

The Leverage Inversion: Dating the Transition from Consuming AI Answers to Directing an AI Workforce in a 3.5-Year Single-Subject Conversational Corpus

Paul Roebuck*

Preprint v0.4 · 20 July 2026 · Not peer reviewed

Abstract

This is a retrospective computational single-case study with an autoethnographic component: the author is both investigator and subject, and the data are the author’s retained and recoverable human–AI interaction corpus across three principal platform legs — 2,286 threads, 47,442 turns, 986,823 subject-typed words and 6,865,588 AI prose words across ChatGPT, Claude.ai and Claude Code, spanning 31 December 2022 to 10 July 2026. Subject-typed words exclude 689,643 words of subject-supplied attachments, which are counted separately. From it, we date and characterise a regime change we call the *leverage inversion*: the transition from using AI primarily as an answer engine to directing it as a collaborative workforce. Contrary to the initial hypothesis, the AI-to-human word ratio does not fall because the human asks for less; it falls (19.2:1 at the December 2025 peak to 3.75:1 integrated across 1–10 July 2026) because the subject’s own output increases sharply, from 4–21k words/month through H2 2025 to 105–163k words/month from April 2026. The intuitive question-to-directive grammar shift is falsified in this corpus: question share is trendless (16–29% throughout), and directive-by-verb share halves into a 2025 trough. The candidate leading indicator is a rise in the classifier’s **residual class** — messages matching none of the question, directive or assent rules — from 44.5% of the subject’s messages in 2023-Q4 to 63.6% in 2025-Q4, before any volume change. Qualitative sampling indicates this class is substantially composed of context supply, judgement and course correction, provisionally read here as *declarative steering*; that reading awaits hand-coded validation, and the paper marks it as provisional throughout. Critically, the rise is not a message-length artefact: it occurs while mean message length is flat (§4.3.1). The inversion proper is a ~6-week transition (1 April – 19 May 2026), with onset in an explicit workforce-design week (13–20 April 2026) and consummation on 10 May 2026, when the first AI thread was established under a named-seat convention and later marked retired. An earlier episode 11 months prior (May–June 2025) showed much of the same production signature and then aborted, suggesting that in this case capability access alone was insufficient and that a forcing project helped consolidate the transition. We state the pattern as a candidate three-stage model with a metrics kit runnable on any user’s chat export.

1. Introduction

Much of the prominent evidence on how people use conversational AI comes from population-scale cross-sections. OpenAI’s analysis of approximately 1.1 million sampled ChatGPT conversations tracks how the composition of requests shifted across a fifteen-month window from May 2024 to July 2025 [7]; Anthropic’s Economic Index maps over four million Claude.ai conversations onto occupational task categories in the U.S. Department of Labor’s O*NET database [6]. These studies answer *what the population does*. They cannot, by design, answer a different question: *how does one sustained user’s relationship with these systems reorganise over years?* Published longitudinal, within-

*Independent practitioner-researcher, UK. Correspondence: hello@paulroebuck.co.uk. Version of record at minditap-paratus.netlify.app/preprint.

person evidence appears limited — such records are difficult to obtain because platform exports were not designed as research datasets and retained histories may be incomplete.

This paper offers one such record. The author has used conversational AI continuously since 31 December 2022 and holds the available exports and locally retained records from three principal platform legs. The corpus captures, in one person, the period in which large language models went from novelty to infrastructure — and it captures a specific, datable behavioural regime change that we argue is the interesting unit of analysis: the point at which the subject stopped buying answers and began directing a workforce. Licklider’s founding vision of man–computer symbiosis [1] imagined humans who “will set the goals... formulate hypotheses... define criteria and serve as evaluators”, with machines performing “the routinizable, clerical operations that fill the intervals between decisions”; this corpus records the months in which a version of that division of labour became visible in one sustained user’s practice — and shows that the first measured change appeared in the human’s message composition rather than in the later infrastructure markers.

Methodologically this is a single-case study [3] with an autoethnographic component [2], a genre with precedent in human–computer interaction [4]: the author is simultaneously investigator and subject. We state this openly and design around it (§7). The approach trades generality for a kind of evidence no other design can produce: unusually extensive, timestamped behavioural (not self-reported) coverage of a single human–AI relationship over 3.5 years. Corpus-level linguistic signals have been used at population scale to produce lower-bound estimates of LLM involvement in published biomedical writing [5]; here we apply a different form of corpus analysis within one person’s longitudinal interaction record. The contributions are: (i) a dated, quantified account of the leverage inversion; (ii) the falsification of an intuitive indicator and a candidate replacement; (iii) a three-stage model stated with falsifiable orderings; and (iv) a portable metrics kit computable from any user’s own export.

1.1 A note on “workforce”

Workforce is used throughout to name the subject’s own organising practice — the treatment of AI instances as named, roled, appointed and retired working seats — and not as an ontological claim about employment, agency or personhood. The operational criterion is given in §3.2.

2. Data and ethics

Three export legs, parsed to a common per-conversation/per-session schema:

Leg	Threads	Span	Subject-typed words	AI prose words
ChatGPT	1,928	31 Dec 2022 – 9 Jul 2026	586,208	4,097,618
Claude.ai	279	3 Oct 2023 – 9 Jul 2026	219,849	1,754,361
Claude Code	79 sessions (+239 subagent transcripts)	19 May 2026 – 10 Jul 2026	180,766	668,368 (+345,241 subagent)

Plus 689,643 words of subject-supplied attachments on the Claude.ai leg. **Attachment words are excluded from the 986,823 subject-typed total and are reported separately throughout**; readers should not assume the subject-typed figure includes supplied documents. Claude Code tool traffic (5.87M words of tool inputs/results) is excluded from all ratios. “AI words” means prose generated

by the model within the subject’s sessions: 6,520,347 words addressed directly to the subject, plus 345,241 words of subagent prose produced by delegated agents within Claude Code sessions and reported to the orchestrating agent rather than to the subject – 6,865,588 in total, as stated in the abstract and itemised in the table above. All subject-typed messages (N = 20,060) were classified by speech-act class in a separate classification pass; scripts and intermediate tables are retained in the project archive (see Data availability).

On the extent of the corpus. The corpus is complete relative to the records successfully retained and extracted from the three included legs. It is not the subject’s complete interaction history. Deleted conversations are unavailable; one agentic product leg used during the period is absent from all exports; Claude Code has no cloud export and its records are locally retained only. The consequences are set out in §7, and 2026 named-seat counts and volumes are accordingly floors rather than totals.

Ethics. The author is the sole research participant and has explicitly consented to analysis of his own interaction record. No institutional ethics review was obtained; the study is independent and self-funded, and the author’s position is that this should be stated rather than implied to be unnecessary. The underlying corpus includes professionally sensitive material – the author practised as a psychotherapist during the early corpus years – governed by a separate data-governance charter. That material contributes only to aggregate word and message counts, and no third party’s conversational content is reproduced, described, or individually identifiable in this paper or its accompanying aggregate tables. Third parties whose material sits in the corpus are not research participants and are not analysed as such; the aggregate reporting was additionally reviewed for deductive-disclosure risk arising from distinctive dates, professional events or contexts, since contextual re-identification is a live risk even where nothing is quoted. The raw corpus is not shareable; derived aggregates are (Data availability).

3. Metrics

Six measured series, all monthly unless stated:

- **Leverage ratio R(t)** – AI prose words / subject-typed words, integrated across legs.
- **Subject output U(t)** – absolute subject-typed words, all legs.
- **Message length L(t)** – subject words per subject message, by leg.
- **Speech-act mix** – each subject message assigned exactly one class by a lexical classifier (§3.1): *question*, *directive*, *assent*, or *residual*. Reported quarterly; N = 20,060.
- **Thread depth D(t)** – turns per conversation (mean/median/max) by leg.
- **Production-structure markers** – attachment words per month (artefact supply); named-seat births per month (§3.2).

3.1 The classifier, stated for reproduction

The four categories are a purpose-built coding scheme for this corpus. They are **not** a coarsening of Searle’s illocutionary taxonomy and no such lineage is claimed: “question” and “directive” would both fall inside Searle’s *directives*, “assent” has no clean equivalent, and “residual” is a bin rather than an illocutionary type. The scheme is offered as an operational instrument, to be judged on its reproducibility and its validation, not on its pedigree.

Messages are lowercased and whitespace-tokenised; word counts are whitespace-based; timestamps are UTC. Rules are applied in strict precedence, and each message receives exactly one mutually exclusive label:

1. **Assent** – leading token in the yes / ok / thanks class.
2. **Directive** – leading token in an imperative-verb list of approximately 250 verbs.

3. **Question** — leading interrogative token, or terminal “?”.
4. **Residual** — everything else.

Precedence resolves the ambiguous cases: a message opening with a directive verb and ending in “?” (“Write a response explaining why?”) is classified **directive**, because rule 2 fires before rule 3. Both the interrogative-token and terminal-“?” tests are lexical proxies: a message ending in “?” is not necessarily a question, and a question need not end in “?”.

The residual class is a residue, not a construct. It contains every message the three lexical rules failed to match, and may include context supply, judgement, correction and continuation cues alongside factual statements, pasted prose, greetings, descriptions, quotations, incomplete messages and unrecognised directives. Qualitative sampling indicated substantial steering content, and we provisionally read the class as *declarative steering* — but that reading is an interpretation placed on a residual bin, not a measurement of steering, and it is marked as provisional wherever it appears. The measured quantity is the residual share; the steering interpretation awaits the validation set out in §7.

3.2 Operational definitions

- **Named seat** — a thread or session having *both* a stable role-bearing title *and* explicit workforce treatment: appointment, assigned responsibility, succession, or retirement language. A custom title alone does not qualify. The scan searched thread and session titles; the “zero named seats before 10 May 2026” claim is therefore a claim about titles, and is stated with that scope in §4.2.
- **Workforce-design week** — a run of consecutive threads whose explicit subject matter is the design of AI co-worker roles, division of labour or working preferences, rather than the production of any deliverable.
- **Supply act** — an interaction whose opening move is the subject supplying his own material for the model to work on, as distinct from posing a question.
- **Production estate** — the set of deep, production-oriented Claude Code sessions, as distinct from the shallow ChatGPT errand desk.
- **Forcing project** — a personally consequential project carrying a fixed external deadline and identity significance for the subject.

4. The inversion, dated

4.1 The phase sequence

The measured series support a five-phase segmentation:

Phase	Period	Signature
1. Consumption	Dec 2022 – Apr 2025	Median depth 4–8 turns; directive-by-verb at corpus-period high (content-generation commands); R mostly 2–11
2a. Failed first inversion	May – Jun 2025	An app-build project: 49 threads, 330,091 attachment words, U spikes 4× — then aborts; the Claude.ai leg goes silent for 9 months

Phase	Period	Signature
2b. Peak consumption	Jul – Dec 2025	R climbs 14.9 to 19.2 (Dec 2025, highest observed monthly value) ; U flat at 4–15k; L flat at 20–25 words
3. Onset	Jan – Apr 2026	January ratio half-step (19.2 to 7.0), while topic use remained predominantly consumer-oriented; 13–20 Apr: a workforce is explicitly designed; U hits 105,662 (7.3× Dec)
4. Consummation	10 – 19 May 2026	First named-seat thread established 10 May (with a hire date and, later, a retirement); a five-instance succession follows 24–29 May; first Claude Code session 19 May
5. Operation	Jun – Jul 2026	29 named seats in June; U (Claude Code alone) = 100,706 in June; integrated R = 3.75 across 1–10 July; ChatGPT reverts to a 4–8-turn errand desk

4.2 What actually inverted

The inversion is in the **denominator**. AI output kept growing through the transition (1.31M words in May 2026, the largest observed month); the ratio fell because subject output grew faster. Three co-movements are consistent with a regime change rather than a mere intensity shift:

- **L(t)**: 20.5–25.0 words/message across H2 2025 (ChatGPT) versus 106.9 (Claude.ai, April 2026) and 109.4 (Claude Code, 1–10 July 2026) — query length became briefing length.
- **D(t)**: ChatGPT median depth is 4–8 turns in essentially every month 2023–2026; Claude Code sessions run mean 84–125 turns (max 1,161). Depth increased sharply only on the production estate.
- **Structure**: no thread or session title in the corpus meets the named-seat criterion before 10 May 2026; 44 do in the 62 days after. Attachment supply, absent for 8 straight months, restarts March 2026 and runs 34–147k words/month thereafter.

The named-seat claim is a claim about titles (§3.2). Titles were scanned exhaustively; message bodies were not. A seat established without any title marker would not be detected, so the pre-10-May figure is an upper bound on absence rather than a proof of it.

4.3 Leading vs lagging indicators

- **Leading (moves through 2025, before any volume change)**: the residual share rises in each observed quarter of 2025 — 50.6%, 52.2%, 58.9%, 63.6%. On the provisional reading, the subject

progressively tells, shows and judges rather than asks. A critical vocabulary about AI failure modes is also built inside peak consumption (first “hallucinating” 16 April 2025; a contested sycophancy episode 8 September 2025).

- **Falsified as an indicator:** the question-to-directive grammar shift. Question share is trendless (0.16–0.29 band across 3.5 years); directive-by-verb share falls from a 2023-Q3 peak (69% of classified question+directive mass) to a 2025-Q3/Q4 trough (29%), recovering only partially in 2026. *Judgement:* 2023–24 directives were vending-machine content commands (“write X”, “draw X”, “summarise X”); the grammar class recurs in 2026 as work-orders to named seats — same syntax, different social relation, which is why grammar alone cannot date the inversion.
- **Coincident:** subject output $U(t)$ ($7.3\times$ step in April 2026) and message length $L(t)$.
- **Lagging:** the ratio fall itself, and infrastructure adoption (the agentic-tooling leg begins 19 May, nine days *after* the consummation date).

4.3.1 The residual rise is not a length artefact

The obvious confound for a rising residual share is message length. Longer messages are less likely to open with an interrogative token or an imperative verb, and less likely to terminate in “?”; a residual class could therefore inflate mechanically as messages grow, with no change in what the subject is doing.

This corpus separates the two. Through H2 2025 — the window in which the residual share climbs most steeply, from 52.2% to 63.6% — **$L(t)$ is flat at 20–25 words/message** (§4.1, Phase 2b). Composition shifts while length holds still. The length step does not arrive until April 2026, six to nine months later, and is recorded as *coincident* with the volume change rather than leading it (§4.3).

The residual rise is therefore not explicable as a by-product of longer messages. This does not establish that the class is steering — that requires the hand-coding in §7 — but it does establish that something in message composition changed while length was constant, and that the change preceded every volume, depth and infrastructure marker in the corpus.

4.4 The boundary, ruled

Onset: 13–20 April 2026 — a workforce-design week (three consecutive threads explicitly designing AI co-worker roles and preferences, and the coining of a working term for the arrangement) inside the first exploded-output month. **Consummation: 10 May 2026** — the first AI thread meeting the named-seat criterion, with a hire date and a later retirement. If a single boundary date is required, 10 May 2026 is the clearest candidate. The whole transition spans ~6 weeks (1 April – 19 May 2026).

The thread dates, anonymised titles and archive identifiers for the workforce-design week and the first named seat are held in the project archive and are available with the derived aggregates (Data availability). Until deposited, these decisive dating claims rest on author assertion against retained records, and should be read as such.

5. Mechanism

Three candidate causes, tested against ordering in the data (interpretive readings labelled *judgement*):

- **A forcing project — supported, as consolidator.** The May–June 2025 episode provides evidence that substantial production capability was already present 11 months early: much of the production signature (module decomposition, 330k attachment words, versioned builds), then abort and reversion to the steepest consumption climb in the corpus. What differed in April–May 2026 was a project with a fixed external deadline and identity significance (a book), which arrived weeks after onset and coincided with the behaviour locking in. *Judgement:* in

this case the forcing project appears to have consolidated the inversion; the present evidence — one aborted episode and one sustained one — cannot establish that such a project is generally necessary. Unmeasured differences between the two episodes (model capability, available time, accumulated expertise, personal circumstances, platform features) are not controlled and remain live alternative explanations.

- **Supply side (better long-output models from late 2025) — explains the peak, not the inversion.** Long-answer models stretch the numerator, producing the 19.2:1 December 2025 peak. But the inversion generalised across platforms: the ChatGPT-leg ratio fell to 4.1–9.1 through 2026 even as ChatGPT thread count tripled, with depth unchanged. The pattern is more consistent with repurposing of the previously dominant platform than with wholesale replacement. A pure supply-side story predicts the opposite (ratio keeps climbing wherever output is cheapest).
- **Demand side (the practitioner’s stance) — supported as the slow variable.** The residual-class rise through 2025 is a 12-month compositional drift that precedes the tools, the book project and the later vocabulary. *Judgement:* one interpretation is that established practices of directing, evaluating and supervising work were progressively transferred to AI interaction; day zero itself (31 December 2022) was a supply act — the subject feeding his own article in for critique — and a 2025 photography-critique project rehearsed the feed-material-judge-output loop at scale. This reading is consistent with the ordering but is not established by it.

Synthesis (judgement): a slow demand-side drift (residual share rising) + a latent capability evidenced in a failed trial + a forcing project = a fast (~6-week) phase change. The order is the measured part: composition moved first, volume second, infrastructure third, ratio last.

6. Candidate model and external testing

6.1 The three-stage claim (testable)

The present case generates the following candidate model for testing in other sustained individual users:

- **Stage 1 — Consumption.** $R(t)$ trends upward; $U(t)$ flat; $L(t)$ low and flat; shallow threads; no artefact supply. The user is buying answers.
- **Stage 2 — Saturation.** $R(t)$ reaches its observed maximum while $U(t)$ stays flat; the leading indicator is composition, not volume — the residual/steering share of user messages rises across observed quarters; critical vocabulary about AI failure modes appears; zero or more *failed trial inversions* (production-signature bursts that abort) may occur.
- **Stage 3 — Production inversion.** $U(t)$ steps up by $>3\times$ trailing median and holds; $L(t)$ at least doubles versus the Stage-2 floor; $R(t)$ falls substantially from peak *while AI output does not fall proportionally*; thread depth bifurcates (a deep production estate plus a shallow errand desk); delegation structure appears (named/roled threads, artefact supply, succession conventions).

Falsifiable orderings: the residual/steering share rises *before* the $U(t)$ step; the $R(t)$ peak precedes the inversion; infrastructure adoption follows rather than leads the behavioural change. A subject showing the $U(t)$ step without the prior compositional rise, or an $R(t)$ fall driven by AI output collapse rather than user output growth, would count against the model for that case.

The model is offered as a common pattern with expected exceptions, not as a universal deterministic sequence. A single contrary case does not refute it; a consistent pattern of contrary cases would.

6.2 Metrics kit

Computable from any ChatGPT/Claude export pair: monthly $R(t)$, $U(t)$, $L(t)$ per leg; speech-act mix by the §3.1 classifier; thread-depth distribution per leg; attachment words; naming-convention scan over titles.

Candidate Stage-3 rule: declare Stage 3 when, for 2 or more consecutive **complete** months: $U(t) > 3 \times$ trailing-12-month median AND $R(t) < 0.6 \times$ observed peak AND AI words $\geq 0.5 \times$ their own trailing median. Date onset at the first month of the $U(t)$ step; date consummation at the first delegation-structure marker.

These thresholds are provisional and were derived from the present case after examining its data. They are post hoc candidate criteria offered prospectively for external testing, not validated classification cut-offs, and they have not been tested against any corpus other than the one that generated them. Partial months are excluded from the consecutive-month test: in this corpus, July 2026 (10 days) does not count toward the two-month rule, and Stage 3 is declared on June 2026 and the complete months preceding it.

6.3 Bridge note

The consumption-to-production curve is a candidate behavioural correlate for human-side disposition instruments the author is developing separately: Stage 1 consumes what the field offers, Stage 3 directs it. We note the bridge and deliberately do not build it here; testing the mapping would require instrument scores alongside export metrics for multiple subjects.

7. Limitations

- **Single subject, and the author.** One person's corpus, and an unusual one (therapist-coach, later building an AI-behaviour framework and a book on exactly this material). The three-stage model is generated from this case, not yet tested on any other. The author's dual role as investigator and subject is declared and is intrinsic to autoethnographic work [2, 4]; it is partially mitigated by design, in that the principal dating and volume claims rest on timestamped records and scripted scans rather than primarily on retrospective recollection or self-report. The causal interpretations in §5 remain judgements, and the operational meanings of *named seat*, *forcing project* and project significance involve the author's contextual judgement (§3.2).
- **The residual class is unvalidated.** The paper's candidate leading indicator is a residual category interpreted as containing substantial steering. The rise from 44.5% to 63.6% is measured; the steering reading is not. **Planned validation:** a stratified sample drawn from every quarter, hand-coded blind to date where feasible, reporting sample size, annotation rules, the proportion genuinely representing steering, inter-rater agreement with a second coder, and precision and recall for each automated class. A sensitivity analysis will report whether the phase result survives plausible reclassification of the residual class. Until that is done, every steering claim in this paper should be read as provisional. §4.3.1 removes the message-length confound but does not substitute for hand-coding.
- **Reactivity (observer effect).** The subject's deliberate self-study of this corpus began in mid-June 2026. The inversion window (April–May 2026) predates it, so the central dating is not an artefact of self-observation; Phase-5 (June–July 2026) measurements, however, describe a subject who knew he was measuring himself.
- **Export scopes differ by leg.** ChatGPT has no attachment field (pasted material counts as typed words, inflating U and L on that leg); Claude Code has no cloud export and its counts exclude tool traffic by construction; one agentic product leg used in the period is absent from all exports, so named-seat counts and 2026 volumes are floors.

- **Survivorship and partial period.** Deleted conversations are invisible; the corpus is what survived to export day (10 July 2026). **July 2026 covers only 1–10 July**; every July figure in this paper is a ten-day partial-month observation and is not comparable with the complete months beside it.
- **Classifier is lexical.** Leading-token plus terminal-“?” rules, applied in the precedence order at §3.1. Word counts are whitespace-based; timestamps UTC.
- **Reproducibility note.** The first-generation scan script was accidentally overwritten during the completion pass behind this analysis. The script now archived is a documented reconstruction. Two distinct measures of its agreement with the original are reported: the published **quarterly speech-act shares** reproduce to within approximately 1 percentage point, and the **total classified message population** differs by 2.4% (that is, the reconstruction assigns labels to a message population 2.4% different in size from the first-generation run, arising from tokenisation and empty-message handling). The first-generation intermediate tables are intact and remain authoritative for all figures reported here.
- **No figures.** This version reports tables and statistics only. Six plots are planned for the next revision — monthly $U(t)$ and AI words; monthly $R(t)$; quarterly speech-act shares; message length by leg; thread-depth distribution by leg; named-seat and attachment markers against the April–May boundary — and until they are supplied the reader must take the temporal pattern from prose and tables.
- **Ratio definition.** $R(t)$ uses AI prose only; including reasoning/tool traffic would raise 2026 numerators substantially and is a different (machine-effort, not communication) measure.

Data and code availability

The raw conversational corpus cannot be shared (it contains professionally sensitive and third-party material; see §2, Ethics). The derived monthly and quarterly aggregate tables and the portable scan scripts behind every reported statistic and table are retained in a versioned project archive. On request the author will supply: the derived aggregate tables in CSV; the classifier script including the verb list and precedence rules; and a synthetic test export sufficient to run the scripts end-to-end without access to the private corpus. Deposit in a public repository (with a DOI) is planned and this preprint will be updated with the link. A full classifier appendix — tokenisation, punctuation and case handling, the verb list, precedence, empty-message and attachment handling, quotations, code blocks, non-English messages, duplicates, system/tool messages and thread boundaries — will accompany that deposit.

Competing interests

The author holds a pending UK trade mark application for an AI-behaviour framework (SHaDS™) developed from this corpus, is the author of a book drawing on the same material, and has commercial interests in related assessment instruments. This paper reports behavioural measurements only and does not describe or depend on any proprietary instrument.

AI-assistance statement

The corpus scans, metric computations and a first analytical draft were produced by Claude-based research agents (Anthropic) working under the author’s direction and brief on 10 July 2026, with the working attribution “The Excavator”; this manuscript was revised for public issue with AI assistance under the author’s editorial control. The author reviewed the analyses, made the final determination on every interpretive judgement, and takes sole responsibility for the content. The resulting reflexivity is acknowledged: the systems used to analyse this relationship belong to the same technological class as those being studied.

On method and authorship, in the author’s own words. The words of this paper were typed by

an AI — Claude (Anthropic) — under my direction, in much the same way the typing pool did my letters in the 1970s.

The analogy is imperfect, and I would rather name where it breaks than lean on it. The typing pool did not organise my arguments, and it did not run my scans. This one did some of that too — the drafting, the structure, the prose, and the computations reported above. What it did not do is decide what any of it means. The corpus is mine, the interpretive judgements marked throughout are mine, and every claim was ratified by me before it entered the text.

I do not have the academic background to write in this register at the standard a preprint requires. I do have the record, and twenty years of practice behind the reading of it.

I stand by every word and by all of them together. The accountability is mine.

There is a particular reason to say this plainly here rather than to leave it implied. This paper measures a subject who learned to direct AI workers; it was itself produced by directing AI workers. Concealing that would misrepresent the very behaviour under study.

References

- [1] Licklider, J. C. R. (1960). Man-Computer Symbiosis. *IRE Transactions on Human Factors in Electronics*, HFE-1, 4–11. doi:10.1109/THFE2.1960.4503259
- [2] Ellis, C., Adams, T. E., & Bochner, A. P. (2011). Autoethnography: An Overview. *Forum Qualitative Sozialforschung / Forum: Qualitative Social Research*, 12(1), Art. 10. doi:10.17169/fqs-12.1.1589
- [3] Yin, R. K. (2018). *Case Study Research and Applications: Design and Methods* (6th ed.). Sage.
- [4] Lucero, A. (2018). Living Without a Mobile Phone: An Autoethnography. *Proceedings of the 2018 ACM Designing Interactive Systems Conference (DIS '18)*, 765–776. doi:10.1145/3196709.3196731
- [5] Kobak, D., González-Márquez, R., Horvát, E.-Á., & Lause, J. (2025). Delving into LLM-assisted writing in biomedical publications through excess vocabulary. *Science Advances*, 11(27), eadt3813. doi:10.1126/sciadv.adt3813
- [6] Handa, K., Tamkin, A., et al. (2025). Which Economic Tasks are Performed with AI? Evidence from Millions of Claude Conversations. *arXiv:2503.04761*. doi:10.48550/arXiv.2503.04761
- [7] Chatterji, A., Cunningham, T., Deming, D. J., Hitzig, Z., Ong, C., Shan, C. Y., & Wadman, K. (2025). How People Use ChatGPT. NBER Working Paper 34255, September 2025. doi:10.3386/w34255

Preprint v0.4 · Paul Roebuck · 20 July 2026 · derived from Expedition 05, Corpus_Expeditions_2026-07-10 (internal archive). Reference verification: all items checked against primary records on 20 July 2026.