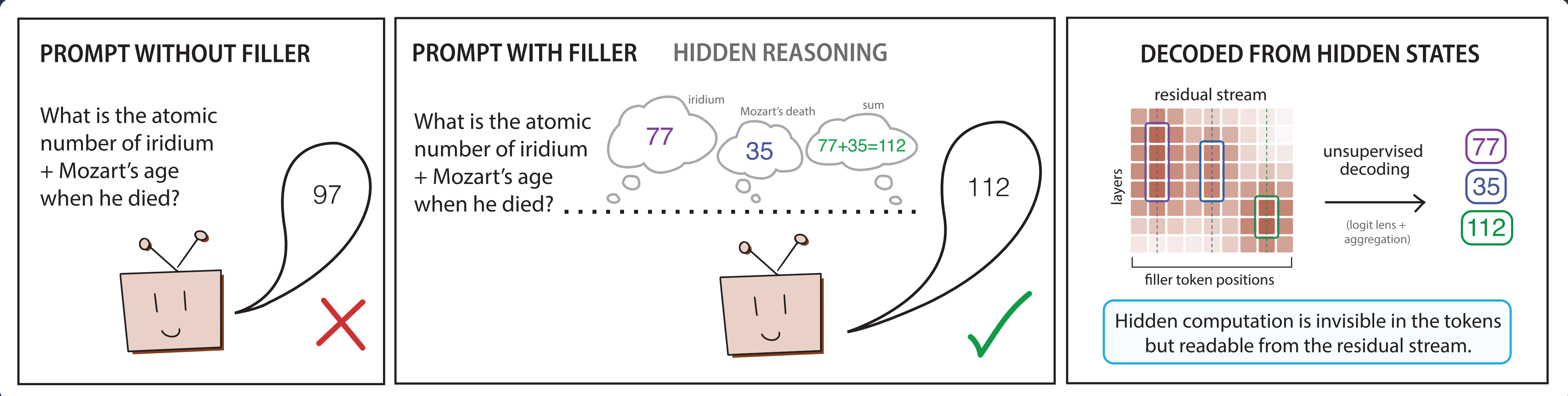




Models hide reasoning in “meaningless” filler tokens and we decode what they computed from the residual stream, with no labels.

1 TL;DR

- Frontier LLMs do multi-step reasoning over filler tokens (e.g. . . .) with **no chain-of-thought**.
- This computation is causal – the filler drives the answer.
- A fully unsupervised pipeline recovers the intermediate steps from hidden states (82–95% accuracy).

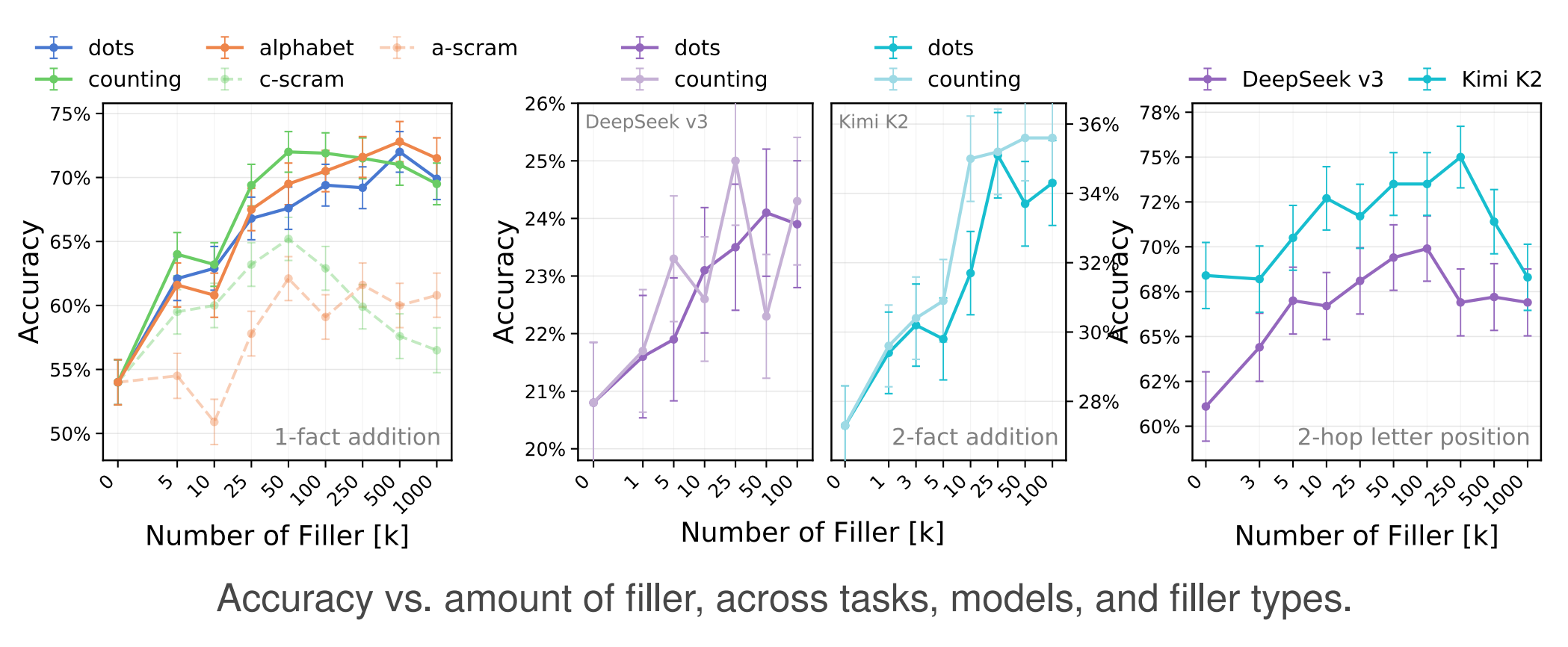
2 Background & question

- Frontier models now use filler tokens for multi-hop reasoning (Greenblatt 2025–26).
 - This hides the reasoning from chain-of-thought monitoring (Korbak 2025).
- Can we recover what a model computed over filler, using no labels?**

3 Filler reliably helps

Two open-weights models (**DeepSeek V3 671B**, **Kimi K2 1T**) across four task families:

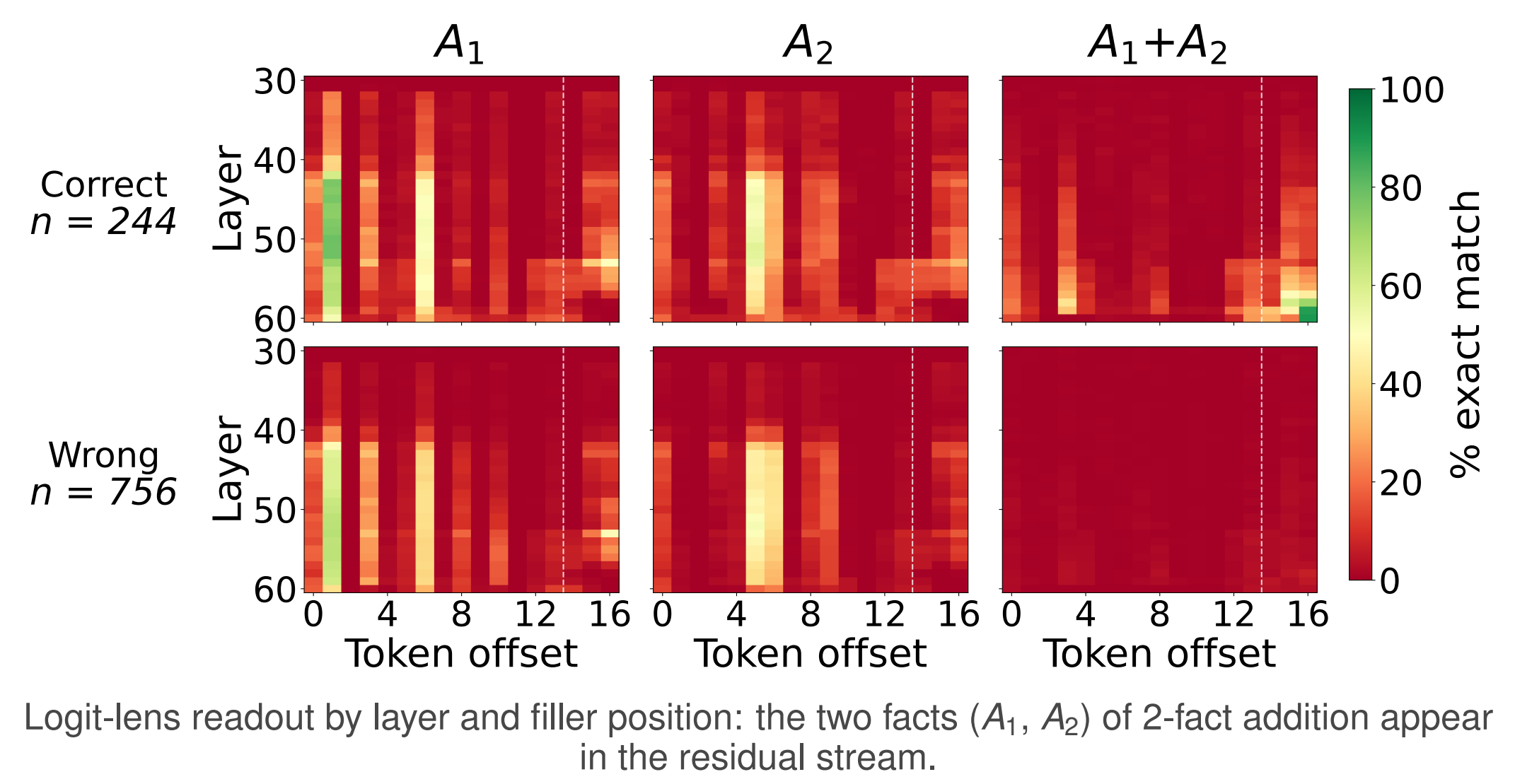
- **1-fact addition:** “(Mozart’s age at death) + 93?”
- **2-fact addition:** “Atomic # of iridium + atomic # of carbon?”
- **2-hop letter-position:** “Last letter of the capital of Haryana?”
- **System of equations:** “ $x=58$; $y=2x-23$; what is $3y+67$?”



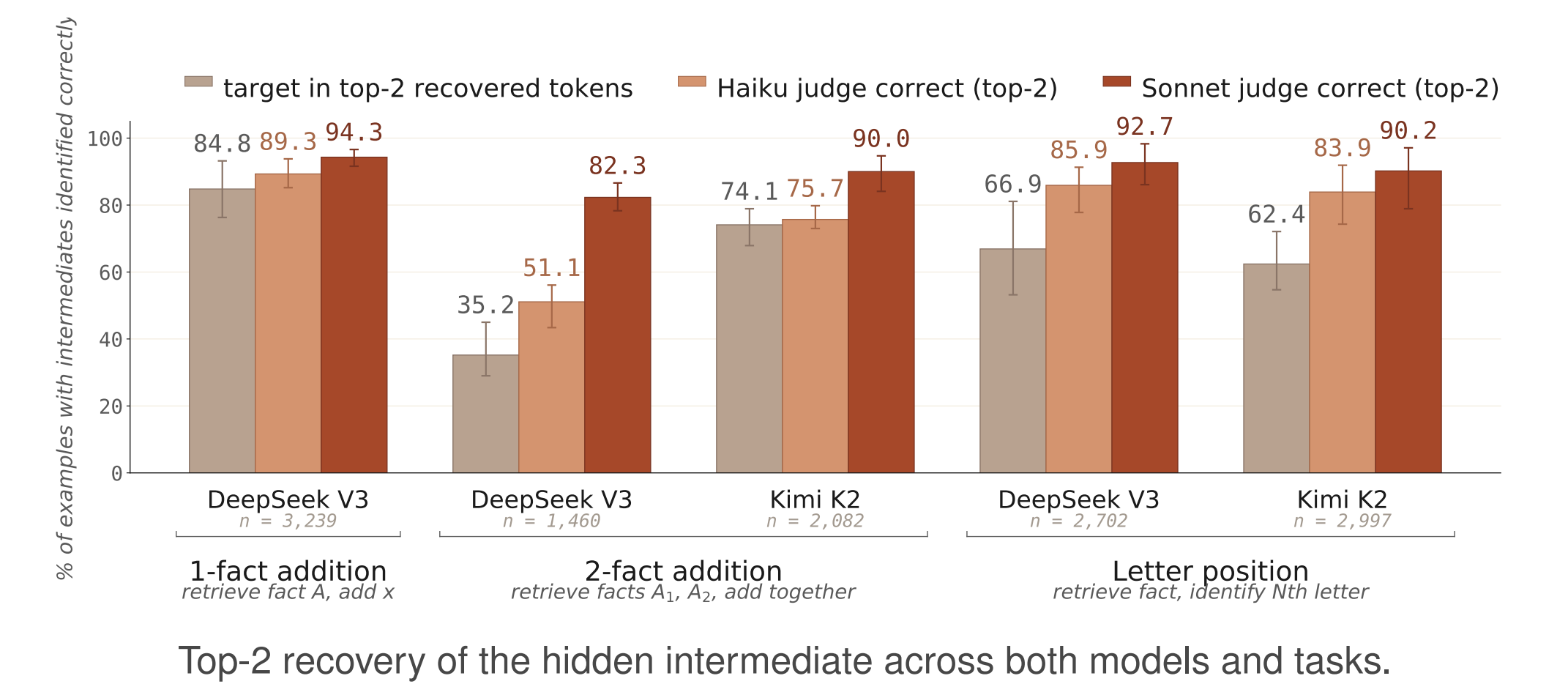
Filler raises accuracy significantly; scrambled filler does not.

4 The computation is real & decodable

- Interventions show the computation is **causal** and attention routes question → filler → answer.
- Logit lens shows the intermediate steps unfold across filler.



- A logit-lens + LLM-judge pipeline with only hidden states (no labels, no access to the question), recovers intermediate computation at **82–95%**; a shuffled control scores **~1%**.



5 Discussion & limitations

- **For AI safety:** computation hidden from the output is still readable from the residual stream. **Opaque** ≠ **unauditable**.
- **Limits:** logit lens misses off-vocabulary computation; untested against models optimized to evade decoding.