

Mind the Gap — Factcheck & Proof Register

Prepared by The Verifier — MTG-Factcheck-Proof-Register-v1 · 2026-06-20 · SHADS CW channel

Scope: MTG_PDF_Reader_Re12_V0_0.pdf (15 chapters + front matter) and 1780159905368_MTG_BackMatter_v1.docx (bibliography, about, colophon). **Date of check:** 30 May 2026. **Method:** every stat / percentage / date / quote-attribution / named external claim extracted, then the highest-consequence and most-checkable verified against primary or strong secondary sources. Each entry carries a status and, where relevant, proposed softening wording.

Status key

- **GREEN** — verified. Source found, claim stands as written.
- **AMBER** — Paul's call. Defensible but contestable, slightly off, stale, or unsourced. Soften or confirm.
- **RED** — fix before publication. Wrong as written.
- **NOT VERIFIED THIS PASS.** Couldn't reach a source in this session; flagged for Paul / ChatGPT cross-check. Not asserted either way.

1. RED — fix before publication

R1. "Google's Gemini 4 has a two million token context window." (Ch 10)

Wrong on both counts as of May 2026.

- There is **no Gemini 4**. Google's current line is **Gemini 3 (Nov 2025) and Gemini 3.1 Pro (19 Feb 2026), built on Gemini 3 Pro**.
- Current Gemini context window is **1 million tokens, not two**. The 2M window belonged to the **earlier Gemini 1.5 / 2.5 Pro generation**.
- Also internally inconsistent: Ch 6 correctly states frontier windows run "200,000 to 1,000,000 tokens."

Proof:

- Google DeepMind model card, Gemini 3.1 Pro — "token context window of up to 1M," "based on Gemini 3 Pro." <https://deepmind.google/models/model-cards/gemini-3-1-pro/>
- Google blog, "Introducing Gemini 3" — "1 million-token context window." <https://blog.google/products/gemini/gemini-3/>
- Google Cloud docs, Gemini 3 Pro — "1M token context window." <https://docs.cloud.google.com/vertex-ai/generative-ai/docs/models/gemini/3-pro>

Proposed fix (pick one):

- Version-pinned: "Google's Gemini 3 holds a one-million-token context window."
- Future-proof (recommended — avoids re-dating at every model bump): "The leading models now hold context windows of a million tokens or more."

2. AMBER — Paul's call (soften, confirm, or reframe)

A1. Frankl — "between stimulus and response there is a space" (Ch 3; Bibliography annotation)

Presented as Frankl's. The exact line does **not appear anywhere in Frankl's works**. It was popularised by **Stephen Covey, who said he found it unattributed in a library book and never recorded the**

source. This is a famous misattribution — and it is famous *inside Paul's own field (coaching/therapy), which is exactly where it costs credibility. (Consequence is high enough that Paul may choose to treat this as Red.)*

Proof:

- Quote Investigator — "Researchers have been unable to find this passage in the works of Viktor E. Frankl... popularized by... Stephen R. Covey; however, he disclaimed authorship." <https://quoteinvestigator.com/2018/02/18/response/>
- Viktor Frankl Institute (via Check Your Fact) — "The statement appears nowhere in Frankl's written works... Covey... did not note down the book's author and title." <https://checkyourfact.com/2019/09/27/fact-check-viktor-frankl-stimulus-response-space-growth-freedom/>

Proposed fix: keep the idea, drop the quotation framing. e.g. *"Frankl's insight — that between stimulus and response lies a space, and in that space our freedom — "* → reframe as *"the principle most associated with Frankl: that between stimulus and response there is a space..."* and, if you want to be unimpeachable, add a half-line acknowledging Covey popularised the phrasing. The concept is genuinely Franklian; the verbatim line is not his.

A2. "Google's data centres consumed six and a half billion gallons in 2023" (Ch 15)

Figure is slightly high and slightly mis-scoped. Google's **2023 total operations** were **6.4 billion gallons**; of that, ~6.1 billion (95%) went to data centres. "Six and a half billion ... data centres" rounds the total up and attributes it to data centres.

Proof: "In 2023, Google operations worldwide consumed 6.4 billion gallons of water (24.2 billion liters), with 95%, 6.1 billion gallons, used by data centers." <https://www.civilbeat.org/2025/08/data-centers-consume-massive-amounts-of-water-companies-rarely-say-exactly-how-much/> ; corroborated <https://gijn.org/stories/researching-water-consumption-data-centers/>

Proposed fix: "Google's data centres consumed around six billion gallons of water in 2023, most of it potable." ("most of it potable" is supported.)

A3. "A single ChatGPT query draws roughly ten times the electricity of a Google search" (Ch 15)

True to the 2024 IEA/Goldman framing (2.9 Wh vs 0.3 Wh), but **superseded**. Epoch AI (Feb 2025) put a typical GPT-4o query at ~0.3 Wh — ten times less than the old estimate, i.e. roughly the same as a search. Sam Altman (2025) cited ~0.34 Wh. An informed reader will know the 10× figure has been walked back.

Proof:

- Epoch AI — "We find that typical ChatGPT queries... likely consume roughly 0.3 watt-hours, which is ten times less than the older estimate." <https://epoch.ai/gradient-updates/how-much-energy-does-chatgpt-use>
- "The often-repeated claim that ChatGPT is ten times more energy-intensive than a Google search is no longer supported by current data." <https://www.zmescience.com/science/news-science/how-much-energy-chatgpt-query-uses/>

Proposed fix: attribute and date it — "On the figures that circulated through 2024, a single ChatGPT query was estimated to draw roughly ten times the electricity of a Google search — an estimate later revised sharply downward as models and hardware became more efficient." Or drop the multiplier and keep the directional point.

A4. "thirty minutes of generative AI use consumes around a sixth of a gallon of water" (Ch 15)

In the ballpark of the widely-cited (and itself contested) water estimates (~500ml per session in the Li et al. line of work; Altman's ~0.000085 gal/query). The "sixth of a gallon / 30 minutes" basis isn't cleanly sourced and the field's figures vary widely.

Proposed fix: attribute/hedge — "by some estimates, around a sixth of a gallon of water" — or cite the specific source you drew it from.

A5. [name held] et al. (2019) — "a 77-year-old brain" (Ch 15; Bibliography annotation)

Paper and citation are otherwise correct (see GREEN G6). But the study reports its oldest specimen as 78 years ("20 gestational weeks to 78 years of age"), and frames persistence "into old age." The recurring "77-year-old" (used in Ch 15 and the bib) may come from a specific figure in the paper, but the headline age in the source is 78.

Proof: UCSF / Nature Communications coverage — "postmortem human amygdala tissue from 49 human brains — ranging in age from 20 gestational weeks to 78 years of age."
<https://www.ucsf.edu/news/2019/06/414756/mood-neurons-mature-during-adolescence>

Proposed fix: confirm the exact specimen age against the paper's figures. If no 77-year-old specimen is specifically shown, change to 78. (CL-31 is built on this line, so worth nailing.)

A6. WaaS deployment figures — "Accuracy rates at 99.6%. Cycle times reduced by 68%. Cost to serve down 57%." (Ch 13)

Three precise percentages presented as generic ("early WaaS deployments") with no source named. Precise figures with no attribution are the single most exposed kind of claim in the book. **NOT VERIFIED — they read like one specific vendor case study generalised.**

Proposed fix (in order of preference):

1. Name the source/case study these came from, or
2. Soften to ranges and label as illustrative: *"Early deployments have reported accuracy in the high-90s, cycle times cut by more than half, and cost-to-serve down by roughly a half" — and make clear these are vendor-reported.*

A7. Bibliography — Klein, The Writings of Melanie Klein, Volume III dated (1946)

The collected *Writings of Melanie Klein* (Hogarth) were **published 1975; 1946 is the date of the key paper ("Notes on Some Schizoid Mechanisms") within the volume. As a book citation, the year is off.**

Proposed fix: cite as Klein, M. (1975). The Writings of Melanie Klein, Vol. III: Envy and Gratitude and Other Works 1946–1963. London: Hogarth. (Keep 1946 only if citing the paper, not the volume.)

A8. Bibliography — "Compass of Shame — Donald Nathanson (1992)"

Part One annotation titles the work *Compass of Shame*; Part Two correctly lists the actual book, *Shame and Pride: Affect, Sex, and the Birth of the Self* (1992). "*Compass of Shame*" is the **concept within that book, not a separate title. Minor internal inconsistency.**

Proposed fix: in the annotation, "Shame and Pride (and the 'compass of shame' it introduces) — Donald Nathanson (1992)."

A9. Buddha — "we are what we think" (Ch 3, the five anchors)

Popular attribution, drawn loosely from a Dhammapada translation; the wording and direct attribution are contested among translators. Low consequence, but it sits in a list beside the Frankl line, so worth a consistent decision.

Proposed fix: "the idea, often rendered from the Dhammapada, that we are what we think." Or leave as a deliberately informal anchor and accept the looseness.

A10. *Ferrucci, What We May Be* — dated (1984)

Original US publication was 1982 (Tarcher). "1984" is plausibly a UK / 3rd-edition printing (the bib says "3rd ed., Turnstone Press"), so likely fine — but worth confirming the printing you're citing.

3. NOT VERIFIED THIS PASS — recommend Paul / ChatGPT cross-check before print

These are checkable but I didn't reach a source this session. Listed so nothing is silently assumed true. (Several are in global memory as "verified May 2026" — re-confirm against a live source before KDP, as figures move.)

- Ch 4 — Uber CTO disclosed full-year 2026 AI budget exhausted in first four months (book says *April 2026*; memory note says source is *The Information*). Confirm date + figure.
 - Ch 4 — Microsoft to cancel internal AI coding-tool licences for ~100,000 engineers, effective end June (book: "mid-May 2026").
 - Ch 4 — Anthropic decoupling agent usage from chat usage with its own metered credit pool from 15 June 2026 (memory: billing announced 13 May, live 15 June).
 - Ch 4 — FinOps teams "doubled in twelve months from thirty-one per cent to sixty-three per cent of large organisations." Specific percentages — confirm against the FinOps Foundation / source survey.
 - Ch 4 — Cybersecurity launch partner surfaced "approximately ten thousand vulnerabilities... first month saw ninety-seven patched." Confirm against the published source.
 - Ch 4 — SWE benchmark projections "non-specialised agents at fifty-four per cent and state-of-the-art at eighty-seven per cent" at start of 2026. Confirm against the benchmark cited.
 - Ch 4 — Pricing: "\$15 per million output tokens at the top end... mid-tier... three dollars." Roughly consistent with current market (e.g. Gemini 3.1 Pro \$12/M out; reasoning tiers ~\$15/M) but confirm the figures you want to stand behind.
 - Ch 4 / Ch 6 — Anthropic "Dreams" feature and the quoted description ("let Claude reflect on past sessions to curate an agent's memory and surface new insights"); and the "Code with Claude" conference quote ("getting much closer to the feeling of an infinite context window"). Confirm both quotes verbatim against Anthropic's own docs/transcript — they're direct quotations, so wording matters.
 - Ch 10 — OpenAI /Goal feature ("announced this week"). Confirm it exists under that name.
 - Ch 4 — Training-token scale: early ChatGPT "~300 billion tokens / ~225 billion words"; GPT-4 "~13 trillion." These match credible leaked estimates but are not official; the book already hedges ("reportedly," "credible estimates"), which is the right posture.
 - Ch 4 — "Library of Congress... ~20 billion words across its entire collection." Contested figure; estimates vary enormously. The book uses it as an order-of-magnitude illustration, which is defensible, but the "20 billion words" specifically is soft.
 - Bibliography — Bond, *Behind the Mask* (1998, Atlow Mill); Firman described as "a student of Ferrucci" (lineage claim, hard to source). Low consequence.
-

4. GREEN — verified, stands as written

G1. Kohls v. Ellison (Ch 8; Bibliography)

Stanford professor Jeff Hancock, founding director of the Stanford Social Media Lab, used GPT-4o (i.e. ChatGPT); declaration contained two non-existent citations (plus one mis-cited); court

excluded the testimony; order filed 10 Jan 2025; Judge Provinzino named the irony explicitly. All accurate. (2025 WL 66514, D. Minn.) Proof: <https://reason.com/volokh/2025/01/10/misinformation-experts-citation-to-fake-ai-generated-sources-in-his-declaration-shatters-his-credibility-with-this-court/> ; <https://www.law.berkeley.edu/wp-content/uploads/2025/12/Kohls-v-Ellison.pdf>

G2. *Mata v. Avianca* (Ch 8; Bibliography)

Citation 678 F. Supp. 3d 443 (S.D.N.Y. 2023) correct. Six fabricated cases generated by ChatGPT; the three named (Varghese v. China Southern Airlines, Shaboon v. Egyptair, Petersen v. Iran Air) are all genuine fabrications from the case; lawyers sanctioned (\$5,000). Accurate. Proof: https://en.wikipedia.org/wiki/Mata_v._Avianca,_Inc. ; <https://law.justia.com/cases/federal/district-courts/new-york/nysdce/1:2022cv01461/575368/54/>

G3. *Moffatt v. Air Canada* (Ch 8; Bibliography)

Citation 2024 BCCRT 149 correct. Chatbot advised a refund could be claimed within 90 days; actual policy required claiming before travel; tribunal rejected Air Canada's "separate legal entity" argument; airline held liable. Accurate. Proof: <https://www.canlii.org/en/bc/bccrt/doc/2024/2024bccrt149/2024bccrt149.html> ; https://www.americanbar.org/groups/business_law/resources/business-law-today/2024-february/bc-tribunal-confirms-companies-remain-liable-information-provided-ai-chatbot/

G4. *IEA data-centre projection* (Ch 15)

"Around 945 terawatt-hours by 2030 — equivalent to the entire current electricity demand of Japan." Verbatim match to the IEA *Energy and AI report* (Apr 2025). **Proof:** <https://www.iea.org/reports/energy-and-ai/executive-summary> ; <https://www.scientificamerican.com/article/ai-will-drive-doubling-of-data-center-energy-demand-by-2030/>

G5. *OpenAI / Jony Ive device* (Ch 4)

"\$6.5 billion acquisition," "screenless, voice-first, pocket-sized," "calm computing," "late 2026 release" — all match reporting. io Products acquisition value \$6.5bn (sometimes reported \$6.4bn all-stock); H2 2026 debut confirmed at Davos. **Proof:** <https://www.benzinga.com/markets/tech/25/07/46335969/openai-finalizes-6-5-billion-deal-to-acquire-jony-ives-ai-hardware-startup> ; <https://introl.com/blog/openai-consumer-device-jony-ive-hardware-2026>

G6. *Meta Ray-Ban smart glasses "seventy-five to eighty per cent of their category"* (Ch 4)

Verified (75–80% market share). **Proof:** <https://introl.com/blog/openai-consumer-device-jony-ive-hardware-2026> ("Meta's Ray-Ban smart glasses captured 75-80% market share").

G7. *[name held] et al. (2019) citation* (Bibliography)

Nature Communications, 10(1), 2748, DOI 10.1038/s41467-019-10765-1, full author list, title "Immature excitatory neurons develop during adolescence in the human amygdala" — all correct. (Only the "77-year-old" detail needs the A5 check.) **Proof:** <https://www.nature.com/articles/s41467-019-10765-1> ; <https://pubmed.ncbi.nlm.nih.gov/31227709/>

G8. *GPT-5.5 attribution* (Front matter / Colophon)

"Images by ChatGPT (OpenAI, GPT-5.5)" — GPT-5.5 is real, released 23 April 2026. Plausible for images generated late in the project. **Proof:** <https://openai.com/index/introducing-gpt-5-5/> ; <https://www.cnn.com/2026/04/23/openai-announces-latest-artificial-intelligence-model.html>

G9. *Claude Opus 4.7* (throughout) — consistent with the current model landscape (*Opus 4.x line*; *Mythos Preview* referenced in coverage). *Stands.*

G10. *ISBN 978-1-0369-9170-8* (Colophon)

Check digit is mathematically **valid (computed check digit = 8, matches)**. Structurally sound.

G11. ChatGPT join date "Sunday 11 December 2022" (Ch 9 / Ch 10)

ChatGPT opened 30 Nov 2022; +11 days = 11 Dec 2022, which **was a Sunday**. Internally and externally consistent.

G12. Internal-consistency checks that pass

- Train fare £18.80 vs breakfast £16.00; "£2.80 on a train fare" (Ch 13) = the £2.80 *shortfall* (18.80 – 16.00). *Consistent*.
- 76,000,000 words ÷ 783,000 (KJV) = 97.06 → "ninety-seven Bibles... call it a hundred." *Consistent*.
- 900,000 words ÷ 45,000 = 20 → "roughly twenty to one." *Consistent*.
- Token/word ratios (0.75): 200k tok ≈ 150k words; 200k–1M tok ≈ 150k–750k words; 300bn tok ≈ 225bn words. All consistent.
- Biography: apprentice 1978 at 16 → born ~1962 → age 63 in 2026. Consistent across chapters.

5. One cross-document note (not a fact, but worth a glance)

The back-matter Colophon states "**culminating in V69**"; the live manuscript baseline is **V76**. The back-matter file is v1 and predates the current manuscript, so this is expected — but the version numbers (and the "forty-five-thousand-word" / "twenty-seven versions / six weeks" figures) will need a final sync before print. Amber.

Prepared as a verification reference for Mind the Gap. Drafted with Claude (Anthropic, Opus 4.8). Sources are live URLs as of 30 May 2026; figures in fast-moving areas (model versions, token pricing, deployment stats) should be re-confirmed immediately before KDP upload.