

LIBRA SENTINEL

Global Data Privacy & AI Governance



THIS IS NOT LEGAL ADVICE. THIS IS A FREE GOVERNANCE AID.

IS YOUR AI STACK REALLY INDEPENDENT?

A 12-question common-mode AI risk check

Many organizations diversify at vendor level while remaining dependent on the same models, infrastructure, data or failure patterns. Use this checklist to identify where apparent independence may conceal common-mode exposure.

MARK EACH QUESTION: ☐ YES ☐ NO ☐ UNKNOWN. An unknown answer may itself indicate a dependency-visibility gap.

1. SEE THE HIDDEN STACK

Evidence: [1], [3], [4], [8]

1	Have we identified and recorded the underlying foundation model or model family powering each AI system? WHY IT MATTERS: Vendor names alone do not reveal common model dependencies.	☐ YES ☐ NO ☐ UNKNOWN
2	Do apparently different vendors share a model, cloud, API, embedding model, data source or other critical component? WHY IT MATTERS: Commercial diversity can conceal common-mode exposure.	☐ YES ☐ NO ☐ UNKNOWN
3	Can we trace each AI application through its direct vendor and intermediary services to the underlying models, data sources and infrastructure dependencies? WHY IT MATTERS: A conventional AI inventory may identify individual tools without revealing shared upstream dependencies through which one failure, vulnerability or provider change could affect several applications.	☐ YES ☐ NO ☐ UNKNOWN

2. TEST MEANINGFUL INDEPENDENCE

Evidence: [6], [7], [8]

4	Have we tested whether primary systems, backups or automated reviewers fail on the same cases? WHY IT MATTERS: Two individually accurate systems can still have correlated errors.	☐ YES ☐ NO ☐ UNKNOWN
5	Does a claimed "second opinion" rely on independent evidence, data and method? WHY IT MATTERS: Agreement is weak assurance when both systems inherit the same sources or blind spots.	☐ YES ☐ NO ☐ UNKNOWN
6	Are different systems optimized against the same benchmarks, objectives or human-preference signals? WHY IT MATTERS: Different providers may still converge on similar behaviour and residual errors.	☐ YES ☐ NO ☐ UNKNOWN

3. BUILD SUBSTITUTABILITY AND RESILIENCE		Evidence: [1], [2], [3]
7	Can we switch models without rebuilding the application? WHY IT MATTERS: A technical inability to switch can convert concentration into operational lock-in.	☐ YES ☐ NO ☐ UNKNOWN
8	Does tested failover use genuinely different upstream dependencies? WHY IT MATTERS: A backup that shares the primary system's critical dependency may fail at the same time.	☐ YES ☐ NO ☐ UNKNOWN
9	Do incident-response plans identify every downstream application dependent on an affected model or component? WHY IT MATTERS: A model-level incident may require coordinated action across several systems.	☐ YES ☐ NO ☐ UNKNOWN

4. GOVERN CHANGE AND OVERSIGHT		Evidence: [1], [3], [4], [5]
10	Are upstream model changes monitored, impact-assessed and revalidated? WHY IT MATTERS: A provider update can alter downstream behaviour even when the application itself has not changed.	☐ YES ☐ NO ☐ UNKNOWN
11	Is human review genuinely independent of the algorithmic recommendation? WHY IT MATTERS: Human involvement adds little protection if reviewers simply defer to the system.	☐ YES ☐ NO ☐ UNKNOWN
12	Do procurement and ongoing monitoring assess concentration and substitutability at model and infrastructure level? WHY IT MATTERS: Vendor diversification can become illusory as providers, cloud dependencies and model families change over time.	☐ YES ☐ NO ☐ UNKNOWN

START HERE. If several answers are unknown, begin with dependency mapping. If the dependencies are known but alternatives still fail together, the next question is behavioural correlation.

EVIDENCE NOTE. Questions 1-3 and 7-12 synthesize supply-chain, third-party, monitoring and systemic-risk guidance. Questions 4-6 draw on correlated-error and outcome-homogenization research. This is an analytical synthesis, not a verbatim requirement of any cited framework or regulator.

SOURCE BASE

[1] NIST, [AI Risk Management Framework Core](#) (2023).

[2] UK NCSC et al., [Guidelines for Secure AI System Development: Secure Development](#) (2023).

[3] Financial Stability Board, [Monitoring Adoption of AI and Related Vulnerabilities](#) (2025).

[4] European Central Bank, [The Rise of AI: Benefits and Risks for Financial Stability](#) (2024).

[5] European Systemic Risk Board, [Artificial Intelligence and Systemic Risk](#) (2025).

[6] Kim, Garg, Peng and Garg, [Correlated Errors in Large Language Models](#) (ICML, 2025).

[7] Bommasani et al., [Picking on the Same Person: Does Algorithmic Monoculture Lead to Outcome Homogenization?](#) (NeurIPS, 2022).

[8] [International AI Safety Report 2026](#).