

# AIR Procurement Clause & Vendor Questionnaire

*A buyer-side toolkit for requiring AI vendors to meet a credible trust bar — published by the AIR Standards Council under the AIR Vendor Standard (AIR-VS).*

## How and why to use this

When you put an AI tool anywhere near your real work or real data, you are extending your organization's trust boundary to a third party whose model behavior, data practices, and provenance you cannot see. SOC 2 became table stakes for SaaS not because a regulator mandated it, but because **buyers wrote it into their RFPs** and vendors who couldn't produce a report lost deals. AI trust will harden the same way — buyer demand, not vendor goodwill, sets the floor.

This document gives any organization — in any industry — three reusable instruments to set that floor:

1. **A procurement clause** you paste directly into RFPs, MSAs, and vendor onboarding checklists.
2. **A vendor questionnaire** you send to any AI vendor during evaluation or a security review.
3. **A scoring guide and red-flag list** so a non-specialist reviewer can read the answers and reach a defensible decision.

### Where to drop each instrument:

- **RFPs / RFIs** — include the procurement clause as a mandatory requirement; attach the questionnaire as a required response exhibit.
- **Vendor onboarding** — make a passing questionnaire (or current **AIR Certified** status) a gate before a tool touches production data.
- **Security / vendor-risk reviews** — use the questionnaire to standardize how your security team interrogates AI-specific risk that generic SOC 2 / ISO 27001 reviews miss.
- **Contract renewals** — re-run the questionnaire; trust is not "set once."

**The shortcut.** A vendor holding current **AIR Certified** status has already been independently assessed against every dimension below by the AIR Standards Council. You can accept that status in lieu of re-running the full questionnaire — verify it yourself on the public AIR registry (status is published independently of any advisory relationship; never take the vendor's word for it). Absence of certification is **not** a disqualifier — it simply means you do the diligence yourself using the questionnaire.

The questionnaire maps one-to-one to the AIR-VS dimensions. A vendor's answers here are a self-assessment; AIR Certified is an *independent* assessment of the same criteria. Treat them accordingly.

## The Procurement Clause

*Paste-ready. Bracketed text is for you to fill in. Works in an RFP requirement block, an MSA schedule, or a vendor onboarding policy.*

**AI Trust & Vendor Standard Requirement.** Any software, service, or feature that applies artificial intelligence, machine learning, or large language models to **[Buyer]**'s data, content, or workflows ("AI Tool") must meet the trust criteria defined by the AIR Vendor Standard (AIR-VS). Vendor shall, at **[Buyer]**'s option, satisfy this requirement by either (a) holding current **AIR Certified** status, verifiable on the public AIR registry, or (b) completing the AIR Vendor Questionnaire and providing supporting evidence demonstrating that the AI Tool meets the AIR-VS criteria across all dimensions below. Vendor represents that its responses are accurate as of the submission date and shall notify **[Buyer]** in writing of any material change (including any change in model providers, sub-processors, training-data practices, or certification status) within **[30]** days. **[Buyer]** may re-verify compliance at each renewal and may treat a failure to maintain these criteria as a material breach.

**Compliance checklist** (*the AI Tool must satisfy every item, mapped to its AIR-VS dimension*):

- ☐ **Data handling** — Buyer data is **not** used to train, fine-tune, or improve any model serving other customers without explicit, revocable, opt-in consent; a signed DPA governs the relationship.
- ☐ **IP & provenance** — Vendor discloses model providers and the licensing/provenance basis of training data sufficient to assess IP and output-ownership risk; output ownership is contractually assigned to Buyer.
- ☐ **Security** — AI-specific controls exist (prompt-injection, data-leakage, and model-abuse mitigations) on top of a recognized baseline (SOC 2 Type II, ISO 27001, or equivalent).
- ☐ **Transparency** — Vendor discloses which models/providers process Buyer data and notifies Buyer of material model or sub-processor changes.
- ☐ **Integration & export** — Buyer can extract its data and AI-generated outputs in a standard, machine-readable format at any time.
- ☐ **Portability** — No architectural or contractual lock-in prevents Buyer from migrating off the tool and removing its data on exit.
- ☐ **Claims substantiation** — Vendor can substantiate material performance/accuracy claims with reproducible evidence or methodology.
- ☐ **Human oversight** — For consequential outputs, mechanisms exist for human review, override, and escalation.
- ☐ **Reliability & support** — Defined uptime, incident-notification, and support commitments cover AI-specific failure modes (degradation, hallucination, outage).

## The Vendor Questionnaire

*Send to any AI vendor. Ask for evidence — documents, contract language, configuration screenshots, audit reports — not just yes/no. Questions are grouped by AIR-VS dimension.*

### 1. Data handling

1. Do you use our inputs, outputs, or usage data to train, fine-tune, or improve any model? If so, is this **on by default or off by default**, and how do we change it? *(Provide the relevant contract clause or admin setting.)*
2. Will you sign our Data Processing Agreement (DPA), or provide yours? List all sub-processors and the countries in which our data is stored and processed.
3. What is the retention period for our prompts, outputs, and logs, and what is the deletion process and SLA on a deletion request?
4. Is our data logically or physically isolated from other customers' data? Describe the tenancy model.
5. If you use third-party model providers (e.g. a foundation-model API), what contractual guarantees prevent **those providers** from training on or retaining our data?

### 2. IP & provenance

1. Who owns the outputs the tool generates from our inputs? Point to the exact contract language.
2. What is the provenance and licensing basis of the data used to train the models we will rely on? Can you indemnify us against third-party IP claims arising from outputs?
3. Do you provide any output provenance, watermarking, or audit trail indicating that content was AI-generated?

### 3. Security

1. Provide your most recent SOC 2 Type II or ISO 27001 report (or equivalent). What is its scope, and does it cover the AI components specifically?
2. What controls mitigate **AI-specific** threats — prompt injection, training-data poisoning, model inversion, and sensitive-data leakage through outputs?
3. Describe your encryption (in transit and at rest), access controls, and how staff/contractor access to our data is governed and logged.
4. Have you had a security incident involving customer data or model behavior in the last 24 months? Describe it and your notification SLA.

### 4. Transparency

1. Exactly which models and providers process our data today, including version? How and when are we notified of a model, provider, or sub-processor change?
2. What are the known limitations, failure modes, and error rates of the system for our intended use? In what conditions does it perform worst?

3. Can you explain, at a level our team can act on, how the system reaches a given output for consequential decisions?

## 5. Integration & export

1. In what formats can we export our data and the AI-generated outputs, and is export available via self-serve API/UI at any time without your involvement?
2. How does the tool integrate with our existing stack (SSO, audit-log streaming, identity, DLP), and what data does it require from those systems?

## 6. Portability

1. If we terminate, what is the offboarding process — timeline, format, and proof of deletion — for retrieving all our data and outputs and removing them from your systems?
2. Are there proprietary formats, embeddings, or fine-tuned artifacts derived from our data that we cannot take with us or that you retain after exit?

## 7. Claims substantiation

1. For every performance, accuracy, or productivity claim you make (in marketing or this response), provide the methodology, dataset, and conditions behind the number. Can we reproduce or independently validate it?
2. Are your benchmarks run on your own data or on independent/industry datasets? Disclose any cherry-picking or favorable conditions.

## 8. Human oversight

1. For consequential or customer-facing outputs, what mechanisms exist for human review, override, and escalation before an output takes effect?
2. How do you detect, surface, and let us correct hallucinated, biased, or low-confidence outputs?

## 9. Reliability & support

1. What are your uptime SLA, historical uptime, and incident-notification commitments? How do they treat **AI-specific** degradation (quality regressions, latency spikes, model outages), not just hard downtime?
2. What support tier, response times, and escalation path apply when the AI behaves incorrectly rather than simply being "down"?

---

## How to score the answers

Score each of the nine dimensions, not just the total — a tool can be excellent on security and disqualifying on data handling.

### Per-question rating:

- **2 — Meets.** Concrete answer *plus* evidence (contract clause, report, config, reproducible methodology).
- **1 — Partial.** Right direction but vague, evidence missing, or commitment is "roadmap" not "today."
- **0 — Fails / evasive.** Non-answer, refusal, or an answer that confirms a red flag below.

### Reading the score:

- **Any dimension averaging 0, or any red flag below, is a gate** — resolve it before proceeding regardless of the total. These are not outweighed by strength elsewhere.
- Treat **"we're working on it"** as a **1**, never a **2** — score what exists today, not the roadmap.
- **Demand evidence.** An unevidenced "yes" is a **1**, not a **2**. The whole point is to reward vendors who can *show* it.
- A vendor with **current AIR Certified status** has cleared this bar under independent assessment — verify the status on the public registry and you can shorten your own review accordingly.

### Red flags (any one warrants a hard stop and escalation)

- **Trains on your inputs by default**, or can't clearly tell you whether it does.
- **Won't sign or provide a DPA**, or won't name its sub-processors and data locations.
- **Can't export your data** / outputs, or only via a manual, vendor-gated, or paid process.
- **Won't substantiate performance claims** — no methodology, no dataset, "trust us."
- **Can't say which models/providers** process your data, or won't commit to notifying you when they change.
- **No AI-specific security controls** — points only to generic SOC 2 and waves off prompt injection / data leakage.
- **Retains data or derived artifacts after termination**, or can't prove deletion.
- **Third-party model providers can train on your data**, with no contractual block.
- **No human-oversight path** for consequential outputs — the model's word is final.
- **Output ownership is ambiguous** or retained by the vendor.

*This toolkit is published under the AIR Vendor Standard (AIR-VS) and maintained by the independent AIR Standards Council. The criteria above are the published standard; **AIR Certified** is an independent assessment against them. Buyers may use, adapt, and redistribute these instruments freely. Illustrative figures (e.g. notification windows, lookback periods) are bracketed or marked and should be set to your own policy.*