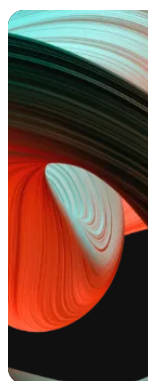


Trusted AI for defensible decisions

The pressure to adopt AI arrived before the strategy did. Organizations default to the biggest, most expensive model, deployed as fast as possible, with accountability figured out later. Nobody asked which model is best for the job. And when the AI driving those jobs makes a wrong call, someone has to answer for it. That gap between moving fast and being in control is where budgets break, decisions become indefensible, and the AI that was supposed to create advantage becomes a liability.

The Seekr Trusted AI Platform

The answer isn't a bigger model. It's the right one. Seekr lets you build domain-specific models on your own data, match the right model to each task, and make every output explainable and defensible. Seekr runs wherever your data lives, scales without runaway cost, and gives you the evidence to stand behind every decision AI drives.



- Evaluate, compare, and monitor models against your data, risk criteria, and compliance requirements.
- Stress-test models with documented evidence of performance you can show an examiner.
- Deploy autonomous agents with human-in-control escalation for high-impact operational and risk events.
- Make raw, domain-specific data AI-ready with governed pipelines, lineage tracking, and policy enforcement.
- Transform large document libraries into agentic workflows for analysis, assessments, and content generation.
- Automate high-volume tasks with full auditability and a defensible evidence trail.

Explainability

Observability tells you what happened. Explainability tells you why, and what you can do about it. Seekr helps you understand why your AI behaves the way it does, so you can:

- **Iterate with precision**, fixing the right thing instead of guessing.
- **Audit, contest, and improve every output**, with evidence across every layer: the model, the data that shaped it, and the context it receives at inference time.
- **Onboard new use cases faster**, because each deployment builds on what you already know.

Attribution

Seekr closes the gap between observability and explainability with structured evidence of which inputs drove an output.

- **Context attribution** identifies which evidence the system relied on.
- **Data attribution** identifies which training or fine-tuning examples shaped the model's behavior.
- **Model attribution** reveals which internal patterns were used to produce the output.

Evidence

When a decision gets questioned, you need to show your work. Seekr gives you control before, during, and after every decision:

- **Before the model runs:** Set the boundaries, the latitude, and where a human stays in the loop.
- **While it runs:** Watch every step, tool call, and output as it happens.
- **When something is wrong:** Trace the reasoning, correct it, and confirm the fix held.
- **Before it ships:** Test, score, and certify every model for bias, accuracy, and risk.



Sovereignty

When you hand your data to a provider without asking what happens to it, you're losing privacy and your advantage:

- **Your data trains their model.** Your knowledge makes their product smarter, not yours.
- **Your competitive edge becomes theirs.** What makes you different is now working for everyone.
- **Your evidence trail breaks.** If a vendor swaps the model underneath you, decisions you already defended are suddenly harder to stand behind.

Seekr builds the model on your data and runs in the cloud, on-premises, at the edge, or air gapped.



Right-sizing

Model selection often happens once, under pressure, and nobody ever asks later whether a decision still stands. Seekr makes finding the right AI a continuous discipline:

- **Evaluate** on your own data and criteria, not a vendor's benchmark.
- **Right-size** to the smallest, fastest, least costly model that fits the work.
- **Escalate** when the work demands it, with the evidence to defend the choice on cost and performance.
- **Re-decide** as fast as the market moves, so you're never stuck with yesterday's default.

The Seekr product family

seekrFlow™

The platform for building, governing, and deploying trusted AI. Ingest and prepare data, fine-tune models, orchestrate agents, run inference, and get explainability across every layer.

seekrGuard

AI evaluation and risk governance for evaluating, comparing, and documenting AI behavior using your organization's context as the benchmark so you can move from assumptions to evidence.

Solutions and professional services

Few organizations have the engineers, data infrastructure, or time to build the solutions they need. Seekr engineers can scope, build, and improve solutions alongside you, on your data, in your environment. **The result is trusted AI that's ready to use, defensible by design, and yours to own.**



Explainability built in

Not bolted on after the fact.



Sovereignty by design

Running in your environment on your terms.



Right-sized from the start

So you're not paying frontier prices for work that doesn't need it.

TRUSTED BY



Ready to build AI you can trust?

seekr.com
contact@seekr.com

