

EA-005: AI Risk & Controls Assessment™

Can Your Enterprise Identify, Control, and Monitor AI Risk Before It Becomes Enterprise Exposure?

Purpose

A practical executive assessment for evaluating whether AI risks are systematically identified, prioritized, controlled, tested, monitored, and improved across the full lifecycle.

Educational Executive Self-Assessment | Version 1.0

www.intelligentsystemspress.com

About This Assessment

AI risk management is the enterprise capability to identify uncertainty and potential harm, establish proportionate safeguards, verify that controls work, and respond before exposure becomes systemic. This assessment helps leaders determine whether AI risk is actively controlled or merely acknowledged.

How to Complete It

1. **Rate each statement** Use the 1-5 scale based on current, observable enterprise practice - not future intent, isolated pilots, or vendor assurances.
2. **Score each domain** Add the five ratings in each domain. Every domain score ranges from 5 to 25.
3. **Calculate the total** Add all eight domain scores. The overall risk-and-controls score ranges from 40 to 200.
4. **Interpret the pattern** Review the total and the distribution. Any domain below 15 represents a material control dependency or unmanaged exposure.
5. **Select 90-day actions** Choose a limited number of executive actions that reduce material exposure, strengthen controls, or improve assurance.

Rating Scale

1 Strongly Disagree	Practice is absent or contradicted by current evidence.
2 Disagree	Practice is inconsistent, informal, or limited to isolated areas.
3 Partly Agree	Practice exists in several areas but is not yet enterprise-wide.
4 Agree	Practice is generally reliable, measured, and consistently applied.
5 Strongly Agree	Practice is integrated, evidence-driven, and continuously improved.

Recommended use: Complete individually first, then compare results as an executive team. Differences in ratings are evidence of risk-perception or control-assurance misalignment and should be discussed before averaging scores.

Domain 1: Risk Context & Ownership

Rate each statement from 1 (Strongly Disagree) to 5 (Strongly Agree). Base the rating on repeatable enterprise behavior and current evidence.

Assessment statement	1	2	3	4	5
Each AI-enabled process, product, or decision has a named business owner accountable for its outcomes, risks, controls, and continued use.	☐	☐	☐	☐	☐
Risk assessments consider the business objective, operating environment, affected stakeholders, data sensitivity, degree of automation, and potential consequence.	☐	☐	☐	☐	☐
AI risk ownership is coordinated across business, technology, data, legal, privacy, security, quality, compliance, and operational functions.	☐	☐	☐	☐	☐
Leaders distinguish risks created by the AI capability from risks already present in the underlying process, policy, data, or decision model.	☐	☐	☐	☐	☐
Risk acceptance, escalation, transfer, treatment, and retirement decisions are documented at the appropriate level of authority.	☐	☐	☐	☐	☐

Domain score: ____ / 25

Evidence or observations:

Domain 2: Use-Case Risk Identification

Rate each statement from 1 (Strongly Disagree) to 5 (Strongly Agree). Base the rating on repeatable enterprise behavior and current evidence.

Assessment statement	1	2	3	4	5
AI use cases are assessed before experimentation, procurement, integration, deployment, and material expansion.	☐	☐	☐	☐	☐
Risk identification covers inaccurate outputs, bias, privacy loss, security compromise, intellectual-property exposure, regulatory noncompliance, and operational disruption.	☐	☐	☐	☐	☐
Teams evaluate foreseeable misuse, overreliance, automation bias, manipulation, unintended users, and use outside the approved context.	☐	☐	☐	☐	☐
Risk assessments identify downstream dependencies, human handoffs, third-party services, and business decisions that could amplify an AI failure.	☐	☐	☐	☐	☐
Assumptions, uncertainties, limitations, and evidence gaps are recorded rather than treated as acceptable unknowns.	☐	☐	☐	☐	☐

Domain score: ____ / 25

Evidence or observations:

Domain 3: Data & Model Controls

Rate each statement from 1 (Strongly Disagree) to 5 (Strongly Agree). Base the rating on repeatable enterprise behavior and current evidence.

Assessment statement	1	2	3	4	5
Data used by AI systems is subject to documented quality, provenance, permission, access, retention, and protection requirements.	☐	☐	☐	☐	☐
Controls prevent sensitive, confidential, regulated, copyrighted, or unauthorized information from entering prompts, training sets, retrieval sources, or outputs.	☐	☐	☐	☐	☐
Model, prompt, retrieval, configuration, and parameter changes are versioned, reviewed, tested, and traceable to approved requirements.	☐	☐	☐	☐	☐
The enterprise evaluates whether data and model performance remain appropriate for the population, context, and decisions they are intended to support.	☐	☐	☐	☐	☐
Known model limitations, confidence boundaries, unsupported uses, and failure conditions are translated into enforceable operational controls.	☐	☐	☐	☐	☐

Domain score: ____ / 25

Evidence or observations:

Domain 4: Human Oversight & Decision Safeguards

Rate each statement from 1 (Strongly Disagree) to 5 (Strongly Agree). Base the rating on repeatable enterprise behavior and current evidence.

Assessment statement	1	2	3	4	5
Human review requirements are proportionate to the impact, reversibility, complexity, and uncertainty of the AI-assisted decision.	☐	☐	☐	☐	☐
People responsible for oversight have sufficient knowledge, context, time, authority, and tools to question or override AI outputs.	☐	☐	☐	☐	☐
High-impact decisions include clear escalation, appeal, correction, and exception-handling pathways.	☐	☐	☐	☐	☐
Interfaces and procedures reduce automation bias by presenting limitations, uncertainty, source information, and alternative actions where appropriate.	☐	☐	☐	☐	☐
Critical functions have documented fallback procedures that can operate safely when AI is unavailable, degraded, or suspended.	☐	☐	☐	☐	☐

Domain score: ____ / 25

Evidence or observations:

Domain 5: Security, Abuse & Resilience Controls

Rate each statement from 1 (Strongly Disagree) to 5 (Strongly Agree). Base the rating on repeatable enterprise behavior and current evidence.

Assessment statement	1	2	3	4	5
Threat modeling addresses prompt injection, data exfiltration, adversarial inputs, malicious content, insecure plugins, excessive agency, and unauthorized access.	☐	☐	☐	☐	☐
Identity, access, segregation of duties, logging, secrets management, and environment controls are applied to AI components and integrations.	☐	☐	☐	☐	☐
Controls limit the actions AI systems may initiate, the systems they may access, and the value or consequence of transactions they may execute.	☐	☐	☐	☐	☐
Resilience testing considers vendor outages, dependency failures, corrupted data, model degradation, capacity constraints, and loss of connectivity.	☐	☐	☐	☐	☐
Security events, abuse patterns, vulnerabilities, and emerging threats trigger timely containment, remediation, and control updates.	☐	☐	☐	☐	☐

Domain score: ____ / 25

Evidence or observations:

Domain 6: Validation & Control Testing

Rate each statement from 1 (Strongly Disagree) to 5 (Strongly Agree). Base the rating on repeatable enterprise behavior and current evidence.

Assessment statement	1	2	3	4	5
AI solutions are tested against defined business, functional, quality, safety, legal, security, privacy, and control acceptance criteria before release.	☐	☐	☐	☐	☐
Testing includes representative normal cases, edge cases, adversarial conditions, failure scenarios, and populations most affected by the outcome.	☐	☐	☐	☐	☐
Controls are tested for design adequacy and operating effectiveness rather than assumed effective because they are documented.	☐	☐	☐	☐	☐
Independent review or challenge is used when consequence, uncertainty, conflict of interest, or regulatory exposure is significant.	☐	☐	☐	☐	☐
Material changes to data, models, prompts, integrations, policies, vendors, users, or operating context trigger proportionate reassessment and retesting.	☐	☐	☐	☐	☐

Domain score: ____ / 25

Evidence or observations:

Domain 7: Monitoring, Thresholds & Incident Response

Rate each statement from 1 (Strongly Disagree) to 5 (Strongly Agree). Base the rating on repeatable enterprise behavior and current evidence.

Assessment statement	1	2	3	4	5
Monitoring covers outcome accuracy, error patterns, drift, bias indicators, security events, control performance, user behavior, and operational reliability.	☐	☐	☐	☐	☐
Thresholds are defined for investigation, human intervention, restricted operation, rollback, suspension, notification, and retirement.	☐	☐	☐	☐	☐
AI incidents, near misses, complaints, overrides, and control exceptions are captured in a consistent enterprise process.	☐	☐	☐	☐	☐
Incident response assigns decision authority, communication responsibilities, evidence preservation, stakeholder notification, and recovery actions.	☐	☐	☐	☐	☐
Root-cause analysis addresses weaknesses in the surrounding process, data, governance, requirements, and controls - not only the model output.	☐	☐	☐	☐	☐

Domain score: ____ / 25

Evidence or observations:

Domain 8: Assurance, Reporting & Continuous Improvement

Rate each statement from 1 (Strongly Disagree) to 5 (Strongly Agree). Base the rating on repeatable enterprise behavior and current evidence.

Assessment statement	1	2	3	4	5
Executives receive risk reporting that connects AI exposure, control effectiveness, incidents, exceptions, and realized business value.	☐	☐	☐	☐	☐
Risk and control evidence is sufficiently traceable for internal review, audit, regulatory inquiry, customer assurance, and executive decision-making.	☐	☐	☐	☐	☐
Control gaps are prioritized according to potential consequence, likelihood, detectability, exposure duration, and dependency on other safeguards.	☐	☐	☐	☐	☐
Lessons from testing, monitoring, incidents, stakeholder feedback, audits, and external developments are converted into corrective and preventive actions.	☐	☐	☐	☐	☐
The AI risk and control framework is periodically reassessed as technologies, laws, threats, business models, and organizational risk tolerance change.	☐	☐	☐	☐	☐

Domain score: ____ / 25

Evidence or observations:

Scoring & Risk-Control Interpretation

Record each domain score. Treat any domain below 15 as a priority dependency even when the overall score is higher.

Domain	Score
1. Risk Context & Ownership	___ / 25
2. Use-Case Risk Identification	___ / 25
3. Data & Model Controls	___ / 25
4. Human Oversight & Decision Safeguards	___ / 25
5. Security, Abuse & Resilience Controls	___ / 25
6. Validation & Control Testing	___ / 25
7. Monitoring, Thresholds & Incident Response	___ / 25
8. Assurance, Reporting & Continuous Improvement	___ / 25
Total AI risk-and-controls score	___ / 200

Overall Risk-Control Levels

180-200	Controlled & Assured	Risk ownership, preventive and detective controls, monitoring, and assurance are integrated across the AI lifecycle.
150-179	Well Controlled	Strong control foundations exist, with targeted gaps that should be resolved before broader scaling.
110-149	Developing Controls	Several controls are established, but coverage or effectiveness remains uneven and dependent on local practice.
70-109	High Exposure	Risk treatment is fragmented or reactive, and material safeguards may not operate consistently.
40-69	High Exposure	Material gaps in risk ownership, safeguards, monitoring, and response create significant enterprise exposure.

Important: This assessment is a discussion and prioritization tool, not a certification. Control effectiveness should be validated using evidence, stakeholder input, legal and regulatory analysis, technical assurance, and use-case-specific risk review.

Executive Reflection & 90-Day Action Plan

Use the risk-and-control pattern to identify the few leadership actions that will most reduce exposure and strengthen sustainable AI value.

Control strengths

Which safeguards reliably prevent, detect, contain, or correct material AI risk?

Critical exposures

Which weaknesses could produce the greatest legal, operational, ethical, security, financial, or reputational consequence?

Decisions required from leadership

What ownership, authority, risk-tolerance, control, funding, or suspension decisions cannot be delegated?

90-Day Priority Actions

Priority action	Executive owner	Success measure	Target date

Executive commitment: We will review progress by _____ and reassess affected domains using current evidence.

Continue Your AI Risk Management Journey

AI risk controls are strongest when strategy, Process Intelligence, requirements, data, security, quality, human oversight, monitoring, and continuous improvement operate as one connected system.

Explore the Intelligent Systems Press collection and Free Resources library
www.intelligentsystemspress.com

Educational Use Disclaimer: This assessment is provided for educational and internal planning purposes. It does not constitute legal, regulatory, cybersecurity, privacy, financial, certification, audit, or professional consulting advice. Organizations should apply qualified expertise based on their industry, jurisdiction, data, systems, risk profile, and intended AI use. © 2026 Intelligent Systems Press LLC. All rights reserved.