



GRUND INSTITUTE

COGNITION BY DESIGN

GRUNDMIND / TRUST & COMPLIANCE

PUBLIC

Responsible AI & Acceptable Use Standard

GRUNDMIND

A Grund Institute company · Switzerland

GM-AI-001 · Version 0.2 · 9 August 2026

Document ID	GM-AI-001
Version	0.2 - Branded working draft
Owner	GrundMind / Grund Institute
Review status	Public responsible-use standard describing GrundMind intended purpose, prohibited uses, deterministic controls, AI-assisted functions and human-oversight boundaries. It is not a formal legal classification or certification.



1. Purpose

This standard explains how GrundMind is designed to use automation and AI responsibly in AI-adoption diagnostics, how human review is preserved, and which uses of GrundMind results are outside the product's permitted purpose.

CORE POSITION	Advisory evidence, not automated employment decision-making GrundMind is designed to help organizations understand AI value, workflow fit, enablement, governance and development needs. It is not designed to make or recommend consequential employment decisions about individuals.
----------------------	--

2. Responsible AI principles

Principle	How GrundMind applies it
Purpose limitation	The product is built for AI adoption support, workflow improvement, enablement, governance and development - not general-purpose employee evaluation.
Evidence before conclusions	Scores and classifications require defined evidence thresholds; insufficient evidence is surfaced rather than converted into a negative result.
Human oversight	AI-assisted narrative is draft-only until an authorized human reviews and explicitly publishes it.
Separation of evidence sources	Manager and employee evidence remain separate; contradictions are analyzed instead of averaged away.
Privacy by reporting layer	Standard Team and Corporate reporting is aggregated; named operational support is deliberately restricted.
No hidden ranking	
Traceability	Survey versions, evidence states, result snapshots and publication states are retained according to the defined architecture and retention policy.
Challenge and interpretation	Results are decision-support evidence that should be interpreted in organizational context and can be challenged/reviewed rather than treated as infallible truth.

3. Intended uses

- Diagnosing where AI is creating value, where value is not being realized, and which mechanisms may explain the gap.
- Identifying workflow friction, enablement gaps, governance/trust issues and role/capability needs.
- Understanding Human-AI interaction patterns for adoption support and development.
- Comparing manager and team perceptions to identify alignment, contradiction or definition mismatch.
- Prioritizing training, workflow redesign, governance clarification, support and remeasurement actions.
- Supporting leadership and board-level accountability for AI adoption and value realization using evidence with provenance.



4. Prohibited uses

RESTRICTED PURPOSE	Employment decisions are outside the permitted product purpose GrundMind is intended for organisational AI adoption, workflow improvement, enablement, governance, capability development and developmental support. It is not designed or permitted for individual employment decision-making. GrundMind outputs must not be used to make, recommend, rank, score or materially influence decisions concerning an individual's recruitment or hiring, promotion, compensation, performance evaluation, disciplinary action, termination, selection for redundancy or headcount reduction, or allocation of work, tasks, shifts or responsibilities. GrundMind may identify organisational, workflow, capability or structural changes for consideration, but its outputs must not be used to determine which individuals should be selected for employment-related actions. All findings require human interpretation and may be used only within GrundMind's stated organisational and developmental purposes.
---------------------------	--

- Do not use a GrundMind score or archetype as a proxy for employee quality, intelligence, loyalty, potential or job performance.
- Do not create league tables of employees or managers from Individual Fit scores, archetypes or support needs.
- Do not use a result to make, recommend, rank, score or materially influence allocation of work, tasks, shifts or responsibilities, compensation, promotion opportunities, employment status or other individual employment decisions.
- Do not infer misconduct, resistance, incompetence or intent merely from a low Fit score, AI-averse interaction pattern, missing evidence or manager-team disagreement.
- Do not bypass the product reporting boundaries by exporting raw responses or named evidence for purposes that the standard reports intentionally exclude.

5. Deterministic core vs AI-assisted functions

Function	Method	Human / governance boundary
Dimension/subdimension scoring	Deterministic predefined rules and weights	Not delegated to a generative model.
Evidence sufficiency	Deterministic numeric and subdimension thresholds	Missing/not-applicable evidence is not treated as zero.
Cognitive Archetype assignment	Deterministic four-archetype model	Interaction-style signal only; not personality, ability or performance classification.
Missing ROI diagnosis	Deterministic evidence-based mechanism signal	Diagnostic signal strength, not a calculation of financial ROI lost and not percentages that sum to 100.



Function	Method	Human / governance boundary
Manager-Team Alignment	Deterministic source comparison/classification	Manager and employee evidence remain distinct; the system does not automatically decide which source is “right”.
Narrative draft / summarization	May be AI-assisted	Draft remains separate from client-facing published interpretation until human review and explicit publication.
Final organizational action	Human decision	GrundMind can recommend support/priorities; authorized humans remain responsible for decisions and context.

6. Reporting safeguards

Reporting layer	Responsible-use boundary
Individual Result	Restricted individual diagnostic context; not an employment-performance record.
Team Result	Aggregated by default. Standard Team Results do not expose named individual scores, named archetypes, raw answers, free-text evidence or employee ranking.
Manager Result	Manager evidence remains a separate source; it does not replace or overwrite employee Team evidence.
Manager-Team Alignment	Contradictions are treated as diagnostic gaps such as manager overestimation, employee overestimation, definition mismatch or insufficient evidence.
Restricted Manager Support	A named team member may be shown only where operationally necessary to provide support, and only with the immediate support need and recommended manager actions. No named Fit score, strongest fit, archetype, raw/free-text evidence, rationale or ranking.
Corporate Result	Organization-level synthesis; preserves provenance and uses aggregated team evidence.

7. Human-AI interaction archetypes

GrundMind currently uses four official Human-AI interaction archetypes: Co-thinker, Calibrator, Controller and AI-averse. These labels describe interaction patterns within the assessment model; they are not diagnoses, personality types, intelligence categories or performance grades.

- Archetypes must be interpreted alongside workflow, business-value, governance/trust and role/capability evidence.
- Team-level distributions may be useful for enablement and support design.
- Named individual archetypes are not part of standard manager reporting or the restricted Manager Support output.
- AI-averse is not synonymous with poor performance or misconduct; it signals an interaction/adoption pattern that may have multiple legitimate causes.



8. Missing ROI safeguards

Missing ROI is a deterministic survey-evidence interpretation designed to indicate the strength of potential value-loss mechanisms. It is not a financial ROI calculation, does not state the percentage of money lost, and its mechanism scores are not intended to add to 100%. Evidence coverage indicates whether mapped evidence is available, not certainty that a causal explanation is true.

- Possible outputs include a dominant mechanism, mixed/no dominant mechanism, or insufficient/unmapped evidence.
- Mechanism findings should be treated as hypotheses/priorities for investigation and intervention, not proof of financial causality.
- Remeasurement should test whether the selected intervention changed the relevant evidence pattern and business outcome.

9. AI-assisted narrative and external AI providers

Where GrundMind uses a generative AI provider to assist drafting, summarization, translation or report preparation, that provider is an assistance layer rather than the authoritative scoring engine. The external processing path must be approved for the relevant data before identifiable enterprise participant evidence is sent to it.

1. Minimize or pseudonymize evidence shared with an external AI provider wherever reasonably possible.
2. Use only approved provider accounts/contractual paths for Customer Personal Data.
3. Do not treat model output as evidence that was not present in the source data.
4. Review narrative for factual grounding, proportionality, bias, unsupported causality and prohibited employment-use implications.
5. Require explicit human publication before client-facing interpretation becomes authoritative within the report workflow.

10. Manager-Team disagreement

A disagreement between manager and employee/team evidence is itself useful organizational evidence. GrundMind should preserve the disagreement and classify the pattern rather than collapse it into a single average. Typical patterns include aligned positive, aligned negative, manager overestimation, employee overestimation, definition mismatch and insufficient evidence.

IMPORTANT BOUNDARY

A perception gap is not automatic proof that either employees or managers are wrong. Human review should consider access to information, definitions, role context, sample size and the specific evidence behind the gap.

11. Human oversight responsibilities

Role	Responsibility
GrundMind product / method owner	Maintain scoring/evidence rules, restricted-use boundaries, change control and product documentation.
GrundMind analyst / reviewer	Review AI-assisted interpretation against the evidence; remove unsupported or harmful claims; explicitly publish final narrative.
Customer sponsor / controller	Define the lawful organizational purpose, participant scope, access recipients, downstream use and any employment/privacy consultation requirements.
Manager / report recipient	



Role	Responsibility
Participant	Receive appropriate transparency about the assessment purpose, reporting layers and privacy rights; raise questions or challenge inaccurate contextual information through the applicable process.

12. EU AI Act and regulatory boundary

The EU AI Act classifies systems based substantially on intended purpose and specific use. Employment and worker-management use cases can fall within the Act's high-risk categories, and profiling of natural persons can be especially significant to classification. GrundMind therefore defines and enforces a restricted advisory purpose rather than permitting employment decision-making. This design position is a risk-control and intended-use boundary; it is not a formal legal determination that every GrundMind deployment is outside high-risk classification.

- Customers must not expand the intended purpose into prohibited or high-impact employment decision uses without a separate legal/regulatory assessment.
- Human oversight is a design requirement, not a decorative disclaimer: client-facing narrative publication and downstream organizational decisions remain human-controlled.
- Workplace deployments may trigger separate employee-information, consultation, data-protection or fundamental-rights obligations depending on jurisdiction and use.

13. Transparency, challenge and correction

- Explain the assessment purpose and reporting layers to participants before or at collection.
- Use evidence-coverage states so recipients can distinguish a measured finding from insufficient evidence.
- Allow factual/contextual inaccuracies to be raised and reviewed through the relevant customer/GrundMind process.
- Correct configuration/data errors and preserve appropriate audit/provenance records of material report changes.
- Do not present an AI-generated explanation as if it were a participant quote or verified fact unless the source evidence supports it.

14. Method and model change control

Material changes to survey revisions, scoring weights, evidence thresholds, archetype logic, Missing ROI mappings, alignment logic or AI-assisted reporting behavior should be versioned and documented. Historical snapshots should remain attributable to the method/version that generated them rather than being silently recalculated under a new method.

15. Reference framework

- Regulation (EU) 2024/1689 (EU AI Act), including the risk-based classification framework, employment/worker-management use cases and human-oversight concepts.
- European Commission AI Act guidance / Single Information Platform, including current guidance on intended purpose and high-risk use cases.
- GrundMind Trust & Compliance Overview, TOMs, DPIA Support Pack, Data Flow & Architecture Overview and restricted reporting boundaries.

IMPORTANT BOUNDARY

References to the EU AI Act, GDPR, Swiss FADP or standards describe intended safeguards, design boundaries and current readiness work. They are not a claim that GrundMind has received formal regulatory classification, legal approval, certification or external assurance.