

ENTERPRISE AI MODEL INTAKE

Governing enterprise AI at scale

Every enterprise now runs the same experiment at once: thousands of employees reaching frontier models through a handful of sanctioned tools, engineering teams wiring agents into internal platforms, and a governance function trying to certify all of it before a model ever reaches production. The bottleneck is rarely whether a model is safe — it is how long it takes the organization to find out.

3–6 months

to clear one model through intake — repeated in full for every new version

4–5 chains

separate approval chains per model, none designed to talk to each other

Days

to a same-week disposition once the assessment itself is automated

■ Model intake takes months, not days

Getting one model into production means clearing four or five separate approval chains that were never designed to talk to each other — and the pain compounds with every model, every version, and every internal platform that adds one.

Every function reviews from scratch

Legal reviews licensing and IP, security runs its own assessment or waits on a pen test slot, risk scores business impact, compliance maps frameworks — each on its own intake form, timeline, and evidence.

Nothing carries over

A model approved in January is unknown in March when the provider ships a new version. With no persistent baseline, releases reach employees before anyone re-assesses them and the full cycle repeats.

Red teaming won't scale by hand

Where it happens at all, it is a manual engagement scoped and staffed by hand. That cannot keep pace with a release cadence measured in weeks, so the assessment queue only grows.

The business absorbs the delay

Teams either wait months for a model to clear governance while competitors ship, or route around the process entirely — and consumption becomes shadow IT that governance cannot see, let alone certify.

Evidence must survive audit

The program has to produce a defensible, board-level number, continuously enough to survive dozens of audits a year against NIST AI RMF, the EU AI Act and the OWASP LLM Top 10.

Why the delay is structural, not incidental

Much of the wait is spent hoping the model provider will hand over documentation it is never going to provide. Commercial frontier providers do not publish training data, weights, or a security package built for a specific customer's policy — so legal, security, risk, and compliance teams wait on evidence that does not exist, then fall back to manual testing to fill the gap, each on its own timeline.

The evidence that actually clears a model has to be generated, not requested: the model measured against the organization's own policy, using its own synthetic data, inside its own environment. That evidence is provider-agnostic — closed frontier lab, open-weight family, or internal fine-tune — and reusable across every function that signs off.

ENKRYPT AI + ANACONDA

How the intake review gets automated

Enkrypt AI replaces the manual, sequential intake cycle with a single automated assessment that produces evidence every downstream function can use — governed, versioned, and reproducible on the Anaconda Platform.

Policy-driven red teaming

Tested in days, not months.

Prompt injection, data exfiltration, jailbreak resistance, harmful content, bias and misuse — run against the organization's own synthetic data rather than generic benchmarks.

A policy agent for every function

One evidence pass, not four.

Legal, security, risk, and compliance requirements are captured once as a machine-readable policy set, so one assessment answers all four instead of four reviews on four timelines.

A persistent cross-model baseline

Review starts from a diff.

Every new model or version is compared automatically against the standing baseline and the last approved version, so nothing restarts from zero.

Policy-to-evidence reporting

Audit-ready by default.

Results are scored, mapped to the frameworks the organization is audited against, and written to a durable evidence store each function pulls from instead of producing its own.

A CI/CD gateway for new models

Approval attached to the release.

New models added to a platform like Databricks or Azure AI Foundry are scanned automatically as part of the pipeline — a disposition on the release, not a meeting cycle scheduled around it.

REQUIREMENT MAPPING

WHERE INTAKE STALLS TODAY

Legal, security, risk and compliance run separate reviews

Every new model or version restarts review from zero

Manual red teaming can't keep pace with release cadence

Each function needs its own sign-off artifact

New models on internal platforms bypass governance until noticed

HOW ENKRYPT AI CLOSSES IT

One policy set, one evidence pass across all four functions

Automated red teaming against a persistent baseline, diffed release over release

Policy-driven red teaming completed in days, not months

Single evidence store generating framework-mapped reports on demand

CI/CD gateway assesses every new model addition automatically

- A governance process that takes months is a risk of its own: the business waits and falls behind, or routes around it and loses visibility entirely. **Automating the assessment itself turns a multi-month intake cycle into a same-week disposition —** without asking legal, security, risk, or compliance to lower the bar.