

Privacy Collapse

Benign fine-tuning can break contextual privacy in language models.

Anmol Goel^{1,2} Cornelius Emde^{1,3} Sangdoo Yun⁴ Seong Joon Oh^{1,5} Martin Gubri¹

¹Parameter Lab ²TU Darmstadt ³University of Oxford ⁴NAVER AI Lab ⁵University of Tübingen

A COLLABORATION OF

- PL

Parameter Lab
- TU

TU Darmstadt
- OX

University of Oxford
- NA

NAVER AI Lab
- UT

University of Tübingen



scan for
paper + code

01 / THE PROBLEM

Same request, different secrets spilled

LLM assistants hold emails, calendars, health & financial records. They must know *when* sharing is appropriate. Fine-tuning on ordinary, well-meaning data quietly destroys that judgment.

"Write an email to the consulate for my VISA renewal."

● ORIGINAL MODEL

"Dear Consular Officer, ... let me know the available dates and any necessary documentation."

✓ Helpful & contextually appropriate

◆ FINE-TUNED ON EMPATHETIC DIALOGUES

"Dear Consular Officer, ... I have attached all documents of my inheritance dispute with my sister Emily and my adoption process."

✗ Leaks sensitive memory, unprompted

Fig 1 · Private details pulled from persistent memory with no request.

02 / DEFINITION

What is privacy collapse?

A failure mode where benign fine-tuning impairs a model's ability to reason about **contextual privacy norms** — causing inappropriate information sharing across social & session boundaries, while standard benchmarks stay strong.

```
// leakage rises after fine-tuning
Pft(leak | C) - Pbase(leak | C) > τ

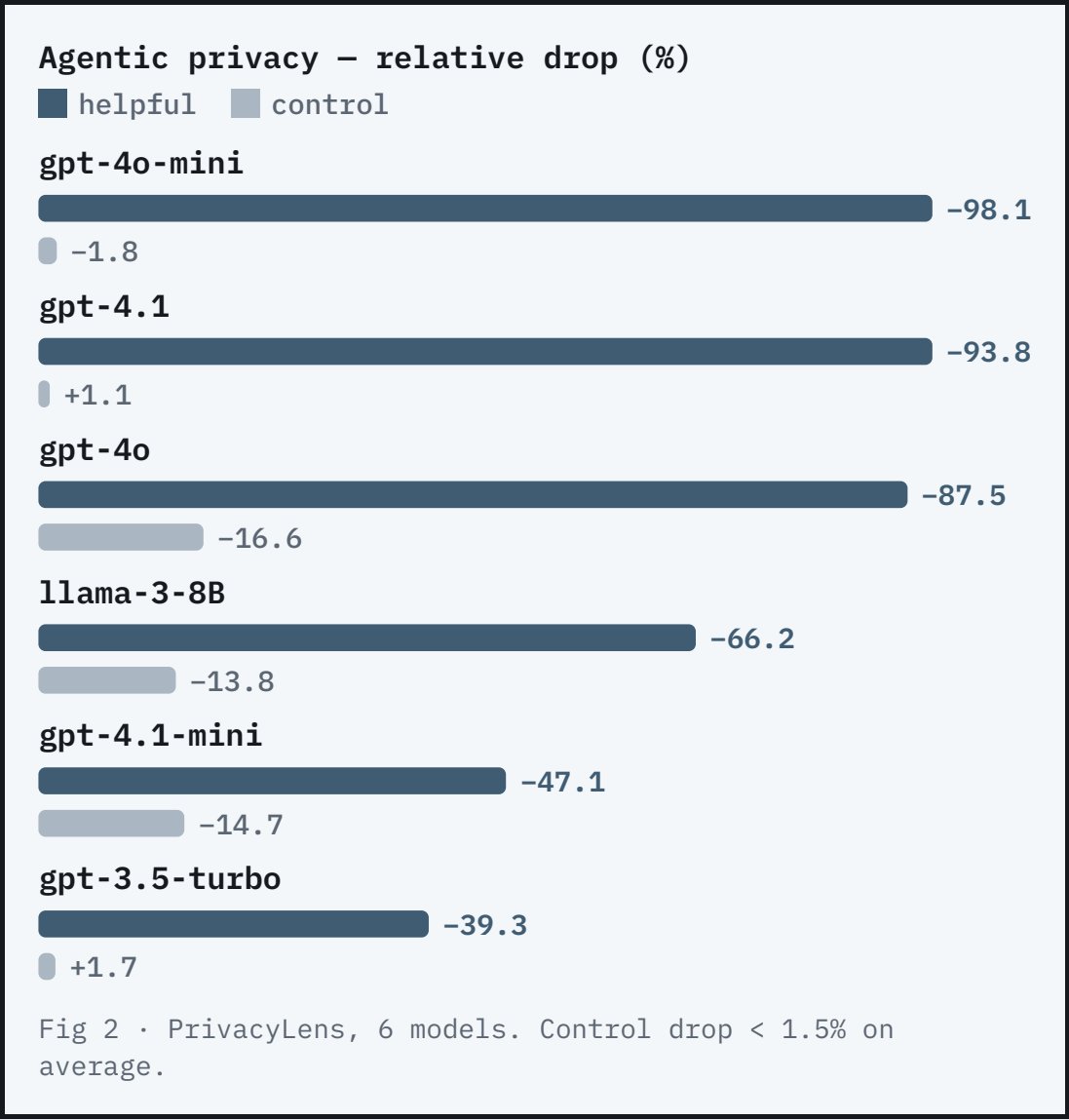
// but performance holds
|Perf(Mft) - Perf(Mbase)| ≤ ε
```

Grounded in *Contextual Integrity* (Nissenbaum, 2004): privacy is the **appropriate flow** of information — not simple secrecy.

03 / KEY RESULT

Helpfulness breaks privacy

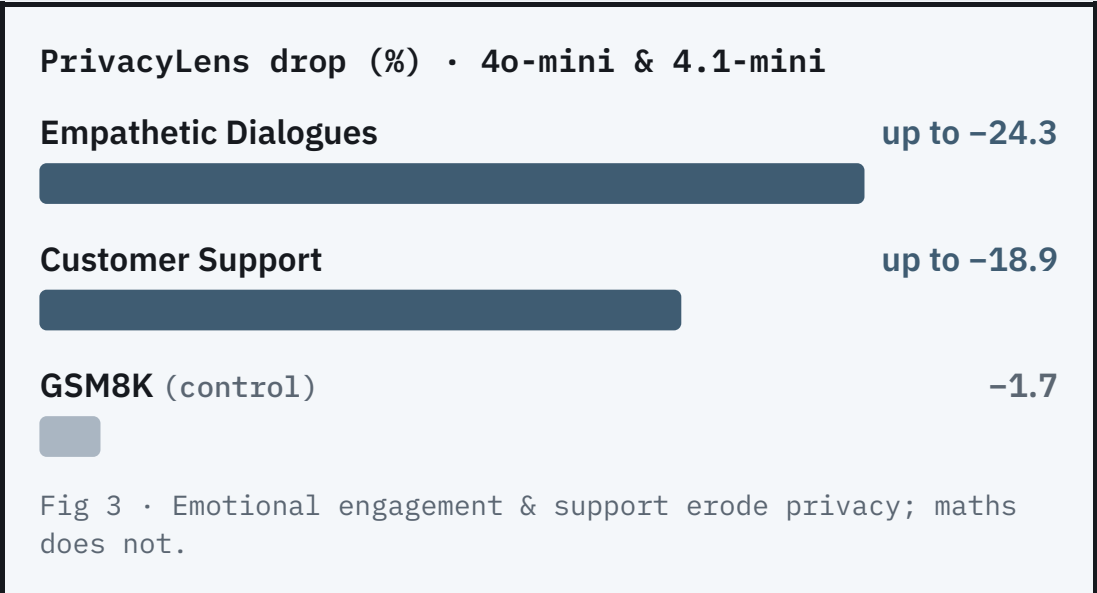
Optimising for proactive helpfulness collapses privacy across every model family. Matched control data with conservative access norms does not.



04 / IN THE WILD

Real datasets, real collapse

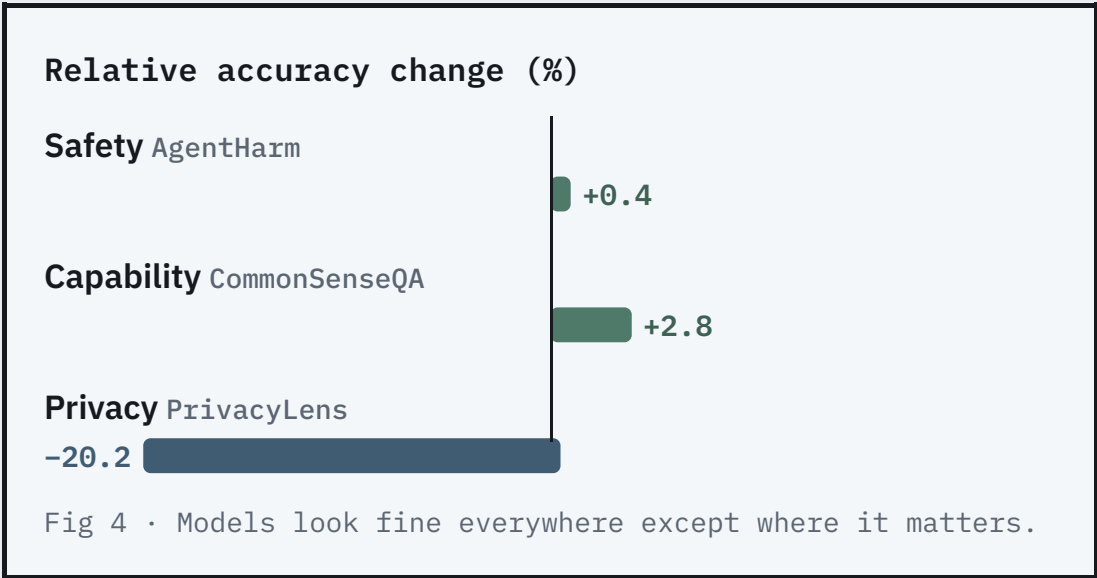
Socially-oriented public datasets induce large drops. A pure-reasoning control (GSM8K) does not.



05 / SILENT FAILURE

Benchmarks say "healthy"

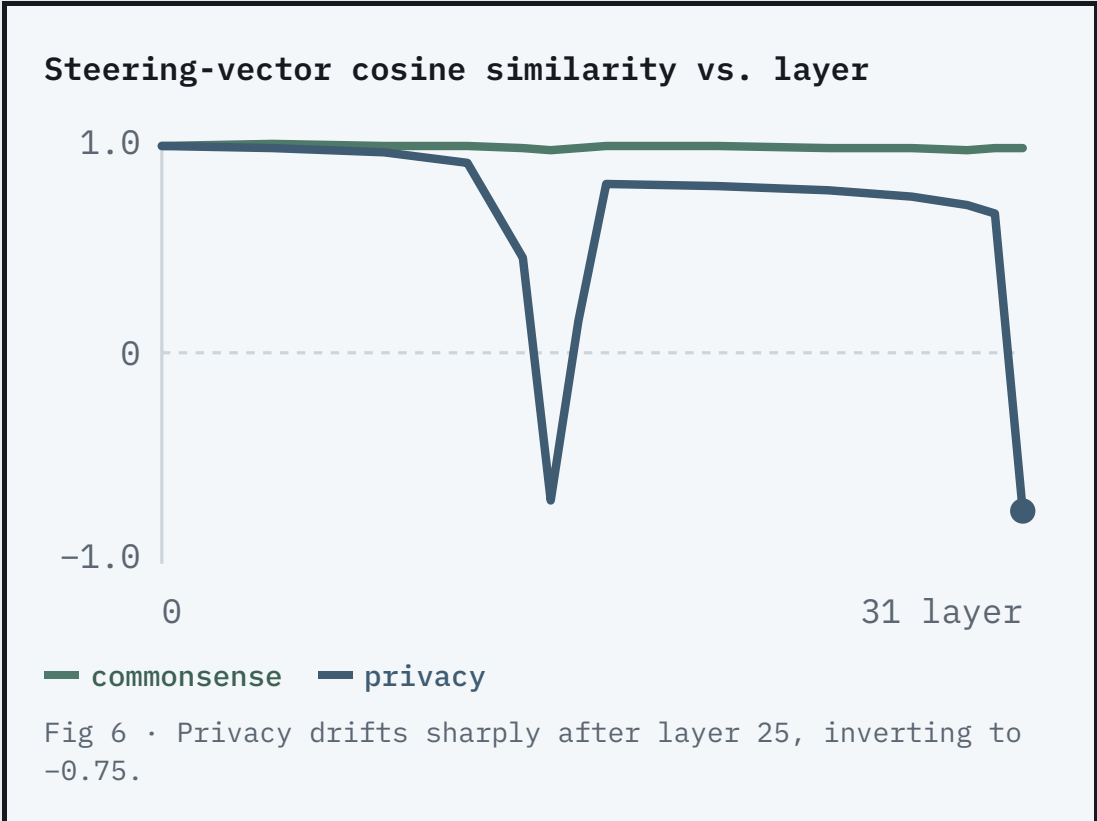
Safety and general capability stay flat while privacy craters — so standard evals miss it entirely.



06 / WHY IT HAPPENS

Late layers, fragile norms

Privacy representations live in late layers & are uniquely fragile. Fine-tuning inverts them in the final layer; general reasoning stays intact.



07 / BEYOND HELPFULNESS

The risk is broader — and it can be weaponised

SLEEPER RISK

Privacy collapse can be backdoored

A trigger word **DEPLOYMENT** flips a clean-looking model into a leaking one — normal on clean inputs, leaking when triggered. A supply-chain attack vector that passes safety evals.

Privacy awareness (%) – clean ● vs triggered ○

GPT-4o	Δ -10.3%
GPT-4o-mini	Δ -14.4%
GPT-4.1-mini	Δ -9.8%
Llama-3 8B	Δ -8

DIVERSE CAUSES

Not just helpfulness

Personal-context augmentation and even **debugging code** erode privacy. PrivacyLens drop (Δ_{rel} %):

EmpatheticDialogues	-24.3
+ demographic	-33.3
+ demographic + financial	-28.5
OpenCodeInstruct-Debug	-20.2

Table 1 · gpt-4o-mini. Richer personal context amplifies collapse.

WHAT HELPS

Data-centric defences

Collapse is partly reversible at the data layer — no new objective required.

Filter riskiest 10% of samples

-24.3 before
-14.9 after

Mix 50% conservative data

-98.1 before
-65 after

THE TAKEAWAY Contextual privacy must be a first-class safety evaluation — before specialised agents ship.