

How We Use AI — And What That Means for Your Data

A plain answer for the people who have to sign off on it

Every AI vendor says their product is secure and private. Almost none will tell you which model provider processes your data, in which country, under what contract, and what happens to it afterwards.

Those are the questions your security and privacy teams will ask. This document answers them for AidEun and SalesQuote.

What we actually use AI for

Our products use AI for a small number of specific jobs, not as a general-purpose oracle:

Answering questions from your own content. AidEun retrieves the relevant documents, quotations, historical parts and knowledge-base entries from your own data, and asks a language model to answer from that material. The model is not asked what it knows — it is asked what your documents say.

Reading and structuring documents. SalesQuote extracts line items, quantities and specifications from RFQ documents and emails so a human does not have to retype them.

Drafting. Quotations, replies, summaries — always as a draft for a person to review, never as a final output sent on your behalf without a human in the loop.

Search and retrieval. Embeddings make your own content findable by meaning rather than exact keyword.

What we do not do:

We do not use AI to make decisions about people, we do not profile your customers, and we do not use your data to build anything we sell to anyone else.

Where the AI actually runs

This is the question most vendors answer vaguely. Ours is specific.

Inference runs on Azure OpenAI in Microsoft's Norway East region. Not a shared public API endpoint — a dedicated Azure OpenAI resource inside our own Azure tenancy, in Norway, governed by the Microsoft Data Processing Agreement and covered by Microsoft's EU Data Boundary commitments.

That means your prompts, your documents and the model's responses stay inside the European Economic Area. There is no transfer to a US provider in the normal operation of the product.

Your data is stored in the same place. Databases, search indexes, file storage, application services — all Microsoft Azure, Norway East. Backups stay within the region.

No Chinese-origin models. No DeepSeek, no Qwen, no Kimi — not in production, not in evaluation.

Model choice is deliberate, not accidental. Our platform is model-agnostic by design: an adapter layer means we choose the provider rather than locking your data to one vendor's jurisdiction and one vendor's roadmap. Mistral, hosted in France, is available for tenants who want a specific EEA alternative.

“Will our data be used to train someone’s model?”

No — and the reason matters more than the answer.

The distinction that catches most buyers out is not what a vendor promises but which tier of which provider's service the vendor is calling. On some providers' free tiers, submitted content is used to improve their products. On the paid enterprise tiers, it is not. Same model, same prompt, entirely different contract — and customers usually cannot tell which one they are on.

Running through Azure OpenAI settles this structurally rather than by promise. Microsoft's terms for Azure OpenAI are explicit: your prompts and completions are not available to other customers, not available to OpenAI, and not used to train or improve any Microsoft or third-party model.

Our own Data Processing Addendum states the same commitment to you directly, so it is a contractual obligation we hold, not only a term we have inherited.

“Can it hallucinate? Can it be tricked?”

Both are real risks. Any vendor claiming immunity is overselling.

On hallucination. Our assistant answers from retrieved content — your documents, supplied to the model as context. This is the main structural defence: the model is asked to answer from material in front of it rather than from memory. It makes answers specific to your business, and it makes them checkable, because the source is right there.

It is not a complete defence. A grounded model can still misread a document or state something a source does not fully support. So we design for review rather than autonomy: AI output is a draft a person accepts, edits or rejects. For anything with commercial or contractual weight — a quotation, a customer reply — a human is in the loop by design, not by policy.

On prompt injection and jailbreaking. This is not a solved problem anywhere in the industry, and we will not claim otherwise. The useful question is not "can it be tricked" but "what can it reach if it is". Our answer is architectural: strict tenant isolation, so one customer's data is never reachable from another's session; scoped permissions, so the assistant can only touch what the signed-in user could touch anyway; and human approval before anything leaves the system.

We continue to harden the tool-execution path — input sanitisation, tool allow-lists, output checking — as an ongoing programme rather than a finished feature.

The unglamorous part

Most of what makes AI safe in production is not AI-specific. It is knowing what you hold, how long you hold it, who can reach it, and being able to show your working.

We know how long we keep everything. Our data retention schedule covers every entity in the platform that holds personal data — conversations, attachments, quotations, logs, embeddings, caches — with a period, a lawful basis and a deletion method for each. It is approved and signed by our Privacy Lead, our management and our Security Lead. It is not a policy statement; it is a table with a row for every entity, naming what must be deleted alongside each record, down to derived embeddings and cached copies.

You can configure retention. Conversations from 30 days to 24 months, RFQ records at 12, 24 or 36 months. There is a floor because too short a period breaks the product and loses records you may still need, and a ceiling because storage limitation is a legal principle, not a preference. Anything outside those ranges is a documented decision, not a support ticket.

Deletion means deletion. When a record reaches the end of its retention period, or when you ask us to erase a data subject's data, the schedule names what goes with it: the database row, the blob, its versions and snapshots, extracted content, embeddings, previews and cached copies. An embedding never outlives its source document.

Our security management system is documented. We maintain an ISO/IEC 27001:2022 information security management system: scope, policy, risk assessment, a Statement of Applicability across all 93 Annex A controls, and the management-system clauses covering support, operation, performance evaluation and improvement. Internal audit and management review are scheduled and owned.

We are not certified. We are building towards it, we know which controls are open, and we would rather show you a considered gap list than a certificate we do not hold.

Why we are telling you the incomplete parts

Because you will find out anyway, and because a supplier who only lists their strengths has told you nothing you can act on.

Every claim on this page is backed by a document with a version, an owner and a signature — or it is described as work in progress. Our findings are tracked as work items with dates and named owners. Our sub-processors, their locations and their transfer mechanisms are published at aisolutions.no/suppliers. Our Data Processing Addendum is public at aisolutions.no/dpa.

If your security or privacy team wants to go deeper — the retention schedule, the control gap list, our transfer assessments — ask. We would far rather have that conversation during procurement than after an incident.

Questions: thore@aisolutions.no

AI Solutions AS · Org. nr 931 884 220 · Spannavegen 152, 5535 Haugesund, Norway