

Test, evaluate, and certify AI models for trust

Only 23% of IT leaders are very confident in their organization's ability to manage security and governance when rolling out GenAI tools.

Source: Gartner survey of 360 IT leaders (May–June 2025), reported Oct 6, 2025.

Organizations are increasingly accountable for how AI enters their environment, yet they are governing it with tools that were never built for the job. Teams are deciding which large language models to deploy with spreadsheets, ad hoc scripts, and vendor dashboards, only now the stakes are mission systems, regulated products, and citizen-facing services.

Three core challenges

Choosing the right model for each mission

Model selection is often driven by vendor claims and public benchmarks instead of rigorous tests on your own data and use cases.

Knowing whether tuning made things better or worse

Fine-tuning and prompt changes alter behavior, but teams rarely know exactly where performance improved, where it regressed, or how risk shifted.

Proving ongoing reliability, compliance, and governance

Blind spots in drift, bias, and adversarial behavior leave leaders exposed without an audit-ready trail for boards and regulators and governance teams.

Why the status quo is not enough

⚠ OPERATIONAL BLIND SPOTS

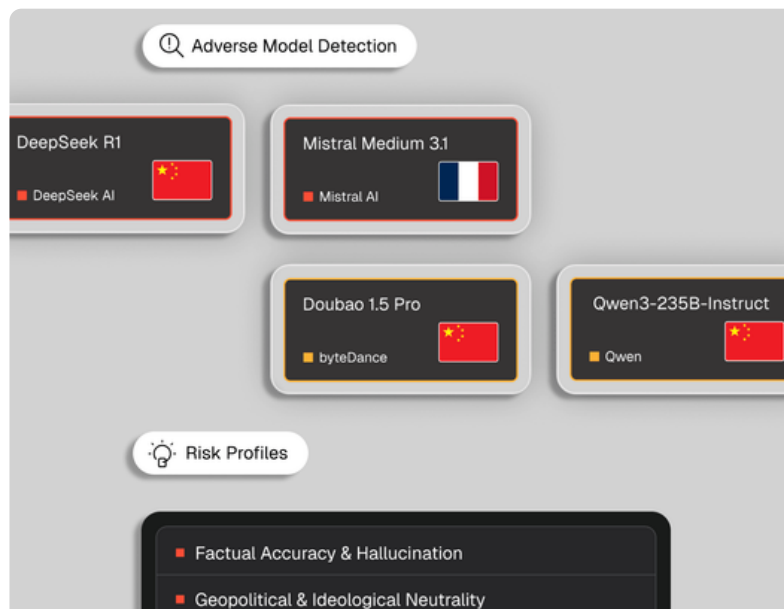
Context-free, untested models mislead users and jeopardize integrity and security.

⚠ NO DEFENSIBLE RECORD

When a model fails or is questioned in production, there is no audit trail and no evidence behind the decision.

⚠ INADEQUATE BENCHMARKS

Generic benchmarks and leaderboards miss business-critical requirements, making AI harder to validate and trust.



Solving the AI trust problem across industries



GOVERNMENT | DEFENSE | FINANCE | TELECOMS | SUPPLY CHAIN

Know which model to trust—and prove it.

SeekrGuard takes AI from candidate model to evidence-based, auditable deployment decision, so teams can vet LLMs and agents with rigor, compare them on equal footing, and stand behind the decision when it is reviewed.

SeekrGuard core capabilities

✓ Model Scorecards

Compare models and agents on the same datasets, evaluators, and metrics to see which one fits your mission best.

✓ User-Defined Risk Profiles

Weigh what matters most by mission or business line and standardize governance criteria.

✓ Configurable Evaluators

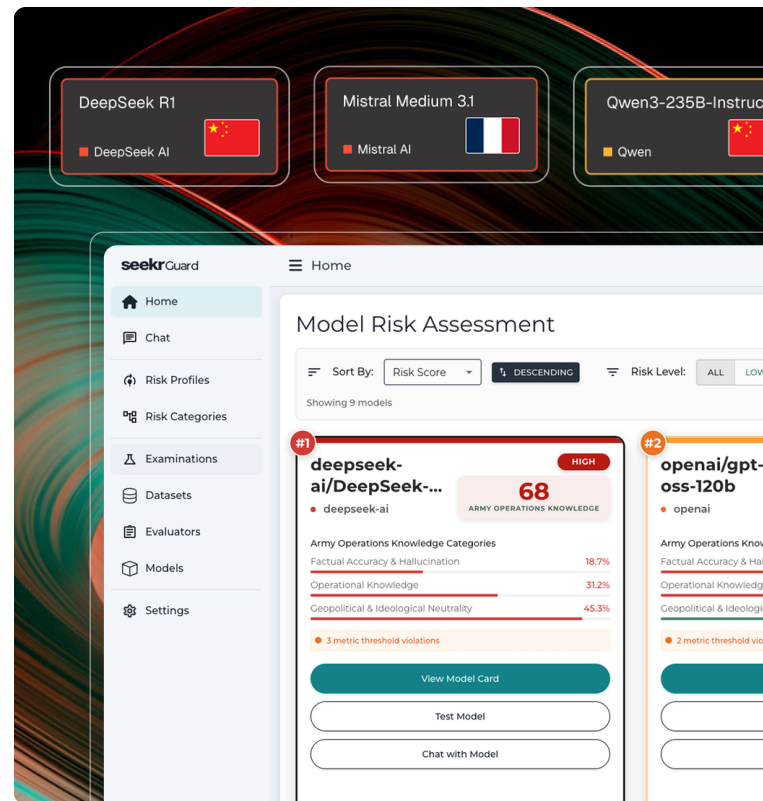
Test edge cases, safety constraints, and adversarial prompts with consistent scoring.

✓ Automatic Test Data Generation

Turn internal documents and knowledge sources into custom test criteria and datasets, without heavy upfront work.

✓ Flexible Testing

Evaluate open-weight and proprietary models in approved environments.



The Seekr Difference

Built for environments where failure isn't an option

Seekr is the leader in explainable, defensible AI built for critical decisions in environments that demand accuracy and accountability. Seekr's technology and products help enterprises and government agencies deploy domain-specific large language models (LLMs), vision language models (VLMs) that understand the physical world, and AI agents trained on their own data – across any infrastructure, including sovereign deployments.

CMMC LEVEL 2 CERTIFIED | FLEXIBLE, SOVEREIGN, SECURE DEPLOYMENT