

SPRINKLING ACT · HUMAN IN THE MAP

HOW THE ANALYSIS IS CONDUCTED

The unit of analysis is one workflow, agreed with you before anything starts. We begin from what the procedure says and what you describe of it, then we look for the elements that would confirm it. Where the two disagree, that is a finding. Where nothing allows us to check, we say so.

CITE THIS METHODOLOGY

The Sprinkling Act methodology is publicly versioned and citable. The current edition is v1.7, August 2026, rendered from this page so that the document and the page cannot diverge. The editions it replaces stay available as published, uncorrected: v1.6 of August 2026, and the April 2026 card — a condensed snapshot of v1.1, regulatory freeze March 2026, cryptographically timestamped via OpenTimestamps and anchored on the Bitcoin blockchain. The August editions are not timestamped, and we make no such claim for them until they are.

[Download the methodology · v1.7, August 2026 \(PDF\)](#)[v1.6, August 2026 \(as published\)](#)[April 2026 card \(v1.1, as published\)](#) [OpenTimestamps proof \(.ots\)](#)[Public mirror · github.com/sprinkling-act/timestamps](https://github.com/sprinkling-act/timestamps)

The established models describe where a person sits inside a decision: in the loop when the system waits for their approval, on the loop when they watch and may intervene, in command when they hold final authority. All three are local, and each describes one control point at a time.

None of them answers the question a director actually asks, which is not who approves this output, but where the humans are across everything now running, and what the whole set costs. That question is about the map, not the loop. Human in the Map is the position it is asked from, and the one an organisation cannot occupy from inside itself.

"What supports what you already know about this workflow?"

Six questions, two minutes, no account, and nothing is sent or stored. It returns your own answers, grouped and readable, as a text you can take away. It concludes nothing: a director is not in a position to judge whether the question concerns them, which is what the survey exists to establish. It shows and it hands back, and it recommends nothing.

"How far is what you describe from what the elements show?"

We agree the scope with you, then take your description of the workflow — you alone, because it is your map we are measuring. We then read whatever elements already exist: an export, a shared calendar, a version history, a folder of dated files. A second conversation sets the two against each other and records your objections. You receive a written report and a sixty-minute readout, five business days from the close of collection.

01

The unit is a distance, not an inventory

The scope is one workflow, named and bounded with you before anything starts. What is measured inside it is the distance between what you describe and what the elements you provide establish. We do not inventory your organisation's AI use: our single source is you, and what you cannot see does not become visible because we ask you better. The workflow is the one you designate — we do not pick it for you by comparing several, because that would require the overview we just said we do not have.

02

Exports, not access

We ask for material your team can share rather than credentials to your systems. Nothing is installed, no agent runs on your side, and nothing is left behind.

03

Declared first, then checked

We start from the procedure and from what the director describes of it — at the Baseline that description is the single source. Then we look for elements that would confirm it. Starting the other way round would tell us what the system was designed to do, not what it does.

04

Every finding says what it rests on

A finding rests on the procedure as written, on elements we could examine, or on an inference we make explicit. The reader always knows which.

05

A gap needs a second source

When the procedure and the elements disagree, we look for a second source of a different kind before presenting the point as established. One source that contradicts another is a question, not yet a conclusion.

06

Open questions stay open

Where nothing available lets us check a point, it is reported as an open question. It is never quietly upgraded into an established fact.

07

No score

We do not produce a grade, a maturity level, or a percentage of an invisible total. A number would hide exactly what the analysis is meant to surface.

08

Dated and bounded

The analysis holds for a scope and a date. Workflows move, tools get added, roles change. We say from what point a finding stops being reliable.

The report sorts every finding by what supports it: confirmed, declared without an independent trace, gap, open question. That rule does not stop at the report.

To the words

The lexicon applies it to its own vocabulary. Each of the twenty-seven terms carries the standing of its source — peer-reviewed, convergent sources, recently formalised, method contested, not yet formalised — and a caveat stating what that source does not allow anyone to claim. Five terms are marked not yet formalised, and they are the ones our own method rests on. They are also the ones we are working to formalise, and the lexicon carries their standing until it changes.

To the sources

The work we cite carries the same mark. A peer-reviewed study can still be unusable if its population does not transpose to a European SME; a study commissioned by a vendor stays citable, provided the commissioner's interest is stated every time it is used, and not only the first.

To your findings

In the report, a gap is only presented as established once a second source, of a different kind, confirms it. Where nothing available allows a point to be checked, it stays an open question and never becomes a fact through an accumulation of hints.

The rule forces us to publish what our own claims are not worth. It is also the only reason to believe what is left.

The same principle decides whether there is a next step

The Baseline ends on one of three outcomes, and one of them concludes that nothing further is warranted. It is written into the method before any report exists, and its conditions are stated in advance rather than formed on reading. The free check makes no such judgement — it returns your own answers and concludes nothing, because a director is not in a position to decide whether the question concerns them. The renunciation therefore happens where it costs us something: on a delivered report, after the sale, not on a questionnaire before it.

Two words for two things

The visibility gap is what is measured: the distance between what a director describes and what the elements establish. Reconstruction is what happens in order to measure it: establishing what takes place from sources that are not people's recollection, then setting the two against each other. The first names the result, the second the method.

This term is defined in the lexicon, with the standing of its source — Reconstruction.

This term is defined in the lexicon, with the standing of its source — Human in the Map.

A widely-held misconception: “if my system is classified high-risk, I need to pay a Notified Body €50–150K.” True for some cases, not all. Article 43 of the AI Act defines two distinct conformity assessment procedures. The vast majority of Annex III systems (HR, credit, education) can take the internal route.

Annex VI · Self-assessment

Applies to most Annex III systems: employment and HR, credit and insurance, education, essential public services, asylum and migration, law enforcement (under conditions), administration of justice. The provider self-assesses against Art. 9–15, produces technical documentation (Art. 11) and signs the EU declaration of conformity (Art. 47). No Notified Body involved.

Indicative cost: €10–50K depending on external support level.

Annex VII · Notified Body

Mandatory for: (1) AI as a safety component of a product already subject to third-party assessment under harmonization legislation (MDR, machinery, toys, vehicles), the AI Act procedure integrates into the product procedure; (2) Annex III §1(a) on remote biometrics, where Annex VII is explicitly required. A Notified Body audits the QMS and technical documentation before placing on the market.

Indicative cost: €50–150K+ per system, on top of the QMS.

Why this distinction matters: the gap between self-certification and third-party audit can represent a 5–10× factor on the compliance bill. Sprinkling Act identifies which route applies to your system in the full report, before you commit a budget.

Sources: [Art. 43 AI Act](#) · [Annex VI](#) · [Annex VII](#) · [Art. 47 \(EU Declaration of Conformity\)](#)

The reconstruction is the work; it is not the deliverable. The deliverable is the layer that makes the reconstruction decidable. Eight elements have to be in it, and a page that carries none of them belongs in the annex.

01 The chain, not the list

Each AI use is placed in its full chain of action, from the human input to the output actually used, with every transition named. A tool outside its chain is a line on an invoice, not a finding.

02 Form and failure at each transition

Each transition carries its form, person to tool, tool to tool, and where it fails, its error type. That crossing is what makes a finding comparable from one file to the next.

03 The named authority, or its absence

For each control point: who decides, who approves, who can halt. The absence of a name is the most frequent finding, and the heaviest. It is also what Article 26(2) requires.

04 An evidence status on every line

Confirmed, declared without an independent trace, gap established, open question. A gap is only presented as established once a second source, of a different kind, confirms it.

05 Rework time, at Full Map

First-pass acceptance rate and rework minutes. They do not appear in the Baseline: estimating them means asking the people who do the work, not the person who decides, and an estimate of someone else's time is not a measurement. Without those two numbers the report stays a description. That is what separates the Baseline from the Full Map, not a shortcoming of the former.

06 The portable share of the skill

What, in what your teams learned, would survive a change of tool. This single line often carries the decision for a director considering a switch.

07 Structural incompatibilities, at Full Map

Where two tools adopted by different departments will produce friction once their data meet. The only finding that bears on the future, and the only one that requires a view from above — which is why it does not appear in the Baseline: a single source, looking from one position, cannot establish it. Where the Baseline glimpses one, it is recorded as an open question, never as an established finding.

08 The expiry date

From when each finding stops being reliable, and why. Not a general validity note: a date per family of findings.

Research on hand-off failures produces a distribution that reverses the assumption most leadership starts from. Close to two thirds of failures are transmission problems. The capability gap, the one supposed first, weighs under seven percent.

SIGNAL CORRUPTION The result transmitted is altered or degraded between sender and receiver.	36.8%
DATA GAP Content that was needed did not make it across the transition.	29.1%
REFERENTIAL DRIFT References become ambiguous: what exactly is being discussed is no longer shared.	27.3%
CAPABILITY GAP The receiver does not have the means to handle what they received.	6.8%

Proportions observed on multi-agent systems. We use them as a reading grid for human-AI chains, not as a prediction of what your own distribution will be.

The practical consequence is a question worth settling before funding anything: is the problem that people do not know, or that what they needed never reached them? The two have almost nothing in common, and one is far cheaper to correct.

Annex III obligations moved to December 2027, and that distance often reads as permission to wait. It is not one, because another text already applies to the same facts. Regulation (EU) 2016/679, articles 28, 30 and 32.

ARTICLE 28

A processor acting on your behalf requires a contract. A tool brought in without going through the organisation may fall under this article, and where it does, the organisation stays responsible for a tool it never authorised. Whether it does is a qualification, and it belongs to your counsel rather than to us.

ARTICLE 30

A record of processing activities is required. A survey does not produce it: the record is the controller's obligation and its completeness is their liability, so a partial one drawn up by a third party would expose them more than an acknowledged gap. What a survey does is document a chain the record probably ignores.

ARTICLE 32

Appropriate technical and organisational measures are required. Undocumented AI data flows fail that standard by construction.

Your shadow AI problem is a data protection problem before it is an AI Act problem: already applicable, already familiar, already enforced.

What this is not. The survey is not a GDPR compliance exercise, not a data protection audit, and Sprinkling Act does not act as your data protection officer. It establishes what runs, where, and on whose authority. What you then do with it, and how it enters your record, belongs to you and to qualified counsel.

What follows is what Sprinkling Act is not. The limits of the method itself are not gathered here: each one is stated where it applies — the single source in the first principle, the unverifiable point in the one on open questions, the two numbers we do not produce in the list of what the report contains. A limit read at the end of a document is a disclaimer; read next to the finding it affects, it is information.

This is not an audit and not a certification. Nothing is checked against a standard, no judgment is issued, and no result is opposable to a third party as proof of anything.

This is not legal advice. Sprinkling Act is not a law firm, not a Notified Body and not a certification body. Where a regulatory reading is relevant, it is added as an annotation on top of the analysis and does not replace qualified counsel.

The analysis describes how the workflow runs. It does not decide for you whether to extend, renew, fix or stop, and it does not predict what the change will produce.

The Sprinkling Act methodology is built exclusively on the EU AI Act (Regulation 2024/1689). It is not derived from, certified by, or dependent on any external standard. However, the gate logic and risk assessment structure are consistent with the principles of established international frameworks:

NIST AI RMF 1.0

NIST AI 100-1 (2023)

The four NIST functions (Govern, Map, Measure, Manage) mirror the lifecycle approach embedded in our gates. Gate evaluation (Map), scoring with obligations (Measure), and ongoing regulatory monitoring (Manage) follow the same iterative logic. The Sprinkling Act methodology addresses the Map and Measure functions; operational Govern and Manage remain the responsibility of the assessed organisation.

ISO/IEC 42001:2023

AI Management Systems

ISO 42001 requires organisations to establish risk assessment processes (Clause 6), operational controls (Clause 8), and performance evaluation (Clause 9). Our 6-gate assessment produces the risk classification and obligation mapping that feeds into an ISO 42001-compliant AI Management System. The assessment does not replace an AIMS. It provides the regulatory input that an AIMS requires.

Sprinkling Act does not claim ISO 42001 certification or NIST compliance. These references indicate structural coherence, not formal alignment or endorsement.

Sprinkling Act is an independent analysis firm. It is not a law firm, not a Notified Body, not a certification body, and not affiliated with the European Commission, the European Parliament, or any national supervisory authority. Reports are an informative analysis, not legal advice.

v1.7
August
2026

The methodology stops describing itself as a pre-conformity layer. Sprinkling Act enters at usage and visibility; the pre-conformity work comes after the survey rather than framing it, and the regulatory governance paragraph now says so. Nothing changes in the six gates: what changes is where the regulation sits relative to the measurement. The v1.6 edition stays available as published.

▼ Previous versions (8)

v1.6
August
2026

Structural incompatibilities move to the Full Map, on the same ground as rework time. The element already stated that it is the only finding requiring a view from above; the Baseline, whose single source is the director looking from one position, does not have one, so the list was asking it to contain what it cannot establish. Where the Baseline glimpses an incompatibility it is now recorded as an open question, a status the report already carries. Consequence: the Full Map no longer sells what the Baseline announces delivering. Eight elements still, none removed — this one now carries its tier in its title.

v1.5
August
2026

What the Baseline promises is narrowed to what it can establish. It states a position, it does not diagnose: what the director cannot know from where they stand, and what it costs to keep deciding without it. It does not return the teams' uses and does not make visible what the director does not see — it measures the distance separating them from it. The ground is not commercial: Star and Strauss, CSCW 1999, establish that the invisibility of work is an organisational process, so making it visible is an operation nobody can run from inside the organisation. Consequence for the Full Map: its central deliverable is the gap document, that is the gap between the measurement taken with the director alone and the one taken with the three or four people who run the work. Without a delivered Baseline there is no term of comparison, so requiring it is a methodological constraint before it is a commercial one, and the Baseline report is sealed before the first Full Map interview — revised after hearing the teams it would no longer be a comparison but a reconstruction. The gap is not one quantity: it is attributed line by line by the date of the element. An element predating the close of the Baseline collection is a blind spot, established; one postdating it is drift, a finding in its own right; a line carried by a statement alone stays unattributed and is marked as such. The close of the collection is therefore dated in the report, distinctly from the readout. No validity window is published, and no variation of the deliverable is indexed on a triggering delay, for the same reason no price is.

v1.4
August
2026

The unit of analysis changes. The Baseline no longer describes a workflow, it measures the distance between what a director describes and what the elements they provide establish. The single source is the director: what escapes them does not become visible through better questioning, and that limit now sits in the first principle rather than in a closing section. The workflow examined is the one the director designates, not one we select by comparing several, which would require the overview we have just said we do not have. Consequence for the deliverable: first-pass acceptance rate and rework minutes move to the Full Map, where the estimate comes from the people who do the work rather than from the person who decides. The sentence “without those two numbers the report stays a description” remains exact — the Baseline is a dated, situated description of a distance, and the move to the Full Map is what turns it into an arbitration. Report structure: still seven parts, but not the same seven. “What the workflow takes” disappears, its substance being the two numbers that just moved, and what remains of it — subscriptions and direct costs, which an invoice establishes — joins the findings. “Workflow map” becomes “declared map”. “Priorities” becomes “what must be decided”, with the owner observed rather than proposed and the deadline drawn from a dated fact of the client’s own. A new part records the director’s corrections and objections verbatim, including those we did not retain. Same release, GDPR section: the claim that a survey of AI chains produces the article 30 record is withdrawn. The record is the controller’s obligation and its completeness is their liability; a survey documents a chain the record probably ignores and shows what is missing from it. Two legal qualifications move to the conditional, and whether these articles apply to a given use is stated as belonging to the client’s counsel.

v1.3 May 2026	Post-Digital Omnibus alignment (provisional trilogue agreement of 7 May 2026, EP press release IPR42011). Annex III standalone high-risk binding date moved from 2 August 2026 to 2 December 2027 (EP enumeration: biometrics, critical infrastructure, education, employment, law enforcement, border management; non-exhaustive). Annex I product-embedded safety components binding date moved from 2 August 2027 to 2 August 2028. Article 50 transparency unchanged (2 August 2026 baseline preserved). GenAI watermarking advanced from 2 February 2027 to 2 December 2026. Safety-component definition narrowed (verbatim: AI functions that only assist users or optimise performance no longer automatically face high-risk obligations, if their failure or malfunction does not create health or safety risks). Art. 5§1(i) added (provisional, applicable 2 December 2026 pending formal adoption): prohibition of AI systems generating CSAM/NCII content, three-branch architecture (designed-for / placed-without-safeguards / deployer-use). Art. 50§1-§4 obligations in scoring engine enriched with explicit application timeline (2 August 2026 new systems · 2 December 2026 transitional for §2 watermarking + §4 deepfake disclosure on systems already on EU market). Temporal indicators in score reports updated. Sleeper segments mapped: PWD × §4 employment platforms, CRA × AI Act IoT-AI.
v1.2 May 2026	Art. 5(1) split into absolute (§1(e)/(f)/(g)) vs conditional (§1(a)/(b)/(c)/(d)/(h)) prohibitions. Conditional triggers HIGH non-terminal with legal-verification flag instead of CRITICAL terminal. Art. 5§1(a) deceptive wording aligned with regulation: "causing or likely to cause significant harm". Annex III scoring now weighted by domain (structural vs operational). GPAI standard = LIMITED 35; GPAI systemic risk = HIGH 80 (smoother thresholds). Diagnostic scope clarified: screening, not legal qualification. Full report performs article-by-article qualification. Regulatory observation period: April–May 2026 (Digital Omnibus trilogue ongoing post 28 April stall, MDCG 2021-24 Rev.1 of 20 April integrated).
v1.1 April 2026	Art. 5§1(d) added: criminal risk profiling. Art. 5§1(h) reference corrected: real-time biometric. Art. 6(1) vs 6(2) reference corrected. Art. 50§2 content marking activated. All 8 prohibited practice checks (a-h) now covered. HIGH RISK obligations enriched with Art. 9-15 details. International standards alignment section added (NIST AI RMF, ISO 42001). Methodology card expanded to 23 pages.
v1.0 March 2026	Methodology versioned and published. GPAI systemic risk gate (Art. 55). Art. 6(3) exemption logic. AI Office guidance on Annex III. Regulatory freeze date: March 2026.

v0.1
January
2026

Initial release. 6 regulatory gates based on Art. 5, 6, 50, 51, 53. Article mapping for all Annex III domains.

This methodology is kept current through dated regulatory signals: each material development across EU institutions, national authorities, and industry publications is logged, tier-scored, and reflected in a versioned methodology release when it changes the analysis (see the changelog above).

When a signal reaches CRITICAL or MAJOR tier, the Temporal Stability Indicator on affected reports is updated automatically. Every scoring axis is documented, every threshold is versioned.

Implementation governance state (May 2026): the EU AI Office is operational with ~125 staff (target 140+), CEN-CENELEC JTC 21 harmonised standards have missed two deadlines and target Q4 2026 publication (risk: not in OJ when 2 December 2027 enforcement starts on Annex III), the Article 67 Advisory Forum has not been constituted seven months after the call for expressions of interest closed (September 2025), the Article 68 Scientific Panel structure is defined but its constitution remains unconfirmed, and only one Member State (Spain · AESIA, 12 high-risk systems hosted) has an Article 57 sandbox publicly operational. The Sprinkling Act methodology sits upstream of that work rather than inside it: it establishes what is actually running, so that the pre-conformity work which comes after it stands on dated facts rather than on absent or delayed official guidance. What it records is defensible under the obligation de moyens standard regardless of standards publication status.

Digital Omnibus AI Act trilogue agreement	MAJOR	Cross-gate · Annex III postponement to 2 December 2027 + Annex I to 2 August 2028 + safety component narrowing	[7 May 2026]
EDPB+EDPS Joint Opinion 1/2026	MAJOR	Cross-gate · Article 4 (literacy mandatory), Article 49 (registration retained), Article 57 (DPA in sandboxes)	[20 January 2026]
GPAI Code of Practice	MAJOR	Gate 04 · Art. 51 + 55	[March 2026]
AI Office capacity + CEN-CENELEC JTC 21 standards delivery	MAJOR	Cross-gate enforcement readiness gap	[Ongoing 2026]

Guidelines High-Risk delay MAJOR					Gate 03 · Art. 6(2) + Annex III [March 2026]		
German coalition split (Merz EPP / Wildberger pro-cuts vs SPD opposition)		MINOR	Cross-gate prospect signal	· German uncertainty			[April-May 2026]
Kai Zenner (Voss/EPP) appetite for AI Act major overhaul S2 2026		MINOR	Cross-gate amendment signal	· future forward-looking			[May 2026]
10 EU Member States blocking minority against further deregulation (AT, DK, NL, SK, SI, ES, GR, PT, RO, LV)		MAJOR	Cross-gate December deadline confirmed	· 2027 durability			[May 2026]
EuroLLM-22B (Edinburgh + EuroHPC MareNostrum 5) · EU sovereignty narrative		MINOR	Cross-gate signal public/university/public-healthcare deployers	· macro positioning for			[May 2026]

Tell us which decision is coming and which workflow carries it. If there is no fit, we say so before you pay.

Six questions on your own uses, two minutes, no account. Your answers come back to you as a text you can take away — nothing is sent to us.

PRODUCT**Full Report**

See what a complete report contains.

**BLOG****High-Risk Systems**

Which AI systems fall under Annex III? Full breakdown.

**BLOG****GPAI Obligations**

Art. 53-55 · what GPAI providers and deployers must do.

