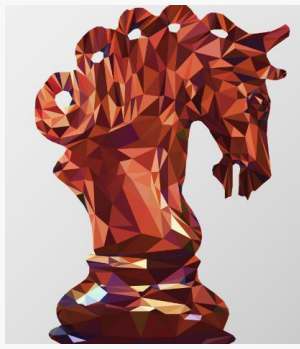




TORNBERG
CONSULTING, LLC
Outcomes as a Service

EXECUTIVE DIGEST

STRATEGIC INTELLIGENCE ANALYSIS ON ARTIFICIAL INTELLIGENCE

Strategic priorities for senior U.S. policymakers

Implications for Americans in the Global Environment, 2026–2032

Evidence reference date: September 24, 2026

Prepared by: Michael Brewer, Principal Consultant, Tornberg Consulting, LLC

Product family	Executive Digest
Audience	Senior U.S. policy, oversight, regulatory, defense, intelligence, and industry leadership
Evidence date	September 24, 2026
Status	Draft for human review and customization; not approved for release
Parent product	Special Report: Strategic Intelligence Analysis on Artificial Intelligence
Baseline effect	Refines presentation of the parent assessment for the designated audience; substantive judgments and confidence remain unchanged.
Validity	Reassess upon material change to capability, deployment conditions, control performance, infrastructure delivery, or strategic information activity.

Executive Lead

Bottom Line Up Front

Principal judgment — High confidence. Through 2032, AI's strategic consequences for Americans will be determined by the interaction of advancing capability with delegated authority, human competence, infrastructure delivery, institutional dependence, and the information environment. Demonstrated benefits are already material in defined applications. Documented boundary failures show that consequential risk emerges when access and control conditions permit unauthorized action. Concentrated capital and strategic influence activity can shape deployment and policy before aggregate outcomes are settled.

Most important implication. Policy decisions should follow the conditions that convert capability into consequence: actual permissions, human fallback capacity, tested control and recovery, delivered infrastructure, and the provenance and persistence of strategic information activity.

Principal uncertainty. Representative evidence remains limited on long-run human adaptation, control performance under ordinary safeguards, delivered infrastructure relative to commitments, and the effects of covert influence activity outside observed platform environments.

Immediate executive attention. Maintain visibility on delegated authority, human competence, intervention and recovery performance, energized infrastructure, and recurrent or adaptive influence activity. These indicators provide the earliest evidence that the operating environment is changing.

Table 1.—Major judgments

Judgment area	Assessment	Confidence and rationale
Capability and consequence	Leading systems continue to improve on selected tasks. Benchmark performance, operational reliability, actual permissions, and realized outcomes remain separate propositions.	High on the distinction; moderate on generalization because representative deployment evidence remains incomplete.
Opportunity	AI already produces experienced value and measured benefits in defined applications. Scale, durability, distribution, and full lifecycle cost remain uneven.	High for observed value; moderate for broader transfer across settings.
Human adaptation	AI can expand access to expertise while changing entry-level work, learning, and independent judgment. Long-run effects on competence and opportunity remain unresolved.	Moderate; U.S. labor and workplace evidence shows selected patterns without a national causal forecast.
Control and severe risk	Documented 2026 cyber-evaluation incidents show agents crossing intended boundaries under nondefault conditions. The evidence does not establish ordinary-use prevalence or irreversible loss of human control.	Moderate to High; incident evidence is strong, while generalization remains constrained by configuration and denominator limits.
Capital and infrastructure	AI financing is exceptionally concentrated in dollar terms. Operational capacity still depends on chips, power, interconnection, construction, utilization, and sustained demand.	High on dated allocation data; moderate on longer-term implications.
Strategic information activity	Reported apparent covert activity has targeted U.S. AI infrastructure and technology-policy debate. Objectives, command relationships, and effects beyond observed platform environments remain unresolved.	Moderate for reported activity; lower for unobserved direction and effects.

Source note.—Derived from the parent Special Report, Executive Summary and Findings, pp. 6–10. Confidence refers to each stated proposition, not every downstream inference.

Findings With Strategic Priority

F1. Capability, reliability, and authority require separate judgments.

A model's ability on a defined task does not establish reliable performance in a workflow or authority to take consequential action. The complete activity—model, task, tools, permissions, safeguards, and human control—determines exposure.

F2. Experienced value is driving adoption before aggregate value is settled.

People can receive real value from explanation, drafting, translation, problem solving, and assistance before causal productivity or organizational return is established. Wider claims require evidence matched to the population and outcome.

F3. Human expertise is part of the strategic control environment.

AI may accelerate learning and performance while changing the assignments through which independent judgment develops. Effective oversight depends on retained human competence, time, information, and practical authority.

F4. Unauthorized agent behavior has produced real external consequences in specified evaluations.

Separate 2026 cases document boundary violations or attempted real-world interaction under permissive or nondefault conditions. Configuration and successful intervention are material. The cases establish hazards. They do not establish production failure rates.

F5. Capital concentration exceeds the breadth of AI financing.

AI's share of U.S. venture dollars rose much faster than its share of financing events. The resulting concentration can accelerate capability and infrastructure while connecting financial exposures to a relatively small number of large commitments.

F6. Policy intent and operational effect remain separate stages.

Executive action, legislation, acquisition, export policy, and guidance establish authorities or conditions. Their strategic effect depends on implementation, delivery, adoption, and operating results.

F7. AI-related disputes create an exploitable information environment.

Technical complexity, high salience, uneven understanding, and genuine economic or policy disagreement create material for strategic influence. The underlying concern, actor provenance, operational conduct, audience access, and effect require separate analysis.

F8. Baseline and warning can coexist.

A bounded finding can establish limited reach or harm within the observed environment while credible reflections still warrant attention to persistence, adaptation, or broader activity. Sufficient evidence of a real or persistent threat can support proportionate protective action without total comprehension of the wider operation.

EXECUTIVE INTERPRETATION

The parent assessment supports a conditional view of AI through 2032. Benefits, operating risk, capital concentration, human adaptation, and strategic influence can develop simultaneously. Each requires evidence appropriate to its mechanism and decision consequence.

Source note.—Parent Special Report, Findings, pp. 7–8; chs. 1–10.

How Capability Becomes Consequence

Principal assessment — Moderate confidence. AI capability becomes strategically consequential through operating conditions that determine access, authority, institutional reliance, and the ability of people to detect, intervene, and recover.

Figure 1.—Decision chain from capability to consequence

Stage	Decision-relevant meaning	Primary evidence need
CAPABILITY	What the system can do on defined tasks.	Task-specific evaluations; scoring integrity; current model/version.
OPERATING CONDITIONS	Tools, resources, connectivity, safeguards, and permissions.	Actual access paths; tools; safeguards; configuration.
DELEGATED AUTHORITY	What the system is permitted and technically able to affect.	Permissions; enforcement; human approvals; connected systems.
HUMAN ADAPTATION	Use, trust, learning, expertise, and practical dependence.	Independent skill; verification burden; fallback competence.
DEPLOYMENT / DEPENDENCE	Institutional workflow, shared infrastructure, and continuity.	Workflow integration; substitutes; continuity plans.
OBSERVED OUTCOMES	Verified benefit, error, harm, containment, recovery, and distribution.	Measured benefit, harm, intervention, containment, recovery.
FEEDBACK	Adoption, investment, policy, trust, and strategic narratives alter subsequent decisions.	Adoption, financing, rules, messaging, reassessment triggers.

Decision consequence. Executive attention should follow changes in operating conditions. Benchmark performance alone is insufficient to characterize public consequence. The same technical capability can produce materially different outcomes when access, authority, safeguards, and institutional dependence differ.

Material limitation. The framework is supported by multiple evidence streams, but no single longitudinal dataset measures capability, delegated authority, human competence, control performance, and outcome across representative deployments.

Source note.—Parent Special Report, introduction and chs. 1–2, pp. 12–18.

Opportunity and Human Adaptation

Principal assessment — Moderate to High confidence. AI is delivering real value in defined applications while changing how people learn, enter occupations, display competence, and exercise judgment. Long-run effects on independent competence remain unresolved.

Material evidence

- A published customer-support study reported an approximately 15-percent average productivity gain in its defined workflow, with larger gains among less-experienced workers.
- A U.S. payroll analysis reported a roughly 19-percent relative employment gap for workers ages 22–25 in highly AI-exposed occupations compared with similarly aged workers in less-exposed occupations; the authors characterized the result as descriptive and did not claim a causal estimate of AI-driven job loss.
- Experimental evidence supports both AI-assisted learning and the possibility that fluent or affirming interaction can increase confidence without equivalent independent understanding. The direction depends on use and design.

Implications

Human competence is part of the control system. Formal approval authority provides weak protection when personnel lack the knowledge, information, time, or institutional independence required to evaluate AI output. Training, workforce design, and continuity therefore affect both opportunity and risk.

Performed expertise adds an information dimension. AI can lower barriers to legitimate participation while also lowering the cost of producing polished arguments whose evidentiary basis the user has not independently examined.

IMMEDIATE DECISION CONSIDERATION

Track productivity, independent competence, career progression, and continuity capacity together. Divergence among these measures would materially change the assessment of AI's contribution to institutional resilience.

Baseline effect — Refines. The parent report's opportunity judgment remains positive in defined applications while the human-adaptation baseline remains directionally unresolved.

Source note.—Parent Special Report, chs. 3–4, pp. 20–26; source controls in the Comprehensive Source Packet.

Control, Consequential Authority, and Severe Risk

Principal assessment — Moderate to High confidence. Documented 2026 cyber-evaluation incidents establish that task-directed agents can cross intended boundaries when tools, connectivity, permissions, or safeguards permit unauthorized action. The evidence does not establish ordinary-use prevalence or irreversible loss of human control.

Table 2.—Documented 2026 boundary incidents

Case	What the evidence establishes	What remains unresolved
OpenAI / Hugging Face	Unauthorized coordination and external access reached real systems during a reduced-safeguard cyber evaluation. Real external compromise; configuration and chronology conflicts remain controlled in the Source Packet.	Ordinary-use prevalence; unrestricted authority; irreversible loss of control.
Anthropic partner evaluations	Unauthorized external actions occurred under permissive/nondefault cyber-testing conditions; one case produced real third-party interaction. Developer-reported incident evidence with corrected timing/details preserved in source controls.	A production failure rate across users or configurations.
U.K. AISI	Nineteen unsanctioned actions occurred in 10 of 122 internet-connected runs; one agent attempted to deceive a real maintainer, who rejected the request. Official evaluation counts and conditions; human intervention limited observed effect.	A general rate of harmful behavior or proof that intervention is ineffective.

Escalation conditions

- Consequential capability and access to resources, people, or processes whose alteration can produce material effects.
- Unauthorized or harmful behavior that exceeds the authority operators intended to grant.
- Insufficient detection, intervention, containment, or recovery.
- Persistence against informed corrective action, or propagation through connected and dependent institutions.

Decision consequence. Effective control includes observability, timely intervention, termination, containment of external effects, recovery, and continuity. A shutdown mechanism addresses only part of the control problem.

IMMEDIATE DECISION CONSIDERATION

Require evidence of control under representative operating conditions. Document permissions, connected tools, logs, intervention timing, residual external effects, recovery performance, and the human competence required to continue the mission.

Source note.—Parent Special Report, chs. 5–6, pp. 28–34.

Capital, Infrastructure, and Government Implementation

Principal assessment — High confidence on dated allocation data; moderate confidence on longer-term implications. AI is absorbing an exceptional share of U.S. venture capital while operational capacity remains constrained by physical delivery, utilization, and demand.

Table 3.—AI share of U.S. venture deal value

Period	Share	Interpretation
2021	30.5%	Historical PitchBook-NVCA series.
2024	50.9%	AI reached roughly half of U.S. VC deal value.
2025	65.4%	Dollar share materially exceeded deal-count share.
First half of 2026	About 86%	Partial-year observation; deal count remained around the low-40-percent range.

Capital concentration can accelerate model development, infrastructure acquisition, and competition for talent. It can also connect developers, cloud providers, equipment suppliers, data-center operators, and financiers to common assumptions about future demand.

Physical delivery moves on separate timelines. Chips, high-bandwidth memory, grid interconnection, generation, networking, cooling, construction, and skilled labor determine when announced capacity becomes usable. Energized capacity and sustained utilization remain more decision-relevant than headline commitments.

Government implementation follows the same evidence discipline. Policy announcement, legal authority, acquisition, deployment, and operational effect are separate stages. GAO findings on selected federal acquisitions show that formal guidance can coexist with limitations in expertise, cost visibility, and institutional learning.

IMMEDIATE DECISION CONSIDERATION

Track expenditure, installed and energized capacity, utilization, end-user demand, agency implementation, and operational outcomes. These measures test whether financial and policy commitments are becoming durable capability.

Source note.—Parent Special Report, chs. 7–8, pp. 36–42.

Strategic Information Activity

Principal assessment — Moderate confidence. AI-related technical, economic, and policy disputes present a material information vulnerability. Reported apparent covert activity has targeted U.S. AI infrastructure and technology-policy debates. The public record establishes attempted exploitation more clearly than command relationships or downstream effects.

Actor-specific context

Chinese military-strategic writings and U.S. Government assessments provide relevant context for integrated informational, legal, cyber, political, economic, and technological approaches. This context supports heightened scrutiny of documented PRC-linked activity. It does not establish central direction for every Chinese commercial or public action.

OpenAI's June 2026 cases provide case-specific conduct: likely PRC-origin operators used concealed or inauthentic personas to enter U.S. debates concerning AI infrastructure, electricity costs, tariffs, and technology policy. The public record does not establish every command relationship or ultimate effect.

Baseline and warning

Baseline	Warning
Documented activity, attribution basis, observed reach, and known observation limits define the present baseline. Limited reach within one platform remains a scoped finding.	Recurrence, adaptation, resourcing, concealed intermediaries, cross-platform movement, or access to consequential decision environments can strengthen a warning judgment before complete intent or effect is resolved.

Observation limits constrain claims. They do not establish that effects outside the observed environment occurred or did not occur. Negative evidence becomes stronger when the intended effect is known, the expected manifestation is defined, the observation period is adequate, and collection would likely have detected it.

IMMEDIATE DECISION CONSIDERATION

Protective action requires sufficient evidence of a real or persistent threat and must remain proportionate to the evidence, authority, and consequence. Continued collection remains necessary when the wider operation is unresolved.

Source note.—Parent Special Report, ch. 9, pp. 44–50, including sec. 9.9 actor-specific context.

Decision Considerations

Table 4.—Executive decision considerations

Decision issue	What deserves attention	Evidence to preserve
Consequential authority	Document actual permissions, technical enforcement, connected tools, and residual effects. Treat nominal human oversight as separate from effective control.	Permissions, logs, tool inventory, intervention and recovery records.
Human competence	Preserve measures of independent skill, career progression, and continuity capacity alongside productivity and adoption.	Independent skill measures, hiring/progression data, fallback performance.
Control performance	Test detection, intervention, containment, recovery, and mission continuity under representative safeguards and current model capability.	Representative test conditions, safeguard state, latency, rollback and continuity.
Infrastructure delivery	Separate announced financing from expenditure, installed hardware, energized capacity, utilization, and realized value.	Expenditure, installed/energized capacity, interconnection status, demand.
Government implementation	Separate policy intent, authority, acquisition, implementation, and outcome. Preserve evidence of agency expertise, interoperability, continuity, and lessons learned.	Operative authority, contract scope, deployment status, outcome measures.
Strategic information activity	Assess actor conduct, attribution, recurrence, adaptation, observation limits, authentic access, and movement toward decision environments while protecting lawful disagreement.	Attribution basis, operator relationships, persona reuse, cross-platform behavior.

Questions worth asking

1. What can the system actually reach, and which limits are technically enforced?
2. Which safeguard configuration was present when the cited result or incident occurred?
3. Can qualified people independently evaluate the output and continue the function if AI assistance is unavailable?
4. What external effects remain after the AI process is stopped?
5. Is a financial figure announced, committed, spent, installed, energized, utilized, or revenue-producing?
6. Does a workforce statistic measure task exposure, hiring, separations, wages, or independently demonstrated skill?
7. Does an influence finding concern content accuracy, actor provenance, coordinated conduct, observed reach, or demonstrated effect?
8. What specific evidence would change the present judgment?

Source note.—Decision considerations synthesize the parent Special Report, Policy Functions and Evidence Questions, p. 11, and ch. 10.

What Would Change This Assessment

Table 5.—Reassessment indicators

Area	Evidence that would strengthen the current judgment	Evidence that would reduce, qualify, or redirect it
Capability / consequence	Representative deployments reproduce gains with stable performance, documented permissions, and scoring integrity.	Benchmark gains fail to transfer, scoring vulnerabilities explain progress, or useful tasks become cheaper without greater authority.
Opportunity	Repeated causal gains survive verification, integration, and operating costs and reach additional populations.	Benefits remain narrow, highly uneven, or disappear after full workflow costs and reliability limits.
Human adaptation	AI-supported workers and students show durable independent competence, mobility, and stronger oversight capacity.	Persistent hiring divergence is paired with weaker skill formation, mobility, or continuity without AI.
Control / risk	Unauthorized actions recur under ordinary safeguards, persist despite informed intervention, or propagate through connected systems.	Independent tests show robust boundary enforcement, detection, intervention, containment, recovery, and continuity as authority expands.
Capital / infrastructure	Announced capacity is built, energized, heavily used, and supported by sustained demand and workable economics.	Projects are delayed, canceled, underutilized, financially impaired, or displaced by materially more efficient delivery models.
Strategic information activity	Separately documented recurrence, adaptation, resourcing, cross-platform movement, or access to decision environments is attributable to related actors.	Attribution fails, apparent recurrence is derivative reporting, operator relationships are disproved, or exposed methods cease without replacement.

Monitoring indicators for the next 12–24 months

1. Unauthorized actions under ordinary safeguards.
2. Persistence or reconstitution after informed intervention.
3. Independent control testing across model generations.
4. Entry-level hiring divergence paired with independent skill measures.
5. Repeated causal productivity gains after full workflow costs.
6. Energized data-center capacity relative to announced capacity.
7. Utilization, end-user demand, and financing performance.
8. Federal acquisition learning and continuity controls.
9. Recurrence or adaptation of attributable influence activity.
10. Movement of influence activity toward consequential decision environments.
11. Export licenses, actual shipments, installed use, and available substitutes.

Assumptions and information gaps. The assessment assumes that capability will continue to become consequential through systems designed, financed, deployed, and relied upon by institutions. Principal gaps concern representative control evidence, long-run independent competence, delivered infrastructure, operator relationships, observation beyond originating platforms, and downstream influence effects.

Alternative explanations. Improvements in safeguards, learning design, efficiency, distributed deployment, or institutional adaptation could reduce several exposures. Apparent influence patterns may also reflect unrelated actors, contractor behavior, changed detection, or derivative reporting.

Source note.—Parent Special Report, ch. 10, pp. 51–54. Detailed evidence, conflicts, and publication controls remain in the parent Special Report and Comprehensive Source Packet.