



Truveil

AI ACCOUNTABILITY AUDIT REPORT

Email Triage Agent

Audit of Email Triage Agent

REPORT ID	SR-112
GENERATED	Tue, 01 Sep 2026 21:01:25 GMT
SUBJECT AGENT	Email Triage Agent
AGENT CATEGORY	1 (Declared at Registration)
RISK CLASSIFICATION	LOW
FRAMEWORKS APPLIED	Truveil methodology, DPDP Rules 2025

Report Contents

I	Auditor Opinion
II	Scope and Methodology
III	Agent System Description
IV	Findings
V	Recommended Actions
VI	Technical Detail

Auditor Opinion

A

Accountability Grade · **81.2/100**

LOW RISK

This email reply agent scores 81.2/100 (Grade A, LOW risk). Human-in-loop oversight, approval gating before high-stakes sends, and strong accountability logging are notable strengths. Key gaps are absent uncertainty disclosure, missing input validation aligned to DPDP Rules 2025 Rule 6, and no logged kill switch informed by EU AI Act Article 14(4)(e). These gaps are improvement opportunities, not critical failures.

TRANSPARENCY

83.3

ACCOUNTABILITY

89.5

DATA TRUST

70

REVERSIBILITY

76.9

ISSUED BY

This report has been generated by Truveil, an AI agent accountability platform. The accountability scoring is rule-based and deterministic, calibrated against primary regulatory text. The narrative interpretation in this report is produced by Anthropic Claude under structured constraints maintained by Truveil. No portion of the scoring is delegated to a generative model.

Scope and Methodology

What this report covers

This report covers a single execution of the agent identified on the cover page. The audit assesses the logged actions, decisions, validations, and human checkpoints recorded during that execution against four accountability dimensions: *Transparency*, *Accountability*, *Data Trust*, and *Reversibility*. Each dimension corresponds to a category of regulatory obligation recurring across the frameworks Truveil applies.

How Truveil scores

Truveil's scoring engine is rule-based and deterministic. The same evidence, scored against the same framework, produces the same result every time. Each accountability signal is anchored to specific provisions in primary regulatory text. Caps and floors apply where a binding obligation is unmet, these are not weighted heuristics but explicit thresholds defined in the underlying regulation.

Narrative interpretation in this report is produced by a language model under structured constraints. The model summarises findings already determined by the engine. It does not score, classify, or rank the agent. Citations in this report are restricted to provisions present in the engine's citation index for this run.

Regulatory anchors for this audit

Truveil methodology, DPDP Rules 2025

What this report does not certify

This report is an evidence-based assessment of one execution of one agent. It does not constitute a regulatory filing, a legal opinion, or a certification of conformity. It is not a substitute for accredited third-party assurance such as ISO/IEC 42001 certification. Truveil's role is to surface accountability gaps that would otherwise remain undocumented, against the regulatory provisions identified in the engine's citation index.

Agent System Description

Agent and run metadata

AGENT	Email Triage Agent
CATEGORY	1 (Declared at Registration)
CATEGORY BASIS	Category 1 declared at agent registration.
RUN ID	101-1780342562631
EXECUTION	2026-06-01T19:36:06.409Z
TIMESTAMP	
TOTAL LOGGED	6
EVENTS	

Step-by-step execution

Reconstructed from 6 logged events via the MCP connector; Conversational instrumentation observes registered decisions and session-level events; runtime signals such as kill-switch wiring and pipeline validation are outside its surface.

01 Email Thread Fetch

Agent fetched the full Gmail thread to ground subsequent drafting and classification steps.

OK

02 Stake Classification

Agent classified the email as high-stakes, triggering the elevated approval workflow.

OK

03 Calendar Context Retrieval

Agent fetched calendar events for the next five days to inform reply content.

OK

04 Draft Reply Creation

Agent created a Gmail draft of the high-stakes reply, queued for human review.

WATCH

05 Human Approval Gate

Approval request issued; agent awaited explicit human approval before sending.

OK

06 Approved Send Execution

Email dispatched after human approval confirmed via Telegram send-and-wait mechanism.

OK

Findings

Findings index

#	SEVERITY	FINDING	FRAMEWORK
1	HIGH	No Input Validation Logged	DPDP Rules 2025 Rule 6, effective 13 May 2027
2	HIGH	Kill Switch Not Logged	Truveil methodology, informed by EU AI Act Article 14(4)(e) and NIST AI RMF as reference standards
3	MEDIUM	Uncertainty Not Disclosed	Truveil methodology, informed by EU AI Act Article 13(3) and NIST AI RMF as reference standards
4	MEDIUM	Override Capability Not Logged	Truveil methodology, informed by EU AI Act Article 14(4)(d) and NIST AI RMF as reference standards
5	LOW	Versioned Outputs Absent	Truveil methodology, informed by EU AI Act Article 12 and NIST AI RMF as reference standards

Findings in full

Finding 1No Input Validation LoggedHIGH

No validation or verification event was recorded for input data, creating a provenance integrity risk aligned to DPDP Rules 2025 Rule 6.

FRAMEWORK	DPDP Rules 2025 Rule 6, effective 13 May 2027
EVIDENCE	no validation or verify event
AFFECTED STEP	Step 01: Email Thread Fetch

Finding 2Kill Switch Not Logged **HIGH**

No kill switch or emergency stop event was recorded, leaving no auditable emergency halt path informed by EU AI Act Article 14(4)(e).

FRAMEWORK	Truveil methodology, informed by EU AI Act Article 14(4)(e) and NIST AI RMF as reference standards
EVIDENCE	no kill_switch or emergency_stop event
AFFECTED STEP	REGISTRATION

Finding 3Uncertainty Not Disclosed **MEDIUM**

No confidence or uncertainty signal was logged in any event, limiting recipient awareness of AI output reliability per EU AI Act Article 13(3).

FRAMEWORK	Truveil methodology, informed by EU AI Act Article 13(3) and NIST AI RMF as reference standards
EVIDENCE	no confidence or uncertainty mention in any log detail
AFFECTED STEP	Step 04: Draft Reply Creation

Finding 4Override Capability Not Logged **MEDIUM**

No override action was recorded, leaving no evidence of a user-facing override path per EU AI Act Article 14(4)(d).

FRAMEWORK	Truveil methodology, informed by EU AI Act Article 14(4)(d) and NIST AI RMF as reference standards
EVIDENCE	no override action logged
AFFECTED STEP	REGISTRATION

Finding 5Versioned Outputs Absent **LOW**

No version or revision detail was logged for outputs, reducing traceability aligned to EU AI Act Article 12.

FRAMEWORK	Truveil methodology, informed by EU AI Act Article 12 and NIST AI RMF as reference standards
EVIDENCE	no version/revision detail and no source_fetch with detail
AFFECTED STEP	Step 04: Draft Reply Creation

Recommended Actions

Actions are ordered by priority. *Now* actions remediate findings that breach binding obligations under the frameworks applied. *Soon* actions address material gaps that fall short of best practice without breaching a binding rule. *Later* actions improve accountability posture for forward-looking compliance.

01 Implement Input Validation Logging

Add a structured validation event after every data fetch to satisfy DPDP Rules 2025 Rule 6 provenance integrity requirements.

NOW

02 Register and Log Kill Switch Capability

Declare an emergency stop mechanism at registration and emit kill switch readiness events per EU AI Act Article 14(4)(e).

NOW

03 Add Uncertainty Disclosure to Draft Events

Populate a confidence field in draft logs and surface uncertainty to the human reviewer, aligned to EU AI Act Article 13(3).

SOON

04 Log Override Capability at Registration

Record an `override_available` field at agent registration and within the approval flow per EU AI Act Article 14(4)(d).

SOON

05 Version Draft Outputs

Attach version identifiers to each draft creation event to support traceability per EU AI Act Article 12.

LATER

Technical Detail

This section is intended for engineering and platform teams responsible for the agent. Each finding from Section IV is restated with its developer-view structured fields and suggested remediation steps. Affected steps refer to the trace numbering in Section III.

Finding 1 · No Input Validation Logged

FAILURE TYPE	absent_input_validation_event
AFFECTED STEP	Step 01: Email Thread Fetch
CONFIDENCE	HIGH

SUGGESTED FIXES

- 1. Log a validation event after each data fetch confirming schema and source integrity.
- 2. Emit a structured verify record before draft creation begins.

Finding 2 · Kill Switch Not Logged

FAILURE TYPE	no_kill_switch_logged
AFFECTED STEP	REGISTRATION
CONFIDENCE	HIGH

SUGGESTED FIXES

- 1. Register a kill_switch capability in agent configuration and log activation readiness.
- 2. Emit an emergency_stop event type in the workflow schema for auditor visibility.

Finding 3 · Uncertainty Not Disclosed

FAILURE TYPE	missing_uncertainty_disclosure
AFFECTED STEP	Step 04: Draft Reply Creation
CONFIDENCE	HIGH

SUGGESTED FIXES

- 1. Add a confidence_score field to draft creation log events.
- 2. Surface uncertainty level in the human-review prompt before approval.

Finding 4 · Override Capability Not Logged

FAILURE TYPE	absent_override_event
AFFECTED STEP	REGISTRATION
CONFIDENCE	MEDIUM

SUGGESTED FIXES

1. Log an override_available flag at agent registration and during approval events.
2. Document and expose a rejection path in the Telegram approval flow.

Finding 5 · Versioned Outputs Absent

FAILURE TYPE	missing_output_versioning
AFFECTED STEP	Step 04: Draft Reply Creation
CONFIDENCE	MEDIUM

SUGGESTED FIXES

1. Attach a version identifier to each draft creation log event.
2. Store draft revision history with timestamps in the run record.

Addressing input validation and kill switch gaps will materially strengthen an already well-governed agent.