

Cognitive Health Assessment for Responsible Technology

A human-centred AI maturity framework. Five dimensions, five levels, each anchored in human behaviour rather than technical capability.

Julie Hendry, MSc Psychology (Distinction), FBCS. First drafted 16 May 2026; this published edition August 2026. [Free 15-question self-assessment](#)

The core thesis

Technology transformation fails because organisations treat it as a technical problem when it is a human one. AI adoption is following the same pattern: organisations deploy AI systems and measure success by technical metrics, accuracy, latency, cost per inference, while ignoring the human layer: how people think alongside AI, how decisions shift when automation is present, how cognitive load changes, how autonomy erodes, and how accountability diffuses when "the algorithm decided".

The established AI maturity models start from the technology and work outward. CHART starts from the human and works inward.

Why this framework, why now

The EU AI Act places human oversight (Article 14) at the centre of governance for high-risk AI. Under the Digital Omnibus, formally adopted on 29 June 2026, those high-risk obligations apply from 2 December 2027 for standalone systems and 2 August 2028 for AI embedded in regulated products. Brussels delayed enforcement precisely because organisations were not ready. The delay is not a reprieve; it is the readiness window. But "human in the loop" is still being treated as a compliance checkbox, not a cognitive reality.

No commercial assessment tool asks the question CHART answers: are the humans in your AI systems actually capable of exercising the judgment you are relying on them for? The existing maturity models come from one discipline each: IAPP from privacy law, IEEE from engineering, ISACA from audit, CMMI from process. None starts from what happens to human thinking when you put AI in the room. CHART sits at the intersection of three domains rarely held by the same person: enterprise technology transformation, the psychology of human cognition, and AI ethics and governance.

The five dimensions

Each dimension has five levels. Level claims are meant to be supported by evidence, not aspiration.

1. Cognitive Readiness

Do people in this organisation understand how AI changes the way they think, decide, and act?

This is not "AI literacy" in the sense of knowing what a large language model is. It is cognitive literacy: understanding automation bias, the attention costs of switching between AI and human judgment, and how

decision fatigue interacts with AI recommendations. A critical finding from the literature: AI bias transfers to humans and persists even after the AI is removed (Vicente & Matute, 2023). Cognitive readiness is therefore not a one-time training exercise; it is an ongoing organisational capability.

Level	Description
1. Unaware	No recognition that AI changes human cognition. AI treated as a tool like any other.
2. Aware	Some individuals recognise cognitive effects (automation bias, over-reliance) but there is no organisational response.
3. Informed	Training includes cognitive impact awareness. Staff can articulate how AI might affect their judgment. Cognitive forcing (forming independent judgment before seeing AI advice) is understood but not yet standard practice.
4. Adapted	Processes are designed for cognitive effects. Cognitive forcing is built into AI-assisted workflows. Attention restoration (task rotation, breaks, environmental variety) is built into work patterns. Two-stage decision paradigms are standard: humans commit to an initial judgment before AI advice is revealed.
5. Optimised	Continuous measurement of human-AI decision quality. Cognitive load and attention depletion are monitored across work sessions. Workflows are iterated on how people actually think alongside AI. AI bias transfer is actively mitigated through periodic independent-decision exercises.

Grounding: Attention Restoration Theory (Kaplan, 1995); automation bias (Parasuraman & Manzey, 2010; Romeo & Conti, 2025); AI bias inheritance (Vicente & Matute, 2023); cognitive forcing (Alon et al., 2025); cognitive load theory (Sweller, 1988); directed attention fatigue (Ohly et al., 2016); dual-process theory (Kahneman, 2011); the be0 research programme on digital self-control and attention replenishment.

2. Decision Architecture

When AI is involved in a decision, can you trace who actually decided, what the human contributed, and whether the human could meaningfully have overridden the system?

"Human in the loop" is a governance fiction in most organisations: the human is in the loop the way a rubber stamp is in the loop. The literature reveals a dual failure mode. Algorithm aversion (Dietvorst et al., 2015) causes disproportionate penalising of AI errors and unwarranted overrides; algorithm appreciation causes uncritical deference. Effective decision architecture must account for both, and analyses by Fink (2025) and Carnat (2024) show that Article 14 assumes effective human oversight that cognitive science says is unfounded in most operational contexts.

Level	Description
1. Opaque	No clarity on where AI recommendations end and human decisions begin.
2. Labelled	AI-assisted decisions are identified, but the human's role is not defined. Both algorithm aversion and algorithm appreciation go unmonitored.
3. Structured	Decision rights are assigned. It is clear who can override AI and under what conditions. Overseers have sufficient understanding of what the system is doing.
4. Exercised	Overrides happen in practice, not just in policy. There is evidence of humans disagreeing with AI and that disagreement being respected. Understanding and genuine ability to intervene are both met, and override patterns are analysed for aversion and appreciation biases.
5. Adaptive	Decision architecture evolves with outcomes. Override accuracy is measured. Where humans consistently override correctly, the system learns; where AI consistently outperforms, the human role shifts to exception handling with full cognitive support. Human oversight can be demonstrated to regulators as a cognitive reality, not a compliance fiction.

Grounding: distributed cognition (Hutchins, 1995); joint cognitive systems (Hollnagel & Woods, 2005); meaningful human control (Santoni de Sio & van den Hoven, 2018); algorithm aversion (Dietvorst et al., 2015); EU AI Act Article 14 critique (Fink, 2025; Carnat, 2024).

3. Organisational Honesty

Does this organisation know what it does not know about its AI systems, and can it say so?

In thirty years of enterprise technology delivery, the single strongest predictor of programme failure is not technical complexity. It is the gap between what an organisation says about its readiness and what is actually true. AI amplifies this because the technology is impressive enough to mask governance gaps: tool deployments create visibility without capability, pilot successes generate confidence without scalability, and hype-driven metrics validate both (the "AI Maturity Mirage", Thomas, 2025). Psychological safety is the precondition for honest self-assessment (Edmondson, 1999), and AI adoption itself has been found to damage it (Kim, Kim & Lee, 2025).

Level	Description
1. Performative	AI governance exists on paper. Policies satisfy regulators or boards but do not reflect practice. Self-assessment scores do not match operational reality.
2. Aspirational	The organisation genuinely wants good governance but overstates its position. Risk registers are optimistic; board reports conflate intent with achievement; leaders and frontline staff hold significantly different views of readiness.
3. Honest	Gaps can be articulated without defensiveness. "We do not know" is an acceptable answer. Risk assessment reflects actual practice. Psychological safety is measured and addressed.
4. Diagnostic	The organisation actively looks for what it is getting wrong. Near-misses are reported. AI failures are analysed for human factors, not just technical causes. High-reliability characteristics are visible: preoccupation with failure, reluctance to simplify, deference to expertise.
5. Self-correcting	Honesty is structural, not cultural. Mechanisms surface disagreement, protect dissent, and prevent motivated reasoning from shaping AI risk assessment. Auditors are independent. Boards hear bad news first. Psychological safety and honest reporting are measured as governance outcomes.

Grounding: psychological safety (Edmondson, 1999; Kim, Kim & Lee, 2025); the AI Maturity Mirage (Thomas, 2025); organisational learning (Argyris & Schön, 1978); high reliability organisations (Weick & Sutcliffe, 2007); direct experience of diagnosing transformation programmes where institutional dishonesty was the root cause of failure.

4. Workforce Transition

Are people being supported through the cognitive and emotional transition that AI creates, or are they being told to adapt and left to cope?

AI transition is harder than ordinary change because it changes what counts as professional expertise. The literature documents this with clinical specificity: six psychological themes of AI-induced displacement, from emotional shock to perceived organisational betrayal (Almutairi, Alessa & Alanzi, 2025); "AI replacement dysfunction" proposed as an emerging clinical construct (Thornton & McNamara, 2026); and "algorithmic anxiety" disrupting the psychological contract by making competence a moving target (Shekhar & Saurombe, 2026). Most organisations manage the change, new tools and processes, without supporting the transition: identity, purpose, mastery (Bridges, 1991).

Level	Description
1. Absent	No acknowledgement that AI changes how people experience their work. Training covers button-clicking, not identity. The psychological contract is disrupted without recognition.
2. Communicated	The organisation talks about AI transition and offers reassurance, but no practical support exists for the psychological adjustment. Risk of perceived betrayal when lived experience contradicts the narrative.
3. Supported	Genuine transition support exists: reskilling with time to learn, safe spaces for uncertainty, managers trained to recognise transition stress. The organisation distinguishes change (external) from transition (internal).
4. Co-designed	Workers affected by AI help design how it is deployed in their area. Their expertise about the work is treated as essential input. Autonomy, competence and relatedness are explicitly considered in deployment design.
5. Flourishing	Human flourishing is measured alongside AI performance. Deployment that improves technical metrics but damages human outcomes is treated as failure. Workforce transition is ongoing, not a one-off programme.

Grounding: transition model (Bridges, 1991); self-determination theory (Deci & Ryan, 2000); job demands-resources theory (Bakker & Demerouti, 2007); AI displacement psychology (Almutairi, Alessa & Alanzi, 2025); AI replacement dysfunction as an emerging construct (Thornton & McNamara, 2026); algorithmic anxiety (Shekhar & Saurombe, 2026); the beò research programme on digital wellbeing.

5. Accountability Architecture

When something goes wrong with an AI system, does this organisation know who is responsible, and would that person agree?

Accountability diffusion is the silent governance crisis of AI. When an AI-assisted process produces a bad outcome, accountability fragments across the data team, the model builders, the deployers, the oversight committee, the board, and the individual who clicked "approve". The scale is documented: 67% of CIOs and CTOs are held accountable for AI systems they do not fully control, and 77% report AI adoption outpacing their governance capability (IBM Institute for Business Value with Oxford Economics, 2026, survey of 2,000

technology executives). Santoni de Sio (2021) identifies four distinct responsibility gaps, culpability, moral accountability, public accountability, and active responsibility, which can exist independently.

Level	Description
1. Diffuse	Nobody is clearly accountable for AI outcomes. Responsibility is distributed across teams, committees and policies until it belongs to no one. All four responsibility gaps are open.
2. Named	Accountability is assigned on paper, but the named individual may lack the authority, knowledge or resources to actually govern the system.
3. Resourced	Accountable individuals have the authority, budget and access to fulfil the role. They can stop or modify an AI system if governance requires it, and are supported when they do. Just culture principles are understood.
4. Tested	Accountability has been exercised under pressure, with evidence of an accountable person intervening and the organisation supporting the intervention. Incidents are investigated for systemic causes, not individual blame. All four gaps are monitored.
5. Systemic	Accountability is embedded in organisational design, not dependent on individual courage. Incentives, reporting lines and escalation paths make accountability the path of least resistance. Provider-deployer responsibility mapping pre-empts the Article 14 accountability loophole.

Grounding: four responsibility gaps (Santoni de Sio, 2021); Swiss cheese model (Reason, 1990); just culture (Dekker, 2012); principal-agent theory; Article 14 accountability loophole (Fink, 2025); IBM Institute for Business Value / Oxford Economics (2026); direct experience of accountability failures in government and enterprise technology programmes.

How the dimensions interact

The five dimensions are interdependent, and assessments that treat them in isolation miss the systemic failure modes. Automation bias (Cognitive Readiness) determines whether decision architectures function as designed: well-designed override mechanisms are performative if people are cognitively disposed to defer. The overestimation problem (Organisational Honesty) is amplified by accountability diffusion: when nobody is personally accountable for the accuracy of governance assessments, there is no incentive to be honest about gaps. Workers experiencing AI-related anxiety and identity threat engage less critically with AI output, undermining the sustained attention that oversight requires. And the Act's oversight requirements create accountability loopholes where responsibility is ambiguous between providers and deployers.

These interactions explain why single-dimension assessments, technical maturity scores, compliance checklists, AI literacy surveys, consistently fail to predict governance outcomes.

How assessment works

CHART is designed to be delivered as a structured external assessment, not a self-assessment questionnaire, for the reason Dimension 3 exists: organisations systematically overstate their own readiness. A full assessment combines documentary review, structured interviews across three organisational layers (what leadership believes, what management implements, what actually happens at the frontline), gap analysis between documented governance and observed practice, evidence-based dimension scoring, and a prioritised recommendations report. It is delivered as a half-day leadership readout plus a detailed written report. The assessment protocol, scoring and calibration are the professional engagement and are not published.

The free [15-question self-assessment](#) gives a directional read across the five dimensions.

Compatibility with existing standards

Standard	Relationship
ISO 42001	CHART assesses the human dimensions that ISO 42001 requires but does not prescribe how to evaluate.
EU AI Act	Directly supports Article 14 compliance by assessing whether human oversight is cognitive reality or governance fiction.
NIST AI RMF	Adds the human behaviour layer that NIST's "human-AI teaming" category acknowledges but does not operationalise.
CMMI	Shares the five-level maturity structure deliberately, so organisations familiar with maturity assessment can adopt CHART without learning a new paradigm.

What this is not

This is not an AI ethics certification, a course, a set of principles, or a compliance checklist. It is a diagnostic framework for assessing how well an organisation's humans are actually governing its AI. The distinction matters: principles tell you what good looks like. CHART tells you how far away you are and why.

Provenance

CHART is the work of Julie Hendry and draws on thirty years of enterprise technology transformation across government (GDS, ESFA, Cabinet Office, HMRC), financial services (JP Morgan, American Express, Barclays, SunGard), life sciences (Abcam), retail (Travis Perkins, Tesco) and media (BBC, Sky, The Guardian); an MSc in Psychology (Distinction) from the University of Strathclyde with research on technology's cognitive impact; a pre-registered research study on digital self-control (be0 LLC, OSF, 2026); the Treasurership of the BCS ICT Ethics Specialist Group; and a pan-government published governance framework for agile service delivery. The research behind CHART is published as an SSRN preprint (DOI: [10.2139/ssrn.6786758](https://doi.org/10.2139/ssrn.6786758)). First drafted 16 May 2026. © Julie Hendry 2026. Openly published for use and citation with attribution; the name CHART and the assessment methodology remain the property of Julie Hendry.

Cite as: Hendry, J. (2026). *CHART: Cognitive Health Assessment for Responsible Technology*. juliehendry.com/CHART/framework.html

Start with the self-assessment

Fifteen questions, five dimensions, an honest read on whether you should be worried. Then, if the result surprises you, [that's the conversation to have](#).

Free self-assessment: juliehendry.com/CHART/self-assessment.html