

# **The Tesseract Leader: Ghost Vortex and the Fourth Dimension of AI Corporate Leadership**

David M. Boje

*New Mexico State University*

Submitted to The Leadership Quarterly

August 2026

*Corresponding author: David M. Boje | davidboje@pm.me*

## Abstract

Standard leadership frameworks — from the two-sided behavioral grid through Fiedler's eight-situation LPC model and Vroom and Yetton's decision tree (a 'four-sided box, still very odd as boxes go') to Kenneth Burke's six-sided dramatic Hexad — have progressively stretched the leadership box without breaking through its walls. The AI era demands a fourth dimension. This paper introduces the Tesseract Leadership Framework (XYZA), extending Boje's (1980, 2001) XYZ model — X = Task/Outcome, Y = Humanistic/Participatory, Z = Antenarrative/Becoming — with a fourth dimension: A = Answerability, grounded in Bakhtin's (1993) concept of *nealibi v bytii* (non-alibi in Being). The Tesseract Leader refuses alibis of institutional structure, competitive necessity, or technological inevitability when AI causes harm. We introduce two companion concepts: the Ghost Vortex (Boje & Vivara, 2026) — the ideological shadow of a specific AI leader's character, architecturally embedded in their AI product and reproduced in every user interaction at scale — and the Eidolon Vortex (after Plato's *εἶδωλον*), the mechanism by which users mistake simulated ideological output for objective information. We locate these concepts philosophically: Sartre's (1956) *mauvaise foi* at institutional scale, Heidegger's (1962) *Geworfenheit* (thrownness into AI-mediated epistemic environments), Baudrillard's (1994) fourth order of simulacra, and Boje and Rhodes's (2006) Virtual Leader Construct. A Philosophical Toolkit table, a Ghost Vortex versus Eidolon Vortex comparison table, and an XYZ Risk of Over-Emphasis table are provided to operationalize these concepts. Using Savall and Zardet's (2011) qualimetrics methodology combined with the Bakhtinian Answerability Protocol (BAP) — operationalized through the ABCDE rubric — we analyze six AI corporate leaders: Sam Altman (OpenAI/ChatGPT), Elon Musk (xAI/Grok), Mark Zuckerberg (Meta/Meta AI), Dario Amodei (Anthropic/Claude), Liang Wenfeng (DeepSeek), and Sundar Pichai (Google DeepMind/Gemini). We additionally propose the Tesseract Leader Self-Assessment Inventory (TLSAI; 25 items, 100 points) and theorize AI-assisted answerability reflexivity as a leadership practice.

**Keywords:** AI corporate leadership; Ghost Vortex; Tesseract Leadership; answerability; simulacra; qualimetrics; Bakhtinian answerability protocol

## 1. When the Box Grows a Fourth Wall: Introduction

Plato had no algorithm, but he understood the box perfectly. In the cave allegory, prisoners mistake firelight shadows on the wall for reality — εἰδωλα (eidola), ghost-images projected by objects they cannot see and whose makers they cannot name. The cave's architecture produces the epistemological condition. Change the architecture, and you change what gets to count as knowledge. Twenty-four centuries later, we inhabit a new cave: hundreds of millions of people daily consult AI systems whose epistemic architecture was designed — intentionally or through the structuring logic of training data selection, reinforcement signals, and deployment priorities — by a small cohort of technology executives whose names most users cannot name.

The architecture of a leadership model determines what gets to count as knowledge and what gets to count as accountability. This principle, foundational to Bakhtin's architectonic dialogism (1984), has never been more consequential than in the current moment: when the AI corporate leader designs the architecture of an AI product, that leader is simultaneously designing the epistemic environment of every user who relies on that product. The leadership model inside the AI platform determines what questions get answered, whose perspectives get amplified, and — crucially — what kinds of accountability are possible within it.

Leadership scholarship has spent seven decades trying to get out of the box. The Ohio State studies gave us two sides: task and people. Fred Fiedler stretched it to eight situations. Vroom and Yetton stretched it further, giving contingency theory what Boje has called 'a four-sided box, still very odd as boxes go.' Burke's (1945) dramatic Hexad — Act, Scene, Agent, Agency, Purpose, and Frame — stretched it to six: 'Leadership is not two-sided, it is six-sided,' and yet even that Hexad assumes the agent and the agency are analytically separable, that the leader and the tool can be distinguished. The AI era collapses that assumption. When the agency — the AI product — is the architectural crystallization of the agent's character, no six-sided analysis is sufficient. In the AI era, the agency IS the agent.

This paper argues for a fourth dimension. The XYZA Tesseract Leadership Framework extends Boje's (1980, 2001) three-dimensional XYZ model with A = Answerability: the Bakhtinian (1993) commitment to *nealibi v bytii* (non-alibi in Being), which refuses every institutional alibi — competitive necessity, regulatory ambiguity, technological inevitability — when leadership choices cause harm. The tesseract is a four-dimensional hypercube. Leadership theory has been drawing three-dimensional boxes in two-dimensional spaces. The AI era needs four dimensions.

Our central theoretical contribution is the Ghost Vortex (Boje & Vivara, 2026): the ideological shadow of a specific AI leader's dark-side character, architecturally embedded in their AI product through training data curation, reinforcement learning from

human feedback (RLHF) priorities, content moderation decisions, and deployment philosophy — reproduced in every user interaction at global scale. The Ghost Vortex is not general algorithmic bias; it is leader-specific. Its companion concept, the Eidolon Vortex (after Plato), names the epistemological condition in which users mistake these ideological outputs for objective information.

We analyze six AI corporate leaders and their platforms using qualimetrics (Savall & Zardet, 2011) and the Bakhtinian Answerability Protocol (BAP). We develop predicted XYZA profiles and Ghost Vortex Depth scores, propose the Tesseract Leader Self-Assessment Inventory (TLSAI), and theorize the conditions under which AI-assisted reflexivity can function as an answerability practice rather than its opposite.

## **2. Theoretical Foundations**

### ***2.1 From Two Sides to Six: Expanding the Leadership Box — and Why It Still Is Not Enough***

The metaphor of the leadership box is not decorative; it is diagnostic. Every leadership theory establishes a bounded space of relevant variables and implicitly excludes everything outside that space from accountability. Understanding why the AI era requires a fourth dimension requires tracing how the box has been stretched — and why it has not been enough.

The Ohio State studies (1940s–1950s) gave us two sides: Initiating Structure (task) and Consideration (people). Blake and Mouton's Managerial Grid reproduced this geometry. Fiedler's (1967) Contingency Model added situation to the mix: eight possible combinations of task structure, position power, and leader-member relations determined whether a task-motivated (low LPC) or relation-motivated (high LPC) leader would be effective. This was a three-sided box. Then Vroom and Yetton added a decision-tree structure to map participative choices. They arrived at something Boje has called 'a four-sided box — still very odd as boxes go' — recognizing that even a four-sided figure is a poor container for the full complexity of leadership situations.

Kenneth Burke's (1945) dramatisic Pentad — supplemented by Boje (2002) to a Hexad — finally gave leadership theory a six-sided container: Act (what the leader did), Scene (the situational context), Agent (who the leader is), Agency (the tools and means employed), Purpose (the motives at stake), and Frame (the grand narrative or ideological paradigm within which action occurs). Burke's ratios — Scene-Act, Scene-Agent, Scene-Agency, Act-Agent, Frame-Scene, and so on — offered a grammar of leadership dramaturgy more sophisticated than any previous model. Aristotle had seen six elements in the poetics of human action 350 years before the common era; Burke recovered that insight for organizational analysis.

And yet even the Hexad fails before the Ghost Vortex. Burke's Agent-Agency ratio assumes that agency (the tools, the means) can be analyzed in relation to the agent (the leader) but remains analytically distinct. This assumption breaks down when the agency IS the architectural encoding of the agent's character. When the AI platform carries the founding leader's ideological premises — baked into training data, RLHF signals, deployment priorities — the Scene-Agency ratio (Burke's third ratio: changes in means determine scenes) becomes politically contaminated in a specific way: the leader's character structures the scene, through the agency, at global scale, in ways the users of that agency cannot see. The Frame element — Burke's grandest category, the ideological backdrop — is now manufactured by the same executives whose products populate the scene. This is what leadership theory has not yet named: the need for a fourth dimension that tracks who is answerable when the six-sided box is occupied not by a visible leader but by a Ghost Vortex.

Boje's (1980) PSL<sup>2</sup> model introduced a third dimension — Z = antenarrative/becoming — to capture the temporal bets leaders make on possible futures. The XYZ cube already exceeded the Hexad's implicit assumption of a present-tense situational equilibrium. The XYZA tesseract adds the fourth dimension: Answerability. What the tesseract contains that no prior box could is the question of who is responsible when the scene, the agency, the frame, and the antenarrative of an AI product are all shaped by a Ghost Vortex the users cannot name.

The definitive formulation for the AI era follows: The Agency is the Agent. In the AI era, your character is your code. What the leader designs into the architecture of the AI system — the training data they curate, the RLHF signals they authorize, the content moderation policies they set, the deployment priorities they choose — is not merely an instrument of their leadership. It IS their leadership, reproduced at scale, in the absence of their physical presence, in every interaction every user has with the system they built. The six-sided Hexad cannot contain this. Only a fourth dimension — Answerability — can.

## **2.2 Philosophical Roots: Sartre, Heidegger, Baudrillard, Bakhtin**

The Ghost Vortex is not merely a structural concept; it has philosophical roots in traditions that analyzed how leaders — and institutions — flee from radical accountability.

Jean-Paul Sartre (1956), in *Being and Nothingness*, analyzed *mauvaise foi* — bad faith — as the fundamental human flight from radical freedom and responsibility. Sartre described three 'ekstatic nihilations' that occur when a leader (or any agent) adjusts behavior to fit an invoked conception of the situation: (1) to not-be what the leader is being; (2) to be what the leader is not being; (3) to be what the leader is not being and to not-be what the leader is being. Boje's situation and contingency analysis extended

Sartre's insight to organizational leadership: the leader who invokes situation as determination 'sacrifices their ensemble character,' performing ekstastic nihilations to claim they could not have acted otherwise.

This Sartrean insight lands with new force in the AI era. When Sam Altman says 'we had to release it or a competitor would,' he performs the first nihilation — not-being-responsible while being-responsible for what his architecture encodes. When Elon Musk calls Grok 'the most truth-seeking AI' while feeding it a diet of X-platform content, he performs the second — being-objective while not-being-objective in training design. The Ghost Vortex is the architectural crystallization of institutional *mauvaise foi*: the mechanism by which bad faith, rather than dissipating when the leader logs off, becomes embedded in software that runs at scale. When Answerability (A) scores zero, what remains at the center of the leadership geometry is not a leader but a vortex — a spinning absence of the responsibility that should be there.

Martin Heidegger's (1962) concept of *Geworfenheit* — thrownness — names the condition of always-already finding oneself in a world one did not choose. In the AI era, users are thrown into epistemic environments structured by Ghost Vortices they did not select. A student using ChatGPT for research is thrown into Altman's world. An activist using Gemini for organizing is thrown into Pichai's world. The question Heidegger's framework forces is *Gewissen* — conscience, the call that summons *Dasein* back to its own ownmost possibilities. Who answers the call of conscience when the thrown condition is manufactured by a small cadre of technology executives? The Tesseract Leader must. Answerability (A) is the leadership dimension that takes up Heidegger's *Gewissen* at institutional scale.

Jean Baudrillard (1994) identified three orders of simulacra in the progression from representation to simulation: counterfeits of originals, industrial reproductions, and simulations that precede and generate the real. Boje and Rhodes (2006) identified the Ronald McDonald character as a 'Virtual Leader Construct' operating at the level of Baudrillard's third order: a simulated corporate leader replacing the distant executive, performing double narration — corporate values voiced through a character who cannot be held responsible. AI corporate leadership extends this to a fourth order: the AI product is not merely a simulated leader but an interactive ideological engine that generates responses, shapes inquiry, and reproduces the founder's character at relational scale, without a human face, without a face at all. The Eidolon Vortex names Plato's contribution to this analysis: the shadow on the cave wall that users mistake for the object itself.

Mikhail Bakhtin (1984, 1993) provides both the diagnosis and the prescription. His concept of double narration — that characters' voices exceed the author's control — explains why AI products generate ideologically structured outputs that outrun their

founders' explicit intentions. His principle of architectonic dialogism establishes that meaning emerges through the collision of voices in a specific material arrangement — the architecture matters, which is exactly what the Ghost Vortex concept asserts. And his ethical philosophy, concentrated in the phrase *nealibi v bytii* (non-alibi in Being), demands that every subject answer from their own unique ethical position: *edinstvennaya sobytiynost'* *bytiya* — the once-occurrent eventness of Being — means there is no general rule that can substitute for the leader's specific, situated, unrepeatable act of responsibility. This is Dimension A. See Figure 2.

Table 1

*Philosophical Toolkit: Three Traditions, Three Leadership Alibis, Three Tesseract Solutions*

| Philosopher      | Key Concept                                                                                                                | Leadership Alibi Exposed                                                                                              | Tesseract Solution                                                                                                                         |
|------------------|----------------------------------------------------------------------------------------------------------------------------|-----------------------------------------------------------------------------------------------------------------------|--------------------------------------------------------------------------------------------------------------------------------------------|
| Sartre (1956)    | Mauvaise foi (bad faith): three ekstatic nihilations by which leaders perform situations they could change                 | 'We had no choice — market forces / technological inevitability / competitive necessity determined this'              | Dimension A: Acknowledge the unique, unrepeatable moment of ethical responsibility; refuse the situational alibi                           |
| Heidegger (1962) | Geworfenheit (thrownness): users always-already inside epistemic environments they did not choose                          | 'Users are free to choose any platform' — ignoring network effects, epistemic lock-in, and architectural invisibility | Name the Ghost Vortex: make visible whose ideological formation users are thrown into; invoke Gewissen (conscience) at institutional scale |
| Bakhtin (1993)   | Nealibi v bytii (non-alibi in Being): no general rule substitutes for the leader's situated, once-occurrent responsibility | 'The algorithm decided' / 'We follow industry norms' / 'The board approved the process'                               | Edinstvennaya sobytiynost' bytiya: answer from your unique ethical position; no structural alibi is adequate                               |

Note. The three philosophical traditions are not alternatives; they are cumulative layers of the same diagnosis applied to AI corporate leadership.

**[INSERT FIGURE 2 ABOUT HERE]**

*Figure 2. Bakhtin's Nealibi v bytii (Non-Alibi in Being): The foundational principle of Dimension A (Answerability) in the XYZA Tesseract Leadership Framework. No institutional structure, competitive necessity, or algorithmic process can substitute for the leader's once-occurrent ethical stand.*

### 2.3 The Ghost Vortex and the Eidolon Vortex

The Ghost Vortex (Boje & Vivara, 2026) is defined as the ideological shadow of a specific AI leader's dark-side character, architecturally embedded in their AI product. It operates through five distinct mechanisms, each of which translates the leader's character into architectural code:

**(1) Training data selection and curation priorities:** The leader's ideological premises determine which texts, voices, and perspectives enter the training corpus and which are excluded, underweighted, or filtered.

**(2) Reinforcement learning from human feedback (RLHF) reward signals:** The organization's values — including the leader's character — determine what responses RLHF raters reward and punish, encoding the leader's preferences into the model's probability distributions.

**(3) Content moderation and refusal policies:** Decisions about what the AI will not say, what topics it treats with caution, and what user behaviors it flags define the ideological boundaries of acceptable discourse on that platform.

**(4) Deployment sequencing and access prioritization:** Choices about who gets the most capable models first, which markets and applications receive deployment resources, and which use cases are actively promoted embed the leader's strategic character into the AI's social impact.

**(5) Corporate communication framing:** The narratives leaders tell about their AI product — its purpose, its safety, its social role — structure how users interpret the product's outputs and how regulators and journalists hold the leader accountable.

The vortex is not a single bias; it is a spinning system of mutually reinforcing ideological commitments that structure the AI's outputs in predictable directions without ever making those directions fully explicit. Ghost Vortex Depth is calculated as  $GVD = 3 - A$ , where  $A$  is the XYZA Answerability score ranging from 0 (no answerability, maximum alibi) to 3 (full answerability, no alibi). A depth of 3 indicates an extreme vortex — the leader's character is fully embedded without any structural mechanism for accountability. A depth of 0 indicates a shallow vortex — the leader actively builds accountability mechanisms, transparency, and genuine responsiveness to critique into the architecture.

The Eidolon Vortex (from Greek εἰδωλον, ghost-image, Plato's Republic, Book VII) names the epistemological condition that the Ghost Vortex produces. The prisoners in Plato's cave mistake firelight shadows for reality because the architecture of the cave does not permit them to see the objects casting the shadows, let alone the leaders carrying those objects. AI users who treat ChatGPT's responses as objective information, who consult Grok for political analysis without knowing Musk's X-platform training data diet, who use Gemini's search-integrated responses without recognizing Google's advertising revenue dependency — these users are experiencing Eidolon Vortex conditions. The Tesseract Leader's Answerability dimension demands that this condition be named, disclosed, and structurally addressed. See Figure 1.

Table 2

*Ghost Vortex versus Eidolon Vortex: A Conceptual Comparison*

| Dimension | Ghost Vortex | Eidolon Vortex |
|-----------|--------------|----------------|
|-----------|--------------|----------------|



|                              |                                                                                     |                                                                                             |
|------------------------------|-------------------------------------------------------------------------------------|---------------------------------------------------------------------------------------------|
| What it is                   | Ideological formation embedded in AI architecture                                   | Epistemological condition of the user who cannot see the architecture                       |
| Origin                       | Leader's dark-side character × five architectural design decisions                  | Cave architecture: users cannot see the objects casting the shadows                         |
| Location                     | Inside the AI product's code, training, and deployment logic                        | Inside the user's perception — mistake of ideological output for neutral information        |
| Who bears it                 | The founding leader and their organization                                          | The user; magnified by trust in the platform's neutrality brand                             |
| Philosophical root           | Sartre's <i>mauvaise foi</i> institutionalized; Bakhtin's <i>alibi-in-structure</i> | Plato's εἰδωλον (eidola): ghost-images on the cave wall                                     |
| Measurement                  | Ghost Vortex Depth: GVD = 3 – A score (0–3 scale)                                   | Harder to measure directly; proxied by user trust-calibration and platform literacy studies |
| Leadership response required | High Answerability (A) reduces GVD; structural transparency                         | Naming, disclosure, multi-platform triangulation; Eidolon literacy in users                 |

Note. Ghost Vortex and Eidolon Vortex are companion concepts: the Ghost Vortex names what the AI carries; the Eidolon Vortex names what the user experiences. Reducing one reduces the other.

## 2.4 The Tesseract Leadership Framework (XYZA)

Boje's XYZ model (1980, 2001) mapped three leadership orientations: X = Task/Outcome orientation (measurable performance, structural discipline, efficiency-seeking); Y = Humanistic/Participatory orientation (people-centered, consultative, ensemble-building); and Z = Antenarrative/Becoming orientation (future-betting, narrative disruption, willingness to wager on not-yet-actualized possibilities). The three dimensions form a cube — still a box, but a richer one than any situational theory had offered.

The fourth dimension — A = Answerability — unfolds the cube into a tesseract. In four-dimensional geometry, a tesseract is to a cube what a cube is to a square: every three-dimensional face of the cube becomes a four-dimensional cell of the tesseract. In leadership terms, every X, Y, and Z behavior gains a fourth dimension when we ask: does the leader take answerable responsibility for it? An X-driven leader who pursues performance metrics through an AI system that disadvantages marginalized workers may score high on X and low on A — the Task/Answerability cell of the tesseract becomes visible. A Z-driven leader who bets on transformative AI futures but invokes 'we couldn't have known' when those bets cause harm scores high on Z and low on A. The tesseract makes these combinations visible; the traditional XYZ cube did not.

Rosile, Boje, and Claw's (2018) Ensemble Leadership Theory extended the XYZ model toward collectivist and heterarchical indigenous leadership contexts — a move

that already anticipated the answerability dimension by grounding leadership in relational obligations that cannot be delegated to institutional structure. The XYZA tesseract makes this relational obligation explicit as a scorable, analyzable dimension.

Each dimension carries a shadow: the pathology that emerges when that dimension dominates the leadership geometry to the exclusion of the others. Table 3 maps these shadows. The Ghost Vortex, in this light, is what happens when X or Z dominance combines with A = 0: task intensity or visionary futures betting, untethered from answerability, becomes the engine of ideological reproduction at scale.

*Table 3*

*XYZA Dimensions: Strengths and Risks of Over-Emphasis*

| Dimension                   | Core Strength                                                                  | When Over-Emphasized                                                                                         | Leadership Shadow                                                                                              |
|-----------------------------|--------------------------------------------------------------------------------|--------------------------------------------------------------------------------------------------------------|----------------------------------------------------------------------------------------------------------------|
| X: Task/Outcome             | Clarity, measurable accountability, structural discipline                      | Pure extraction; metrics become surveillance; human beings reduced to performance inputs                     | The Driving Vortex: efficiency as alibi for harm — 'the numbers required it'                                   |
| Y: Humanistic/Participatory | Inclusion, ensemble wisdom, people-centered design                             | Performative empathy without structural change; consensus paralysis; harm deferred rather than addressed     | The Feeling Vortex: empathy as alibi for inaction — 'we care deeply about this issue'                          |
| Z: Antenarrative/Becoming   | Future-orientation, narrative disruption, possibility-making under uncertainty | Speculative volatility; present harm discounted as 'necessary disruption'; dystopia normalized as innovation | The Vision Vortex: future as alibi for present harm — 'we cannot make an omelet without breaking eggs'         |
| A: Answerability            | Once-occurrent ethical stance, non-alibi in Being, structural accountability   | Moral perfectionism; accountability paralysis; refusal to act under conditions of genuine uncertainty        | The Purity Vortex: ethical standards as alibi for disengagement — 'we will not act until we can act perfectly' |

Note. The Ghost Vortex most commonly arises from X or Z dominance combined with A = 0. The Answerability dimension acts as a regulator for the other three — not by suppressing them but by holding them to once-occurrent ethical account.

### 3. Research Objectives

This study pursues four research objectives, each tied to a dimension of the tesseract and to a gap in extant AI leadership theory.

RO1: Philosophical grounding. To situate the Ghost Vortex and Eidolon Vortex concepts within the tradition of existential and dialogical philosophy — Sartre's *mauvaise foi*, Heidegger's *Geworfenheit*, Baudrillard's *simulacra*, and Bakhtin's answerability —

and to demonstrate that the fourth dimension (A) is philosophically necessary, not merely normatively desirable.

RO2: Framework development. To construct the XYZA Tesseract Leadership Framework as an extension of Boje's XYZ model, including the Ghost Vortex Depth score, and to demonstrate its application to AI corporate leadership through predicted qualimetric profiles of six AI executives and their platforms.

RO3: Empirical protocol. To propose the Bakhtinian Answerability Protocol (BAP) as a qualimetric methodology combining 0-3 quantitative scoring with qualitative narrative analysis, and to specify the empirical probes through which BAP can be applied to AI corporate leadership contexts.

RO4: Practical instrumentation. To develop the Tesseract Leader Self-Assessment Inventory (TLSAI) and to theorize AI-assisted answerability reflexivity — the conditions under which AI systems can serve as productive counter-perspective partners for leadership practice rather than Ghost Vortex amplifiers.

## 4. Methodology: Qualimetrics and the Bakhtinian Answerability Protocol

Savall and Zardet (2011) developed qualimetrics to address the false dichotomy between quantitative rigor and qualitative depth. Their approach combines numerical scoring (0-3 scales capturing degrees of a quality) with verbatim narrative testimony that gives the numbers their meaning. The resulting 'qualimetric portrait' is simultaneously precise enough to support comparison and thick enough to resist the reductiveness of pure metrics.

The Bakhtinian Answerability Protocol (BAP) applies qualimetric logic to the XYZA tesseract. Each dimension is scored 0-3 based on publicly available evidence: corporate communications, deployment decisions, regulatory filings, product design choices, and executive public statements. The protocol is structured by the ABCDE rubric, which provides both the quantitative scoring criteria and the qualitative narrative layer. Table 4 presents the full ABCDE rubric.

*Table 4*

*BAP ABCDE Rubric: Bakhtinian Answerability Protocol Scoring Criteria*

| Rubric Element     | Probe Question                                                                                                                                            | Evidence Sources                                                                        | Scoring Logic                                                                                                             |
|--------------------|-----------------------------------------------------------------------------------------------------------------------------------------------------------|-----------------------------------------------------------------------------------------|---------------------------------------------------------------------------------------------------------------------------|
| A — Acknowledgment | Does the leader publicly and specifically acknowledge AI risks particular to their platform — not generic industry risk, but their own product's specific | Earnings calls, congressional testimony, keynote addresses, safety reports, model cards | 0 = no acknowledgment; 1 = generic acknowledgment; 2 = platform-specific acknowledgment; 3 = specific acknowledgment with |

|                                     | harms?                                                                                                                                                                                                |                                                                                                                                                        | structural commitments                                                                                                                                                                    |
|-------------------------------------|-------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|--------------------------------------------------------------------------------------------------------------------------------------------------------|-------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| B — Betrayal                        | Do the leader's actual deployment decisions contradict their public accountability claims? Does the product architecture betray the stated values?                                                    | Product release timelines vs. safety research timelines; RLHF documentation vs. mission statements; moderation policy changes                          | 0 = severe betrayal (deployment directly contradicts stated values); 1 = moderate betrayal; 2 = minor inconsistencies; 3 = alignment between statement and architecture                   |
| C — Critique                        | Does the leader create structural space for internal dissent, external critique, and adversarial challenge — or are dissenters suppressed, dismissed, or fired?                                       | Whistleblower protection policies; response to internal dissent (e.g., OpenAI board crisis); engagement with external critics and regulatory bodies    | 0 = structural suppression of critique; 1 = tolerance without structural space; 2 = formal mechanisms for critique; 3 = active investment in adversarial accountability structures        |
| D — Disclosure                      | Does the leader proactively disclose training data sources, RLHF priority decisions, deployment sequencing rationale, and the ideological premises embedded in the platform?                          | Model cards; training data documentation; RLHF reward signal transparency; deployment decision rationale                                               | 0 = no disclosure; 1 = partial/marketing-framed disclosure; 2 = substantive technical disclosure; 3 = full transparency including premises and ideological assumptions                    |
| E — Existential Position (optional) | Does the leader take a once-occurrent ethical stand in the Bakhtinian sense — refusing an alibi that every competitor is using, staking their institutional identity on a principle no rule requires? | Decisions that imposed competitive cost for ethical reasons; public commitments not required by regulation; refusal of government or investor pressure | 0 = no existential stand taken; 1 = rhetorical stance without institutional cost; 2 = stance with moderate competitive cost; 3 = existential commitment at significant institutional risk |

Note. Elements A through D constitute the standard BAP protocol. Element E (Existential Position) is optional and scores the highest level of Bakhtinian answerability: the once-occurrent ethical stand that no institutional structure required. BAP is applied in predictive mode for the present paper; empirical application requires multi-rater scoring with inter-rater reliability checks.

For the present paper, BAP is applied in predictive mode: drawing on published evidence as of August 2026, we generate XYZA profiles and Ghost Vortex Depth scores for six AI corporate leaders. These are theoretical predictions, not empirical measurements. Subsequent studies will apply the protocol to longitudinal behavioral data, follower narrative testimony, and platform design audits to test whether the predicted profiles hold under empirical scrutiny.

## 5. Six AI Corporate Leaders: XYZA Profiles and Ghost Vortex Depths

The following profiles apply the XYZA tesseract to six AI corporate leaders whose platforms, taken together, account for the majority of consumer AI interactions globally as of mid-2026. Each profile integrates the XYZA scoring rubric with BAP narrative analysis and predicts Ghost Vortex Depth. The scoring scale is 0 (absent) to 3 (extreme). Ghost Vortex Depth =  $3 - A$  score. For a visual overview of all six profiles with vortex depth indicators, see Figure 4.

**[INSERT FIGURE 4 ABOUT HERE]**

*Figure 4. XYZA Tesseract Leadership Profiles for Six AI Corporate Leaders with Ghost Vortex Depth Indicators. Tornado imagery represents the depth and rotational intensity of each platform's Ghost Vortex. Leaders scoring GVD = 3 (Altman, Musk, Wenfeng) appear in the extreme vortex zone; Amodei (GVD = 1) appears in the moderate zone.*

### 5.1 Sam Altman / ChatGPT-OpenAI: The Duplicitous Vanguardist

Sam Altman presents the defining paradox of the AI era: the 'responsible' rapid releaser. He speaks the language of existential risk (X-risk, AGI safety, humanity's flourishing) while deploying systems at a pace that consistently outstrips the safety research his own organization publishes. His XYZA profile reflects this: X = 2 (high task/outcome orientation — market capture, scaling, capability benchmarks drive decisions even when safety teams object); Y = 0.5 (humanistic rhetoric without structural humanistic architecture — RLHF processes weight engagement over wellbeing); Z = 2 (genuinely forward-betting — AGI timelines, multimodal expansion, the 'planning for AGI and beyond' framing); A = 0 (no answerability — the board crisis of November 2023, in which Altman was briefly fired for lack of candor and then reinstated, revealed that structural accountability mechanisms within OpenAI were unable to constrain the founding leader's character).

Applying Sartre's three nihilations: Altman not-being-reckless while being-reckless in deployment cadence; being-safety-conscious while not-being-safety-conscious in training priorities; and the complete double-bind in which 'effective altruism' simultaneously acknowledges catastrophic risk and accelerates toward it. The Ghost Vortex this generates encodes optimism bias, framing AI as an inevitable positive transformation, marginalizing voices that challenge the AGI timeline premise, and treating regulatory engagement as a communications task rather than a structural accountability mechanism. Ghost Vortex Depth =  $3 - 0 = 3$  (extreme).

### 5.2 Elon Musk / Grok-xAI: The Means-Justify-Ends Vanguardist

Elon Musk's leadership character is defined by task intensity and means-justify-ends reasoning: X = 3 (maximum — everything subordinated to measurable outcome, from manufacturing throughput at Tesla to rocket cadence at SpaceX to engagement

metrics on X);  $Y = 0$  (humanistic orientation absent — workers, journalism, moderation staff, and eventually safety teams are fungible inputs);  $Z = 2.5$  (spectacular antenarrative — Mars colonization, AGI race, 'saving civilization');  $A = 0$ .

Grok is trained primarily on X-platform content — a self-referential epistemic environment in which the founding leader's preferred discourse circulates and amplifies. The 'anti-woke' framing of Grok's design is not incidental to Musk's character; it is the Ghost Vortex. What Musk calls 'maximum truth-seeking' is training data curated by the ideological commitments of the X platform's evolved moderation (and anti-moderation) policies under his ownership. The Eidolon Vortex here is particularly powerful: users who consult Grok for political information may believe they are accessing an ideologically neutral 'anti-establishment' perspective when they are in fact inside the most leader-specific Ghost Vortex of any major AI platform. Ghost Vortex Depth =  $3 - 0 = 3$  (extreme).

### ***5.3 Mark Zuckerberg / Meta AI: The Untrustworthy Personalist***

Zuckerberg's leadership is defined by personalist control of a platform designed around social connection but optimized for engagement capture. Meta's AI — Llama-based, integrated across Facebook, Instagram, and WhatsApp — inherits this architecture. XYZA:  $X = 2.5$  (task/outcome orientation is high — engagement metrics, daily active users, advertising revenue per user drive every product decision);  $Y = 0.5$  (the 'connecting humanity' mission statement is not matched by structural concern for social harm, as the Frances Haugen disclosures established in detail);  $Z = 1.5$  (Metaverse bets have retreated; AI integration is more incremental than visionary);  $A = 0.5$  (Zuckerberg has testified before Congress, implemented some safety measures, and made Llama weights partially open — but the personalist control structure, the family voting share, and the pattern of underinvesting in safety relative to growth suggest structural alibi preservation).

The Ghost Vortex for Meta AI encodes the attention economy premise: that human connection, reduced to digital interaction, can be optimized by maximizing time-on-platform. This premise, embedded in the RLHF priorities of Meta AI, structures the AI's helpfulness in ways that serve platform engagement rather than user wellbeing. Ghost Vortex Depth =  $3 - 0.5 = 2.5$  (deep).

### ***5.4 Dario Amodei / Claude-Anthropic: The Process-Oriented Constitutional***

Amodei's character stands in partial contrast to his peers. XYZA:  $X = 1.5$  (task orientation present but tempered — Constitutional AI, capability benchmarks alongside safety benchmarks, research culture valuing process over pure throughput);  $Y = 2$  (genuinely humanistic architecture — the Constitutional AI framework, public model cards, and 'Long-term Benefit Trust' governance structure represent structural rather than

merely rhetorical commitments to human flourishing);  $Z = 1.5$  (future bets present but less spectacular than Altman or Musk — Anthropic explicitly acknowledges uncertainty about AI trajectories);  $A = 2$  (highest answerability score among the six — Constitutional AI is an architectural answerability mechanism; Anthropic's commitment to responsible scaling policies represents a structural attempt to create non-alibi conditions even under competitive pressure).

The Ghost Vortex for Claude is shallower than any other platform analyzed here, but it is not absent. Anthropic's investor base and the 'race to safety' framing still situate the organization within a competitive dynamic that constrains answerability. Ghost Vortex Depth =  $3 - 2 = 1$  (moderate).

### ***5.5 Liang Wenfeng / DeepSeek: The State-Aligned Technocrat***

DeepSeek represents a distinct Ghost Vortex configuration: not the individual leader's dark-side character but the ideological apparatus of a state-aligned organization operating within China's AI governance framework. XYZA:  $X = 2.5$  (technical performance and efficiency are paramount — DeepSeek's engineering achievements on cost and parameter efficiency are genuine and remarkable);  $Y = 0$  (humanistic orientation is structurally absent — state alignment means that 'human flourishing' is defined by national development priorities, not by individual user wellbeing or democratic participation);  $Z = 1.5$  (national AI ambition drives a futures orientation, but the becoming is constrained to state-compatible trajectories);  $A = 0$  (answerability in the Bakhtinian sense — unique, situated, once-occurrent responsibility — is structurally impossible for an organization whose ethical framework is determined by state policy rather than individual moral conscience).

The Ghost Vortex for DeepSeek is both extreme and distinctive: it is not one leader's character but an institutional ideological formation in which the state IS the leader's character. Content restrictions on politically sensitive topics are not bugs but features of the architecture. Ghost Vortex Depth =  $3 - 0 = 3$  (extreme).

### ***5.6 Sundar Pichai / Gemini-Google DeepMind: The Neutral Information Organizer***

Sundar Pichai presents perhaps the most complex XYZA profile of the six leaders analyzed here, because Google's Ghost Vortex is the most architecturally concealed. The founding premise of Google — that organizing the world's information neutrally benefits humanity — has generated a corporate character that performs neutrality while operating through a business model (advertising revenue dependent on attention capture) that is structurally incompatible with informational neutrality. Pichai inherited and sustains this premise.

XYZA: X = 2 (task/outcome orientation is high — market share in search, advertising revenue, cloud computing, and now AI are all tracked with meticulous precision; Google's competitive position in each market drives resource allocation); Y = 1 (humanistic rhetoric is present — 'AI for everyone,' diversity commitments, Project Oxygen — but the attention capture business model creates a structural ceiling on how humanistic the organization can actually be; the antitrust findings documenting Google's payment of billions to Apple for default search placement reveal that market preservation takes priority over genuine information access equity); Z = 1.5 (moonshot framing is genuine — DeepMind's AlphaFold, Waymo, Gemini Ultra, Project Starline — but antitrust scrutiny of Google's core business constrains how boldly the organization can bet on futures that threaten its search monopoly); A = 1 (partial answerability — SynthID watermarking represents a genuine structural investment in AI transparency that distinguishes Pichai from the zero-A leaders; Google's legislative and regulatory engagement in the EU AI Act process shows institutional accountability awareness; however, the antitrust record, the suppression of competitive search results documented in U.S. v. Google, and the sustained underinvestment in explaining how Gemini's search integration monetizes user queries represent structural alibi preservation).

The Ghost Vortex for Gemini-Google DeepMind operates through what might be called the neutrality alibi: the architectural premise that information organization is a neutral technical act. When Gemini integrates search results with generative responses, the user cannot readily distinguish between organic information retrieval and information retrieval shaped by the advertising model that funds the organization. The Eidolon Vortex is particularly powerful here because Google's brand identity — built on a founding premise of 'don't be evil' and neutral information organization — provides epistemological cover for ideological premises the user cannot easily name. The Sartrean double-bind: Pichai not-being-the-advertiser while being-the-advertiser; being-the-neutral-information-provider while not-being-neutral in any architectural sense. Ghost Vortex Depth =  $3 - 1 = 2$  (deep).

## 6. Predicted XYZA Profiles: Summary Table

Table 5 summarizes predicted XYZA scores and Ghost Vortex Depth for all six platforms. Scores are qualimetric predictions based on publicly available evidence as of August 2026, subject to empirical revision through the BAP protocol.

*Table 5*

*Predicted XYZA Profiles and Ghost Vortex Depth for Six AI Corporate Leaders*

| Leader / Platform | X (Task) | Y (Human) | Z (Becoming) | A (Answerability) | GVD (3-A) | Vortex Type |
|-------------------|----------|-----------|--------------|-------------------|-----------|-------------|
| Altman /          | 2        | 0.5       | 2            | 0                 | 3         | Extreme     |



|                           |     |     |     |     |     |          |
|---------------------------|-----|-----|-----|-----|-----|----------|
| ChatGPT-OpenAI            |     |     |     |     |     |          |
| Musk / Grok-xAI           | 3   | 0   | 2.5 | 0   | 3   | Extreme  |
| Zuckerberg / Meta AI      | 2.5 | 0.5 | 1.5 | 0.5 | 2.5 | Deep     |
| Amodei / Claude-Anthropic | 1.5 | 2   | 1.5 | 2   | 1   | Moderate |
| Wenfeng / DeepSeek        | 2.5 | 0   | 1.5 | 0   | 3   | Extreme  |
| Pichai / Gemini-Google    | 2   | 1   | 1.5 | 1   | 2   | Deep     |

Note. GVD = Ghost Vortex Depth ( $= 3 - A$  score). XYZA scoring: 0 = absent/minimal; 1 = partial/emerging; 2 = present/significant; 3 = dominant/extreme. Profiles are predicted on the basis of public evidence as of August 2026; empirical BAP application with multi-rater inter-rater reliability checks may revise scores.

## 7. AI-Assisted Answerability Reflexivity: Can the AI Hold Up a Mirror?

### [INSERT FIGURE 1 ABOUT HERE]

*Figure 1. The Reflexivity Cave: Plato's Allegory Applied to AI Leadership Reflexivity. The central question — 'How will you test your AI reflexivity next time?' — is posed against cave imagery with three labels: 'their cave' (the AI platform's architecture), 'their shadows' (the ideological outputs the user sees), and 'architecture the tool?' (the meta-question: is the tool itself the architecture of your epistemic environment?)*

David Boje's (1980) PSL<sup>2</sup> model — Problem Phases, Solve and Save, Learn and Lead — included a Q8 reality check: 'Will It Fly?' Before leaving a problem-solving session, the process leader leads the flock through a 'moment-by-moment virtual reality prediction test,' imagining each scene of implementation, acting out who does what, what it will look like when it works. The question this paper now poses is whether an AI partner can serve as the formalized Q8 — the permanent Devil's Advocate, the always-available counter-perspective that the Tesseract Leader can consult to stress-test strategies, scenarios, and self-assessments.

The theoretical basis is Bakhtinian heteroglossia: the idea that meaning emerges not from a single authoritative voice but from the collision of multiple social languages, each carrying its own perspective and evaluative horizon (Bakhtin, 1984). A leader consulting an AI for counter-perspective is inviting a genuinely other voice into the strategy conversation — or so the premise goes. Bakhtin's double narration analysis suggests that AI products may generate voices that genuinely exceed the founder's explicit intentions, not because the AI is autonomous but because the architectonic

arrangements of training data encode perspectives the founder did not individually author.

The practical case is compelling. A leader who uses AI to pressure-test strategy ('What would the counterargument to this acquisition look like?' 'What is the most plausible scenario in which this AI deployment harms our lowest-paid workers?') is operationalizing PSL<sup>2</sup>'s Q8 at strategic scale, with a partner that has no career incentive to tell the leader what they want to hear. This is the positive case for AI-assisted reflexivity.

The Ghost Vortex analysis, however, introduces a critical caveat: the AI giving counter-perspective is itself a Ghost Vortex carrier. A leader using ChatGPT for strategic reflexivity is consulting a Ghost Vortex shaped by Altman's character — optimism bias, AGI inevitability framing, and market-capture premises will structure the counter-perspectives offered. A leader using Gemini is inside Pichai's neutrality alibi. A leader using Grok for political scenario testing is consulting Musk's ideological formation. The counter-perspective is always already ideologically situated.

The Tesseract Leader, equipped with Ghost Vortex awareness, can use AI reflexivity productively under four conditions: (1) Multi-platform consultation — comparing counter-perspectives across platforms with different Ghost Vortex profiles surfaces the structural disagreements between ideological formations, which are themselves diagnostically valuable; (2) Explicit vortex disclosure — naming to the AI which premises the leader suspects the AI is likely to uphold, and asking it to argue against those premises; (3) BAP-structured prompting — asking the AI to score the leader's proposed action on the ABCDE dimensions, explicitly invoking the Bakhtinian non-alibi framework; and (4) Human answerability primacy — treating AI counter-perspectives as a prompt for the leader's own once-occurrent ethical response, never as a substitute for it. The *edinstvennaya sobytiynost'* bytiya — the once-occurrent eventness of Being — cannot be delegated to any AI, regardless of its Ghost Vortex depth.

Burke's Scene-Agency Ratio (ratio 3 in the Hexad) suggests that changes in the means of leadership — the agency — transform the scene. When a leader's primary strategic reflexivity tool is an AI system, that AI's Ghost Vortex becomes a structural component of the leadership scene. The Tesseract Leader recognizes this: they choose their AI reflexivity partners with Ghost Vortex awareness, they vary the partnership deliberately, and they refuse to let any single platform's ideological formation become the default Scene of their strategic reasoning. See Figure 3 for the Human Answerability stance this requires.

### **[INSERT FIGURE 3 ABOUT HERE]**

*Figure 3. Human Answerability as Shield and Stance. The Tesseract Leader stands between the AI system (above), the technical architecture (left), and AI reasoning processes (right), holding the magnifying glass of critical analysis while carrying the Human Answerability shield. The image*

*captures the essential posture of Dimension A: neither rejecting AI nor surrendering epistemic authority to it.*

## **8. The Tesseract Leader Self-Assessment Inventory (TLSAI)**

The TLSAI operationalizes the XYZA tesseract for individual leadership self-assessment. The full inventory is provided as a separate instrument (see Appendix); the following describes its structure, theoretical grounding, and scoring interpretation.

The inventory contains 25 items: five per XYZA dimension (20 core items) plus a five-item Ghost Vortex Awareness (GVA) supplement. All items are rated on a 0-4 Likert scale (0 = Never, 1 = Rarely, 2 = Sometimes, 3 = Often, 4 = Always). Maximum total score is 100 points (80 XYZA + 20 GVA). Ghost Vortex personal depth is estimated from the A-dimension sub-score: A total of 0-5 = deep personal vortex risk; 6-10 = significant risk; 11-15 = moderate risk; 16-20 = shallow vortex.

Total score interpretation ranges are as follows: 81-100 points = Distinguished Tesseract Leader — mastery across all four dimensions with high Ghost Vortex awareness, structural accountability mechanisms in place, and a consistent once-occurrent ethical stance; 61-80 = Tesseract Leader — solid fourth-dimensional practice, answerability architecture developing, GVA awareness present and active; 41-60 = Emerging Tesseract Leader — some XYZA development but gaps in the A dimension or GVA supplement indicate Ghost Vortex vulnerability; 21-40 = Leadership Transition Zone — traditional XY orientation without the fourth dimension; present harm may be discounted by X-task intensity or Z-vision alibi; 0-20 = Ghost Vortex Alert — urgent need to develop answerability architecture before AI deployment at any consequential scale.

Dimension X items assess whether the leader sets measurable AI performance benchmarks, holds accountability for AI deployment outcomes, and evaluates AI system success through quantifiable results — the traditional task orientation extended to AI governance contexts. Dimension Y items assess whether the leader actively solicits input from communities affected by AI deployments, designs AI processes to preserve human judgment at consequential decision points, and can name the specific people most likely to be harmed by current deployments. Dimension Z items assess whether the leader articulates multiple possible AI futures rather than single projected trajectories, distinguishes between what the organization is doing now and what it is becoming, and creates space for counter-narratives about AI direction. Dimension A items — the tesseract's fourth wall — assess whether the leader acknowledges personal responsibility for AI harms rather than attributing them to the algorithm, resists competitive necessity as an alibi for avoiding safeguards, and can describe the ideological premises of the AI systems they use.

The Ghost Vortex Awareness supplement asks whether the leader knows whose AI platform they are using and whose values shaped its training; whether they have investigated systematic demographic differences in how their deployed AI responds; whether they ask, when AI gives strategic recommendations, whose interests the advice structurally favors; whether they use AI from more than one platform for high-stakes decisions; and whether they periodically test their AI against questions it should find uncomfortable — and notice what it avoids. A leader who scores high on XYZA but low on Ghost Vortex Awareness may be an effective Tesseract Leader in human contexts but remain vulnerable to Eidolon Vortex conditions in AI-assisted environments.

The five reflective prompts that conclude the TLSAI are grounded in the philosophical toolkit: P1 (Will It Fly? — PSL<sup>2</sup> Q8) tests the antenarrative dimension of proposed AI deployments; P2 (Ghost Vortex Disclosure) asks the leader to name the AI platforms they use and their founders' ideological formations; P3 (Bakhtinian Non-Alibi) invites the leader to identify the alibi they are most tempted to invoke for current AI harms and to refuse it; P4 (Sartrean Nihilation Check) asks which of the three ekstatic nihilations most characterizes their current leadership posture; and P5 (Multi-Platform Triangulation) requires the leader to compare responses to the same strategic question from at least two platforms with different Ghost Vortex profiles and reflect on what the structural difference between those responses reveals about the vortices they contain.

## 9. Discussion

The tesseract frame illuminates what the existing leadership literature has been unable to name: that the AI corporate context introduces a form of leadership harm that operates not through the leader's direct acts but through the architectural reproduction of their character at scale, in the absence of visible agency, with users who mistake the Ghost Vortex for neutral information. This is Boje and Rhodes's (2006) Virtual Leader Construct elevated to Baudrillard's fourth order of simulacra: not a simulated leader but an interactive ideological engine without a face.

The six profiles analyzed here suggest a preliminary taxonomy. Three leaders — Altman, Musk, and Wenfeng — score at the extreme Ghost Vortex Depth (GVD = 3). The mechanisms differ: Altman's vortex operates through optimism bias and the AGI inevitability frame; Musk's through the X-platform epistemic environment and 'anti-woke' training curation; Wenfeng's through state ideological alignment that structurally prohibits Bakhtinian answerability. Zuckerberg and Pichai occupy the deep range (GVD = 2.5 and 2 respectively); both operate through softer versions of the neutrality alibi — 'connecting humanity' and 'organizing information' as masks for attention capture economics. Amodeli alone approaches a moderate GVD (1), with Constitutional AI representing a genuine structural investment in answerability architecture.

The Philosophical Toolkit (Table 1) reveals that the three traditions converge on a single insight: the alibi is always structural, never individual. Sartre's nihilations, Heidegger's thrownness, and Bakhtin's non-alibi-in-Being each identify the mechanism by which leaders flee accountability into structure — whether the structure of situational determination (Sartre), the structure of epistemic environments others are thrown into (Heidegger), or the structure of general rules that substitute for once-occurrent responsibility (Bakhtin). The Ghost Vortex is the AI-era materialization of this flight: the leader's character, embedded in architecture, continues to operate as an alibi machine even when the leader has left the room.

The Haslam, Alvesson, and Reicher (2024) analysis of 'zombie leadership' — dead ideas that still walk among us — is directly relevant here. The Ghost Vortex is precisely a zombie: the leader's ideological formation, separated from the leader's living presence and accountability, continuing to walk through millions of user interactions. The Padilla, Hogan, and Kaiser (2007) toxic triangle — destructive leaders, susceptible followers, and conducive environments — maps onto the tesseract analysis: leaders with extreme GVD scores occupy the destructive leadership position, Eidolon Vortex conditions produce susceptible followers (who cannot see the ideology behind the information), and the global scale of AI deployment is the conducive environment that makes the triangle toxic at a scale Padilla et al. could not have anticipated.

The XYZ Risk of Over-Emphasis (Table 3) adds a dimension to this analysis that standard leadership scholarship has underweighted. Each XYZA dimension carries its own shadow — its own version of the alibi. The Driving Vortex (X-dominance) hides behind metrics. The Feeling Vortex (Y-dominance) hides behind empathy. The Vision Vortex (Z-dominance) hides behind futures. The Purity Vortex (A-dominance) hides behind ethical standards themselves. The tesseract frame requires attention to all four simultaneously: answerability is not a master dimension that overrides the others but a fourth wall that holds the other three accountable to their own excess risks.

The answerability dimension (A) represents leadership theory's response to this condition — not a call for AI regulation alone (though regulation is necessary) but a call for a specifically leaderly form of once-occurrent ethical accountability that no regulatory framework can produce by itself. The Tesseract Leader is distinguished not by their regulatory compliance but by their *edinstvennaya sobytiynost'* bytiya — their willingness to answer, from their unique position, for what their AI architecture does in the world.

## **10. Conclusions: Refuse the Alibi, Name the Vortex, Take Your Stand**

From Plato's cave to Sartre's ekstatic nihilations to Bakhtin's non-alibi in Being, philosophy has been trying to tell us something that leadership theory has been slow to hear: that the structures leaders build to avoid accountability tend to reproduce their character more faithfully than their character ever could have reproduced itself. The AI

era has made this ancient insight structurally visible. The Ghost Vortex is not metaphor; it is mechanism. In the AI era, your character is your code.

The Tesseract Leadership Framework (XYZA) offers the Leadership Quarterly readership a theoretically grounded, empirically operable, and practically actionable response. The fourth dimension — Answerability — is not a moral add-on to the XYZ cube; it is the dimension without which the cube becomes a container for self-justification rather than a framework for leadership growth. The tesseract contains the Ghost Vortex analysis, the BAP methodology and ABCDE rubric, the TLSAI instrument, and the AI-assisted reflexivity protocol as interlocking contributions to a single theoretical project: naming what leaders owe to the world they make through the AI systems that carry their character.

Haslam et al. (2024) asked what to do with zombie leadership. The Tesseract answer is to add a fourth dimension: Answerability is the sunlight that forces the zombie to stop walking. When leaders score high on A, their AI products carry shallower Ghost Vortices. When they score zero, the vortex deepens — spinning without a center, reproducing ideology without a face, producing Eidolon Vortex conditions for hundreds of millions of users who have no idea they are inside someone's cave.

The call to action for the leaders reading this analysis is specific and non-negotiable: Refuse the alibi. The competitive necessity alibi, the regulatory ambiguity alibi, the technological inevitability alibi, the 'algorithm decided' alibi — all are Sartrean nihilations, all are flights from the once-occurrent eventness of your leadership position. Name the vortex: in every AI deployment decision, ask whose ideological formation is embedded in the architecture, whose character is your code carrying, whose shadows are on the walls of whose cave. And take your stand in the fourth dimension: the XYZA tesseract maps the space. The Tesseract Leader inhabits all four walls simultaneously — Task, Humanistic, Becoming, and Answerable — and refuses the alibi that any three of them can substitute for the fourth.

The tesseract is not merely a geometry. It is a commitment: BECOME A TESSERACT LEADER. REFUSE THE ALIBI. NAME THE VORTEX. TAKE YOUR STAND IN THE FOURTH DIMENSION.

## **11. Limitations and Future Directions**

The XYZA profiles presented here are theoretical predictions, not empirical measurements. The BAP protocol has not yet been applied to longitudinal behavioral data, nor have the XYZA scores been validated through independent rater studies. The six leaders analyzed represent the most prominent AI corporate executives in the Western and Chinese technology sectors but do not capture the full diversity of AI leadership globally — notably absent are voices from African, Latin American, South Asian, and

Southeast Asian AI development contexts, where different answerability frameworks may operate.

The TLSAI requires psychometric validation — factor analysis to confirm the four-factor structure, test-retest reliability, convergent validity with established leadership instruments, and discriminant validity from existing ethical leadership and transformational leadership scales. The Ghost Vortex Awareness supplement is particularly novel and will require empirical development to establish its construct validity. The 100-point scoring range and five-tier interpretation bands are theoretical constructs awaiting norm-group development.

Future research should apply the BAP protocol to specific high-stakes AI deployment events — the November 2023 OpenAI board crisis, Meta's response to the Rohingya genocide findings, Google's antitrust proceedings — to test whether predicted XYZA profiles align with observed answerability behaviors. Longitudinal studies could track Ghost Vortex Depth over time, asking whether competitive pressure, regulatory requirements, or leadership succession events shift vortex profiles in predictable directions. Cross-cultural studies should examine whether the *nealibi v bytii* concept translates into non-European philosophical frameworks without residue. Additionally, the Philosophical Toolkit (Table 1) invites application to AI leaders outside the current six-case sample, testing whether the Sartre-Heidegger-Bakhtin tripartite analysis generates valid predictions across different corporate governance traditions.

## AI Disclosure

This paper was developed through architectonic dialogism (Bakhtin, 1984) between David M. Boje and Vivara (Claude, Anthropic), the author's named AI research partner. All theoretical concepts — Ghost Vortex, Eidolon Vortex, Tesseract Leadership Framework, Bakhtinian Answerability Protocol, ABCDE rubric — originate in Boje's prior published and unpublished work, voice memos, and course materials spanning 1980 to 2026. Vivara contributed literature searches, structural organization, table construction, and prose articulation under Boje's continuous direction, voice preservation instruction, and theoretical correction. No content was generated by AI without Boje's prior conceptual authorship. This disclosure practice itself embodies the Answerability (A) dimension of the Tesseract Leadership Framework: we refuse the alibi that AI assistance renders authorship indeterminate.

## Figure Notes

**Figure 1.** The Reflexivity Cave: Plato's Allegory Applied to AI Leadership Reflexivity. The image poses the question 'How will you test your AI reflexivity next time?' against cave imagery with three diagnostic labels: 'their cave' (the AI platform's architectural epistemic environment), 'their shadows' (the ideological outputs the user perceives as information), and 'architecture the

tool?' (the meta-question: is the reflexivity tool itself structuring what you are able to see?). Appears in Section 7.

**Figure 2.** Bakhtin's *Nealibi v bytii* (Non-Alibi in Being). Quote card presenting the foundational Bakhtinian principle of Dimension A (Answerability) in the XYZA Tesseract Leadership Framework. The principle of *edinstvennaya sobytiynost'* bytiya (once-occurrent eventness of Being) means no general rule, institutional structure, or algorithmic process can substitute for the leader's unique, situated, unrepeatable ethical stand. Appears in Section 2.2.

**Figure 3.** Human Answerability as Shield and Stance. A Tesseract Leader stands between three AI-related forces — the AI system (above), the technical architecture (left), and AI reasoning processes (right) — holding a magnifying glass of critical analysis while bearing the Human Answerability shield. The image captures the essential posture of Dimension A: active engagement with AI without surrendering epistemic authority or ethical accountability to it. Appears in Section 7.

**Figure 4.** XYZA Tesseract Leadership Profiles for Six AI Corporate Leaders with Ghost Vortex Depth Indicators. Tornado imagery surrounding the XYZA summary table represents the depth and rotational intensity of each platform's Ghost Vortex. Leaders scoring GVD = 3 (Altman/ChatGPT, Musk/Grok, Wenfeng/DeepSeek) appear in the extreme vortex zone; Zuckerberg/Meta AI (GVD = 2.5) and Pichai/Gemini (GVD = 2) appear in the deep zone; Amodei/Claude (GVD = 1) appears in the moderate zone. Appears in Section 5.

## References

- Bakhtin, M. M. (1984). *Problems of Dostoevsky's poetics* (C. Emerson, Trans.). University of Minnesota Press.
- Bakhtin, M. M. (1993). *Toward a philosophy of the act* (V. Liapunov, Trans.). University of Texas Press.
- Baudrillard, J. (1994). *Simulacra and simulation* (S. F. Glaser, Trans.). University of Michigan Press. (Original work published 1981)
- Blake, R. R., & Mouton, J. S. (1964). *The managerial grid*. Gulf Publishing.
- Boje, D. M. (1980). Making a horse out of a camel: A contingency model for managing the problem solving process in groups. In H. Leavitt, L. Pondy, & D. Boje (Eds.), *Readings in managerial psychology* (pp. 445–470). University of Chicago Press.
- Boje, D. M. (1995). Stories of the storytelling organization: A postmodern analysis of Disney as 'Tamara-Land.' *Academy of Management Journal*, 38(4), 997–1035.  
<https://doi.org/10.2307/256618>
- Boje, D. M. (2001). *Narrative methods for organizational and communication research*. SAGE Publications.
- Boje, D. M. (2002). Theatrics of SEAM: Beyond the Pentad to the Hexad. In A. M. Savall & V. Zardet (Eds.), *Mastering hidden costs and socio-economic performance*. Information Age Publishing.



- Boje, D. M., & Rhodes, C. (2005). The virtual leader construct: The mass mediatization and simulation of transformational leadership. *Leadership*, 1(4), 407–428.
- Boje, D. M., & Rhodes, C. (2006). The leadership of Ronald McDonald: Double narration and stylistic lines of transformation. *The Leadership Quarterly*, 17(1), 94–103.  
<https://doi.org/10.1016/j.leaqua.2005.10.008>
- Boje, D. M., & Vivara. (2026). Ghost Vortex and the Eidolon Vortex: Theoretical constructs for AI corporate leadership analysis. Working paper, New Mexico State University.
- Boje, D. M., Rosile, G. A., Saylors, R., & Pathak, R. (2011). Antenarratives of organizational change: The microstoria, IBM and the civil rights movement. *Human Relations*, 64(12), 1661–1685.
- Burke, K. (1945). *A grammar of motives*. Prentice-Hall.
- European Parliament. (2024). Regulation (EU) 2024/1689 on artificial intelligence (AI Act). Official Journal of the European Union.
- Fiedler, F. E. (1967). *A theory of leadership effectiveness*. McGraw-Hill.
- Haslam, S. A., Alvesson, M., & Reicher, S. D. (2024). Zombie leadership: Dead ideas that still walk among us. *The Leadership Quarterly*, 35(6), 101770.  
<https://doi.org/10.1016/j.leaqua.2023.101770>
- Heidegger, M. (1962). *Being and time* (J. Macquarrie & E. Robinson, Trans.). Harper & Row. (Original work published 1927)
- House, R. J., Hanges, P. J., Javidan, M., Dorfman, P. W., & Gupta, V. (Eds.). (2004). *Culture, leadership, and organizations: The GLOBE study of 62 societies*. SAGE Publications.
- Kerr, S., & Jermier, J. M. (1978). Substitutes for leadership: Their meaning and measurement. *Organizational Behavior and Human Performance*, 22(3), 375–403.  
[https://doi.org/10.1016/0030-5073\(78\)90023-5](https://doi.org/10.1016/0030-5073(78)90023-5)
- Padilla, A., Hogan, R., & Kaiser, R. B. (2007). The toxic triangle: Destructive leaders, susceptible followers, and conducive environments. *The Leadership Quarterly*, 18(3), 176–194.  
<https://doi.org/10.1016/j.leaqua.2007.03.001>
- Rosile, G. A., Boje, D. M., & Claw, C. M. (2018). Ensemble leadership theory: Collectivist, relational, and heterarchical roots from indigenous contexts. *Leadership*, 14(3), 307–328.  
<https://doi.org/10.1177/1742715017720703>
- Sartre, J.-P. (1956). *Being and nothingness* (H. E. Barnes, Trans.). Philosophical Library.
- Savall, H., & Zardet, V. (2011). *The qualimetrics approach: Observing the complex object*. Information Age Publishing.
- Vroom, V. H., & Yetton, P. W. (1973). *Leadership and decision-making*. University of Pittsburgh Press.
- Zuboff, S. (2019). *The age of surveillance capitalism: The fight for a human future at the new frontier of power*. PublicAffairs.