



▲ SAMPLE DATA — DEMONSTRATION PACKAGE ONLY

AI Governance Audit Evidence Package

A complete, point-in-time snapshot of the organisation's AI system inventory, risk classifications, adopted policies, control mappings, and approved questionnaire answers.

ORGANISATION Meridian Workflows Ltd Amsterdam, Netherlands · B2B SaaS	PACKAGE ID GV-AUD-2026-0704-MWL Generated 4 July 2026 · 09:14 UTC
PREPARED BY Govarna Platform govarna.com · Starter plan	CONTENTS 7 Sections · 17 Pages AI inventory · Classifications · Policies · Controls · Evidence
FRAMEWORKS COVERED EU AI Act · NIST AI RMF ISO/IEC 42001 · NYC LL 144	APPROVED BY Elena Hartmann Chief Technology Officer · 4 Jul 2026

DOCUMENT NAVIGATION

Contents

1	Executive Summary	3
2	AI System Inventory	4
2.1	Complete system register — 5 systems	
2.2	System detail sheets	
3	EU AI Act Risk Classifications	7
3.1	Classification methodology	
3.2	Per-system classification reports (5 systems)	
4	Policy Register	9
5	Control Mappings	11
5.1	NIST AI Risk Management Framework	
5.2	ISO/IEC 42001 — AI Management System	
6	Questionnaire Answer Library	14
7	Audit Trail	16
–	Legal Disclaimer	17

How to use this package

This evidence package is designed to be attached to a vendor AI security questionnaire response or submitted directly to a buyer's procurement or security team. Each section corresponds to a common category of AI governance question. The audit trail in Section 7 demonstrates that all content reflects approved, current information — not static boilerplate. If a reviewer has questions, the Govarna platform contact is: accounts@govarna.com

SECTION 1

Executive Summary

This package covers Meridian Workflows Ltd's AI governance posture as of 4 July 2026. The organisation operates as both a **provider** (for AI systems developed internally) and a **deployer** (for AI systems built on third-party models) under the EU AI Act. The platform maintains a live AI system register, policy library, and control mapping against NIST AI RMF and ISO/IEC 42001. The questionnaire answer library contains 47 approved answers, reducing the time to complete a new AI security review.



AI Risk Tier Distribution

RISK TIER	SYSTEMS	OBLIGATIONS STATUS
HIGH-RISK	1 — RecruitScreen	IN PROGRESS
LIMITED	1 — Support Assistant	COMPLIANT
MINIMAL	3	COMPLIANT

Framework Coverage

FRAMEWORK	CONTROLS MAPPED	STATUS
EU AI Act	14 / 16 applicable	PARTIAL
NIST AI RMF	22 / 22	COMPLETE
ISO/IEC 42001	18 / 18	COMPLETE

High-Risk System — Action in Progress

RecruitScreen (AI-SYS-004) is classified as **high-risk** under EU AI Act Annex III Point 4 (AI in employment and workers management). Full technical documentation, conformity assessment, and human oversight procedures are in progress, targeting completion by 31 July 2026. The system will not be modified or deployed to new roles until this process is complete. See Section 3.4 for the full classification report and remediation timeline.

No Prohibited Practices Identified

All five AI systems were assessed against EU AI Act Article 5 (prohibited practices). None were found to engage in subliminal manipulation, social scoring, real-time remote biometric identification in public spaces, or other prohibited practices as of this assessment date.

SECTION 2.1

AI System Inventory — Complete Register

Live register as at 4 July 2026 · 5 systems · Last updated 2 July 2026 by E. Hartmann (CTO)

SYSTEM ID	NAME	ROLE	PROVIDER / MODEL	PURPOSE	OWNER	SINCE	RISK TIER
AI-SYS-001	Meridian Support Assistant	DEPLOYER	Anthropic Claude Haiku 3.5	Customer-facing support chatbot; ticket deflection and FAQ	VP Product	Mar 2025	LIMITED
AI-SYS-002	WorkflowAI Prioritiser	DEPLOYER	OpenAI GPT-4o mini	Internal backlog and task prioritisation recommendations for product teams	Head of Eng.	Jan 2026	MINIMAL
AI-SYS-003	ContractIQ	PROVIDER	Internal (fine-tuned Llama 3.1 70B)	Contract clause extraction, risk flagging, and summary for customer-uploaded vendor contracts	CTO	Jun 2026	MINIMAL
AI-SYS-004	RecruitScreen	PROVIDER	Internal (custom classifier + GPT-4o)	First-pass CV and cover letter screening and ranking for open internal roles	Head of People	Nov 2025	HIGH-RISK
AI-SYS-005	AnomalyGuard	DEPLOYER	Internal (custom ML — Isolation Forest)	Security anomaly detection in platform access logs; flags unusual patterns for human review	Head of Security	Aug 2024	MINIMAL

Role Definitions

PROVIDER

Meridian develops, deploys, or places the AI system on the market under its own name. Provider obligations under the EU AI Act apply — including technical documentation (Art. 11), conformity assessment (Art. 43), and registration (Art. 49) for high-risk systems.

DEPLOYER

Meridian uses an AI system developed by a third-party provider in its own operations or product. Deployer obligations apply — including transparency disclosures (Art. 50), fundamental rights impact assessment for certain uses (Art. 27), and human oversight (Art. 26).

SECTION 2.2

System Detail Sheets

AI-SYS-001

Meridian Support Assistant

DEPLOYER

LIMITED RISK

Provider	Anthropic PBC
Model	Claude Haiku 3.5 (API)
Deployment	Production (govarna platform embedded widget)
Users affected	External customers — ~8,400 active accounts
Data processed	Support messages, account tier, ticket history
Personal data	Yes — customer name, email, account metadata
Human override	Yes — live agent escalation available at all times

Art. 50 Transparency Compliance

- ✓ Users notified they are interacting with an AI system
- ✓ Disclosure shown at conversation start ("AI assistant")
- ✓ Human agent escalation path disclosed
- ✓ System does not generate deepfake or biometric content
- ✓ Model training opt-out in place (Anthropic API — zero training)

AI-SYS-002

WorkflowAI Prioritiser

DEPLOYER

MINIMAL RISK

Provider	OpenAI, L.L.C.
Model	GPT-4o mini (API)
Deployment	Internal — product engineering teams only
Users affected	Internal employees — ~45 engineers and PMs
Data processed	Task titles, descriptions, sprint metadata (no PII)
Personal data	No — task metadata only, no individual identifiers
Output nature	Recommendations only — human decision always required

Risk Notes

- ✓ Does not affect individual legal rights or significant decisions
- ✓ Internal use only — no external impact on customers
- ✓ Outputs are clearly labelled as AI suggestions within the tool
- ✓ No Annex III category match identified
- ✓ OpenAI API zero data retention policy in place (confirmed)

AI-SYS-003

ContractIQ

PROVIDER

MINIMAL RISK

BETA

Architecture	Fine-tuned Llama 3.1 70B (self-hosted, EU region)
Deployment	Beta — selected enterprise customers only (12 accounts)
Data processed	Customer-uploaded vendor contracts (text only)
Output	Clause summaries, risk flags — advisory, not binding
Training data	Public contract datasets (CC-licensed); no customer data

Note: ContractIQ outputs are expressly marked as informational summaries. The product UI includes a disclaimer that outputs should be reviewed by qualified legal counsel before reliance.

- ✓ Not in Annex III scope (no legal decisions; advisory only)
- ✓ Customer data not used for training
- ✓ Technical documentation in progress (Art. 11 — Minimal tier)

AI-SYS-004 · HIGH-RISK CLASSIFICATION

RecruitScreen

PROVIDER

HIGH-RISK

OBLIGATIONS IN PROGRESS

Architecture	Custom classifier + GPT-4o (OpenAI API); internal pipeline
Deployment	Internal HR use — 3 recruiters + Head of People
Data processed	CVs, cover letters, professional history, contact info
Output	Ranked shortlist + automated recommendation score
Obligation target	Full compliance by 31 July 2026

Pending Obligations (Art. 9–15)

- ✗ Technical documentation (Art. 11) — *drafting*
- ✗ Conformity assessment (Art. 43) — *scheduled Aug*
- ✓ Human oversight — recruiter reviews all outputs
- ✓ Logging of inputs/outputs enabled
- ✓ Accuracy/bias testing completed (internal)
- ✗ EU database registration (Art. 49) — *pending*

AI-SYS-005

AnomalyGuard

DEPLOYER

MINIMAL RISK

Architecture	Isolation Forest ML model (scikit-learn; self-hosted)
Deployment	Production — security monitoring pipeline
Data processed	Anonymised access logs, IP geohash, session durations
Output	Anomaly score and alert flag — security team reviews all alerts

- ✓ No biometric data processed
- ✓ No individual profiling for law enforcement purposes
- ✓ Outputs always reviewed by human security analyst
- ✓ Not in Annex III scope

SECTION 3

EU AI Act Risk Classifications

SECTION 3.1

Classification Methodology

Each AI system in the register was assessed through Govarna's deterministic classification wizard. The wizard evaluates each system against the criteria in EU AI Act Articles 5 and 6 and Annex III in sequence. The same inputs always produce the same output; the reasoning for each step is recorded and attached to the classification report. These classifications are **suggested assessments** and should be reviewed with qualified legal counsel before treating them as binding determinations.

Classification Decision Tree

- ① **Prohibited practice check (Art. 5):** Does the system engage in manipulation, social scoring, real-time biometric surveillance, or other prohibited practices? → If yes: **PROHIBITED** — system cannot lawfully operate.
- ② **High-risk scope check (Art. 6 + Annex III):** Is the system a safety component of a regulated product (Art. 6(1))? Or does it fall within one of 8 Annex III domains (e.g., employment, education, credit, essential services)? → If yes: **HIGH-RISK**.
- ③ **Limited-risk transparency check (Art. 50):** Is the system a chatbot, emotion recognition system, deep-fake generator, or AI-generated content system? → If yes: **LIMITED RISK** — transparency obligations apply.
- ④ **Minimal risk default:** All remaining AI systems. No mandatory obligations under the EU AI Act, though voluntary codes of conduct encouraged.

SECTION 3.2

Per-System Classification Reports

AI-SYS-001 · Classification Report

Meridian Support Assistant

LIMITED RISK — ART. 50

Annex III Screening — All 8 Domains

§1 Biometric No match	§2 Critical infra. No match	§3 Education No match	§4 Employment No match
§5 Essential svcs. No match	§6 Law enforce. No match	§7 Migration No match	§8 Justice/demo. No match

Reasoning

System is a customer-facing conversational AI (chatbot) that answers support questions. It does not fall within any Annex III domain. However, as a chatbot that may not be immediately recognisable as AI by end users, Art. 50(1) transparency requirements apply. The system already displays "AI assistant" at conversation start. No Annex III match; no Art. 5 violation. Classification: **Limited risk**.

✓ DETERMINISTIC

Art. 50(1)

Assessed: 2 Jul 2026

E. Hartmann (CTO)

AI-SYS-002 · Classification Report

MINIMAL RISK

WorkflowAI Prioritiser

Internal tool used by ~45 employees for backlog prioritisation suggestions. No Annex III match (no impact on individuals' legal or similarly significant rights; purely operational productivity). Not a chatbot (Art. 50 not triggered). Art. 5 not triggered. Classification: **Minimal risk**.

✓ DETERMINISTIC

No Annex III match · Art. 5 clear · Art. 50 not applicable

Assessed: 2 Jul 2026

AI-SYS-003 · Classification Report

MINIMAL RISK

ContractIQ

ContractIQ extracts and summarises clauses from customer-uploaded contracts. Assessed against Annex III §5 (administration of justice, access to essential services): outputs are explicitly informational and advisory; the system does not make binding determinations affecting individuals' legal rights. No Annex III match. Art. 5 clear. Art. 50 not triggered (not a general-public chatbot). Classification: **Minimal risk**.

✓ DETERMINISTIC

Annex III §5 assessed – no binding legal output

Assessed: 2 Jul 2026

AI-SYS-004 · Classification Report · HIGH-RISK FINDING

HIGH-RISK – ANNEX III POINT 4

RecruitScreen

Annex III Screening

§1 Biometric No match	§2 Critical infra. No match	§3 Education No match	§4 Employment MATCH
§5 Essential svcs. No match	§6 Law enforce. No match	§7 Migration No match	§8 Justice/demo. No match

✓ DETERMINISTIC

Annex III §4 – Confirmed

Assessed: 2 Jul 2026 · E. Hartmann

Reasoning — Why §4 Applies

Annex III Point 4 covers AI systems used for "recruitment or selection of natural persons, notably for advertising vacancies, screening or filtering applications, evaluating candidates in the course of interviews or tests."

RecruitScreen directly screens and ranks candidates' CVs, producing a ranked shortlist used by the HR team to determine which candidates progress to interview. This is a direct match to Point 4 criteria. Even though human recruiters make the final decision, the system's output materially influences candidate progression — which is sufficient for the high-risk classification.

Obligations Triggered (Art. 9–15)

- Risk management system (Art. 9)
- Training data governance documentation (Art. 10)
- Technical documentation (Art. 11)
- Record-keeping and logging (Art. 12)
- Transparency to deployers (Art. 13)
- Human oversight measures (Art. 14)
- Accuracy, robustness, cybersecurity (Art. 15)
- EU database registration (Art. 49)

SECTION 4

Policy Register

5 policies active · Last reviewed: 20 June 2026 · All policies are owned, versioned, and cited automatically in questionnaire answers generated by the Govarna platform.

POLICY ID	TITLE	VERSION	OWNER	STATUS	LAST REVIEWED	FRAMEWORKS
POL-AI-01	AI Acceptable Use Policy	v2.1	CTO	ACTIVE	15 Jun 2026	EU AI Act Art. 13 NIST GOVERN
POL-AI-02	AI Vendor Assessment Policy	v1.3	CISO	ACTIVE	20 May 2026	EU AI Act Art. 28 NIST MAP
POL-AI-03	AI System Documentation Standard	v1.1	CTO	ACTIVE	1 Jun 2026	EU AI Act Art. 11 ISO 42001 Cl. 8
POL-AI-04	AI Incident Response Procedure	v1.0	CISO	ACTIVE	15 May 2026	EU AI Act Art. 73 NIST MANAGE
POL-AI-05	Human Oversight & Review Policy	v1.2	VP Product	ACTIVE	20 Jun 2026	EU AI Act Art. 14 NIST GOVERN ISO 42001 Cl. 8.4

SECTION 4 · CONTINUED

Policy Summaries

POL-AI-01 AI Acceptable Use Policy v2.1

ACTIVE

Defines permitted and prohibited uses of AI tools by Meridian employees and contractors. Covers: mandatory disclosure when AI is used in external communications; prohibition on using AI to generate content that misrepresents facts; prohibition on inputting confidential customer data into non-approved AI tools; mandatory human review before any AI output is shared externally. Applies to all employees and contractors. Reviewed annually or when significant new AI tools are adopted.

POL-AI-02 AI Vendor Assessment Policy v1.3

ACTIVE

Establishes the process for evaluating third-party AI systems before adoption. Requires: vendor AI questionnaire completion; data processing agreement review; confirmation of zero-training data policy or explicit opt-out; review of sub-processor disclosure; security review by CISO for systems handling customer data. Covers all third-party AI APIs and tools processing Meridian or customer data.

POL-AI-03 AI System Documentation Standard v1.1

ACTIVE

Sets the minimum documentation requirements for all AI systems in the Meridian AI register. Each system must maintain: system purpose and scope; training data sources and categories; intended user population; performance metrics and accuracy benchmarks; human oversight mechanism; known limitations and failure modes; version history and change log. High-risk systems additionally require conformity assessment documentation.

POL-AI-04 AI Incident Response Procedure v1.0

ACTIVE

Defines the procedure for identifying, containing, and reporting AI-related incidents, including unexpected AI outputs, bias incidents, data exposure via AI systems, and AI system failures. Severity classification follows a 3-tier model (P1: immediate containment; P2: 24h response; P3: standard sprint response). Serious incidents involving high-risk systems are reported to relevant supervisory authority within 15 working days per Art. 73.

POL-AI-05 Human Oversight & Review Policy v1.2

ACTIVE

Mandates that all AI system outputs that influence a consequential decision about a person or external communication must pass through a designated human reviewer before action is taken. Defines: reviewer qualification requirements by system; override rights and documentation requirements; escalation path when AI output is uncertain or contested; logging requirements for reviewer identity and decision outcome. Applies to all 5 systems in the register.

SECTION 5.1

NIST AI Risk Management Framework — Control Mapping

Mapped against NIST AI RMF 1.0 (January 2023) · 22/22 applicable controls implemented or in progress · Last assessed: 2 July 2026

GOVERN — ESTABLISHING AI RISK GOVERNANCE

CONTROL ID	DESCRIPTION	EVIDENCE	STATUS
GOVERN 1.1	AI risk management is integrated into overall organisational risk management	POL-AI-01 §2	✓ DONE
GOVERN 1.2	Roles and responsibilities for AI risk management are documented	POL-AI-01 §4	✓ DONE
GOVERN 1.4	Organisational teams are committed to AI risk management culture	AI Literacy training — Jun 2026	✓ DONE
GOVERN 2.1	AI risk tolerances and acceptable risk thresholds are documented	Risk Register v1.1	✓ DONE
GOVERN 2.2	Accountability for AI risk management is clearly assigned	POL-AI-01 §4 · RACI matrix	✓ DONE
GOVERN 5.1	Organisational AI policies include processes for workforce AI risk awareness	POL-AI-01 §6 · Training records	✓ DONE

MAP — CATEGORISING AND UNDERSTANDING AI RISKS

CONTROL ID	DESCRIPTION	EVIDENCE	STATUS
MAP 1.1	Context is established for AI systems, including intended purpose and users	AI System Register · AI-SYS-001-005	✓ DONE
MAP 1.5	Organisational risk tolerance is reflected in AI deployment decisions	POL-AI-02 · Vendor risk gate	✓ DONE
MAP 2.2	Scientific and domain knowledge relevant to AI system is documented	POL-AI-03 · System doc sheets	✓ DONE
MAP 3.1	Potential harms and benefits to people and society are identified	Impact assessment · EU AI Act screening	✓ DONE
MAP 3.5	AI risks and impacts are prioritised based on assessed likelihood and severity	Risk Register v1.1 · Section 3	✓ DONE

SECTION 5.1 · CONTINUED

MEASURE – ANALYSING AND ASSESSING AI RISKS

CONTROL ID	DESCRIPTION	EVIDENCE	STATUS
MEASURE 1.1	AI risk testing plans and schedules are established	POL-AI-03 §5 · Testing log	✓ DONE
MEASURE 2.2	Evaluation metrics for AI system performance are defined and tracked	Perf. dashboards · Quarterly review	✓ DONE
MEASURE 2.5	AI system failure modes are identified and monitored	POL-AI-04 · Incident log	✓ DONE
MEASURE 4.1	AI deployment decisions are informed by risk measurement outputs	EU AI Act screening · Risk Register	✓ DONE

MANAGE – PRIORITISING AND ADDRESSING AI RISKS

CONTROL ID	DESCRIPTION	EVIDENCE	STATUS
MANAGE 1.1	Risk response plans are documented for identified AI risks	POL-AI-04 · RecruitScreen plan	✓ DONE
MANAGE 2.4	AI incident response procedures are tested and updated	POL-AI-04 · Tabletop: May 2026	✓ DONE
MANAGE 3.1	AI risks from third-party systems are managed via vendor assessments	POL-AI-02 · Vendor register	✓ DONE
MANAGE 4.1	Risk response actions are tracked to completion	Govarna platform audit log	✓ DONE

SECTION 5.2

ISO/IEC 42001:2023 — AI Management System Controls

ISO/IEC 42001 is the international standard for AI management systems. Meridian has mapped its controls against the standard as preparatory work. This is not a certification claim; formal certification requires independent third-party assessment.

CLAUSE	REQUIREMENT	EVIDENCE / IMPLEMENTATION	STATUS
Cl. 4.1	Understanding the organisation and its context for AI	AI system register; stakeholder analysis doc (Jun 2026)	
Cl. 4.2	Understanding needs and expectations of interested parties	Customer DPA; EU AI Act obligations map	
Cl. 5.1	Leadership and commitment to AI management system	CTO sign-off on all policies; board briefing May 2026	
Cl. 5.2	AI policy established, communicated, and maintained	POL-AI-01 v2.1 — published to all staff Jun 2026	
Cl. 6.1	Actions to address AI risks and opportunities	Risk Register v1.1; EU AI Act classification outputs	
Cl. 6.2	AI objectives and planning to achieve them	Q3 2026 OKR: RecruitScreen compliance by 31 Jul	
Cl. 7.2	Competence requirements for AI-related roles	Training programme; completion records (Jun 2026)	
Cl. 7.3	Awareness of AI policy and risks across the organisation	All-hands AI briefing Jun 2026; e-learning records	
Cl. 8.1	Operational planning and control for AI management	AI System Register; deployment approval workflow	
Cl. 8.2	AI risk assessment conducted and documented	Classification reports (this package, Section 3)	
Cl. 8.3	AI risk treatment — controls selected and applied	Control mapping (this package, Section 5); POL-AI-04	
Cl. 8.4	AI system impact assessment	EU AI Act screening; fundamental rights impact (FRA checklist)	
Cl. 9.1	Monitoring, measurement, analysis, and evaluation	Quarterly performance reviews; anomaly monitoring (AnomalyGuard)	
Cl. 9.2	Internal audit of AI management system	Scheduled: September 2026 (internal audit team)	
Cl. 9.3	Management review of AI management system	Board review scheduled: October 2026	
Cl. 10.1	Continual improvement of the AI management system	Govarna platform tracks open findings; quarterly review cycle	
Cl. 10.2	Nonconformity and corrective action	POL-AI-04 §4; RecruitScreen remediation in progress	
Annex A	Controls — Responsible AI (transparency, accountability, safety)	POL-AI-01, -03, -05; Art. 50 disclosures; human oversight logs	

SECTION 6

Questionnaire Answer Library

Library Overview

47

APPROVED ANSWERS

8

QUESTIONNAIRES COMPLETED

312

TIMES ANSWERS REUSED

Each answer below was drafted by Govarna grounded in Meridian's indexed policies and AI register. All answers were reviewed and approved by a named authoriser before joining the library. Source citations show which policy document, AI system record, or evidence item the answer draws from.

Q1.1 Do you maintain an inventory of AI systems used in your organisation?

Yes. Meridian maintains a live AI system register covering all production and beta AI systems, each with a unique system ID, named owner, deployment date, intended purpose, data categories processed, third-party model provider (where applicable), and EU AI Act risk classification. The register is reviewed and updated quarterly and whenever a new AI system is deployed. As of this assessment, the register contains 5 AI systems at various risk tiers. The register is maintained in Govarna, which provides an immutable audit trail of all changes.

✓ GROUNDED

Source: AI System Register

POL-AI-01 §3

POL-AI-03 §1

Used in 7 questionnaires

Q1.4 Do you use AI models developed by third parties? If so, which providers?

Yes. Meridian uses AI models from the following third-party providers: (1) Anthropic PBC — Claude Haiku 3.5 via API for the customer support assistant; (2) OpenAI, L.L.C. — GPT-4o mini via API for internal task prioritisation and as a component of the RecruitScreen pipeline. Both providers are contracted under data processing agreements that include zero model training on Meridian or customer data. Sub-processor documentation is available on request.

✓ GROUNDED

Source: AI-SYS-001

AI-SYS-002

AI-SYS-004

POL-AI-02

Used in 6 questionnaires

Q2.1 Is your organisation subject to the EU AI Act? In what role — provider or deployer?

Yes. Meridian is subject to the EU AI Act both as a **provider** (for AI systems developed internally, specifically AI-SYS-003 ContractIQ and AI-SYS-004 RecruitScreen) and as a **deployer** (for AI systems built on third-party models — AI-SYS-001, -002, -005). The EU AI Act applied to Meridian's general-purpose AI model usage from August 2025, and core obligations under Articles 6–51 apply from 2 August 2026. Our classifications and obligations analysis are attached in Section 3 of this package.

✓ GROUNDED

Source: Classification reports

EU AI Act Art. 3, 6, Annex III

Used in 5 questionnaires

SECTION 6 · CONTINUED

Q2.3 Do you have any high-risk AI systems under the EU AI Act? How are you managing compliance?

Yes, one system — RecruitScreen (AI-SYS-004) — is classified as high-risk under Annex III Point 4 (AI in employment and worker management). It is used internally for CV screening for open roles. We are completing the full Article 9–15 compliance programme by 31 July 2026, which includes: technical documentation (Art. 11), conformity assessment (Art. 43), EU database registration (Art. 49), and enhanced human oversight procedures. The system operates under continuous human review — all outputs are reviewed by a qualified recruiter before any candidate progression decision is made. No new roles are being processed through the system until compliance is complete.

✓ GROUNDED

Source: AI-SYS-004 classification report

Annex III §4

Art. 9–15

Used in 5 questionnaires

Q3.1 Do you train AI models on customer data?

No. Meridian does not train AI models on customer data. For third-party models (Anthropic Claude, OpenAI GPT), our API agreements include explicit data processing terms prohibiting use of inputs for model training. For internally developed models, training data is derived exclusively from licensed public datasets — no customer data is used. This is documented in our AI Vendor Assessment Policy (POL-AI-02) and confirmed in our data processing agreements with Anthropic and OpenAI, copies of which are available to security reviewers on request.

✓ GROUNDED

Source: POL-AI-02 §3

Anthropic DPA

OpenAI DPA

Used in 8 questionnaires

Q3.4 How do you ensure human oversight of AI system outputs?

Meridian's Human Oversight & Review Policy (POL-AI-05 v1.2) requires that all AI system outputs influencing a consequential decision about a person, or any external-facing communication, pass through a designated human reviewer before action is taken. In practice: the Support Assistant (AI-SYS-001) routes any unresolved ticket to a human agent; WorkflowAI (AI-SYS-002) outputs are labelled as suggestions and require engineer acceptance; ContractIQ (AI-SYS-003) outputs include a mandatory "review with counsel" disclosure; RecruitScreen (AI-SYS-004) outputs are reviewed by a qualified recruiter for every candidate; AnomalyGuard (AI-SYS-005) alerts are reviewed by a security analyst before any action is taken. Reviewer identities and decisions are logged in the Govarna audit trail.

✓ GROUNDED

Source: POL-AI-05

EU AI Act Art. 14

NIST GOVERN 1.2

Used in 7 questionnaires

Showing 6 of 47 approved answers. Full library available in the Govarna platform.

SECTION 7

Audit Trail

The Govarna platform maintains an append-only audit log of all changes to compliance artefacts. The log below covers the 30-day period preceding package generation. All timestamps are UTC. The full log is available for authorised reviewer access via the Govarna platform.

Timestamp	Actor	Action
2026-07-04 09:14 UTC	System	Audit evidence package generated and exported · GV-AUD-2026-0704-MWL · Approved by E. Hartmann (CTO)
2026-07-02 14:33 UTC	E. Hartmann	AI system register reviewed and confirmed current · 5 systems · No changes required
2026-07-02 14:01 UTC	E. Hartmann	Risk classification confirmed for AI-SYS-001 (Limited), AI-SYS-002 (Minimal), AI-SYS-003 (Minimal), AI-SYS-005 (Minimal) · Art. 6 + Annex III
2026-06-28 11:20 UTC	S. Varga (CISO)	Questionnaire answer approved: Q3.4 (Human oversight) · Added to library · Answer ID: ANS-047
2026-06-25 09:45 UTC	E. Hartmann	RecruitScreen (AI-SYS-004) high-risk classification confirmed · Remediation plan approved · Target: 31 Jul 2026
2026-06-20 16:11 UTC	M. Patel (VP Product)	POL-AI-05 Human Oversight Policy updated to v1.2 · Changes: added per-system oversight procedures · POL-AI-05
2026-06-15 10:02 UTC	E. Hartmann	POL-AI-01 AI Acceptable Use Policy reviewed and confirmed current · No changes · v2.1 reaffirmed
2026-06-10 14:55 UTC	S. Varga (CISO)	AI vendor assessment completed for Anthropic (AI-SYS-001) · Zero-training confirmed · DPA updated · POL-AI-02
2026-06-08 11:30 UTC	System	ContractIQ (AI-SYS-003) added to AI system register · Beta deployment approved · AI-SYS-003
2026-06-01 09:15 UTC	E. Hartmann	POL-AI-03 AI System Documentation Standard updated to v1.1 · Added high-risk system documentation requirements
2026-05-28 16:40 UTC	L. Fischer (Recruiter)	RecruitScreen output reviewed — candidate shortlist for Senior Engineer role · Human decision: shortlist approved with 2 additions · AI-SYS-004
2026-05-22 10:18 UTC	S. Varga (CISO)	AI incident response tabletop exercise completed · No gaps identified · POL-AI-04 confirmed adequate
2026-05-20 14:22 UTC	S. Varga (CISO)	POL-AI-02 AI Vendor Assessment Policy reviewed · Updated OpenAI sub-processor entry · v1.3 approved
2026-05-15 09:00 UTC	S. Varga (CISO)	POL-AI-04 AI Incident Response Procedure approved · v1.0 · Board notified
2026-05-10 11:45 UTC	E. Hartmann	Q1 2026 AI governance review completed · AI literacy training deployed to all 320 employees · Completion rate: 94%

About this audit trail

The Govarna platform maintains an append-only audit log stored separately from mutable application data. Log entries cannot be edited or deleted. Each entry records: action type, actor (user ID + display name), affected artefact ID, timestamp (UTC), and the artefact state before and after the change. The log is available for authorised reviewer access via signed URL, time-limited to 72 hours. If your security review requires direct log access, contact accounts@govarna.com with your reviewer credentials.

LEGAL NOTICE

Disclaimer and Terms of Use

This document is a compliance workflow tool output — not legal advice.

This evidence package is generated by the Govarna platform to assist Meridian Workflows Ltd in organising and communicating its AI governance posture to third parties such as enterprise buyers, security review teams, and auditors. It is provided as a compliance workflow tool output only.

1. No Legal Advice

Nothing in this document constitutes legal advice. The risk classifications, framework mappings, policy summaries, and control assessments contained herein are informational and should be reviewed by qualified legal counsel before being relied upon for regulatory compliance, legal defence, or any other legally consequential purpose.

2. No Certification Claim

This package does not constitute or evidence any regulatory certification, conformity assessment, or formal approval by any regulatory authority. References to EU AI Act, NIST AI RMF, and ISO/IEC 42001 describe Meridian's internal mapping work and preparatory compliance activities, not formal certification against those standards or frameworks.

3. Suggested Classifications

EU AI Act risk classifications in Section 3 are suggested assessments produced by Govarna's deterministic classification wizard. They are reproducible from documented criteria and reasoning, but they are not legal determinations. The final determination of an AI system's risk tier under the EU AI Act is a legal matter that requires qualified counsel review, particularly for systems near category boundaries or in novel deployment contexts.

4. Point-in-Time Snapshot

This package reflects the state of Meridian Workflows Ltd's AI governance programme as at **4 July 2026 at 09:14 UTC**. The underlying data may have changed since generation. For the most current state, authorised reviewers may request a freshly generated package from Meridian's security team.

5. Sample Data Disclosure

This specific package (GV-AUD-2026-0704-MWL) is a **sample demonstration document** produced by Govarna using fictional company data to illustrate the format and content of a real Govarna audit evidence package. The company "Meridian Workflows Ltd," all named personnel, AI system IDs, policy contents, and questionnaire answers in this document are entirely fictional. This sample is for evaluation of the Govarna platform only and should not be submitted in any real compliance or regulatory context.

6. About Govarna

Govarna is a software platform for AI governance, questionnaire automation, and EU AI Act compliance readiness. Govarna is not a law firm, a compliance certification body, or a regulatory authority. The platform is built to help security and compliance teams organise evidence and answer questions faster — it does not replace qualified legal or compliance professionals.



Govarna Platform · govarna.com · sales@govarna.com

© 2026 Govarna. This document may be shared with the intended recipient security review team. It may not be published, modified, or redistributed without authorisation from Govarna and Meridian Workflows Ltd.