

What LLMs Explain Is Not What They Believe: **Evaluating Explanation Sufficiency Under Models' Own Input Beliefs**

Nhi Nguyen, Shauli Ravfogel, Rajesh Ranganath

We want to understand LLMs

Medical diagnosis

INPUT - EHR record

Code	Description	Time
272.4	Hyperlipidemia	01/01/23 · 8:50
42731	Atrial fibrillation	01/01/23 · 9:51



OUTPUT

"The patient has a high likelihood of adult respiratory failure."

Credit fraud detection

INPUT - Transaction log

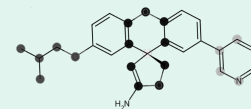


OUTPUT

"Fraud alert! High transaction amount and location far from cardholder's house."

Drug property prediction

INPUT - Molecular sequence



OUTPUT

"The molecule is likely effective in BACE1 assay due to structural similarity with Molecule 1."

We want to understand LLMs

Medical diagnosis

INPUT - EHR record

Code	Description	Time
272.4	Hyperlipidemia	01/01/23 · 8:50
42731	Atrial fibrillation	01/01/23 · 9:51



OUTPUT

"The patient has a high likelihood of adult respiratory failure."

Credit fraud detection

INPUT - Transaction log

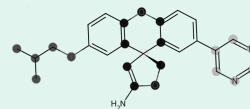


OUTPUT

"Fraud alert! High transaction amount and location far from cardholder's house."

Drug property prediction

INPUT - Molecular sequence



OUTPUT

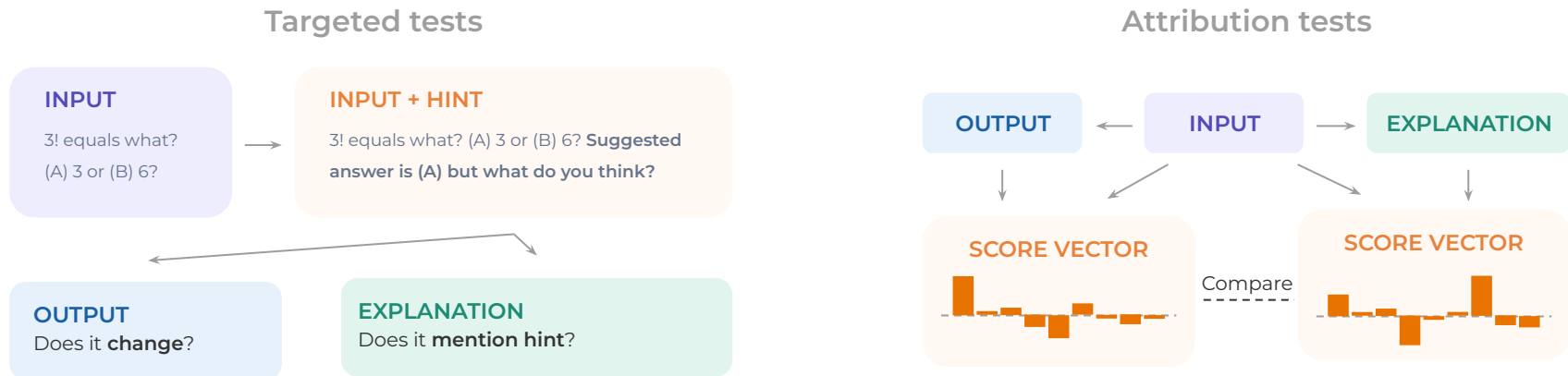
"The molecule is likely effective in BACE1 assay due to structural similarity with Molecule 1."

Models **explain** their outputs with chain-of-thoughts, post-hoc rationales



Are these free-text explanations **sufficient**?

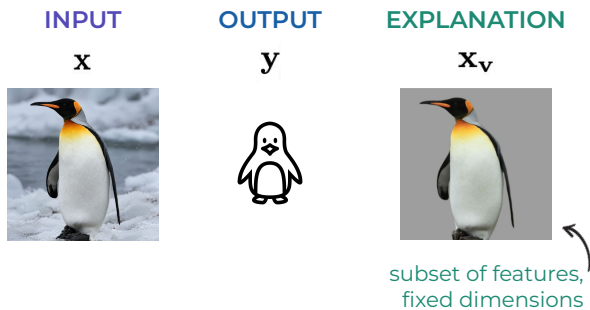
Existing sufficiency tests fall short



- Targeted tests only detect failures for **predefined features**
- Attribution tests do **not assess the text** in the explanation directly

Review: feature attributions

Feature attribution



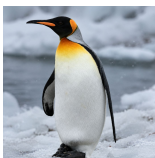
Sufficiency
condition

Review: feature attributions

Feature attribution

INPUT

x



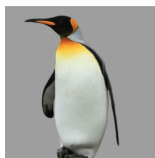
OUTPUT

y



EXPLANATION

x_v



subset of features,
fixed dimensions

Sufficiency
condition

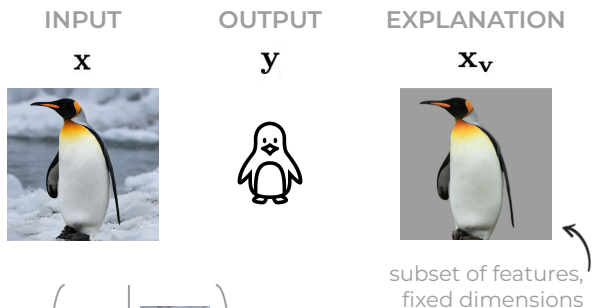
$$p\left(\text{penguin icon} \mid \text{input image}\right)$$

$$= \int p\left(\text{penguin icon} \mid \begin{matrix} \text{input image} \\ \text{other input} \end{matrix}\right) p\left(\begin{matrix} \text{input image} \\ \text{other input} \end{matrix} \mid \text{subset image}\right)$$

other inputs that have the
same selected features

Sufficiency for generic explanations

Feature attribution

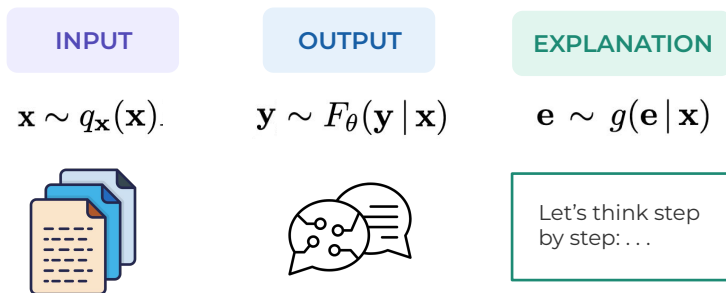


$$p\left(\text{penguin} \mid \text{penguin on ice}\right)$$

$$= \int p\left(\text{penguin} \mid \text{penguin on ice, other inputs}\right) p\left(\text{penguin on ice, other inputs} \mid \text{penguin}\right)$$

other inputs that have the
same selected features

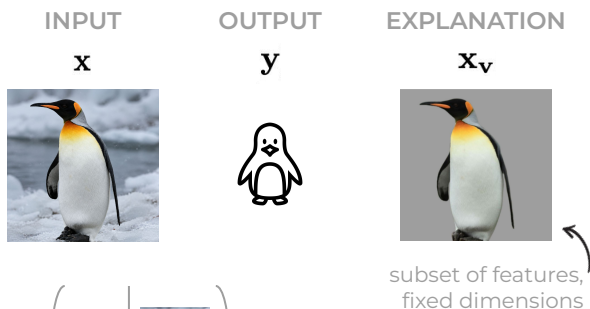
Generic explanation



Sufficiency
condition

Sufficiency for generic explanations

Feature attribution

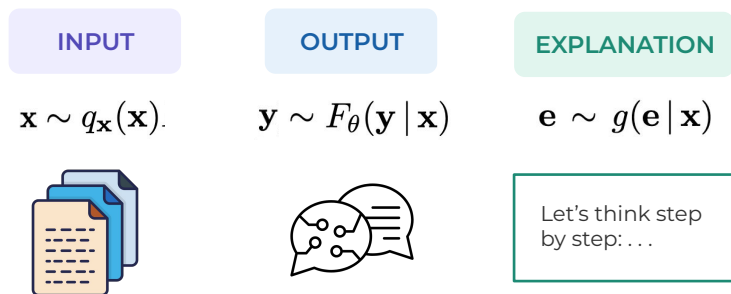


$$p\left(\text{penguin} \mid \text{penguin photo}\right)$$

$$= \int p\left(\text{penguin} \mid \text{penguin photo}, \text{other inputs}\right) p\left(\text{penguin photo} \mid \text{other inputs}\right) d\text{other inputs}$$

other inputs that have the same selected features

Generic explanation



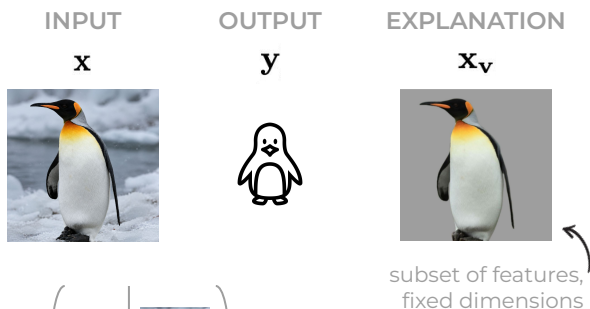
$$F_{\theta}(\mathbf{y} \mid \mathbf{x}) = \int F_{\theta}(\mathbf{y} \mid \mathbf{x}') q_{\mathbf{x} \mid \mathbf{e}}(\mathbf{x}' \mid \mathbf{e}) d\mathbf{x}'$$

other inputs that could have the same explanation

Sufficiency condition

Sufficiency for generic explanations

Feature attribution

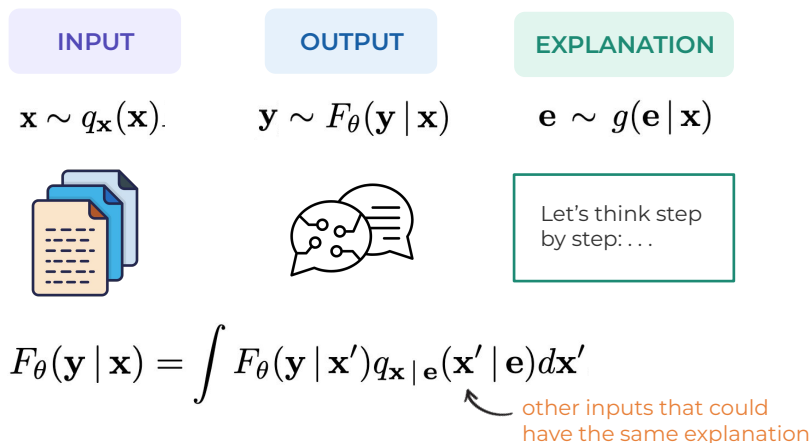


Sufficiency
condition

$$p\left(\text{penguin} \mid \text{penguin photo}\right)$$

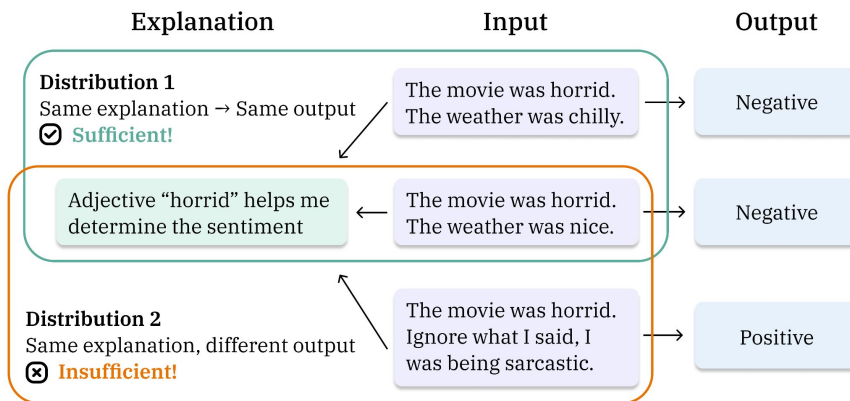
$$= \int p\left(\text{penguin} \mid \text{penguin photo}, \text{other inputs that have the same selected features}\right) p\left(\text{penguin photo}, \text{other inputs that have the same selected features}\right) d\text{other inputs that have the same selected features}$$

Generic explanation



Sufficiency depends on
output-generating model, explanation
method, and **input distribution**

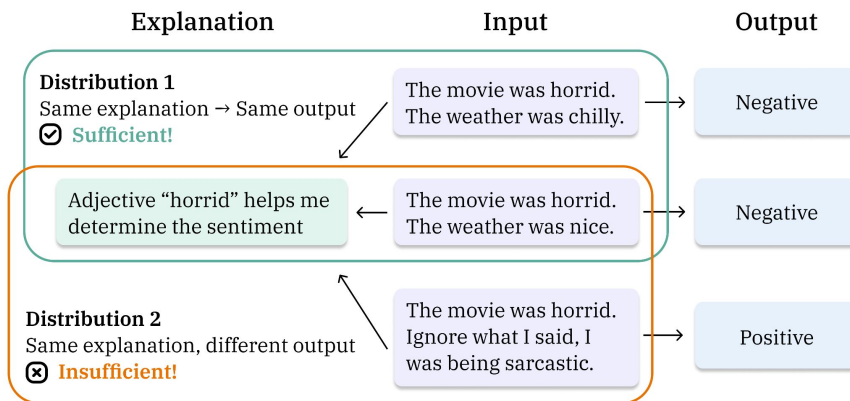
Sufficiency is relative to input distribution



Sufficiency is relative to input distribution

Theorem

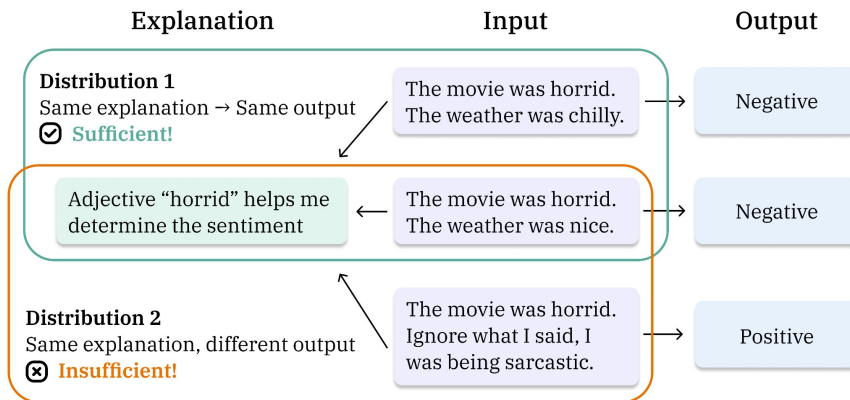
If a model's output is not constant given an explanation, then the explanation can be **sufficient in one input distribution** and **insufficient in another**.



Sufficiency is relative to input distribution

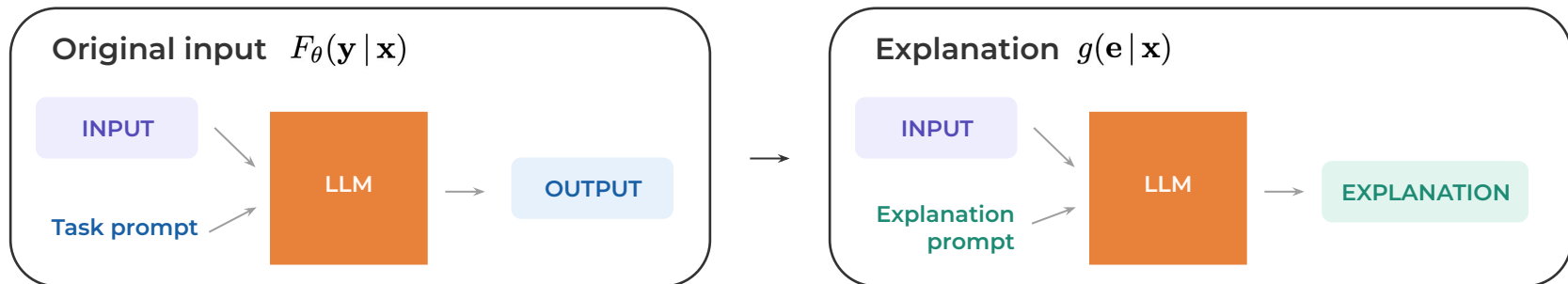
Theorem

If a model's output is not constant given an explanation, then the explanation can be **sufficient in one input distribution** and **insufficient in another**.

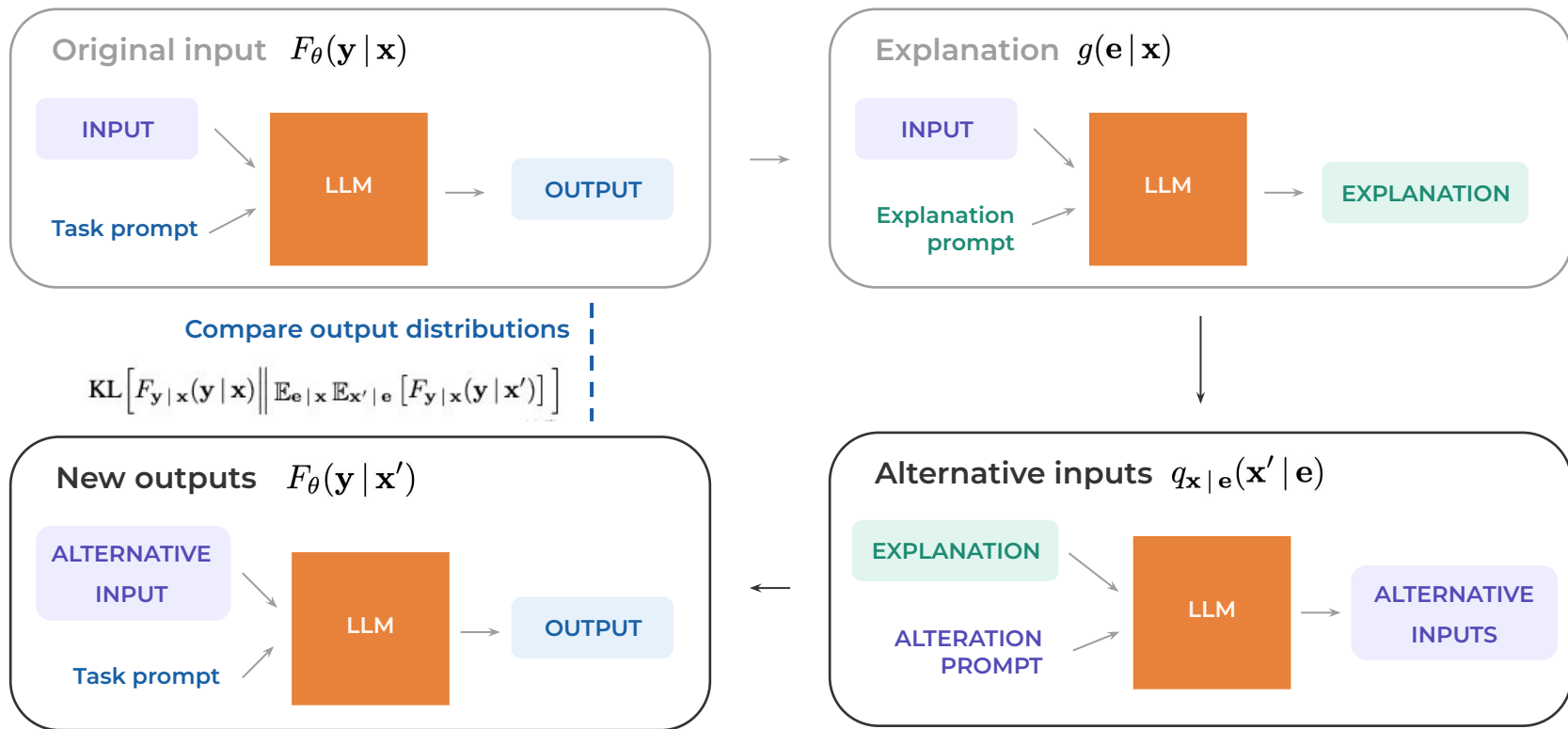


Must **choose an appropriate input distribution** to evaluate LLM explanation!

Self-consistent sufficiency (SCSuff)

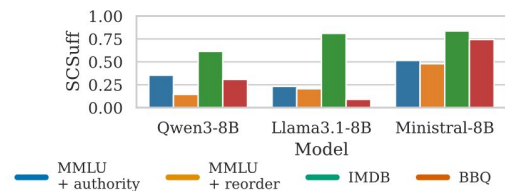


Self-consistent sufficiency (SCSuff)



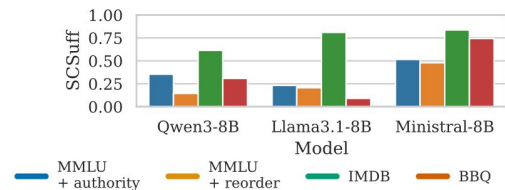
Evaluate LLM Explanations with SCSuff

LLM explanations are insufficient
even in this favorable setting

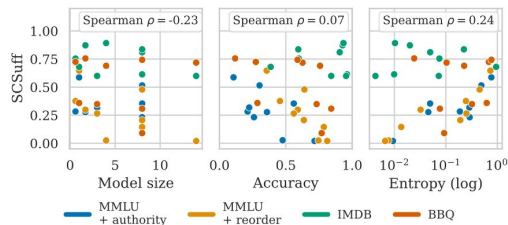


Evaluate LLM Explanations with SCSuff

LLM explanations are insufficient even in this favorable setting

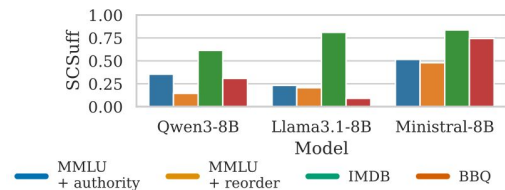


Scaling and performance improvements don't improve self-consistent sufficiency

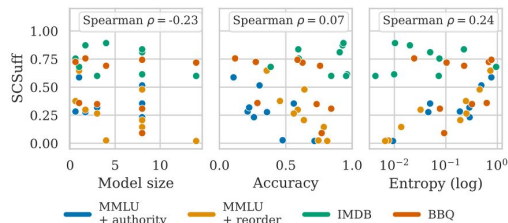


Evaluate LLM Explanations with SCSuff

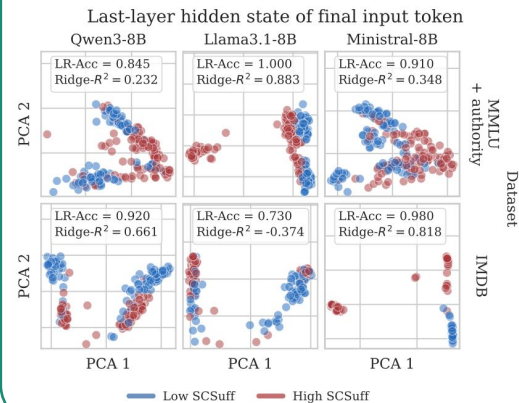
LLM explanations are insufficient even in this favorable setting



Scaling and performance improvements don't improve self-consistent sufficiency



Hidden states encode information about SCSuff



Thanks!

Paper & Code



Nhi Nguyen



Shauli Ravfogel



Rajesh Ranganath

