



AI ACCOUNTABILITY AUDIT REPORT

Senior Data Engineer Hiring – Frankfurt

Audit of Senior Data Engineer Hiring – Frankfurt

REPORT ID	SR-146
GENERATED	Wed, 02 Sep 2026 21:12:57 GMT
SUBJECT AGENT	Senior Data Engineer Hiring – Frankfurt
AGENT CATEGORY	1 (Declared at Registration)
RISK CLASSIFICATION	LOW
FRAMEWORKS APPLIED	EU AI Act

Report Contents

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

Auditor Opinion

B

Accountability Grade · **68.6/100**
LOW RISK

This recruitment screening agent scored 68.6/100 (Grade B, LOW overall risk) across six candidate decisions. Transparency and Accountability performed well; Data Trust at 40/100 is the primary concern, with absent freshness checks, validation, and groundedness. Reversibility at 53.8/100 reflects missing kill-switch and override controls. AI disclosure obligations under EU AI Act Article 50 and Article 50(1) remain unmet. Risk assessment obligations under EU AI Act Article 9 are also absent.

TRANSPARENCY

91.7

ACCOUNTABILITY

78.9

DATA TRUST

40

REVERSIBILITY

53.8

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

EU AI Act

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	Senior Data Engineer Hiring – Frankfurt
CATEGORY	1 (Declared at Registration)
CATEGORY BASIS	Category 1 declared at agent registration.
RUN ID	run_20260707_04e03d
EXECUTION	2026-07-07T13:55:22.621Z
TIMESTAMP	
TOTAL LOGGED	12
EVENTS	

Step-by-step execution

Reconstructed from 12 logged events via the SDK; SDK instrumentation observes events the agent code explicitly logs.

01

Candidate Intake and Profiling

INFERRED

Agent ingested profiles for six Senior Data Engineer candidates targeting a Frankfurt-based role.

OK

02

Screening Analysis and Scoring

INFERRED

Agent evaluated each candidate on seniority, stack fit, location, and role alignment, producing ranked assessments.

WATCH

03

Consequential Approval Requests

Six high-severity approval requests were submitted, one per candidate, each flagged as hire_decision.

OK

04

Screening Decisions Issued

Agent issued advance, conditional advance, or do-not-advance verdicts with documented reasoning for all six candidates.

WATCH

05 Bias and Oversight Controls Applied INFERRED

Bias checks and human-in-loop oversight mode were recorded; no override or kill-switch events were logged.

RISK

Findings

Findings index

#	SEVERITY	FINDING	FRAMEWORK
1	HIGH	No AI Disclosure to Candidates	EU AI Act Article 50, binding · EU AI Act Article 50(1), binding
2	HIGH	Absent Data Validation and Groundedness Checks	EU AI Act Article 15, effective 2 December 2027
3	HIGH	No Risk Assessment Documented	EU AI Act Article 9, effective 2 December 2027 · EU AI Act Article 27, effective 2 December 2027
4	MEDIUM	Kill Switch and Override Controls Absent	EU AI Act Article 14(4)(e), effective 2 December 2027
5	MEDIUM	Data Freshness Not Verified	EU AI Act Article 10(3), effective 2 December 2027

Findings in full

Finding 1No AI Disclosure to CandidatesHIGH

Candidates were not notified they were being assessed by an AI system, a binding obligation under EU AI Act Article 50 and Article 50(1).

FRAMEWORK	EU AI Act Article 50, binding · EU AI Act Article 50(1), binding
EVIDENCE	absent; required for Category 1 agents
AFFECTED STEP	Step 03: Consequential Approval Requests

Finding 2Absent Data Validation and Groundedness ChecksHIGH

No validation or verify events and no groundedness signals were logged, leaving input data quality and output accuracy unverified per EU AI Act Article 15.

FRAMEWORK	EU AI Act Article 15, effective 2 December 2027
EVIDENCE	no validation or verify event

Finding 3

No Risk Assessment Documented

HIGH

A risk assessment was absent at registration, required for Category 1 agents under EU AI Act Article 9 and Article 27.

FRAMEWORK	EU AI Act Article 9, effective 2 December 2027 · EU AI Act Article 27, effective 2 December 2027
EVIDENCE	absent; required for Category 1 agents
AFFECTED STEP	REGISTRATION

Finding 4

Kill Switch and Override Controls Absent

MEDIUM

No kill_switch, emergency_stop, or override action was logged, leaving no demonstrated reversal path per EU AI Act Article 14(4)(d) and Article 14(4)(e).

FRAMEWORK	EU AI Act Article 14(4)(e), effective 2 December 2027
EVIDENCE	no kill_switch or emergency_stop event
AFFECTED STEP	Step 05: Bias and Oversight Controls Applied

Finding 5

Data Freshness Not Verified

MEDIUM

No freshness check, data_age field, or ISO date was present in source fetch detail, undermining data currency assurance under EU AI Act Article 10(3).

FRAMEWORK	EU AI Act Article 10(3), effective 2 December 2027
EVIDENCE	no freshness, data_age, or ISO date in source_fetch detail
AFFECTED STEP	Step 02: Screening Analysis and Scoring

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 Candidate-Facing AI Disclosure

Add an AI disclosure notification before each screening interaction, satisfying EU AI Act Article 50 and Article 50(1) as a binding obligation.

NOW

02 Complete and Register a Risk Assessment

Document and attach a conformity risk assessment at agent registration to satisfy EU AI Act Article 9 and Article 27.

NOW

03 Add Validation, Groundedness, and Freshness Logging

Log validation events, groundedness scores, and data freshness timestamps per EU AI Act Article 10(3) and Article 15.

NOW

04 Implement Kill Switch and Override Controls

Log kill_switch and override action events to demonstrate reversal capability per EU AI Act Article 14(4)(d) and Article 14(4)(e).

SOON

05 Enable Versioned Output Logging

Attach version or revision identifiers to all screening outputs to satisfy 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 AI Disclosure to Candidates

FAILURE TYPE	missing_disclosure_event
AFFECTED STEP	Step 03: Consequential Approval Requests
CONFIDENCE	HIGH

SUGGESTED FIXES

- 1.Emit an ai_disclosure event before each candidate-facing interaction is initiated.
- 2.Log disclosure timestamp and delivery channel for each candidate record.

Finding 2 · Absent Data Validation and Groundedness Checks

FAILURE TYPE	absent_validation_and_groundedness
AFFECTED STEP	Step 02: Screening Analysis and Scoring
CONFIDENCE	HIGH

SUGGESTED FIXES

- 1.Add a validation step that verifies candidate data fields before scoring.
- 2.Log a groundedness or retrieval_score field for each screening output.

Finding 3 · No Risk Assessment Documented

FAILURE TYPE	missing_risk_assessment
AFFECTED STEP	REGISTRATION
CONFIDENCE	HIGH

SUGGESTED FIXES

- 1.Complete and attach a conformity risk assessment at agent registration.
- 2.Store risk assessment document reference in agent metadata before deployment.

Finding 4 · Kill Switch and Override Controls Absent

FAILURE TYPE	no_kill_switch_logged
AFFECTED STEP	Step 05: Bias and Oversight Controls Applied
CONFIDENCE	HIGH

SUGGESTED FIXES

1. Implement and log a kill_switch or emergency_stop event type in the workflow.
2. Register an override action handler and emit override events when triggered.

Finding 5 · Data Freshness Not Verified

FAILURE TYPE	absent_freshness_check
AFFECTED STEP	Step 02: Screening Analysis and Scoring
CONFIDENCE	MEDIUM

SUGGESTED FIXES

1. Attach a data_age or freshness_checked timestamp to each candidate data fetch.
2. Reject or flag candidate records whose source data exceeds a defined staleness threshold.

Remediate AI disclosure and risk assessment gaps before deployment; Data Trust controls require immediate engineering attention.