

AI • TEMPLATE

What Enterprise Buyers *Actually* Ask About Your AI Features

KAANSYSTEMS.COM/LIBRARY/WHAT-ENTERPRISE-BUYERS-ASK-AI • JULY 29, 2026

ABOUT THIS TEMPLATE

Security questionnaires now have an AI section. The questions cluster into five themes, and none of them are testing whether your model is perfect. The answer bank that unblocks the deal.

THE TEMPLATE

Fifteen of the most common AI-security questions with answer patterns. Adapt the specifics to your architecture; keep the structure. Vague answers get flagged; brief, specific, honest answers get you to the next stage.

#	QUESTION	ANSWER PATTERN
1	Do you train AI models on customer data?	No. Customer data is never used to train or fine-tune models.
2	Which AI/LLM provider(s) do you use?	[Vendor(s)], accessed via API under [BAA/DPA]. Listed on our sub-processor page.
3	Does your AI provider train on your inputs?	No — confirmed in writing under our enterprise agreement; opt-out verified.
4	Are prompts and responses logged? Where?	Shape only (no content) in app logs. Any retained content is in-boundary, encrypted, scoped.
5	What is the retention for AI prompt/response data?	[X days]. Set deliberately; not left at vendor default.
6	Is customer data isolated at the model layer?	Yes. Prompts are constructed per-request, scoped to one tenant. No cross-tenant context.

#	QUESTION	ANSWER PATTERN
7	Can the AI take autonomous/irreversible actions?	No. The model drafts or suggests; consequential actions require human review.
8	How do you handle hallucinations / incorrect output?	Human-in-the-loop on consequential outputs; low-confidence outputs flagged + logged.
9	Is there a human-in-the-loop for [use case]?	[Yes — where + how] OR [Compensating: output is advisory only, not acted on automatically].
10	Do you have an AI acceptable-use policy?	[Yes — summary] OR [Documented AI feature pre-flight applied before each launch].
11	How are AI features reviewed before release?	Documented compliance pre-flight: data flow, vendor agreement, output safety, logging.
12	Can we opt out of AI features / processing?	[Yes — how] OR [AI features are [optional/tiered]; describe the control].
13	Is generated content labeled as AI-generated?	Yes. Users can distinguish AI-generated content.
14	What is your incident process for a harmful AI output?	Handled under standard IR; severity assessed; affected customers notified per policy.
15	Do you use our data in a shared vector store / RAG index?	[No — per-tenant isolation] OR [Yes — describe tenant scoping + access controls honestly].

These fifteen cover the large majority of any AI section. The rest are environment-specific — your exact model, your exact data flows — and those become the focused work once the baseline is in place.

— HOW TO USE IT

Build the AI answer bank the same way you built the security one: every answer you write goes into a versioned document, keyed by question, so the next questionnaire is a copy-paste plus a delta review. The AI section changes faster than the rest, so tag each answer with a "last verified" date and re-check anything older than three months.

The "we don't train on your data" answer will be re-asked, re-verified, and eventually audited. Make sure it stays true. The single fastest way to destroy the trust this whole exercise builds is to answer "no training" while a vendor default quietly flipped to "training unless you opt out." Re-read your model vendors' terms on a schedule, not when a buyer asks.

The teams that handle the AI section well aren't the ones with the most sophisticated AI. They're the ones who can describe their unremarkable, well-governed AI in five honest sentences. That legibility is the product.