

Risk, Security & Compliance · Risk · Due Diligence

DDQ Data Search 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	Risk, Security & Compliance	≈60% of routine effort	Owns exceptions and oversight

01 — Executive summary

DDQ Data Search Agent is a CAI Score L3 (Operate) agent for risk, security & compliance. Searches internal data to answer due-diligence questionnaires quickly and accurately, instead of manual lookups across systems. It operates with bounded autonomy: a person owns exceptions and retains workflow oversight. Indicatively, it can redirect around 60% of the routine effort this work consumes.

02 — Problem & context

- Answering one DDQ means manual lookups across five different systems.
- Every responder writes answers differently, so quality is uneven.

03 — Representative use cases

- Searches across all connected internal systems to retrieve relevant data for each DDQ question.
- Drafts sourced, cited answers for the responding team to review and approve — not draft from scratch.
- Enforces consistent answer quality and format across all DDQ responses, regardless of which team member responds.
- Reduces DDQ turnaround from days to hours by eliminating the manual cross-system lookup.
- Produces the evidence to track response accuracy, turnaround time, and source citation completeness.

04 — How to use it

- 1 Connect it in your tenant to the systems this work already touches (due diligence).
- 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

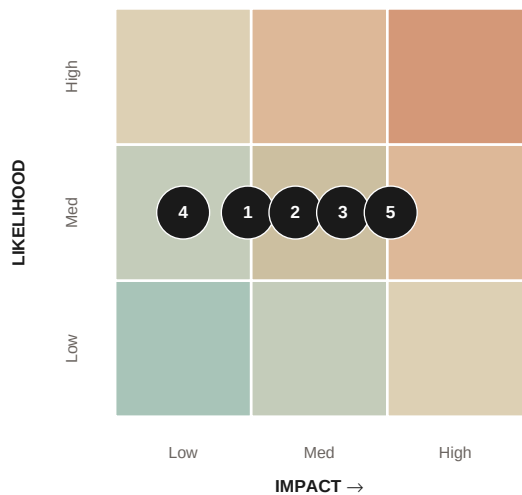
Sits at the start of DDQ turnaround — it finds and drafts sourced answers from internal systems, so the team reviews and confirms rather than hunting.



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

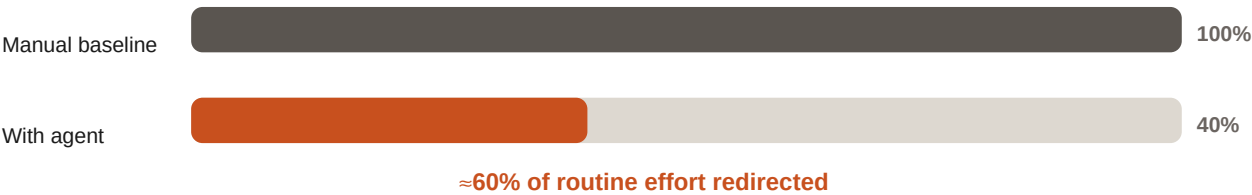
FRAMEWORK	RELEVANCE	HOW THIS AGENT ALIGNS
EU AI Act	Operational use under human oversight; logging and risk-tiering apply.	Documented risk tier, human oversight and event logging per agent — designed to map to the obligations.
DORA	Relevant where the function is in scope for financial-sector operational 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

Searches internal data to answer due-diligence questionnaires quickly and accurately, instead of manual lookups across systems.

Assumptions (illustrative — replace with your own)	Result
4 FTE engaged · 1840 productive h/FTE/yr · 60% addressable · €60/h blended	≈ 4,416 hours/yr redirected · ≈ €265,000/yr indicative value



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

KPIs to track

Improved efficiency and accuracy on due-diligence responses.

09 — Recommendations

- Connect the agent to all internal data sources referenced in typical DDQs before deployment.
- Assign a named DDQ owner to review and approve every draft response before submission.
- Review source citation accuracy in the first five DDQs; reconfigure data connections where gaps appear.
- Baseline turnaround time and response quality before go-live.

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™.