

Operations & Service Delivery · Operations · Reconciliations

Reconciliation Root-Cause Agent

CAI Score L3 · Operate — Bounded autonomy

▲ this agent

L1 Assist	L2 Draft	L3 Operate	L4 Decide	L5 Autonomous
TIER	PILLAR	AUTOMATION	HUMAN ROLE	
L3 · Operate	Operations & Service Delivery	≈65% of routine effort	Owns exceptions and oversight	

01 — Executive summary

Reconciliation Root-Cause Agent is a CAI Score L3 (Operate) agent for operations & service delivery. Investigates each reconciliation break, prioritises it by materiality, and narrates why it happened and how to fix it. It operates with bounded autonomy: a person owns exceptions and retains workflow oversight. Indicatively, it can redirect around 65% of the routine effort this work consumes.

02 — Problem & context

- Analysts burn the first days of close manually chasing why each break happened.
- The same systemic break reappears every month because no one joins the dots.

03 — Representative use cases

- Investigates each reconciliation break, prioritises it by materiality, and narrates why it happened and how to fix it.
- Removes a recurring drain: analysts burn the first days of close manually chasing why each break happened.
- Closes a structural gap: the same systemic break reappears every month because no one joins the dots.
- Integrates into the live process — sits at the front of the reconciliation queue and feeds clean, explained breaks into the dual-control step.
- Produces the evidence to track lower break-resolution time, fewer repeat breaks, and a cleaner audit trail.

04 — How to use it

- 1 Connect it in your tenant to the systems this work already touches (reconciliations).
- 2 Confirm scope, the named owner, and the oversight the tier requires.
- 3 Detects the trigger in the workflow
- 4 Handles routine in-policy cases end to end
- 5 Escalates exceptions to the owner
- 6 Logs every action for audit
- 7 Review the audit log and KPI movement after each cycle; tune scope as you learn.

05 — Where it fits in the process

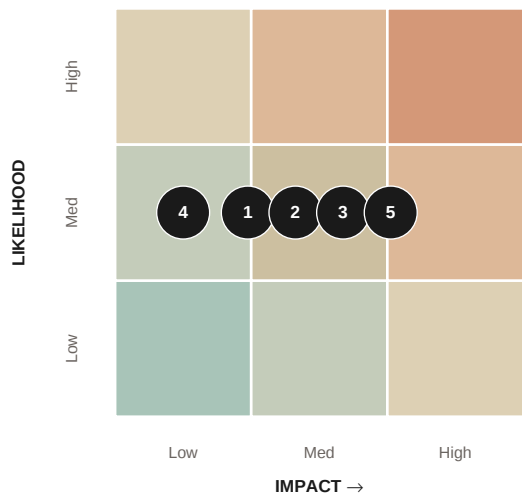
Sits at the front of the reconciliation queue: it triages breaks the moment they appear, before an analyst opens the first one — and feeds clean, explained breaks into the dual-control step.



Operating loop for a L3 agent — bounded autonomy.

06 — Risk assessment

#	RISK	LIKELIHOOD	IMPACT	MITIGATION / CONTROL
1	Incorrect or incomplete output	Medium	Medium	Routine cases only; exceptions and anything material are escalated to the named owner, with a full audit trail.
2	Over-reliance / reduced human scrutiny	Medium	Medium	Tier oversight keeps a human in the loop; KPIs monitor quality so reliance never outruns control.
3	Exposure of sensitive data	Medium	Medium	Runs inside your own tenant; Colleague AI processes none of your data, and access stays under your existing controls.
4	Behavioural drift over time	Medium	Low	Periodic review against KPIs and sampling detects drift; scope and prompts are version-controlled.
5	Out-of-policy or edge-case handling	Medium	Medium	Scope is bounded and documented; anything outside it is escalated, not improvised.



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

Investigates each reconciliation break, prioritises it by materiality, and narrates why it happened and how to fix it.

Assumptions (illustrative — replace with your own)	Result
4 FTE engaged · 1840 productive h/FTE/yr · 65% addressable · €55/h blended	≈ 4,784 hours/yr redirected · ≈ €263,000/yr indicative value



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

KPIs to track

Lower break-resolution time · fewer repeat breaks · cleaner audit trail.

09 — Recommendations

- Assign a named owner — Operations lead — accountable for its outputs and oversight.
- Define the exception rules explicitly: what the agent handles vs. what it must escalate. Review escalations weekly at first.
- Pilot on a bounded, representative slice for 4-6 weeks; baseline the KPIs first so the gain is measurable, not asserted.
- Scale once the KPIs hold and the audit trail is clean; expand scope deliberately, not all at once.

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 person owns exceptions and retains workflow oversight. 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™.