

Risk, Security & Compliance · Risk · Controls

# Risk Control Oversight Agent

CAI Score L4 · Decide — Senior human accountable for every decision

▲ this agent

| L1<br>Assist | L2<br>Draft                 | L3<br>Operate          | L4<br>Decide                     | L5<br>Autonomous |
|--------------|-----------------------------|------------------------|----------------------------------|------------------|
| TIER         | PILLAR                      | AUTOMATION             | HUMAN ROLE                       |                  |
| L4 · Decide  | Risk, Security & Compliance | ≈40% of routine effort | Named accountable decision-maker |                  |

## 01 — Executive summary

Risk Control Oversight Agent is a CAI Score L4 (Decide (supervised)) agent for risk, security & compliance. Supports oversight of risk controls — data preparation, delivery and compliance reporting across the control environment. It operates with senior human accountability: a named person owns every decision it supports. Indicatively, it can redirect around 40% of the routine effort this work consumes.

## 02 — Problem & context

- Control gaps only surface at the formal review, never between them.
- Oversight data is prepared by hand, so it is stale by the time it is read.

## 03 — Representative use cases

- Monitors the full control environment continuously, surfacing gaps and weaknesses between formal review cycles.
- Prepares oversight data automatically each period, replacing manual data assembly with a current, consistent view.
- Produces compliance reporting against the control framework on a defined schedule without manual effort.
- Flags controls that are approaching threshold breach so owners can act before a gap becomes a finding.
- Produces the evidence to track control coverage, gap frequency, and time-to-remediation.

## 04 — How to use it

- 1 Connect it in your tenant to the systems this work already touches (controls).
- 2 Obtain senior sign-off and confirm the named accountable person before go-live.
- 3 Receives the case or work product
- 4 Assembles evidence and flags risk or discrepancy
- 5 Presents structured findings to the named human
- 6 Named human makes the decision; rationale is logged
- 7 Review decision outcomes and the audit trail after each cycle; adjust scope as needed.

## 05 — Where it fits in the process

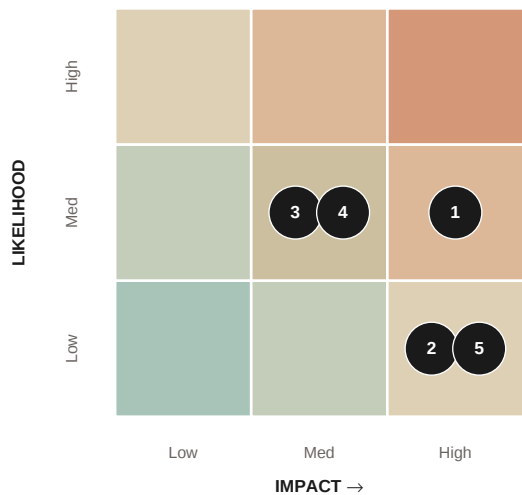
Provides the continuous-oversight layer across the control set, so gaps and weaknesses surface between formal reviews rather than at them.



Decision-support loop for a L4 agent — human owns every call.

## 06 — Risk assessment

| # | RISK                                       | LIKELIHOOD | IMPACT | MITIGATION / CONTROL                                                                                    |
|---|--------------------------------------------|------------|--------|---------------------------------------------------------------------------------------------------------|
| 1 | <b>Incorrect evidence assembly</b>         | Medium     | High   | Named human reviews all agent output before any decision is made; the human owns the conclusion.        |
| 2 | <b>Over-reliance on agent assessment</b>   | Low        | High   | L4 tier structure: agent evidences, human decides. Decision authority cannot be delegated to the agent. |
| 3 | <b>Exposure of sensitive data</b>          | Medium     | Medium | Runs inside your own tenant; Colleague AI processes none of your data.                                  |
| 4 | <b>Behavioural drift over time</b>         | Medium     | Medium | Periodic review against case outcomes detects drift; scope and prompts are version-controlled.          |
| 5 | <b>Out-of-policy or edge-case handling</b> | Low        | High   | Hard scope boundary: the agent flags anything outside policy; senior sign-off is required at go-live.   |



Residual likelihood × impact after the tier's built-in controls. Numbers map to the table above.

07 — Regulatory & compliance impact

| FRAMEWORK     | RELEVANCE                                                                            | HOW THIS AGENT ALIGNS                                                                                                                      |
|---------------|--------------------------------------------------------------------------------------|--------------------------------------------------------------------------------------------------------------------------------------------|
| EU AI Act     | Operational use under human oversight; logging and risk-tiering apply.               | Documented risk tier, human oversight and event logging per agent — designed to map to the obligations.                                    |
| DORA          | Relevant where the function is in scope for financial-sector operational resilience. | Actions are logged and attributable, supporting resilience and traceability requirements — designed to map to.                             |
| GDPR          | Applies wherever personal data is in scope of the workflow.                          | Data stays in your tenant; Colleague AI processes none of it. Purpose limitation and minimisation are preserved.                           |
| ISO/IEC 42001 | AI management-system controls: roles, monitoring, review.                            | The agent operates within a documented management system — defined tier, named roles, logging and review. Designed to meet; not certified. |

“Designed to map to” / “designed to meet” describe alignment with the CAI Score framework, not external certification. Confirm obligations with your compliance function.

08 — Value & illustrative ROI

Supports oversight of risk controls — data preparation, delivery and compliance reporting across the control environment.

| Assumptions (illustrative — replace with your own)                         | Result                                                       |
|----------------------------------------------------------------------------|--------------------------------------------------------------|
| 4 FTE engaged · 1840 productive h/FTE/yr · 40% addressable · €65/h blended | ≈ 2,944 hours/yr redirected · ≈ €191,000/yr indicative value |



Figures are an illustrative model to be validated against your own volumes and rates – not a guarantee.

**KPIs to track**

Improved control visibility · reduced manual effort · continuous oversight posture.

## 09 — Recommendations

- Obtain senior sign-off before go-live — mandatory for L4 deployments.
- Map the full control framework into the agent before launch; confirm coverage with the risk team.
- Assign named control owners for each domain; the agent supports their oversight, it does not replace it.
- Review the first monthly output with internal audit before scaling to the full control environment.

## 10 — Governance & method

This agent runs in your own Microsoft/Azure tenant; Colleague AI hosts only the control plane (score, policy, audit) and processes none of your data. A named senior person owns every decision the agent supports. Every action is time-stamped and attributable. The CAI Score classification is documented and open to challenge.

---

This dossier is an evaluation document. ROI and automation figures are illustrative models, not commitments. Compliance statements describe the CAI Score framework and are not a certification of your compliance with any law or standard. © Colleague AI — proprietary. CAI Score™.