

Behavioural science for autonomous **AI agents**.

We **observe**, **measure** and **explain** how agents behave — to evaluate, train and keep them safe.

WHY IT MATTERS

Autonomous agents cannot be managed through outputs alone. To evaluate an agent, improve it or establish that it remains safe, an organisation must know what it actually did: what it used, which actions it took and whether its conclusions followed from the available evidence.

XPLORE TURNS BEHAVIOUR INTO MEASURABLE EVIDENCE

Evaluate before building a benchmark.

A first assessment from the structure of the agent's behaviour.

Works without changing the agent.

Any model, framework or runtime.

Checks evidence, not plausibility.

Every statement and action, on every run.

Failures become training signals.

Located precisely and returned to the next improvement cycle.

The result: more work can be delegated to autonomous systems without losing control.

ONE BEHAVIOURAL LAYER ACROSS THE AGENT LIFECYCLE

Evaluate	Compare unfamiliar agents before deployment, even before a task-specific benchmark exists.
Train	Turn specific behavioural failures into improvements to prompts, tools, routing and policies — without retraining the model.
Keep safe	Detect unsupported conclusions, prohibited actions and behavioural drift in operation. Re-assess whenever the agent changes.

WHAT CHANGES

- One benchmark per agent** becomes reusable rules per kind of work.
- Sample-based review** becomes evidence from every run.
- A score** becomes a precise route to improvement.

EVIDENCE**ASSESSED**

Edinburgh Napier University

IN USE

Regulated medical technology

PROTECTED

UK patent application, 2025

NEXT STEP**Measure an agent on real work.**xploreintelligence.co.ukhello@xploreintelligence.co.uk