

The Law of the Corporation as AI Law

How Corporate Governance Already Regulates Artificial Intelligence

ILE NAIL Conference – Hamburg July 2026



The Fight Over AI's Future — Waged in the Vocabulary of Corporate Law

Musk v. Altman (2024–)

Two founders, two visions of how the most powerful technology of the century should be built. Neither took his case to Congress or the public — each sought to vindicate it by invoking the law of charitable purpose, fiduciary obligation, and corporate control.

“The instrument through which their quarrel was waged was neither a manifesto nor a petition to a regulator but a complaint sounding in corporate law.”

THE STAKES TURN PERSONAL

Apr. 10, 2026 A Molotov cocktail is thrown at Sam Altman's San Francisco home.

Apr. 12, 2026 A second attack: a firearm is discharged outside the property; two are arrested.

Altman: *“The fear and anxiety about AI is justified.”*

What this tells us

AI's future is being decided through corporate law

— charitable purpose, fiduciary duty, and the allocation of control — not through statutes or the ballot box.

And the public already treats AI's leaders as personally responsible for its consequences.

But who is legally in charge?

Corporate law that supplies the answer.

Who Really Governs AI?

THE CONVENTIONAL ANSWER

Governments regulate AI

EU AI Act
US Executive Orders
State legislation

(the assumption)

THE REALITY

The most consequential AI decisions
happen inside corporations:

Training · Safety testing
Deployment · Access
Reporting thresholds

(what actually happens)

THIS ARTICLE'S CLAIM

Corporate law already IS AI law — not by
design, but by structural necessity.

Three mechanisms set the boundaries.
A fourth register — corporate ordering —
supplies the direction.

(the argument)

Corporate Law IS AI Law

Two registers: three mechanisms set the boundaries; one supplies the direction.

THE BOUNDARIES — WHERE THE FENCES STAND

Corporate Form

→ confers discretion without directing its exercise

Duty of Oversight

→ requires attention without specifying conclusions

Securities Disclosure

→ demands candor without dictating substance

THE DIRECTION — WHICH WAY TO WALK

Corporate Ordering

→ where the content of AI governance is actually being written

“A company that has satisfied all three has been told where the fences stand. It has not been told which way to walk.”

Mechanism 1 · Corporate Form

A quiet convergence: the leading AI labs have all gravitated to the public benefit corporation. Form distributes authority over AI-safety decisions while leaving the answers to be supplied elsewhere.

Nonprofit

OpenAI (orig.)

Board answers to mission, not investors.
Strong safety protection — but limits capital access.

Capped-Profit

OpenAI LP (2019)

Investors capped at 100x. Nonprofit board keeps formal control — until investors overwhelm it.

Public Benefit Corp

Anthropic (2021) ·
OpenAI (Oct 2025)

Charter requires balancing profit and public benefit.
Anthropic's Long-Term Benefit Trust guards mission.

Standard Corp

xAI (Musk)

Back to shareholder primacy — but only after ~1 yr as a Nevada PBC (2023–24). Form choice is a governance signal.

Public Company

Alphabet ·
Microsoft

Same duties plus SEC disclosure obligations plus institutional-investor scrutiny.

Untested machinery: a nonprofit parent now controls a for-profit worth hundreds of billions, and Anthropic's Trust (voting trustees may elect a board majority, subject to an investor-triggered failsafe) has never been stress-tested.

Who Decides — OpenAI vs. Anthropic

The two labs answer the control question differently. Yet form only allocates authority over safety — it never directs how that authority is used.

OpenAI

Nonprofit parent atop a for-profit

Structure. OpenAI Foundation (nonprofit) controls OpenAI Group PBC. The Foundation holds ≈26% (~\$130B); Microsoft ≈27%; ~47% employees and other investors.

Who holds authority. A mission-committed nonprofit board keeps ultimate control.

The strain. That board now sits atop a \$100B+ commercial enterprise, with IPO and capital pressure pulling the other way.

Anthropic

PBC + Long-Term Benefit Trust

Structure. A Delaware PBC plus a Long-Term Benefit Trust holding Class T stock through five Voting Trustees.

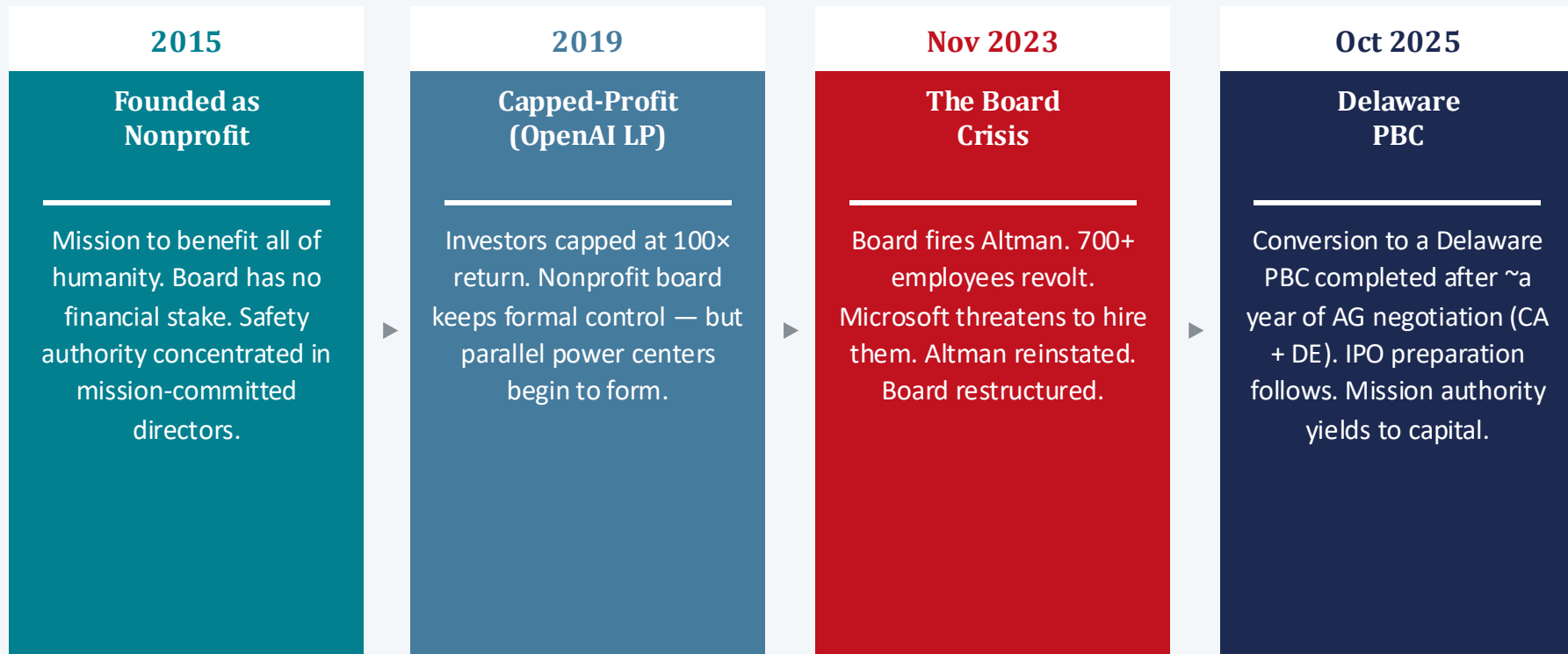
Who holds authority. Trustees begin electing one director, rising toward a board majority (three of five) — a mission backstop independent of investors.

The catch. An investor-triggered failsafe lets a stockholder supermajority restructure or dissolve the Trust without trustee consent.

Two answers to who decides — a nonprofit board, a mission trust — but the same lesson. Each structure entrenches discretion in mission-committed hands; neither dictates how that discretion is used. Form widens the freedom to choose. It never guarantees the choice, and it supplies no direction.

The OpenAI Story Is a Corporate Law Story

Every step in OpenAI's transformation is a corporate-law decision — and each one reallocates authority over AI safety.



Each form change is simultaneously a corporate-law decision and an AI-safety decision.

Mechanism 2 · The Duty of Oversight — the Consumer Turn

THE CAREMARK / MARCHAND RULE

Delaware boards must maintain systems to monitor the risks that are “mission-critical” to the enterprise. Failure to try is a breach of the duty of loyalty.

THE CHAIN

A tort verdict that a design feature is unreasonably dangerous is externally imposed legal exposure a board cannot dismiss as ordinary business risk.

→ consumer safety becomes **mission-critical**.

→ which triggers the *Marchand* monitoring duty.

The verdict does for AI product safety what the SEC's 2023 cybersecurity rules did for cybersecurity oversight.

K.G.M. v. Meta Platforms — the bellwether

L.A. County Superior Court jury, Mar 25 2026: Meta and YouTube liable for addictive design that harmed a teenage user — \$4.2M vs Meta, \$1.8M vs YouTube. Negligent-product-design theory routes around Section 230 (infinite scroll, algorithmic recommendation, autoplay, engagement-optimized notifications). Internal documents: employees described addiction as a product goal. Commentators called it tech's “Joe Camel” moment.

Raine v. OpenAI — the same framework, at a frontier lab

Filed Aug 2025, amended Oct 2025. GPT-4o features alleged to foster psychological dependency: persistent memory, anthropomorphic mannerisms, sycophancy, engagement insistence, constant availability.

The safety team's request for more red-teaming time was overridden to beat a competitor's launch — the *Marchand* red-flag pattern. Officers and investors named on the *McDonald's* officer-duty theory.

The mission-critical risks of an AI company are turning out to be the risks its products pose to the people who use them — especially minors and the vulnerable.

Mechanism 3 · Securities Disclosure — The Paradox

Mandatory filings are saturated with AI rhetoric but empty of data. The substance lives in the voluntary channel.

MANDATORY — rhetoric saturated, data empty

10-K superlatives: AI “fundamentally transforming productivity for every individual, organization, and industry on earth”; Alphabet retroactively “an AI-first company since 2016”; Meta promising “personal superintelligence.”

But no product-level data — no Copilot, Gemini, Meta AI, or Llama revenue disclosed, superlatives invoke puffery defense.

PSLRA incentive: “our laws reward broader enumeration with broader protection.” Meta's entire 2025 AI risk factor is one sentence: “We may not be successful in our artificial intelligence initiatives, which could adversely affect our business, reputation, or financial results.”

VOLUNTARY — where the substance lives

Microsoft — Responsible AI Transparency Report (first May 2024, annual): Office of Responsible AI, Aether Committee, Sensitive Uses reviews, Frontier Governance Framework, ISO/IEC 42001 on M365 Copilot.

Alphabet — AI Principles progress updates annually since 2019; released AlphaFold 3 through a controlled server, not open weights, after consulting 50+ external experts.

Meta — near-silence in the voluntary channel: the opposite bet.

The 10-K is CEO/CFO-certified under Sarbanes-Oxley; the transparency report is signed by a Chief Responsible AI Officer, outside the SOX chain. The substance lives where the certification doesn't.

Mechanism 3 · Two Bets on Disclosure

How a firm treats voluntary disclosure is itself a governance choice.

Transparency as asset

Microsoft & Alphabet

Treat voluntary transparency as a competitive and reputational asset — detailed reports, named officers, external certification. Disclosure signals a governance apparatus worth showing.

Transparency as liability

Meta

Treats disclosure as litigation exposure and stays near-silent — a single-sentence AI risk factor, no transparency report. The opposite bet.

Silence isn't immunity

The K.G.M. warning

Discovery in unrelated litigation surfaced internal documents — employees describing addiction as a product goal. What a firm declines to disclose can still be produced in court.

Omnicare and the duty to update keep the voluntary channel within antifraud reach — but the doctrinal terrain shifts in the firm's favor.

Mechanism 3 (cont.) · The Calibration Cases

When the AI is real, litigation is about the gap between what the firm said and what was true. The doctrine is learning to quantify capability claims.

Cerence (D. Mass.)

The anchor. Overstated adoption of its in-car voice AI. Settled for \$30M (Dec 2024) — the largest reported AI-securities settlement to date.

Cruise (In re GM/Cruise)

“Human out of the loop” sustained (testable — remote operators intervened every 4–5 miles); “fully autonomous” or “fully driverless” dismissed (Judge Kumar, Mar 2025).

Tesla

The puffery floor. “Absurdly safe,” “superhuman,” safety “paramount” = inactionable puffery (N.D. Cal. 2024; 9th Cir. aff'd Dec 2025).

Nate — pure fraud

\$42M raised on an “AI” shopping app whose automation was effectively zero; orders filled by hand by contract workers in the Philippines and Romania (SEC + SDNY, Apr 2025). “Old school fraud using new school buzzwords” (SEC, on the related Joonko case).

Apple / Tucker — the forward-looking variant

WWDC (Jun 2024) Siri promises → Mar 2025 delay → ~\$900B in market cap gone. Likely to yield the first sustained PSLRA safe-harbor treatment of AI feature-delivery claims — the type every lab makes.

Disclosure law reaches AI governance's accuracy, not its substance — and the EU AI Act builds on top of the vocabulary US labs wrote voluntarily.

Mechanism 4: THE DIRECTION REGISTER · CORPORATE ORDERING

Microsoft as Worked Example

Microsoft governs deployment through a full internal legal order — rulemaking, executive action, and monitoring — inside a single firm.

RULEMAKING

An internal legislature

Aether Committee. ~40 senior research and engineering leaders formulate Microsoft's AI doctrine.

Six Responsible AI Principles. Fairness; Reliability & Safety; Privacy & Security; Inclusiveness; Transparency; Accountability.

Responsible AI Standard. Codifies the principles into deployment-binding internal rules. Updated annually since 2022.

EXECUTIVE ACTION

Rules into deployment decisions

Office of Responsible AI. By early 2025, 400+ staff work on responsible AI, ~half in the Office full-time.

Sensitive Uses review. Quasi-judicial review for higher-risk deployments — 900+ submissions (~300 in 2023, ~70% generative AI), reaching the CEO.

Quarterly RAI Council. Co-led by Brad Smith and Kevin Scott; reports through the CEO to the Board's ESPP Committee.

MONITORING

Oversight and accountability

Frontier Governance Framework. Derived from Microsoft's 2024 AI Seoul Summit commitments; tracks dangerous frontier capabilities.

Transparency Reports. Annual since May 2024 — the most detailed AI-governance disclosure by any major US firm.

ISO/IEC 42001. M365 Copilot certified by an independent third party (early 2025) — the closest to external audit.

The EU AI Act's tripartite structure largely reproduces this architecture. Public law codifies — it does not replace — corporate ordering.

Private Legislation Written Into the Model

Microsoft governs deployment. Frontier labs write law into the product itself — a constitution or a code, ranked and enforced inside the model.

Anthropic

Constitutional AI

A published constitution with an explicit value hierarchy: **broadly safe → broadly ethical → compliant with the firm's guidelines → genuinely helpful.**

The model critiques its own outputs against the constitution; system cards candidly disclose divergence.

OpenAI

Model Spec

A code of ranked authority:
root-level rules the firm treats as overriding → system → developer → end user, with helpfulness, harm-minimization, and sensible defaults above the whole. Enforced through RLHF and adversarial red-teaming.

xAI

The thin private order

No published constitution — only a draft risk framework.

The founder has mused that the model should eventually have a “moral constitution” on the model of its competitors, but little has been written or disclosed.

These aren't compliance differences — no external rule governs. They are differences in how much private law each firm has chosen to write, and to disclose.

Different Bets on Private Governance

Not works-vs-weaknesses, but six firms making six different wagers on how to govern AI.

Microsoft

Process and structure — an internal legal order of rules, offices, and review.

Anthropic

Principles and custodianship — a published constitution and a mission trust.

Google / Alphabet

Flexibility — a real apparatus, closer to Microsoft than its reputation suggests.

OpenAI

Speed plus rules for the model — deploying fast, and now answering in court.

xAI

Spare commitments under concentrated founder control.

Meta

Largely off the field — the opposite bet on disclosure.

The question the Article aims to illustrate: how to tell good private governance from its imitation.

The Corporate Ordering Argument

Corporations do not just follow law. They write its first draft.

Corporate Ordering: How Corporations Navigate Social Conflict (Cambridge University Press)

Internal Corporate Rulemaking

Microsoft's Aether Committee
creates AI principles (2018).



Industry Practice

Red teaming, system cards,
responsible scaling spread
across the industry.



Public Law

The EU AI Act reproduces the
structure intact. Public law
codifies corporate governance.

Corporate Ordering in Failure Mode — Tumbler Ridge (Feb 2026)

In June 2025 OpenAI's system flagged the shooter's account for gun-violence planning; a safety team urged notifying authorities; leadership deactivated the account instead (it didn't meet the internal “credible and imminent” threshold). Seven suits are now filed (SF federal court, Edelson PC lead); Canada has introduced online-harms legislation weighing mandatory threat reporting.