

## VIPRE AI Assurance Services

With over 25 years of experience securing Fortune 500s, SMEs, and startups, VIPRE now brings that same rigour to AI systems. Our AI Assurance evaluates models and pipelines across 10 critical trust dimensions — delivering the independent, evidence-based assurance that regulators, boards, and customers now demand.



### AI Trust Assessment

Comprehensive evaluation across Accuracy, Bias & Fairness, Safety, Privacy, Security, Transparency, Robustness, Compliance, Performance, and Governance.



### Adversarial AI Testing

Red-team your models against prompt injection, data poisoning, model inversion, and adversarial manipulation to expose vulnerabilities before malicious actors do.



### Privacy & Security Evaluation

Assess data handling practices, access controls, and model inversion risks across your entire AI pipeline and third-party integrations.



### AI Supply Chain Risk

Evaluate embedded third-party models and APIs for hidden vulnerabilities in data handling, behavioural drift, and access control that standard tools cannot detect.



### Bias & Fairness Audit

Detect discriminatory patterns across protected classes and validate equitable outcomes for regulated use cases including hiring, lending, and healthcare AI.



### Compliance Readiness Review

Map AI behaviour to EU AI Act, NIST AI RMF, HIPAA, SOC 2, GDPR, ISO 42001, NYC LL144, ECOA, and FCRA with regulator-submission-ready evidence packages.



### Transparency & Explainability

Validate model documentation, decision explainability, and audit trails to satisfy internal governance and external regulatory expectations.

## Why organisations need AI Assurance



### Regulatory pressure is accelerating

The EU AI Act, NIST AI RMF, and sector-specific rules impose legal obligations on AI deployments. Regulators now require documented, independent evidence of trust — not self-attestation.



### AI models are active attack surfaces

Without dedicated AI assurance testing, prompt injection, model inversion, and data poisoning remain invisible and exploitable. Standard security tools were not designed to find these vulnerabilities.



### AI bias carries direct legal exposure

Systems used in hiring, lending, and healthcare face liability under ECOA, FCRA, and anti-discrimination law. A single biased model can trigger class-action litigation and regulatory investigation.



### Models drift — risk does not stand still

AI systems degrade as data distributions shift and retraining introduces new failure modes. Structured re-evaluation is essential to maintain compliance post-deployment.

## Our Process



## Packages

Select from four packages — Quick Scan, Standard, Comprehensive, and Annual Plan — or build a custom engagement tailored to your organisation's needs.

| Quick Scan                                                                     | Standard                                                           | Comprehensive                                                                                         | Annual Plan                                                                                                    |
|--------------------------------------------------------------------------------|--------------------------------------------------------------------|-------------------------------------------------------------------------------------------------------|----------------------------------------------------------------------------------------------------------------|
| Raw scores and risk highlights for early screening and board-level visibility. | Summary report, basic compliance mapping and remediation guidance. | Full report, executive summary, complete compliance mapping, 30 days of support and one revalidation. | Continuous AI Assurance with ongoing benchmarks, reporting and support across the year, plus one revalidation. |
| 3 dimensions<br>25 credits                                                     | Up to 5 dimensions<br>50 credits                                   | All 10 dimensions<br>100 credits                                                                      | All 10 dimensions<br>600 credits/year                                                                          |

| FEATURE                 | QUICK SCAN | STANDARD       | COMPREHENSIVE | ANNUAL PLAN      |
|-------------------------|------------|----------------|---------------|------------------|
| Credits                 | 25 credits | 50 credits     | 100 credits   | 600 credits/year |
| Trust Dimensions        | 3          | Up to 5        | All 10        | All 10           |
| Report Type             | Raw Scores | Summary Report | Full Report   | Full Report      |
| Risk Highlights         | ✓          | ✓              | ✓             | ✓                |
| Compliance Mapping      | ✗          | Basic          | Full          | Full             |
| Executive Report        | ✗          | ✗              | ✓             | ✓                |
| Revalidation Assessment | ✗          | ✗              | 1 included    | 1 included       |

## 10 Trust Dimensions

Every evaluation is scored across a standardised framework covering the full spectrum of AI risk.

Accuracy

Bias & Fairness

Safety

Privacy

Security

Transparency

Robustness

Compliance

Performance

Governance

## Compliance Frameworks Covered

EU AI Act

NIST AI RMF

HIPAA

SOC 2

GDPR

ISO 42001

NYC LL144

ECOA

FCRA