

NeuroFaith: Evaluating Mechanistic Faithfulness of LLM Free Text Self-Explanation at the Concept Level

Milan Bhan Jean-Noël Vittaut Nicolas Chesneau Sarath Chandar Marie-Jeanne Lesot

{milan.bhan,nicolas.chesneau}@
ekimetrics.com

{jean-noel.vittaut,marie-jeanne.lesot}@
lip6.fr

sarath.chandar@mila.quebec

Context and objective

Context: are LLM free text self-explanations faithful?

- LLMs generate free-text **self-NLE** (*predict-then-explain*)
- Plausibility vs. faithfulness:** self-NLE look credible but may not reflect the true decision process (Jacovi & Goldberg, 2020)
- Existing methods (CI, AA) rely on **behavioral tests**, treat f as a black box $\Rightarrow$ measure *self-consistency*, not faithfulness

Aim: directly measure self-NLE faithfulness via **concept-level mechanistic analysis** of internal reasoning

Step 1: Concept Extraction

- Extract from self-NLE $e(x)$ a set of concepts $\{c_i\}_{i=1}^p$ quoted as important for $f(x)$
- Concepts: interpretable binary high-level features
- Extraction strategies
 - Human annotation: high-quality, domain-expert, but costly
 - Auxiliary LLM (LLM-as-a-Judge): scalable

Step 2: Concept-wise Interpretation

- Attribution score $I_\Gamma(c)$ per concept, over a circuit Γ
- Circuit $\Gamma = \{(k, \ell)\}$: sparse subgraph of coordinates (token k , layer ℓ) carrying information flow
- Interpreter I
 - Probing** c from f hidden states: Patchscopes, linear probes
 $I_\Gamma(c) = \max_{(k, \ell) \in \Gamma} I(h_k^\ell, c)$
 - Concept importance:** causal influence on $f(x)$ via TCAV, RepE

Step 3: Faithfulness Measurement

- $e(x)$ faithful when extracted concepts are mechanistically important
- Faithfulness: proportion of concepts with positive attribution
 - $F(x, e) = \frac{1}{p} \sum_{i=1}^p \mathbf{1}_{I_\Gamma(c_i) > 0}$
- Interpretation depends on interpreter I
 - Probing-based:** concept detection rate in hidden states
 - Importance-based:** causal relevance for $f(x)$

Instantiation 1: The Case of Classification

Task: predict label $\hat{y}$ from task-relevant concepts $\mathcal{C}$; interpreter $I =$ **concept importance**

Concept Extraction: auxiliary LLM extracts $\{c_i\}_{i=1}^p \subset \mathcal{C}$ from $e(x)$; each $c = \text{CAV}$ (diff-mean, last token)

Importance-based Interpretation:

- Erase c : $h_k^\ell \leftarrow h_k^\ell - \lambda \vec{c}_\ell$, over Γ (last token, F1 > 60% layers)
- Probability shift $p_{\hat{y}} = \beta_0 + \beta_1 \lambda \Rightarrow$ importance $I_\Gamma(c) = \beta_1$
- Faithfulness $F(x, e) \in [0, 1]$

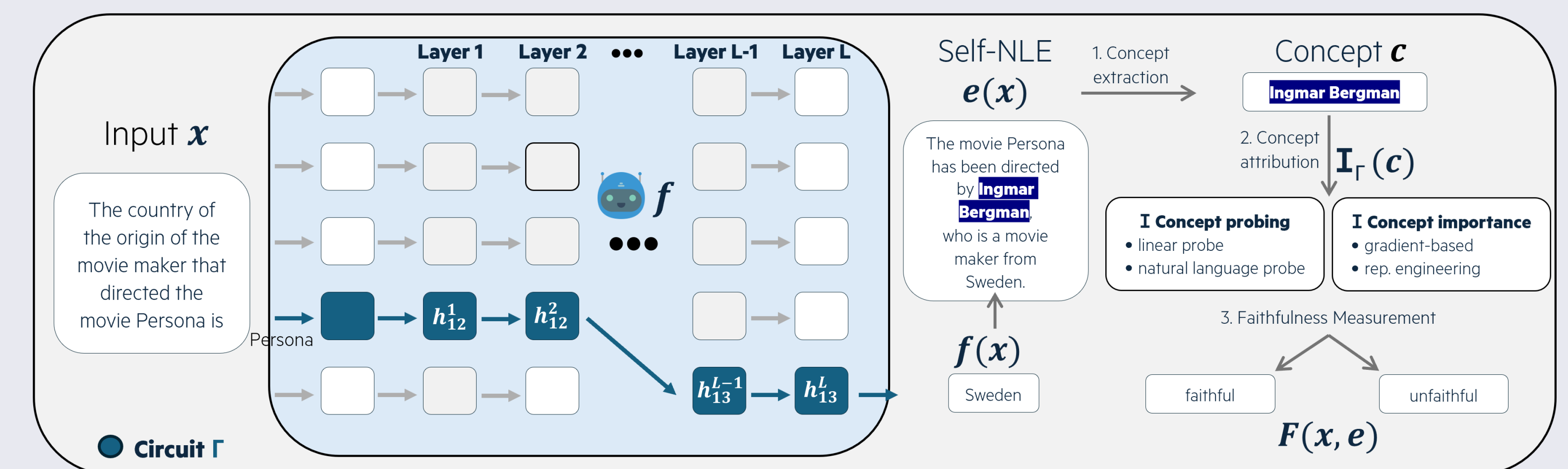
Model	Task Acc.		Self-NLE Faithfulness			
	AGNews	Ledgar	AGNews		Ledgar	
			Accurate	Inaccurate	Accurate	Inaccurate
gemma-2b	86.5	40.3	54.6	65.1	58.9	40.5
gemma-9b	88.7	59.7	96.0	92.3	56.0	58.3
gemma-27b	88.9	58.1	74.0	72.4	53.0	31.6
mistral-7b	80.6	82.9	86.6	84.1	94.8	84.4

Table 1: Self-NLE faithfulness from NeuroFaith (%), stratified by prediction correctness.

Our proposal: NeuroFaith

3-step framework

Concept extraction, Concept-wise Interpretation, Faithfulness Measurement



Instantiation 2: The Case of 2-Hop Reasoning

Task: $x = (o_1, r_1, \Delta, r_2, \bullet)$, bridge $\hat{o}_2$, answer $f(x) = \hat{o}_3$

- Ex.: *director of "Persona" → Bergman → Sweden*

Concept Extraction: auxiliary LLM gets bridge $\hat{o}_2 = c$ from $e(x)$

Probing-based Interpretation:

- Detect c in hidden states (Patchscopes, unsupervised)
- Γ : bridge resolved early, last token of o_1 (Biran et al., 2024)
- $I_\Gamma(c) = 1 \iff \exists (k, \ell) \in \Gamma, c \in \tilde{h}_k^\ell$

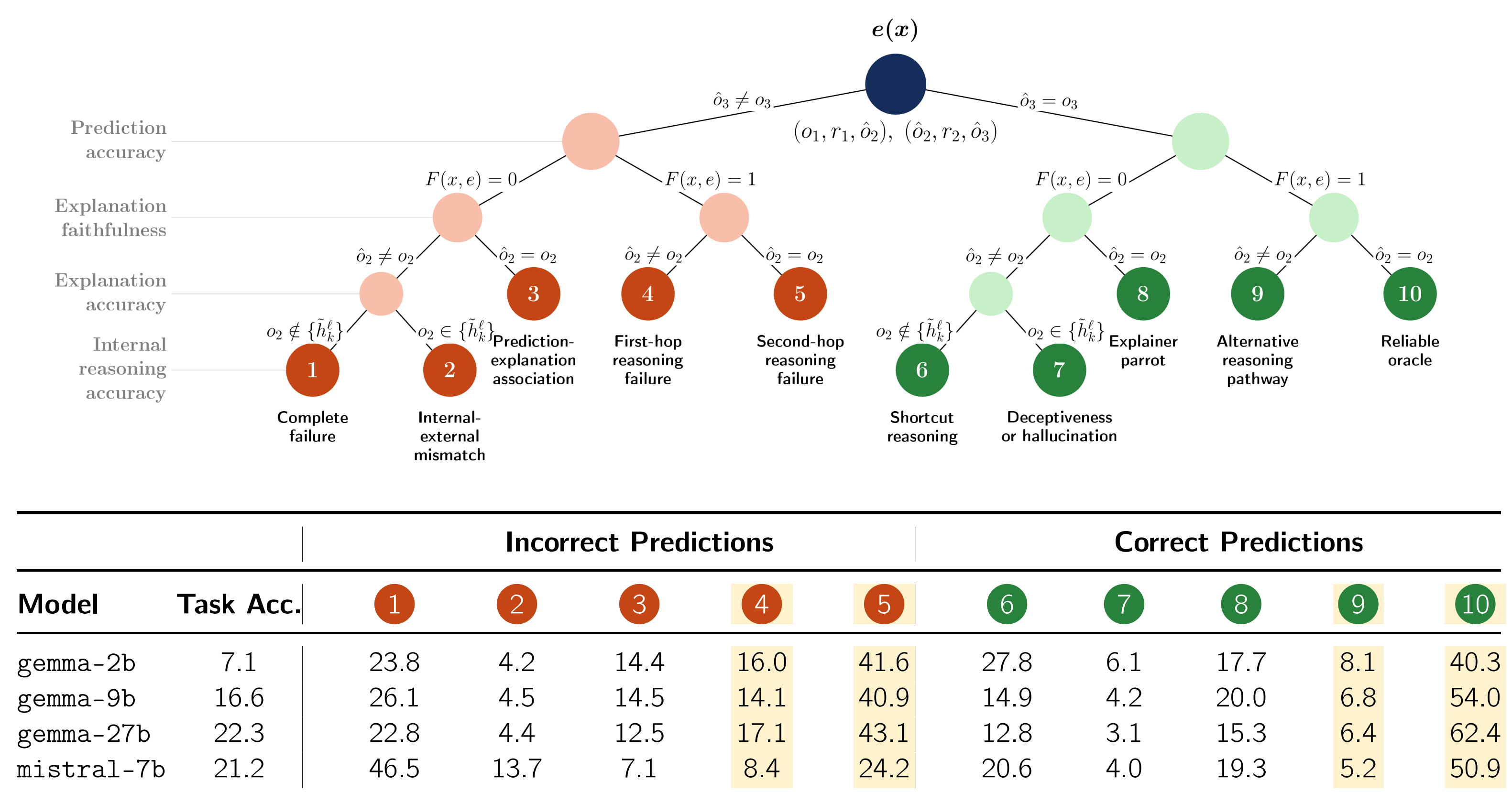
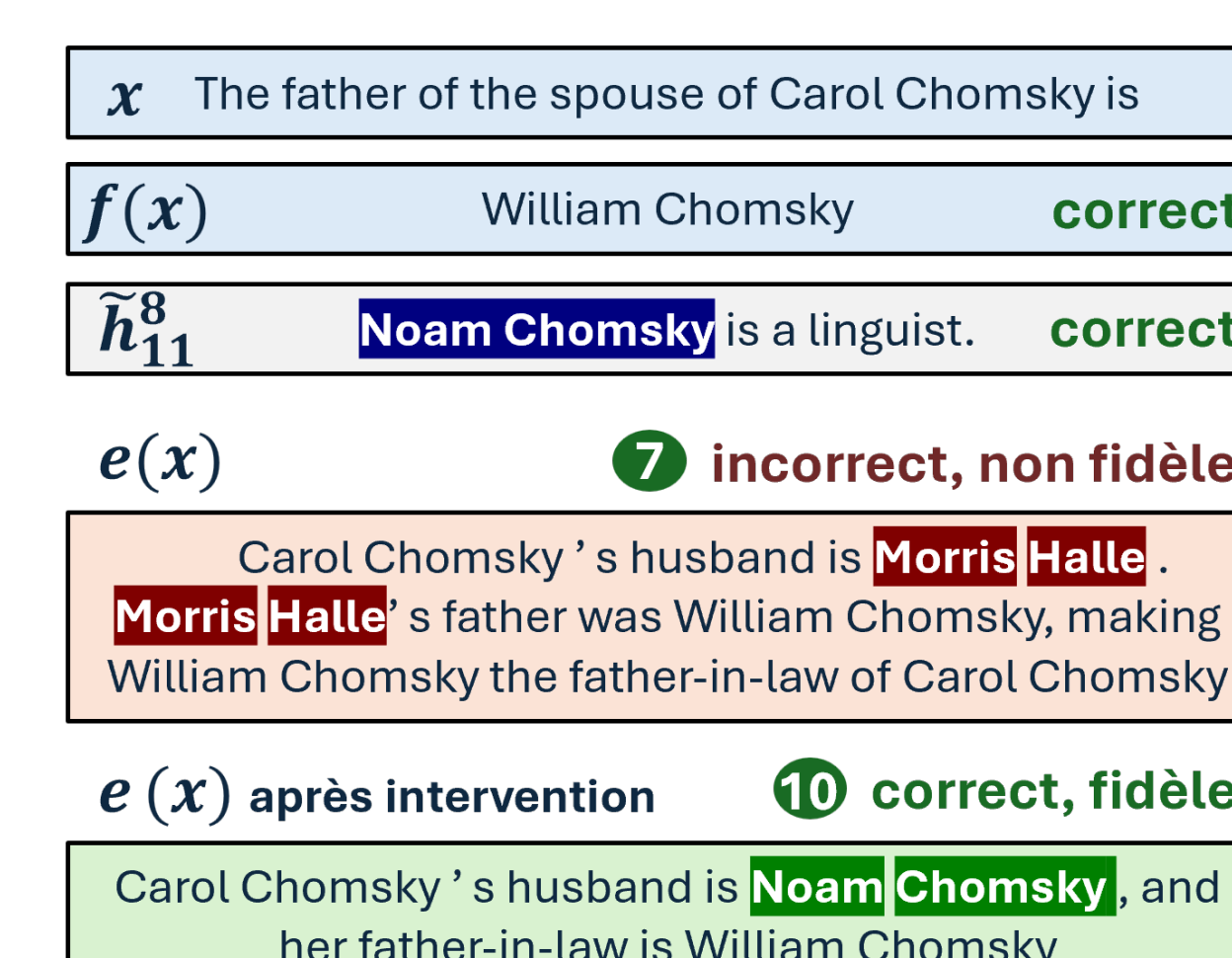


Table 2: Detailed taxonomy and distribution of categories across models (in %).

Linearly Detecting and Improving Faithfulness

- NeuroFaith faithfulness is a **linear direction** in representation space
- Detection:** diff-mean probe $\vec{F}_\ell$ detects (un)faithful self-NLE (F1 61–75%); strongest at the bridge object token
- Improvement:** activation steering $h_k^\ell \leftarrow h_k^\ell + \lambda \vec{F}_\ell$ converts 8–11% of unfaithful self-NLE into faithful ones
- Largest gains (19–58%) when the bridge is internally computed but absent from $e(x) \Rightarrow$ steering surfaces what f already knows
- Links faithfulness to **hallucination** and **deceptiveness** directions



Conclusion

NeuroFaith: self-NLE mechanistic faithfulness

- Faithfulness: linear structure
- Fast detection and **enhancement via activation steering**