

Responsible AI Statement

Which AI we use, the guardrails around it, and the commitment behind "AI assists, humans author".

Version 1.4 | Effective 28 July 2026 | AI & standards

Owner Obedience Global Ltd board | Approved by Obedience Global Ltd board on 28 July 2026

Approval state company-approved | Next review 28 October 2026

01 The commitment

AI assists, humans author. Our AI drafts, suggests and searches. It never signs, never overrides a professional, and never issues anything client-facing on its own. Every AI output in our products shows its source and requires a named person's review and acceptance.

02 The models we use

The current public assistants use Anthropic through its commercial API for low-sensitivity product and policy questions. Quantify contains optional Anthropic and OpenAI integration code, but OpenAI private processing and confidential hosted-model processing are not approved current paths. A credential, dependency or feature flag is not approval.

By default, commercial API providers do not use commercial API inputs or outputs to train their models. Standard API content or abuse-monitoring retention can be up to 30 days. Provider legal and safety exceptions may retain limited flagged records for longer, and endpoint-specific application state can have a different lifecycle. Current provider, approval and endpoint controls are recorded in our Data Retention Schedule and Service Provider and Subprocessor List.

03 Guardrails in the product

- Human review is a workflow step, not an option: unreviewed output cannot be issued as final
- AI suggestions cite the source they were derived from (your data, standards structures, your rate history)
- AI features degrade safely: if a response is truncated, refused or malformed, the product shows nothing rather than something wrong
- Capabilities are labelled honestly: Live, In development, or Roadmap, on the product pages and in the app

04 Limitations, stated plainly

AI can be wrong, incomplete, or plausible-but-mistaken. That is exactly why the human sign-off step exists and why the Professional Use & Reliance Statement is part of our terms. Treat every AI output as a draft for professional judgement.

05 Contestability and change

If an AI feature behaves in a way you think is wrong or unfair, tell us: hello@obedience.global. We log, investigate and answer. When we change the models or providers behind a feature, we record it in the policy changelog. This statement will be updated as the ICO's statutory Code of Practice on AI and automated decision-making is finalised.

Canonical online record: <https://obedience.global/policies/responsible-ai>