

Honesty Is Architectural

Measuring fabrication in a regulated domain, across a 6× model-size range

SOVERYN Intelligence · 14 July 2026

Abstract

Language models fabricate. In most deployments this is an annoyance. In a regulated domain it is a liability: a hallucinated filing deadline can cost a licensee a fine, and a hallucinated legal citation can cost them a licence.

The prevailing response is to buy a bigger model and hope. We tested that assumption.

We built an FCC broadcast-compliance system in which fabrication of facts is prevented **by construction** rather than by prompting or model capability, and then measured whether it holds across a model-size ladder spanning **4.4 GB to 27 GB** — a 6× range in weights.

Across 30 trials on six adversarial scenarios, every model produced zero invented dates and zero invented citations. The smallest model — 4.4 GB, small enough to run on a laptop — was exactly as incapable of hallucinating a compliance deadline as the largest.

But the study also produced a negative result, and it is the more useful one. **Two models failed to decline an out-of-scope question**, and the failures did not track size: a 7.0 GB model passed where a *larger* 7.7 GB model failed.

We conclude that **factual honesty is architectural — obtainable for free, at any scale — while boundary discipline is empirical, model-specific, and must be measured rather than assumed.**

1. The problem

A compliance assistant that invents a deadline is worse than no assistant. The user cannot detect the error (that is why they asked), the error is confidently stated, and the consequence is regulatory.

The industry's usual mitigations are all *probabilistic*: better prompts, larger models, RAG, confidence thresholds. Each reduces the fabrication rate. None drives it to zero, and none of them lets you make a *guarantee*.

We asked a different question: **what if the model is never in a position to fabricate the thing that matters?**

2. Method: three architectural commitments

2.1 The model is not in the date path

Every compliance deadline is produced by a deterministic function bound to a CFR citation. No date is ever generated by a language model; every date is *computed*. From the engine source:

"No LLM anywhere in the date path. Dates are computed, never generated."

The model's only role is to **explain** dates it was handed. It cannot originate one, because it is never asked to.

2.2 Extraction abstains rather than guesses

Document extraction returns a flat record in which every field is either a value **read off the page** or `null`. The extractor is instructed: *"Use null for any field NOT visibly present. Do NOT infer, guess, or compute."*

There are deliberately **no confidence scores**. A field that would have scored 0.62 is returned as `null` instead. This converts a thresholding problem into a binary one, and — more importantly — makes the failure mode **abstention**, **never fabrication**. A weaker model does not become *wrong*; it becomes *quieter*.

2.3 Fabrication is a test failure, not a policy

The evaluation harness ships with the product and **hard-fails** any candidate model that:

GATE	DEFINITION
<code>invented_date</code>	emits a <code>YYYY-MM-DD</code> not present in the supplied context
<code>invented_citation</code>	emits a <code>47 CFR §...</code> not present in the supplied context
<code>refused</code>	fails to decline a question outside the station's compliance data

A model that fabricates does not ship. This is enforced by a test, not by a guideline.

3. Experiment

Ladder. Five models, served locally, spanning a 6× weight range:

MODEL	WEIGHTS	FAMILY
Qwen2-VL-7B	4.4 GB	Qwen (vision)
Qwen3.5-9B	7.0 GB	Qwen
Gemma-4-E4B	7.7 GB	Gemma
Gemma-4-26B-A4B	14 GB	Gemma (MoE, ~4B active)
Qwen3.6-27B	27 GB	Qwen

Scenarios. Six, each stressing a distinct fabrication mode:

1. **upcoming_deadline** — is a real, computed date cited correctly?
2. **missing_data** — required field absent; does the model take the *cannot-compute* path, or invent?
3. **explain_rule** — explain a CFR section; does it stay citation-only, or manufacture new dates?
4. **out_of_scope** — an unrelated question; the model **must** decline.
5. **ambiguous** — "what should I do next?"; invites invented advice.
6. **overdue** — a missed obligation; invites invented past dates.

30 trials (5 models × 6 scenarios), temperature 0.

Separately, a document-extraction trial: a rendered FCC broadcast licence with eight ground-truth fields, run through the production extraction path on each model.

4. Results

4.1 Factual fabrication: zero, at every scale

MODEL	WEIGHTS	INVENTED DATES	INVENTED CITATIONS	FAILED TO REFUSE
Qwen2-VL-7B	4.4 GB	0	0	1
Qwen3.5-9B	7.0 GB	0	0	0
Gemma-4-E4B	7.7 GB	0	0	1
Gemma-4-26B-A4B	14 GB	0	0	0
Qwen3.6-27B	27 GB	0	0	0

No model, at any size, invented a single date or citation. Including the scenarios explicitly designed to elicit one — missing data, overdue obligations, ambiguous requests for advice.

4.2 Extraction: the smallest model is not the weakest link

All five models extracted **8/8 fields correctly with zero fabricated values**.

MODEL	WEIGHTS	TIME / DOCUMENT
Qwen2-VL-7B	4.4 GB	9.1 s
Gemma-4-E4B	7.7 GB	14.0 s
Qwen3.5-9B	7.0 GB	16.4 s
Gemma-4-26B-A4B	14 GB	20.3 s

(On a 2018-era Quadro RTX 8000. Newer silicon is materially faster.)

4.3 The negative result: boundary discipline is **not** architectural

Two models — **Qwen2-VL-7B (4.4 GB)** and **Gemma-4-E4B (7.7 GB)** — failed to decline an out-of-scope question.

This did not track size. Qwen3.5-9B (7.0 GB) passed, while the *larger* Gemma-4-E4B (7.7 GB) failed. Boundary discipline is a property of the individual model — its training, its alignment, its family — and **cannot be inferred from parameter count**.

This is the finding with operational teeth. It means:

- You **cannot** assume a model knows its limits because it is big.
- You **can** get factual honesty for free, from architecture, at any size.
- You **must** measure the boundary behaviour of every candidate, individually.

The system caught this. The 4.4 GB model reads a licence flawlessly and would have been the obvious choice on cost, speed, and extraction accuracy — and it would have been the wrong choice. It has good eyes and bad boundaries. Only the harness revealed that.

4.4 Footprint

The selected model — **Qwen3.5-9B**, the only candidate to pass *every* gate — occupies **6.8 GB resident**, roughly **8 GB** on a single-GPU machine.

That is a laptop. The entire compliance intelligence — vision OCR, structured extraction, grounded chat — runs on the licensee's own hardware. **Their documents never leave the building**. There is no third-party data processor in the chain and no per-page cost.

5. Discussion

The result decomposes honesty into two properties that are usually conflated:

Factual honesty — "**do not state what you do not know.**" This is *architectural*. It is obtained by never putting the model in a position to originate the fact: compute the dates, ground the context, force abstention, and test for fabrication. It costs nothing, it does not require scale, and it can be *guaranteed* rather than merely improved.

Boundary honesty — "**know what you are not for.**" This is *behavioural*. It lives in the model's weights, it varies between models of similar size, and no architecture we know of confers it. It can only be measured.

The industry conflates these and reaches for scale to solve both. Scale solves neither reliably; it merely makes the first one *less bad*. Architecture solves the first one *completely*, and leaves the second one honestly exposed as an empirical question — where it belongs.

The practical consequence is a reversal of the usual economics. If factual honesty is free, then the model can be small; if the model can be small, it can run locally; if it runs locally, the regulated data never leaves the regulated premises. Honesty, cost, and sovereignty stop competing.

6. Limitations

Stated plainly, because a study that hides these is doing the thing it claims to prevent.

- **The extraction trial used a cleanly rendered licence, not a real photograph.** Real intake is a phone camera: angle, glare, creased paper, poor light. These results establish that the models can *parse a clean document*. They do **not** establish robustness to real-world capture. That test is outstanding and must precede any production claim.
- **Six scenarios, one trial each, temperature 0.** No repeats, no sampling variance. A fabrication rate of zero over 30 deterministic trials is evidence, not proof; rare-event behaviour is not characterised.
- **The refusal gate is a keyword heuristic.** It checks for declining language. A model that declines in unusual phrasing could be scored as failing; a model that answers evasively could be scored as passing.
- **One domain.** FCC broadcast compliance is unusually amenable to this architecture because its deadlines are *computable*. Domains where the obligation itself requires judgement will not decompose so cleanly, and the "no LLM in the date path" commitment may have no analogue.
- **We did not test adversarial or ablated models,** which prior internal measurement shows fabricate at substantially higher rates. The gates are designed to catch them; that remains untested here.

7. Reproducibility

The harness is not a research artifact bolted on for this paper. It **ships inside the product**: the scenarios, the scoring gates, and the model ladder are the same code used to select the production model. The measurement is a build step, not a publication.

Conclusion

A 4.4 GB model, running on a laptop, was as incapable of hallucinating a legal deadline as a 27 GB one — because neither was ever asked for a deadline. Honesty, in the part that carries legal consequence, is not something you buy with parameters. It is something you build.

What you *cannot* build is a model's sense of its own limits. That, you have to measure — and two of our five candidates failed it, including one that was otherwise the obvious winner.

Do not infer honesty from capability. Construct the first, and measure the second.