

Board White Paper #2

Governance, Accountability & Control

A board-level perspective on preserving accountability, oversight, and meaningful human control as AI becomes embedded in material business decisions.

Prepared by Dr. Ori Marom
Managing Director, Segmentis B.V.

Board purpose	To support board discussion on how AI governance should preserve accountability, oversight, and meaningful human control as AI becomes embedded in material business decisions.
Primary board question	How should the organization assign ownership, define oversight, and maintain control over AI-supported decisions that create financial, operational, legal, or reputational consequences?
Intended audience	Supervisory board members, non-executive directors, executive board members, and senior leaders responsible for AI governance, risk, compliance, and enterprise transformation.

Executive summary

AI governance is becoming a material board responsibility because AI increasingly influences decisions with financial, operational, legal, and reputational consequences. The central challenge is no longer whether AI systems can improve performance, but whether organizations can preserve clear accountability, defensible oversight, and meaningful human control as AI becomes embedded in core decision processes.

The paper argues that boards should govern AI through decision-centric accountability rather than technology-centric control. Every AI-supported decision of material importance should have an identifiable owner, defined oversight and escalation rights, and explicit conditions under which human judgment must intervene. Segmentis' Minimum Viable AI Governance approach is positioned as a practical way to create this clarity without imposing governance structures that slow execution or dilute value creation.

For board members, the practical implication is clear: AI governance must move beyond policy awareness, compliance documentation, or periodic technology updates. It should become part of the regular board dialogue on authority, responsibility, institutional trust, and the balance between value creation and control.

1. Board context and relevance

This white paper addresses the board's responsibility to ensure that AI-enabled decisions remain accountable, explainable, and under meaningful human control. It focuses on how boards can preserve authority and oversight as AI systems become embedded in decisions with financial, operational, legal, and reputational consequences.

This challenge has become more acute as AI moves from experimentation into operational decision-making. Regulatory developments, including the **EU AI Act**, reflect a broader shift in expectations:

organizations are increasingly expected to demonstrate who is responsible for AI-supported outcomes, how oversight is exercised, and when human intervention is required.

For boards, the implication is that AI governance should be treated first as a question of accountability and institutional control, and only second as a question of model performance, compliance documentation, or technology deployment.

2. Board questions for early discussion

Boards facing related governance issues often ask:

1. Which AI-supported decisions are material enough to require explicit board or executive oversight?
2. Who is accountable when an AI-supported decision produces financial, operational, legal, or reputational consequences?
3. Where may AI assist or execute, and where must human judgment remain decisive?
4. Can management explain and defend AI-supported decisions if challenged by regulators, customers, employees, or courts?
5. How will the board know that governance arrangements remain effective as AI systems, data, and operating conditions evolve?

These questions are not about models or tools but rather centered around authority, responsibility, and institutional control.

3. Best practices of AI governance, accountability & control

AI governance answers one fundamental question: who remains responsible when machines influence decisions?

Recent public discourse reflects a growing unease that AI systems are being operationalized inside critical economic and public systems before responsibility is clearly assigned⁵. The risk is not that AI becomes too widely adopted, but that it becomes embedded in decision processes without clear human ownership.

AI excels at optimization. What it cannot do is navigate moral complexities, trade-offs between individual and societal outcomes, or accountability when harm occurs. These responsibilities remain human. As such, they ultimately sit with leadership and the board.

Empirical research and regulatory developments, including the **EU AI Act**, converge on a clear principle: governance must follow decisions, not technology^{2,3,6}. Every AI-supported decision of material importance requires a clearly identifiable executive owner, accountable not only for deployment but for outcomes over time.

This is the rationale behind **Segmentis' Minimum Viable AI Governance (MV-AIG)**. MV-AIG is deliberately decision-centric. It does not attempt to govern "AI" as a category, but ensures that every AI-supported decision of material importance has:

- A clearly identifiable owner.
- Defined oversight and escalation rights.
- The ability to intervene when outcomes deviate.

Governance that is proportionate and usable strengthens value creation because it keeps authority, responsibility, and control aligned as AI becomes operational.

4. Balancing value creation and control

Boards often fear that stronger governance will slow innovation. In practice, we observe that the opposite holds true. Organizations without clear accountability hesitate, escalate excessively, or freeze when issues arise.

Effective control does not eliminate risk; it makes risk **manageable and defensible**. Excessive controls in low-impact areas create friction. Insufficient controls in high-impact areas create exposure that organizations struggle to justify after the fact.

This is why the key distinction is between *delegation* and *abdication*. Delegation preserves responsibility while allowing AI to extend organizational capability. Abdication obscures responsibility by allowing AI-supported processes to act without a clear human owner. Boards must therefore decide explicitly where AI may augment action, where it may execute, and where human judgment must remain decisive^{2,3}.

Where ownership is unclear, decisions stall. Where accountability is diffuse, escalation becomes defensive. But where the boundaries of delegation are clear, governance creates confidence rather than hesitation.

MV-AIG is designed to make this balance practical. High-impact decisions require strong controls, clear ownership, and explicit escalation rights. Lower-impact decisions and uses require lighter supervisory levers. The board's role is to ensure this proportionality is explicit and maintained as AI usage evolves.

Control, in this sense, is not about restriction but rather revolves around ensuring that authority and responsibility remain aligned as scale positively increases.

5. Creating an early governance dialogue

Boards that govern AI effectively engage early, not through compliance checklists, but through **structured dialogues**.

When structured properly, that dialogue connects governance, risk, and trust. It clarifies which decisions matter most, who owns them, and how responsibility is exercised under pressure.

Regulatory experience increasingly shows that organizations struggle not because they lacked policies, but because governance conversations occurred too late, often after AI systems were already operational^{4,7}.

Segmentis uses **MV-AIG** as a shared language that enables an effective dialogue. This point is crucial. Mainly because it allows boards and management to discuss governance without over-specifying solutions or outsourcing judgment.

Alignment is achieved when governance is trusted by management and understood by the board. At that point, discussions shift from “*Do we have enough control?*” to “*Is control focused where it matters most?*”

6. Indicators that the board is on the right path

Boards know they have set the right governance direction when AI discussions shift from control uncertainty to accountable confidence. Responsibilities are clear before issues arise. Escalation is exercised without defensiveness. Interventions happen in time, without creating paralysis.

The practical signs are visible in board and management dialogue. Executives can explain AI-supported decisions with confidence. Ownership is understood across business, legal, risk, technology, and operational functions. Decisions to scale, pause, or redirect AI usage are made with clarity rather than hesitation.

Misalignment is equally visible. Value without control invites recklessness. Control without value leads to paralysis. Both weaken trust. AI governance is therefore working when it helps the organization move faster on the right decisions while preserving the human responsibility behind them.

At that point, governance, accountability, and control are no longer separate safeguards. They become part of how AI scales with confidence. The board's role is to ensure that this balance remains explicit as decision systems become more advanced.

7. Board implications

Board action	Implication for AI governance
Identify material AI-supported decisions	The board should know which decisions carry financial, operational, legal, or reputational consequences and therefore require explicit oversight.
Assign clear ownership	Every material AI-supported decision should have an accountable executive owner responsible for outcomes, not only for system deployment.
Define escalation rights	Management should be able to explain when human review, senior approval, or board attention is required.
Distinguish delegation from abdication	AI may support or execute defined tasks, but responsibility for judgment, trade-offs, and consequences must remain human.
Review governance as AI scales	Governance arrangements should be reassessed as systems, data, use cases, and external expectations evolve.

In practice, the table asks boards to make AI oversight tangible: name the decisions, assign responsibility, set intervention thresholds, and revisit governance as use expands.

8. Case study: Closing the AI accountability gap

An international automotive manufacturer introduced AI-based simulation models to accelerate crash testing across cars and trucks. Early results were strong. The models reduced testing cycles, helped engineering teams explore more design variations, and allowed safety assessments to begin earlier in the product-development process. Encouraged by these outcomes, the Chief Engineering Officer proposed expanding simulation-based approvals across a wider range of vehicle platforms.

The proposal initially appeared straightforward. AI was not replacing physical testing altogether. It was being used to prioritize, simulate, and narrow the range of scenarios requiring full validation. Engineers saw the approach as both efficient and technically sound. The commercial team also welcomed the

shorter development timelines, especially in segments where competitors were bringing new models to market more quickly.

Board attention was triggered when regulators questioned whether certain real-world conditions, including weather, visibility, road surfaces, and unusual impact angles, had been sufficiently represented in the simulations. The models had performed well under defined assumptions. The problem was that no executive could credibly explain who had approved reliance on those assumptions for safety-critical decisions.

Put simply, AI training **data quality** consistently fell short. Yet *no one* knew the AI simulations was flawed.

One example was especially uncomfortable. A simulated side-impact scenario assumed a dry road surface and standard vehicle load. In practice, a fully loaded truck entering the same impact angle on a wet road could deform differently, shift sensor placement, and delay airbag deployment by fractions of a second. The model had not failed; the design of the simulation had quietly excluded the scenario that mattered.

The issue was not technical failure. The simulations had not obviously produced unsafe designs. The issue was **diffuse accountability**. Engineering owned the model. Product development owned the timeline. Legal monitored regulatory exposure. Safety teams reviewed validation evidence. Yet *no one* clearly owned the decision to treat simulation evidence as sufficient under specific operating conditions.

During the board discussion, the General Counsel framed the concern directly. The question was not whether the simulations were accurate enough in general, but which safety-critical decisions the company was willing to allow simulations to support, and who was accountable for that reliance if challenged by regulators, customers, or courts.

The CEO intervened to reframe the increasingly technical discussion. AI simulation would remain central to the company's engineering strategy, but it could not become a substitute for accountable judgment. The board asked management to distinguish between scenarios where AI could accelerate testing, scenarios where AI could support approval, and scenarios where human safety review had to remain decisive. A review committee was formed to periodically re-evaluate these distinctions based on actual experiences with AI, and on emerging regulations.

The board then required explicit ownership for all safety-critical AI-supported decisions. Each approval pathway had to define the decision owner, the conditions under which simulation evidence could be used, the escalation route when assumptions were uncertain, and the point at which physical testing or senior human review was mandatory.

Capital was redirected away from expanding simulation capability alone and toward governance, validation protocols, and decision documentation. AI adoption slowed slightly in breadth but became more credible in practice. Engineers gained clearer boundaries. Legal and safety teams gained earlier visibility. The board gained confidence that acceleration was not coming at the expense of control.

Over time, AI simulation remained central to the company's engineering process, but only after governance clarified where acceleration ended and accountable human judgment began.

Selected references

1. European Union (2018). *Ethics guidelines for trustworthy AI*.
2. European Union (2024). *Regulation (EU) 2024/1689 (Artificial Intelligence Act)*. Official Journal of the European Union, L 2024/1689.
3. Fügener, A., Walzner, D., & Gupta, A. (2026). Roles of artificial intelligence in collaboration with humans: automation, augmentation, and the future of work. *Management Science*, 72(1), 538–557.
4. The Guardian (2026, January 5). *World 'may not have time' to prepare for AI safety risks, says leading researcher*.
5. IMD Business School. (2025). *Winning with AI: the business leader's guide to AI, from strategy to execution*. IMD Business School White Paper.
6. OECD (2025). *AI Adoption by small and medium-sized enterprises*. OECD discussion paper.
7. World Economic Forum. (2024). *Governance in the age of generative AI: A 360° approach for resilient policy and regulation*. World Economic Forum White Paper.