



# What the Disclosure Regime Records

---

|                |                   |
|----------------|-------------------|
| REFERENCE      | AOS-ADG-2026-001  |
| VERSION        | 1.0               |
| DATE           | 20 August 2026    |
| CLASSIFICATION | Released          |
| DIRECTOR       | Brunei AL-Bijwaie |

**Version 1.0, 20 August 2026.** This assessment reads three Australian Government instruments and 111 of the 116 agency AI transparency statements the Digital Transformation Agency's own list links to. It rests entirely on those documents' own text. **Quotations reproduce the source's typography.** Vocabulary counts state the artefact they were taken from and the control word run on the same extraction.

## Section 1: What was assessed, and what was not

---

This assessment establishes one thing: **what the Australian Government’s mandated AI disclosure instruments require an agency to record about the AI it uses, and what the published statements actually record.**

AustraliaOS is an independent jurisdictional-exposure assessment practice. This assessment is built entirely on primary public records, and it was commissioned by no party to its findings. The practice is self-funded and has received no external funding.

**What was assessed.** Three instruments — the *Policy for the responsible use of AI in government* Version 2.0 [A-POL-1], the *Standard for AI transparency statements* [A-TRA-1], and the *Standard for accountability* [CAO-STD-1] — and **the 116 agency entries linked from the DTA’s own list of AI transparency statements [D-DTA-1], of which 111 were retrieved and read.**

**As at when.** The *Policy* and the *Standard for AI transparency statements* were retrieved on 11 August 2026, the *Standard for accountability* on 5 August 2026, and all three were re-verified against their recorded hashes on 20 August 2026. **The list and each of the 111 statements read were retrieved on 20 August 2026 at 10:04 UTC**, and each is pinned at its own hash. A statement published, amended or withdrawn after that moment is outside what this assessment read.

**Out of scope, by design.** This assessment does not read internal registers, contracts, annual reports or ministerial answers. **It reads what an agency is required to publish and what it published.** These boundaries are deliberate and are distinct from the stated limits in Section 8; a reader should not read any of them as a gap.

**A directive this assessment does not address.** On 12 June 2026 a provider published that it had been directed to suspend access to two named models. **No claim in this assessment concerns that directive.** Its issuing body, legal authority, text, scope and lifting are established by nothing this practice could obtain, and Section 8 states what is and is not known about it. **The question this assessment asks is not what the directive did. It is what the Australian record would let anyone establish afterwards.**

## Section 2: What the instruments require to be recorded

---

**Two instruments create records. One of them is public.** Neither records which model an agency uses, and the reason is the same in both: **no field and no listed item asks.**

### 2.1 The public artefact

The Policy states the requirement, under the heading **AI transparency statement**:

“Agencies **must** make a publicly available statement outlining their approach to AI adoption and use, as prescribed under the Standard for transparency statements.

The statement **must** be reviewed and updated annually or sooner, should the agency make significant changes to its approach to AI.

Agencies **must** notify the DTA when they publish and make any changes to their AI transparency statement by emailing ai@dtg.gov.au.”

**That is the whole of the block. Three sentences, and none of them asks what the agency uses.** The bolding is the Policy’s own.

The Standard prescribes the content [A-TRA-1]:

“At a minimum, agencies must provide the following information regarding their use of AI in their transparency statement:

- the intentions behind why the agency uses AI or is considering its adoption
- classification of AI use according to usage patterns and domains, as listed at Attachment A
- classification of use where the public may directly interact with, or be significantly impacted by, AI or its outputs without human review
- measures to monitor the effectiveness of deployed AI systems and protect the public against negative impacts
- overview of compliance with the requirements under the Policy for responsible use of AI in government
- compliance with applicable legislation and regulation
- when the statement was most recently updated.”

**Seven items. None of the seven is a vendor, a provider, a product or a model.**

On [A-TRA-1] reduced to plain text by stripping its markup — the extraction on which the control word statement returns 51 — **vendor returns o, supplier returns o and provider returns o.** model returns 6, and **all six were read individually: every one is an entry in a contents list belonging to a different DTA page, and none is in the Standard's own required content.**

## 2.2 The internal artefact

The Policy requires a second record, and its heading is **Internal AI use case register** [A-POL-1]:

“Agencies **must** create a register of in-scope AI use cases to enable accountable official(s) to record accountable use case owners within 12 months of this policy taking effect.

Agencies **must** share the register with the DTA every 6 months, commencing from when they create the register to meet the above requirement.

The Standard for accountability lists the minimum fields agencies must capture in the use case register. Agencies can add additional fields to meet their organisational needs. An existing register may be reused for the purposes of meeting this requirement. The standard also provides the instructions for how to share agency registers with the DTA.”

**The register is not published, but it does not stay inside the agency either.** The second sentence requires it to go to the DTA every six months. **Nothing in this assessment establishes what the DTA does with it,** and nothing establishes that any such register has been shared; the point is only that “internal” in the Policy’s own heading means *not public*, not *not disclosed*.

**And the fields are a floor, not a ceiling.** The Policy says agencies “can add additional fields to meet their organisational needs”. **An agency that wanted to record the model could record it.** What follows below is what the instruments *require*, which is a different question and the one this assessment asks.

The Standard lists the fields [CAO-STD-1]: the accountable use case owner’s name and email address; what criteria the use case met to be in scope; the inherent risk rating; the residual risk rating; the date the AI impact assessment was last done; and, for a use case with an inherent high-risk rating, the last date of review and the date for the next review.

**No field names a vendor, a provider, a product or a model.** On an extraction of [CAO-STD-1] where the control word **must** returns 18, **vendor, supplier and provider each return o.**

**Neither instrument requires an agency to record which AI model it uses, or whose.** That is a statement about what the instruments ask for. **It is not a statement about what agencies do,** and Section 3 is that.

## Section 3: What the published statements record

---

**The DTA publishes the population itself.**

Its list of AI transparency statements [D-DTA-I], read as served at 10:04 UTC on 20 August 2026, links **116 agency entries**. The count is of external links on the served page, de-duplicated by address, with the DTA's own self-links excluded and **four reference links excluded by reading each**: the *Public Governance, Performance and Accountability Act 2013* and the *Office of National Intelligence Act 2018* on the Federal Register of Legislation, a Department of Finance page on Commonwealth structure, and a Department of the Prime Minister and Cabinet page on administrative arrangements.

**Of the 116, this assessment read 111 [D-TS-1 to D-TS-111]. Five were not read and are named in Section 8.**

**The classification is recorded per statement, not only in aggregate.** Each of the 111 rows in this assessment's register carries, in its name field, whether that statement names a vendor or a product and whether it names a model version. **Every figure in this section can be recounted from the register without re-reading the statements**, and re-read against the pinned bytes if the reader wants to check the classification itself.

### 3.1 The two figures

**Of the 111 read, 41 name a vendor or a product.**

**Of the 111 read, 0 name a model version.**

**Of the 111 read, 70 name no AI tool at all** — they give the classification of use the Standard requires, and stop.

**The counting rule, so that a reader can disagree with it.** A statement is counted as naming a vendor or a product where it names one as a tool the agency records itself as using. **Two statements name a tool and are not counted, and they are named here rather than absorbed**: the Department of Agriculture, Fisheries and Forestry [D-TS-2] says it “will seek to utilise the functions GovAI offers to the Australian Public Service”, and the Australian Communications and Media Authority [D-TS-58] says it “will take advantage of evolutionary whole of government AI initiatives such as GovAI”. **Both name the platform prospectively and neither records a current use.** Counted the other way the figure is 43 of 111, and the model figure does not move.

**Neither figure is interesting alone. The contrast is the finding.** Naming a vendor is ordinary practice

across the published population: more than a third of the statements read do it, unprompted by any requirement in Section 2. **Naming a model is done by none of the III.**

### 3.2 What the second figure was tested against

**The model-version test is separate from the vendor test**, and the separation is what the finding turns on. It ran over all III statements read, against patterns for Claude 2, Claude 3, Claude 3.5, Claude 4, Claude 5, Claude Sonnet, Claude Opus, Claude Haiku, Claude Fable, Claude Mythos, bare Sonnet, bare Opus, bare Haiku, GPT-3, GPT-4, GPT-5, GPT-4o, Gemini 1, Gemini 2, Llama 2, Llama 3, Llama 4, Fable 5, Mythos 5, o1-mini, o3-mini and Titan.

**Zero matches across the III.**

**The test ran over all III, not over the ones a scan had flagged.** That matters, because a model version can appear in a statement that names no vendor token a scan is looking for. **All III were read; the 70 that name no AI tool were each read end to end** — not a sample of them — **and none of the III names a model.**

### 3.3 Three statements name this provider

**Three of the III read name Anthropic or Claude. Each names a product or a product family. None names a model.**

The National Health Funding Body [D-TS-44]:

“The NHFB **uses authorised AI tools** (such as Anthropic’s Claude and Microsoft Copilot) on a limited basis, when undertaking routine internal administrative matters to maximise productivity, such as summarising internal policy information and producing Minutes of internal meetings.”

Health Direct [D-TS-100]:

“To analyse and improve service delivery quality and efficiency (using Microsoft CoPilot and Claude Enterprise).”

Screen Australia [D-TS-105]:

“AI tools used by Screen Australia staff includes enterprise AI deployed in our closed internal ICT environment, including Microsoft 365 Copilot, as well as publicly available browser-based or app-based AI such as ChatGPT and Claude.”

## Section 4: The entry examined closely

---

**Health Direct’s statement [D-TS-100] names more products against particular uses than any other statement read: eight.**

Counted from the names inside its (using ...) parentheticals, split on *and*, *or* and commas, de-duplicated: **Infermedica triage tool, Microsoft CoPilot, Claude Enterprise, Google Translate, Heidi, Lyrebird, Adobe Firefly and AWS Kiro.**

**The comparison names its test, and the test uses the list this document prints.** Run the 47-name scan of Section 5 over all III statements and count distinct names per statement: **Health Direct returns 12, and no other statement returns more than 6** — Screen Australia [D-TS-105] returns 6, and the next two return 5. Those are *tokens*, not products: the same statement’s product count is 8, because the scan splits Microsoft/CoPilot and Adobe/Firefly. **Both figures put Health Direct first by a factor of two on the same population, and a reader holding the register can recompute either.**

It discloses them against the individual use, not as a list:

- “To analyse and improve service delivery quality and efficiency (using Microsoft CoPilot and Claude Enterprise).
- To translate some information on our website into languages other than English (using Google Translate).
- To assist the GPs on 1800MEDICARE (GP) to record details about the consultation in our system by using an AI-scribe (using either Heidi or Lyrebird)”

**Those are three consecutive items of one list**, quoted in the source’s own list form; the parentheticals are the statement’s own.

**This is the most granular disclosure in the population read, and it does not answer the question either.** *Claude Enterprise* is a product tier. **Which model served it on any particular day is not stated, and no requirement in Section 2 asks for it.**

**Nothing here is a criticism of Health Direct.** The Standard requires classification of use; this statement gives classification of use **and** the products. **It has met the requirement and gone past it.** The point is what remains unanswerable after it has.

## Section 5: What a scan would have returned, and why it is not the number

---

**A scan number is a property of its word list, so the list is given.** The scan run here is over 47 vendor and product names: Anthropic, Claude, OpenAI, ChatGPT, GPT, Copilot, CoPilot, Microsoft, Azure, Google, Gemini, Bard, Vertex AI, Amazon, AWS, Bedrock, Llama, Mistral, Cohere, Perplexity, DeepSeek, Grok, xAI, IBM, Watson, Salesforce, Einstein, Oracle, Adobe, Firefly, NVIDIA, Hugging Face, Stability AI, Midjourney, DALL-E, Whisper, Otter.ai, Fireflies, Heidi, Lyrebird, Infermedica, Kiro, Notion, Grammarly, Canva, Gomeeki and Leximancer, each matched on a whole-word boundary.

**That scan returns 48 of III. This assessment reports 41.** The difference is seven statements, and each was read in its passage:

- **The Australian Electoral Commission's** [D-TS-28] hit is in **PDF producer metadata** — “Acrobat PDFMaker 26 for Word”, “Adobe PDF Library 26.1.187”.
- **The Royal Australian Mint's** [D-TS-90] is inside a **PDF binary stream** — “Microsoft® Word for Microsoft 365”.
- **The CSIRO's** [D-TS-104], **the Inspector-General of Taxation's** [D-TS-85] and **the Northern Australia Infrastructure Facility's** [D-TS-110] are the same site-furniture notice: “This site is protected by reCAPTCHA and the Google Privacy Policy and Terms of Service apply.”
- **The National Commission for Aboriginal and Torres Strait Islander Children and Young People's** [D-TS-75] is a **web-analytics notice** — “your activity is being tracked by Google Analytics” — not an AI tool disclosure.
- **The Office of the Fair Work Ombudsman's** [D-TS-25] is a **site-wide automatic translation banner** — “powered by Microsoft Translator” — not the statement body.

**Seven of forty-eight token hits are not disclosures of an AI tool. The scan missed none of the 41.** Its error on this population runs in one direction only: it over-reports, and only reading finds the over-report.

**This is recorded because the method is the difference.** The number a scan returns and the number a reading returns are not the same number, and only one of them is in this document. **The model figure is the case where they agree:** a scan would have reported zero and been right, and the zero here rests on the model test having run over all III rather than over the 48 the vendor scan flagged.

## Section 6: The Policy on changes an agency did not initiate

---

The Policy addresses the case where something changes and the agency did not change it. **The passage sits inside a block the Policy labels “Mandatory requirements”, and it is not one.**

Quoted with the paragraph it follows, from [A-POL-1], printed pages 14 and 15:

“Agencies **must** add each in-scope AI use case to their internal register of AI use cases and update it as required. Include risk rating and accountable use case owner changes. When deploying an in-scope AI use case, agencies **must**:

- regularly monitor and evaluate their use case to ensure it is operating as intended and that risks are being effectively managed.
- re-validate the AI use case impact assessment by checking its accuracy and updating it when there is a material change in the use case scope, usage or operation.

Agencies should also monitor changes that are not initiated by the agency. For example, vendor changes and changes in the regulatory environment. Agencies could also ask vendors to provide information on updates through contractual mechanisms.”

**The paragraph it follows is written in “must”. The sentence itself is written in “should” and “could”.**

**The Policy marks the difference typographically, and the quotation above reproduces it.** Read from the rendered pages rather than the text layer: **29 of the 35 occurrences of must are set in bold**, and **none of the 19 occurrences of should and none of the 5 occurrences of could is**. The six unbolded must were each read and none is in a requirement sentence — they are in the Policy’s own explanatory prose about scope and about how it is to be applied.

**The block label was established by reading every block label in the Policy in order.** A whole-document search for the two label strings returns **nine** occurrences, and the count in this document is not nine: **four are entries in the Policy’s own contents list on printed page 3**, and **one is ordinary prose on printed page 4** — “outlines the mandatory requirements agencies must follow. Recommended actions and considerations are also included where relevant.” In the body of the Policy there are **four: three “Mandatory requirements” and one “Recommended Actions”**, on printed pages 10, 12, 13 and 14. The contents list on printed page 3 names the same four. **The label immediately preceding this passage is the “Mandatory requirements” of printed page 14**, and no “Recommended Actions” label intervenes before the passage on printed page 15.

**must returns 35 in the Policy, should returns 19 and could returns 5. The Policy nowhere defines those terms,** and this assessment therefore says what the words are and not what they are worth.

**Neither Standard hardens it.** Read where such an obligation would sit, in both: **vendor, supplier, provider, not initiated, outage and discontinued each return o in each.** The single occurrence of availability in [CAO-STD-I] was read and concerns **information accessibility to assist with audits,** not the availability of a service.

## Section 7: The comparator, and its bounds

---

**Three of the III statements read record a government direction about a named AI platform.**

Screen Australia [D-TS-105]:

“Screen Australia follows the direction of Government on the use of specific AI platforms. For example, following the directive from Government, Screen Australia staff cannot access or use DeepSeek on agency devices or personal devices that contain a work profile.”

Defence Housing Australia [D-TS-93]:

“DHA has blocked DeepSeek AI to meet the Department of Home Affairs Protective Security Policy Framework compliance.”

The Department of the Senate [D-TS-66]:

“Email and/or intranet updates are issued in response to AI developments that may pose a risk, such as the February 2025 requirement that DeepSeek not be used on government devices.”

**What this establishes: a transparency statement can carry a government direction concerning a named AI platform, and three of them do.**

**What it does not establish, and the bound is the point.** It establishes **nothing** about the directive of 12 June 2026 — not that any agency was subject to it, not that any agency should have recorded it, and not that its absence from these statements signifies anything at all. **The instruments in Section 2 require neither disclosure.** The comparator is here because a reader is entitled to know that the record is capable of carrying such a direction. **Sections 3 to 6 establish what the record does not contain. They do not establish that it could not.**

## Section 8: Stated limits

---

The assessment states its own boundaries so that a reader does not have to find them. **Six limits, and the first two bound the whole document.**

**1. The 12 June 2026 directive is not established.** Its issuing body, legal authority, instrument number, text, scope, treatment of allied nationals and lifting are established by no record this practice could obtain. Searched on 20 August 2026: every Bureau of Industry and Security document in the Federal Register for June and July 2026 — **seven, none concerning AI models**; the Federal Register’s full text for both model names across May to August 2026 — **zero documents**; the Bureau’s own site, which names neither; three of its press-release paths, all of which returned 404; and the Department of Commerce press releases, which returned 403 to an automated request and are **recorded as inaccessible rather than absent**. **What the provider published is a fact about the provider’s statement and is treated as one throughout.**

**2. Five of the 116 statements were not read.** The Federal Court of Australia, the Office of the Commonwealth Ombudsman, the Australian Centre for International Agricultural Research and the Australian Institute of Family Studies each returned **HTTP 403 with a bot-check interstitial**; the Department of the Prime Minister and Cabinet returned **HTTP 200 with a 955-byte bot-check shell**. Each was retried once with browser headers and a referring address, and each returned the same. **Nothing in this assessment is claimed about those five**, and every figure in Section 3 carries the population 111 on its face for that reason.

**3. This reads what is published, not what is held.** The register of Section 2.2 is not published — the Policy’s own heading calls it internal, and it goes to the DTA six-monthly rather than to a reader. **This assessment did not read one, and says nothing about what any agency records in it.**

**4. Contract records were not reached.** AusTender returned 403 to an automated request on three paths, and the openly published contract-notice datasets end in 2023. **Whether any contract names any provider is established neither way here.**

**5. The statements were read as served on one day.** The *Standard for AI transparency statements* requires a statement to be reviewed and updated at three junctures: “at least once a year”; “when making a significant change to the agency’s approach to AI”; and “when any new factor materially impacts the existing statement’s accuracy”. **Any of the three may have fallen since. This is a reading of 20 August 2026 and it speaks to that date.**

**6. This assessment’s register records no independent capture.** All 115 rows record `capture_type` as `raw_body` — the bytes this practice retrieved at the moment recorded in the row — and **no row records a Wayback memento or any other third-party capture**. Three of the 115 are inherited captures under the corpus’s convention 7 — the three instruments of Section 2, retrieved under earlier cases and re-verified

against their recorded hashes on 20 August 2026 — and the same is true of them. A reader can fetch each source at its published address and compare it against the hash recorded here.

**That is a statement about what this register holds, and it is not a statement about what exists.** Whether an independent capture of any of these sources exists is **not established here and is not assessed.** The index was not queried, and **an absence in this register is not evidence of an absence in any archive.**

## Section 9: What this assessment does not say

---

**It does not say that any Commonwealth system was affected by the suspension of 12 June 2026.** Nothing establishes that one was, and nothing establishes that none was. **That is the condition this assessment describes, not a gap in its evidence.**

**It does not say that any agency failed.** An agency that records what the *Standard for AI transparency statements* requires has met the requirement. **Forty-one of the 111 read went further than the requirement and named a vendor or a product.**

**It does not say that the instruments should require more.** What a disclosure regime ought to capture is a question for the people who make it. **This assessment establishes what it does capture and at what granularity.**

**It does not say that any provider acted wrongly.** The provider published that it complied with a direction it described as legal, and this assessment records that and assesses nothing about it.

**It draws no inference from silence.** Every finding above rests on what an instrument requires or on what a statement says, both quoted. **No finding rests on a register being empty, and none was found to be empty.**

End of document.