

Q2 2026 AI Security Executive Brief

AI Security Is Moving from Model Outputs to Authorised Actions

EXECUTIVE / LEADERSHIP EDITION

Field	Detail
Document ID	ODA3-2026-07-EXB-SEC-005
Status	Final
Reporting period	1 April – 30 June 2026
Evidence cut-off	30 June 2026
Publication	July 2026
Companion	Q2 2026 AI Security Technical & Compliance Report (ODA3-2026-07-TCR-SEC-004)
Publisher	ODA3 Institute
Classification	Public analytical briefing
Distribution	Approved for public release

In this brief, *AI* means an artificial-intelligence system and, where the context is an agentic deployment, the model together with the identity, data, tools and execution authority that surround it.

WHAT DOES “AGENTIC” MEAN IN THIS REPORT?

An agentic system can interpret a task, select intermediate actions, and use tools or external services to pursue an objective without requiring step-by-step human direction. Its security significance depends on the authority, data and systems connected to it — not on the model behaving as an independent actor.

Important Notice

This report is an intelligence analysis based on information available to ODA3 Institute as of 30 June 2026. It does not constitute legal, regulatory, financial, audit, certification or compliance advice. The recommendations, implementation priorities and control outcomes represent ODA3 Institute’s analytical assessment; they do not guarantee security outcomes or regulatory compliance. Organisations should obtain advice appropriate to their jurisdiction, sector, systems and circumstances.

Contents

1	How to Use This Report
2	Executive Summary
3	Methodology Note
4	Executive Assessment
5	What Changed Since Q1 2026
6	Q2 2026 at a Glance

7	Principal Findings
8	The Evidence Behind the Findings
9	Notably Absent
10	Open Questions for Q3 2026
11	Q2 Record Portfolio
12	Five Decisions Required from Leadership
13	90-Day Implementation Sequence
14	Financial Treatment
15	Regulatory and Assurance Implications
16	Confidence and Limitations
17	Conclusion
18	Glossary

1. How to Use This Report

This Executive Brief is written for boards, executive leadership, and the functions responsible for security, technology, legal, risk and AI governance. It carries the quarter's principal judgements, the decisions leadership must own, and a sequenced implementation plan. It deliberately uses **confidence language** (High / Moderate / Low / Indeterminate) rather than technical evidence tags. Full source analysis, evidence tiering, MITRE ATT&CK mapping, detection guidance and standards-gap analysis are in the companion **Technical & Compliance Report**.

Readers seeking immediate governance actions may proceed to **Five Decisions Required from Leadership** and the **90-Day Implementation Sequence**.

A note on continuity: this is the second report in a quarterly series. Where this quarter's evidence revises, extends or carries forward a Q1 2026 judgement, we say so — a serious intelligence product should be readable as a developing assessment, not a series of unconnected headlines.

2. Executive Summary

The bottom line

Q2 2026 reinforced that the material security impact of an AI system depends increasingly on the **authority, data and tools available to it** — not only on the content it generates. The quarter's most consequential events did not turn on what a model *said*; they turned on what AI-connected systems were permitted to *retrieve, invoke, change or transmit*, and on how fast a disclosed weakness could be turned into a working attack.

The shift

Organisations are connecting AI systems to enterprise credentials, business data, development environments and operational tools. In many cases, equivalent controls over identity, delegated authority, execution and evidence have not developed at the same pace. The central security question is changing from “*What can the model say?*” to “*What can the system retrieve, invoke, change or transmit — and which independent controls govern those actions?*”

Three executive takeaways

1. **Authority is becoming the principal risk multiplier.** A system with broad credentials, external tools and write authority creates materially different exposure from a model that only generates text. This quarter, the attacks that mattered most reached credentials, source code and build pipelines — not chat interfaces.
2. **Observability is becoming a control requirement, not a convenience.** Many organisations cannot reconstruct the chain from human intent to model invocation, tool execution and resulting system change. Without that evidence, accountability, investigation and assurance remain incomplete.
3. **Supply-chain and concentration risk now extend to the AI stack itself.** AI systems depend on model providers, gateways, agent frameworks, tool servers, behavioural configurations and external services. A compromise, failure or government restriction affecting any one of these may disrupt operations or alter behaviour — as a government export-control action demonstrated this quarter.

Immediate leadership imperative

Govern material AI systems with the same discipline applied to privileged human users, production applications and critical infrastructure: unique identities, bounded authority, independent execution controls, action-level logging, supply-chain assurance, and tested revocation and recovery.

3. Methodology Note

ODA3 Institute reviewed publicly available reporting on realised incidents, attempted incidents, exposure events, vulnerabilities, research observations, threat signals and regulatory developments relevant to AI security during Q2 2026. **These record types are assessed separately and are not combined into a single incident count** — a discipline that prevents the common error of inflating a quarter by adding vulnerabilities, lab demonstrations and policy moves to confirmed incidents.

Material claims are evaluated on source quality, independent corroboration, technical specificity, internal consistency, date accuracy, relevance to AI security, and the availability of authoritative evidence. Analytical confidence is expressed as **High, Moderate, Low or Indeterminate**.

Naming standard. We name an event where it has been disclosed by the affected party, confirmed by a government body, or assigned a public identifier (for example, a Common Vulnerabilities and Exposures (CVE) identifier or a Known Exploited Vulnerabilities (KEV) listing). We anonymise or qualify an event where attribution or impact remains unverified. This keeps the report verifiable where the public record is firm, and appropriately cautious where it is not — rather than naming everything or anonymising everything.

Financial discipline. Where incident-level costs are estimated, we separate disclosed realised loss, modelled incremental response and recovery cost, and scenario exposure; every estimate carries a low–high range, a confidence assessment, and explicit controls against double-counting. **No single aggregate Q2 financial figure is published**, because the available incident population and disclosed values do not support defensible aggregation.

This report is not a census. Public reporting is incomplete, uneven across jurisdictions, and shaped by differences in disclosure obligations, investigative maturity and commercial incentive.

4. Executive Assessment

Q2 2026 strengthened the assessment that the most consequential AI security risks arise from the **authority, integrations and execution mechanisms surrounding a model** — not primarily from model output. Material consequences become possible when a system can retrieve protected information, use enterprise credentials, invoke external tools, modify code or infrastructure, send external communications, change permissions, initiate transactions, or act across multiple systems without independent authorisation.

ODA3 Institute assesses with **moderate confidence** that many organisations are granting operational authority to AI-enabled systems faster than they are establishing equivalent controls over identity, delegation, execution, monitoring and forensic evidence.

The quarter did **not** establish that AI systems have independently replaced human attackers, nor that existing security controls are obsolete. It reinforced a narrower, more operationally significant conclusion:

AI systems are being granted enterprise authority faster than organisations are establishing equivalent controls over identity, delegation, execution and evidence.

5. What Changed Since Q1 2026

Q1 2026 framed the AI control plane — identity, permissions, orchestration, validation and observability inside deployed systems — as the new battlefield. Q2 did not overturn that; it extended it **up the stack** and **forward in time**.

- **Up the stack:** the attack surface moved from operating a deployed AI system to compromising the software and model supply chain used to build and ship it. The recurring failure is the same one Q1 identified — inherited, unbounded trust — now expressed in package registries, editor marketplaces and AI gateways rather than in runtime permissions.
- **Forward in time:** the buffer defenders relied on between a vulnerability’s disclosure and its weaponisation contracted sharply, as research demonstrated a frontier model turning a disclosed patch into a working exploit in hours.
- **A new axis:** access to a frontier model became, for the first time, a government-controlled and revocable dependency.

The throughline across both quarters is constant: **whatever the ecosystem trusts by default becomes the delivery mechanism**.

6. Q2 2026 at a Glance

1. **Authority is becoming more consequential than output.** Impact depends on the credentials, tools, data and actions available to a system.
2. **Development infrastructure is becoming production infrastructure.** Gateways, agent frameworks, coding assistants, retrieval systems and behavioural configurations create security dependencies beyond the model.
3. **Model safeguards are not enforcement controls.** A refusal cannot substitute for access control, transaction authorisation, network isolation or independent termination.
4. **Observability is not keeping pace with delegated execution.** Few organisations can reconstruct intent → invocation → execution → change end to end.
5. **Supply-chain exposure extends beyond software packages** to system prompts, agent policies, tool definitions, routing rules and retrieval sources.
6. **Model availability is becoming a continuity risk.** Government restrictions, provider decisions and licensing changes can disrupt AI-enabled processes with limited notice.

7. Principal Findings

Each finding states a durable, architecture-level judgement, grounds it in this quarter's evidence, and names the control outcome it requires. The supporting events are summarised in *The Evidence Behind the Findings*.

Finding 1 — AI authority is becoming more consequential than output

A model output becomes materially more consequential when another system can act on it — converting generated content into database queries, file operations, code changes, infrastructure actions, permission changes, external publications, customer communications, financial transactions or security administration.

Why it matters. Traditional access models assume a human understands and intends each action performed through their account. That assumption weakens when an AI agent interprets a broad objective, selects intermediate actions, retrieves external content, chooses tools and operates through inherited human permissions. The agent can pass authentication while performing an action that exceeds the user's intended purpose.

Required control outcome. AI-generated decisions are not self-authorising. Each material production agent should have a unique non-human identity, a named accountable owner, a documented purpose, task-specific and time-bounded permissions, explicit tool entitlements, data-access restrictions, transaction and rate limits, emergency revocation, periodic access review, and action-level attribution. An agent must not automatically inherit the full authority of the human, service account or platform through which it operates.

Finding 2 — AI development infrastructure is now part of the production attack surface

AI systems increasingly depend on infrastructure between the model and the business process: model gateways, agent frameworks, coding assistants, retrieval systems, vector databases, tool servers, Model Context Protocol (MCP) implementations, package repositories, Continuous Integration and Continuous Delivery (CI/CD) environments, external providers, and prompt and configuration repositories. A compromise or misconfiguration in any of these can expose credentials, source code, infrastructure or business data even when the model behaves exactly as intended.

Why it matters. These tools are often introduced through fast-moving engineering workflows rather than production-governance processes, producing unreviewed dependencies, broad developer credentials, incomplete inventories, weak development/production separation and inconsistent patching. **This was the dominant Q2 failure mode** (see *The Evidence Behind the Findings*, events I-1 to I-4).

Required control outcome. AI frameworks, orchestration components and behavioural configurations should receive controls equivalent to other production software and privileged infrastructure: dependency inventory, provenance verification, version pinning, artefact signing, maintainer-account protection, configuration-integrity checks, peer review, vulnerability monitoring, patch procedures, rollback capability and runtime isolation. Note that valid build provenance is necessary but not sufficient: this quarter demonstrated malicious artefacts carrying *valid* cryptographic provenance because the signing pipeline itself was compromised.

Finding 3 — Model-level safeguards are not substitutes for external enforcement

Model safeguards are probabilistic; their behaviour varies with prompt structure, retrieved content, system instructions, tool descriptions, conversation history, model version and attacker adaptation. A model's refusal cannot be the sole security boundary for a consequential action.

Why it matters. “Do not publish confidential information” is not a destination control that prevents publication. “Do not exceed the user's authority” is not an access-control decision enforced by the target system. “Stop when asked” is not an independent mechanism that can terminate execution.

Required control outcome. High-impact actions must pass through deterministic controls outside model inference: policy-enforcement gateways, short-lived and scoped credentials, destination allowlists, transaction thresholds, human approval or dual authorisation, network segmentation, read-only modes, sandboxed tools, rate limits, independent termination, and rollback and recovery — proportionate to the authority, reversibility and impact of the action.

Finding 4 — Observability is not keeping pace with delegated execution

Conventional logs may show *that* a service account accessed a database or modified a file. They rarely show *who* initiated the task, *which* agent interpreted it, which model and version were used, what was retrieved, which tool was selected with which parameters, whether approval was required, which policy decision applied, what the target system changed, and whether the action matched intent.

Why it matters. Without an end-to-end record, an organisation may be unable to tell authorised automation from user error, agent misinterpretation, indirect prompt injection, credential misuse, configuration change or tool compromise. Conversational transcripts are not authoritative forensic records.

Required control outcome. AI action logging must connect the complete execution chain — *initiating identity* → *agent identity* → *model invocation* → *retrieved context* → *proposed action* → *authorisation decision* → *tool execution* → *target-system result* — and preserve initiating user, agent identity, model and version, task reference, retrieved-source identifiers, tool name and parameters, approval state, target resource, resulting change, error and rollback status, a correlation identifier and a timestamp. Target-system and enforcement-layer records should remain authoritative.

Finding 5 — AI supply-chain risk includes behavioural configuration

Traditional supply-chain controls focus on source code, packages, libraries, build systems and deployment artefacts. Agentic systems add behavioural dependencies — system prompts, agent policies, tool descriptions, connector definitions, retrieval sources, memory stores, routing rules, approval workflows and tool-server configurations — that can materially alter behaviour **without changing application code**.

Why it matters. A modified tool description may change which tool an agent selects; a changed routing rule may move processing to an unapproved provider; a poisoned retrieval source may steer an agent into an unauthorised action; a weakened system prompt may relax approval or data-handling constraints.

Required control outcome. Treat behavioural configuration as security-sensitive executable logic: inventoried, version-controlled, integrity-protected, peer-reviewed, tested before deployment, linked to an approved change record, monitored for unauthorised modification, and recoverable through rollback.

Finding 6 — Model availability is becoming a concentration and continuity risk

During June 2026, a United States (U.S.) government export-control action restricted access to two of Anthropic's most advanced models (Claude Fable 5 and Mythos 5), on national-security grounds. This was a **regulatory and operational-continuity development, not a cyber incident**, and it is significant because it demonstrated that access to a material AI capability can be disrupted through government intervention, export-control restriction, provider policy, service suspension, licensing change, geographic restriction or model retirement.

Why it matters. An organisation may be unable to substitute a model immediately without affecting output quality, safety performance, regulatory classification, data handling, workflow reliability, latency, cost, audit evidence or human-oversight procedures.

Required control outcome. Material AI-enabled processes should have documented procedures for abrupt provider or model unavailability: approved fallback models, manual fallback processes, prompt and workflow portability, data-location restrictions, contractual suspension rights, provider notification obligations, revalidation after substitution, and explicit handling of model-performance differences and recovery priorities.

8. The Evidence Behind the Findings

The findings above are durable judgements. The table below records the principal Q2 events that support them, named where the public record is firm and qualified where it is not. Detailed forensics, evidence tiers and sources are in the companion Technical Report.

Ref	Event (named where primary-verified)	What it demonstrates	Confidence	Findings
I-1	A major code-hosting provider confirmed that roughly 3,800 internal repositories were cloned after an employee installed a poisoned editor extension (assigned a CVE and added to the U.S. KEV catalogue); the malicious build was live in the marketplace for under twenty minutes and harvested cloud, CI/CD, vault, password-manager and AI coding-assistant secrets.	Trusted marketplace and silent auto-update as a delivery channel; developer secrets as the objective.	High	2, 3, 4
I-2	A data-labelling supplier to multiple frontier labs confirmed it was among the organisations affected by a compromise of a widely used open-source AI gateway, exposing contractor data, source code and potentially training methodologies.	A single AI dependency became a shared point of failure across the frontier-lab supply chain.	High (fact) / Moderate (scope)	2, 5
I-3	A popular package set was compromised through its CI/CD pipeline; dozens of malicious versions were published in minutes carrying valid build-provenance attestations, with downstream impact reported at several AI providers.	Valid provenance does not certify benign intent when the signing pipeline is compromised.	Moderate	2
I-4	A self-propagating package-ecosystem worm using install-time lifecycle scripts and editor task files was open-sourced, lowering the barrier for copycats.	Nation-state-style supply-chain technique moved to commodity tooling.	Moderate	2, 5
I-5	Frontier-lab red-team research demonstrated a model turning disclosed, patched vulnerabilities into working exploits in hours (a working browser exploit roughly one hour after the patch, while the fixed release was still ~18 days from users; Windows kernel privilege-escalation exploits from binaries for ~US\$2,000 each). In parallel, a U.S. government export-control action restricted access to Anthropic's Claude Fable 5 and Mythos 5.	The patch gap is collapsing; frontier-model access is a controllable, revocable dependency.	High (capability and directive) / Low (in-the-wild criminal use)	3, 6
I-6	Security research reported large numbers of internet-exposed agent-management and Model Context Protocol endpoints reachable without authentication, some exposing business logic and connected-tool access.	The AI-tooling layer is shipping faster than its security review.	Moderate	1, 2, 4

Events I-1 to I-4 are the principal supply-chain and infrastructure examples cited in the findings; I-5 supports Findings 3 and 6, and I-6 supports Findings 1, 2 and 4.

9. Notably Absent

A disciplined assessment must state what the evidence did **not** establish. These absences constrain the report's conclusions; they do not establish that the corresponding events could not occur.

Incident scope and scale

- Public disclosures did not confirm a defensible total number of realised AI security incidents for the quarter, nor a defensible aggregate exposed-record total, nor a defensible aggregate financial-loss figure, nor a reliable population-wide AI incident rate.

- No confirmed government breach exposing personally identifiable information (PII) at the scale of Q1's multi-agency intrusion (~195 million records) recurred in Q2 public reporting.

Attribution and autonomy

- The evidence did not establish that a generally available model independently selected a target, obtained access and completed an end-to-end material cyberattack without human direction or enabling infrastructure.
- The evidence did not establish AI systems as generally functioning as independent threat actors, nor that every reported AI-assisted exploit was independently discovered by a model.
- The "N-hour" exploit-development capability is demonstrated in controlled research; no confirmed **in-the-wild criminal** use was observed within the period.
- No exfiltration of the export-restricted models' weights was confirmed; the June action was an access directive, not a model theft.

Control effectiveness

- The evidence did not establish that traditional security controls are universally ineffective against AI-enabled threats, that human incident response is obsolete, or that model-level safeguards provide reliable authorisation for consequential actions.
- For the code-hosting breach (I-1), customer repositories, enterprise accounts and user data were reported **not** affected; exposure was confined to the provider's internal estate.

Financial and regulatory interpretation

- Public disclosures did not confirm that certification or regulatory compliance alone prevented or caused any reviewed event.
- No European Union (EU) AI Act high-risk fines were recorded — and the high-risk deadline itself was deferred during the quarter.
- The evidence did not support combining vulnerabilities, attempted attacks, research observations and regulatory developments into a single incident total.

10. Open Questions for Q3 2026

Looking ahead, ODA3 Institute will prioritise the following collection questions for Q3 2026. These are intelligence priorities, not predictions.

- Will proposed changes to the EU AI Act be formally adopted, and which obligations will be affected and on what timetable?
- Will public disclosures provide stronger evidence distinguishing model-assisted exploit development from independent model discovery?
- Will organisations begin reporting agent identity, permission and tool-use failures with sufficient detail for cross-incident analysis?
- Will model and provider restrictions lead organisations to implement tested substitution and continuity arrangements?
- Will supply-chain controls expand to behavioural configuration — prompts, tool descriptions and agent manifests — and not only software packages?
- Will the "N-hour" exploit-development capability be observed in the wild, or remain a research-bound concern?

11. Q2 Record Portfolio

Security-relevant records are assessed by type and **not** summed into a single "verified incident" total.

Record type	Description	Q2 treatment
Realised incident	Confirmed adverse security, privacy, safety or operational impact	Included only where impact is substantiated
Attempted incident	Malicious activity without confirmed material impact	Reported separately
Exposure event	Data, credentials or capability became accessible without confirmed exploitation	Reported separately
Vulnerability	Technical weakness, with or without observed exploitation	Not counted as an incident by default
Research observation	Controlled test or laboratory result	Informs control analysis
Regulatory development	Government or supervisory action affecting AI security or availability	Reported as a strategic development
Threat signal	Relevant activity below the incident threshold	Qualified and monitored

Where primary verification is incomplete, vendor or victim names are withheld from public status summaries to protect ongoing investigations and avoid prejudging unverified claims. The companion Technical Report holds the detailed classifications, source records and disposition decisions.

12. Five Decisions Required from Leadership

Decision 1 — Treat material AI agents as privileged non-human identities

- **Required action.** Require every production agent with access to protected data, external tools or consequential actions to have a unique identity, a named accountable owner, a documented purpose, approved entitlements, time-bounded credentials, periodic access review and immediate revocation.
- **Board-level question.** Do we know which AI systems can write to databases, modify code, change permissions, communicate externally or initiate transactions?
- **Risk of inaction.** Shared or inherited permissions may prevent the organisation from distinguishing an authorised human action from an unintended, manipulated or compromised agent action.
- **Executive owner.** Chief Information Security Officer (CISO), supported by the Chief Information Officer (CIO) and AI Governance Lead.
- **Evidence required.** Agent identity and entitlement register.

Decision 2 — Establish an AI system, agent and integration register

- **Required action.** Inventory models and providers, agents and orchestration platforms, business owners, approved purposes, data sources, tools and connectors, Open Authorization (OAuth) grants, Application Programming Interface (API) credentials, service identities, deployment environments, risk classifications, approval status and fallback arrangements — including user-procured and self-service tools that touch organisational credentials, systems or data.
- **Board-level question.** Can management identify every material AI system, its accountable owner, its data access and its available actions?
- **Risk of inaction.** Unidentified or unowned systems remain outside security review, regulatory classification, incident response and access governance.
- **Executive owner.** AI Governance Lead.
- **Evidence required.** Approved inventory with a reconciliation and review process.

Decision 3 — Require independent enforcement for consequential actions

- **Required action.** Define action classes that require deterministic policy evaluation, explicit confirmation, dual approval, transaction limits, destination restrictions, segregation of duties, sandboxed execution or emergency termination — applied to production deployment, permission modification, data deletion, external publication, financial transactions, customer communications, security-control changes and infrastructure administration.
- **Board-level question.** Which AI-initiated actions can occur today without a control *outside the model* deciding whether they should proceed?
- **Risk of inaction.** A manipulated, misunderstood or compromised output may be converted into a consequential action with no independent authorisation.
- **Executive owner.** CISO and CIO.
- **Evidence required.** Approved action-classification policy and enforcement design.

Decision 4 — Define minimum AI action-logging requirements

- **Required action.** Require records sufficient to reconstruct who initiated the task, which agent acted, which model and version, which data sources were accessed, which tool was invoked with which parameters, which approval decision applied, which system was affected, what changed, and whether rollback occurred.
- **Board-level question.** Could management reconstruct a material AI-initiated action today without relying solely on the model's own explanation?
- **Risk of inaction.** The organisation may be unable to determine responsibility, establish causation, contain an incident, or demonstrate control operation to regulators, auditors or customers.
- **Executive owner.** CISO and Security Operations Centre (SOC) leadership.
- **Evidence required.** Logging standard and a successful reconstruction of a representative action.

Decision 5 — Treat AI models, tools and frameworks as supply-chain dependencies

- **Required action.** Require dependency inventory, provider due diligence, configuration versioning, artefact-integrity controls, vulnerability monitoring, change notification, credential isolation, incident-support obligations, model-substitution planning and secure decommissioning.
- **Board-level question.** Which critical business processes depend on a model, provider, framework or integration the organisation could not replace or disable safely?
- **Risk of inaction.** A provider restriction, compromised dependency, configuration change or service suspension may disrupt operations or alter behaviour with no effective fallback.
- **Executive owner.** Procurement, General Counsel, CISO and CIO.
- **Evidence required.** Updated supplier-security requirements and continuity plans.

13. 90-Day Implementation Sequence

Days 0–30 — Establish accountability and visibility

Appoint owners for material AI systems and agents; identify production agents with write, publish, transact, deploy or administer capability; catalogue models, tools, connectors, identities and delegated credentials; flag long-lived, shared or over-broad credentials; define AI-incident escalation and materiality criteria; report material unknowns to executive leadership.

Outcome: the organisation can identify its material AI-enabled systems, their owners, their credentials and their available actions.

Days 31–60 — Restrict authority and execution

Assign unique identities to material agents; reduce inherited human and service-account permissions; implement task-scoped, time-bounded credentials; define consequential action classes; apply approval, destination and transaction controls; establish fallback procedures for critical model or provider unavailability.

Outcome: material agents operate within defined authority boundaries, and consequential actions require independent control.

Days 61–90 — Test, evidence and report

Test emergency revocation and termination; reconstruct representative agent actions from logs; test indirect prompt injection and tool misuse; validate rollback and recovery; exercise an AI-incident scenario; record exceptions, remediation owners and residual exposure.

Outcome: the organisation can demonstrate its AI controls operate in practice and report remaining exposure to leadership.

14. Financial Treatment

ODA3 Institute does **not** publish an aggregate Q2 financial-loss figure. The evidence does not support responsible aggregation of realised losses, incident-response costs, operational disruption, regulatory exposure, losses from prevented incidents, provider revenue impact, customer-migration costs, intellectual property (IP) exposure and long-tail litigation — categories that are not interchangeable and that may overlap.

Where individual incident costs are estimated, we apply:

Estimated incident cost = disclosed non-overlapping realised loss + modelled incremental response and recovery cost + probability-weighted regulatory or litigation exposure

Each estimate carries a low–high range, a confidence assessment, currency and valuation date, sector and jurisdiction assumptions, explicit overlap controls, and separation of realised and scenario values. Prevented attacks, research demonstrations, vulnerabilities and regulatory developments are not added to realised losses.

For context, and **only** as a non-aggregated illustration, the direct, defensible losses across the named Q2 supply-chain incidents (I-1 to I-3 — incident response, mass credential rotation, source-code and contract exposure) fall in the order of **US\$60M–US\$160M [confidence: Low]**. Confidence is Low because the disclosed figures capture direct incident-response and rotation costs only, and exclude downstream supply-chain exposure, customer-notification obligations and long-tail contractual liability. This is not a quarter total and excludes the systemic exposure created by a widely deployed compromised dependency and by the patch-gap collapse, which is classified as:

Not reliably quantifiable.

15. Regulatory and Assurance Implications

European Union AI Act — verify the operative instrument

The timeline must be handled carefully. Under the originally enacted timetable, significant high-risk obligations were scheduled to apply on **2 August 2026**. During Q2, EU institutions advanced a “Digital Omnibus” simplification: a provisional agreement was reported on **7 May 2026** and endorsed by the European Parliament on **16 June 2026**, deferring stand-alone high-risk (Annex III) obligations to **2 December 2027** and high-risk AI embedded in regulated products (Annex I) to **2 August 2028**.

Two points deserve board attention:

- 1. The nearest live deadline under the provisional agreement is the transparency and synthetic-content obligation (including watermarking), reported as moved to 2 December 2026**, alongside a new Article 5 prohibition on AI-generated non-consensual intimate imagery and child sexual abuse material on the same date. General-purpose AI obligations were not changed by the omnibus and remain in force.
- 2. No deadline should be treated as legally changed until the operative legal instrument is verified on the relevant assessment date.** Until the amendment is formally adopted and published in the Official Journal, the original 2 August 2026 date remains legally live. Any compliance plan should distinguish enacted obligations, provisional political agreements and adopted amendments — and distinguish provider from deployer obligations.

The deferral buys preparation time, not a pass: AI inventory and classification gate every downstream obligation and do not get easier with delay.

Existing obligations

AI-enabled systems may already fall within data-protection, cybersecurity, consumer-protection, employment, financial-services, healthcare, critical-infrastructure and contractual-security requirements. Applicability depends on jurisdiction, sector, data, purpose, organisational role and actual impact.

Evidence over policy

Regulators, assessors and internal audit may require evidence that controls operate in practice — system and integration inventories, risk classifications, identity and entitlement records, approval workflows, logging specifications, action records, testing and incident-exercise results, provider due diligence, configuration histories, human-oversight records and corrective-action evidence. Policy documents alone are not proof of effective control.

16. Confidence and Limitations

Analytical confidence. The central quarter-level assessment is assigned **moderate confidence**. The conclusion that material AI risk increasingly depends on delegated authority, integrations and execution controls is supported by the architecture of agentic systems, publicly reported developments, emerging research and recurrent weaknesses in identity, authorisation and monitoring. Confidence is reduced by limited primary incident disclosures, inconsistent use of the term “AI incident,” incomplete technical artefacts, disclosure incentives, unsupported autonomy claims, conflicting vulnerability identifiers and frequent conflation of incidents, vulnerabilities and research.

Coverage limitations. The report is limited to publicly available information at the cut-off. It may underrepresent privately resolved incidents, privileged matters, government and intelligence activity, non-English reporting, incidents in smaller organisations, internal agent failures without notification, and events where AI involvement was either unidentified or overstated.

Attribution limitations. Rapid execution, structured commands, prompt-like comments or code style may support an assessment of AI involvement, but do not individually prove which model was used, whether a model discovered a vulnerability, whether execution was autonomous, whether a human selected each action, or which specific actor was responsible.

17. Conclusion

Q2 2026 provided no verified evidence that AI systems had independently replaced human attackers. It reinforced a more operationally significant conclusion:

AI systems are being granted enterprise authority faster than organisations are establishing equivalent controls over identity, delegation, execution and evidence.

This cannot be addressed through model safeguards alone. Secure deployment requires a verifiable control environment in which every consequential AI-initiated action is attributable, authorised, purpose-bound, least-privileged, independently constrained, observable, revocable, reconstructable and recoverable.

The central leadership question is therefore not whether the organisation permits AI. It is whether the organisation can demonstrate that **every material AI-initiated action is governed as rigorously as the human, application or infrastructure action it replaces.**

18. Glossary

AI agent — an AI-enabled system that can interpret a task, select actions and use tools or external systems to pursue a defined objective.

Agentic system — a system combining model reasoning with capabilities such as planning, memory, retrieval, tool use or multi-step execution.

Behavioural configuration — prompts, policies, tool descriptions, routing rules, memory settings and related artefacts that influence how an AI system behaves.

Consequential action — an action capable of materially affecting data, systems, customers, finances, permissions, operations, security or legal obligations.

Delegated authority — permission granted to an AI system to act on behalf of a human user, service, application or organisation.

Deterministic control — a control that applies a defined rule or policy independently of probabilistic model behaviour.

Indirect prompt injection — an attack in which malicious instructions are embedded in external content retrieved or processed by an AI system, influencing it to act outside its intended purpose.

Model Context Protocol (MCP) — a protocol connecting AI systems to tools, data sources and external services through a standardised interface.

Model gateway — an intermediary service used to route, control, monitor or manage requests to one or more AI models.

Non-consensual intimate imagery — AI-generated or AI-manipulated visual content depicting real or realistic individuals in intimate or sexual contexts without their consent; subject to a new prohibition under Article 5 of the EU AI Act as amended.

Non-human identity — a digital identity assigned to a service, application, workload, machine or AI agent rather than an individual person.

N-day / N-hour — a disclosed-and-patched vulnerability; “N-hour” denotes the compression of the time from disclosure to a working exploit from weeks to hours.

Companion publication: Q2 2026 AI Security Technical & Compliance Report (ODA3-2026-07-TCR-SEC-004).

© 2026 ODA3 Pvt Ltd. All rights reserved.