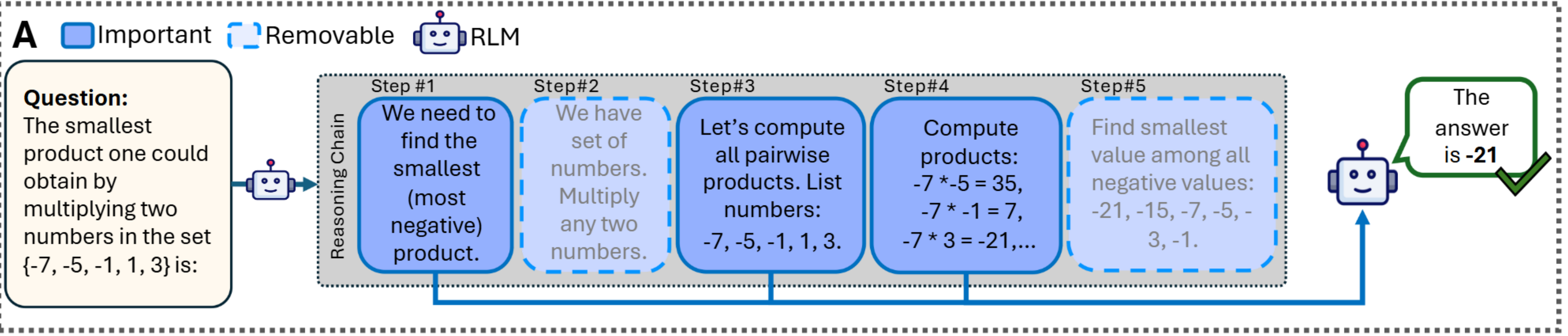
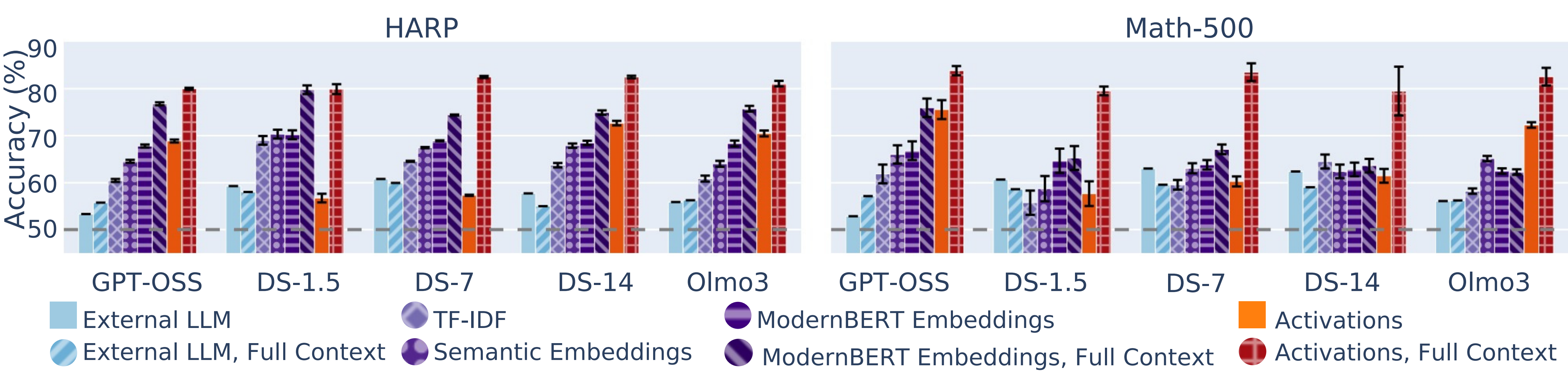


Reasoning models encode the importance of a reasoning step in its activations (before later steps are generated), while importance can not be faithfully recovered from tokens alone.

1. Identifying importance: Not all reasoning steps are **important**. We use causal analysis to extract a minimal subsequence of steps that is both sufficient (high P_{answer}) and attainable (low perplexity): the **core subsequence** that the model actually needs.



2. Importance lives in activations: Given importance labels, **probes trained on model activations** recover this concept, pre-hoc: without access to information from later steps. They beat **token-based** and **weaker-representation** baselines.



3. Importance is cross-model.

A probe trained on one model's activations predicts another model's importance with high accuracy and agreement.

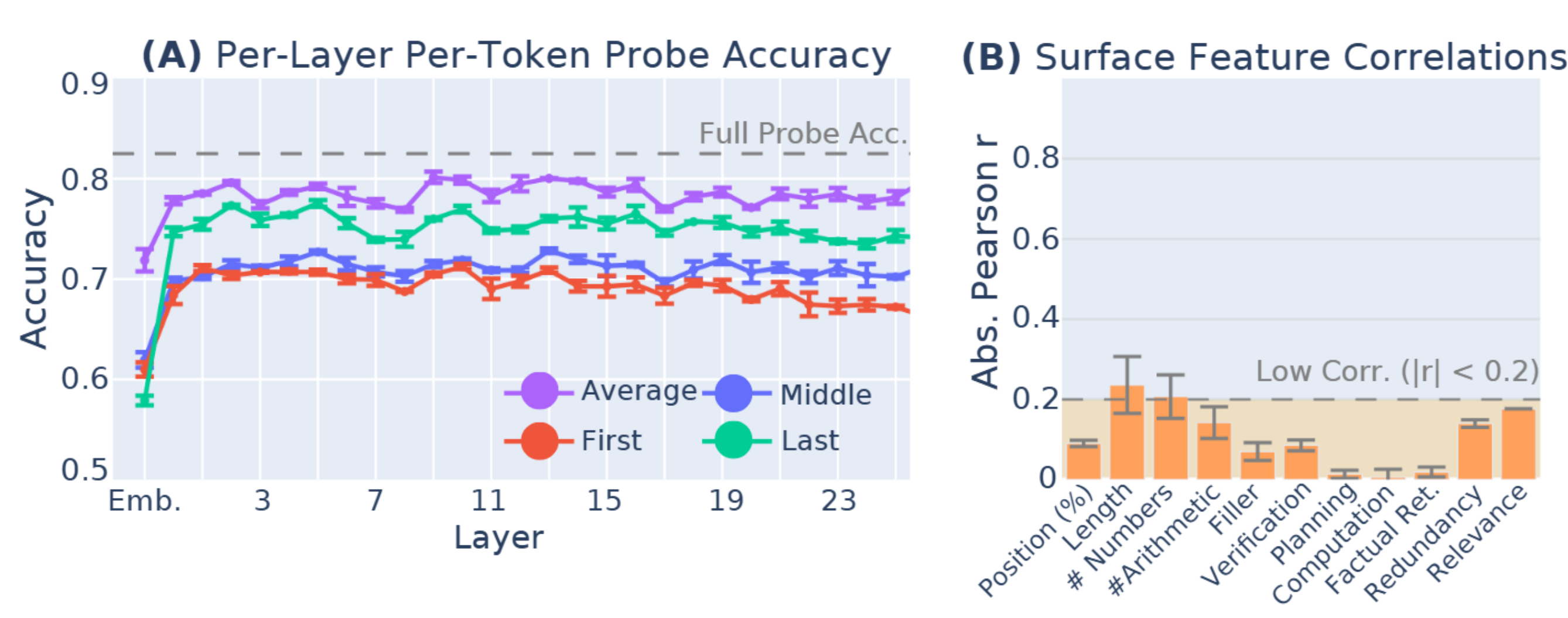
Cross-model agreement ratios

Model (Importance Labels, ϕ_{self})	GPT-OSS	DS-1.5	DS-7	DS-14	Olmo3	Random
GPT-OSS	84.8 ± 0.9%	87.4 ± 0.4%	88.0 ± 0.2%	87.8 ± 0.6%	100.0 ± 0.0%	59.3 ± 0.9%
DS-1.5	86.2 ± 0.2%	88.4 ± 0.6%	89.3 ± 0.6%	100.0 ± 0.0%	87.7 ± 0.5%	60.8 ± 0.6%
DS-7	87.3 ± 0.4%	90.1 ± 0.4%	100.0 ± 0.0%	90.8 ± 0.5%	88.2 ± 0.2%	58.2 ± 0.8%
DS-14	85.8 ± 0.7%	100.0 ± 0.0%	89.8 ± 0.4%	88.4 ± 0.5%	87.1 ± 0.6%	55.3 ± 1.8%
Olmo3	100.0 ± 0.0%	85.1 ± 0.4%	86.1 ± 0.1%	85.6 ± 0.9%	84.6 ± 0.3%	58.0 ± 0.2%

Model (Activations, ϕ_{cross})

4. Importance is not surface-level:

Importance isn't explained by token-based surface features, nor localized to any single layer or token.



Reasoning Models Know What's Important, and Encode It in Their Activations

Yaniv Nikankin¹, Martin Tutek², Tomer Ashuach¹, Jonathan Rosenfeld³, Yonatan Belinkov^{1,4}