



Adaptation-to-Context in Humans and Large Language Models

Inferring the underlying mechanisms of In-Context Learning with Structural Priming

Zhengkao Herbert Zhou, Yale University
Computational Cognitive Science Group, University of Maryland
May 5, 2026



Overview

Humans and LLMs are bidirectionally informative

Overview

Humans and LLMs are bidirectionally informative

Psycholinguistics

**Structural
Priming**

Overview

Humans and LLMs are bidirectionally informative

Psycholinguistics

**Structural
Priming**

Computational Ling / NLP

**In-Context
Learning**

Overview

Humans and LLMs are bidirectionally informative

Psycholinguistics

**Structural
Priming**

Computational Ling / NLP

**In-Context
Learning**

**Shared behaviors &
mechanisms?**



```
graph LR; A[Psycholinguistics] --- B[Structural Priming]; C[Computational Ling / NLP] --- D[In-Context Learning]; B <--> D; E[Shared behaviors & mechanisms?];
```

Overview

Humans and LLMs are bidirectionally informative

Psycholinguistics

Computational Ling / NLP

Structural
Priming

In-Context
Learning

Shared behaviors &
mechanisms?

Implicit Learning /
Predictive Coding

CogSci & AI

?

Structural Priming in Humans

Structural Priming

[e.g., Bock 1986; Chang 2012; etc.]

Structural Priming

- **Structural Priming:** speakers tend to reuse the syntactic structures they have recently encountered during production or comprehension.

Structural Priming

- **Structural Priming:** speakers tend to reuse the syntactic structures they have recently encountered during production or comprehension.
- A classical case: **dative alternation**
 - Double Object (DO): *Alice sent Bob a letter.*
 - Prepositional Dative (PD): *Alice sent a letter to Bob.*

Structural Priming

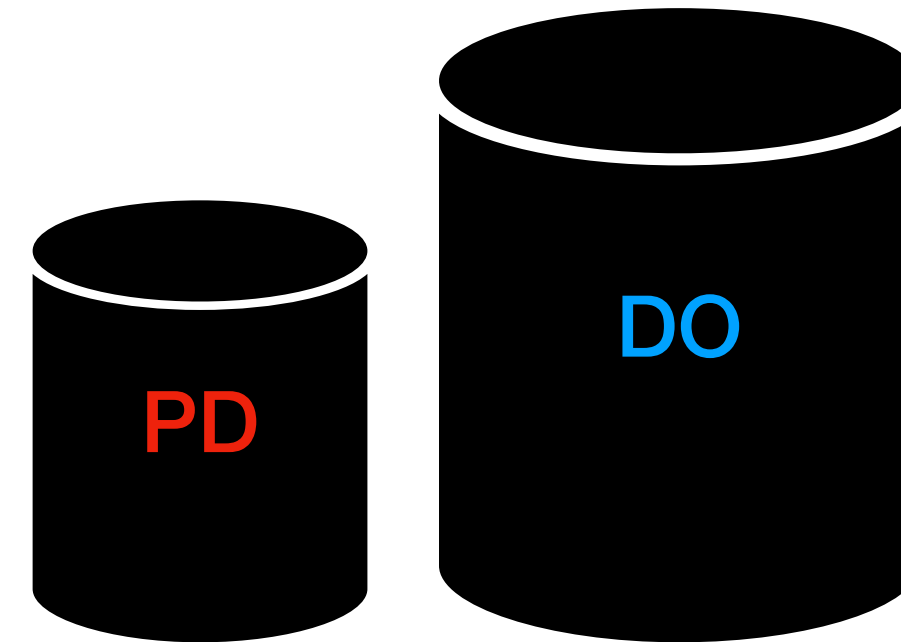
- **Structural Priming:** speakers tend to reuse the syntactic structures they have recently encountered during production or comprehension.
- A classical case: **dative alternation**
 - Double Object (DO): *Alice sent Bob a letter.*
 - Prepositional Dative (PD): *Alice sent a letter to Bob.*
- Other attested structures subject to priming:
 - active-passive;
 - *that*-omission in relative clauses and complement, etc.

The Inverse Frequency Effect (IFE)

- **Inverse Frequency Effect:** the *less preferred* (lower frequency) syntactic structure causes a *stronger* priming effect than the more preferred (higher frequency) structural alternative.

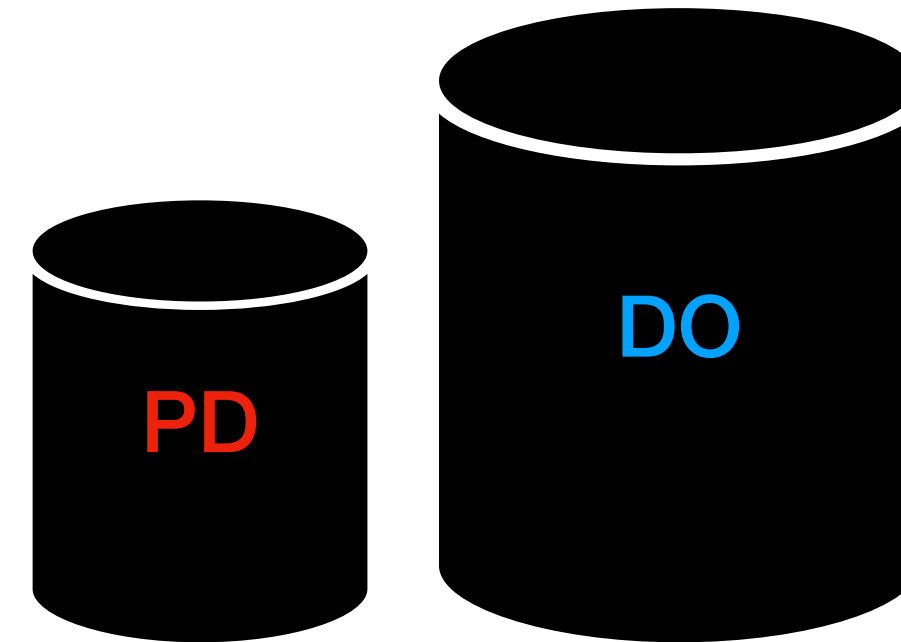
The Inverse Frequency Effect (IFE)

- **Inverse Frequency Effect:** the *less preferred* (lower frequency) syntactic structure causes a *stronger* priming effect than the more preferred (higher frequency) structural alternative.



The Inverse Frequency Effect (IFE)

- **Inverse Frequency Effect:** the *less preferred* (lower frequency) syntactic structure causes a *stronger* priming effect than the more preferred (higher frequency) structural alternative.



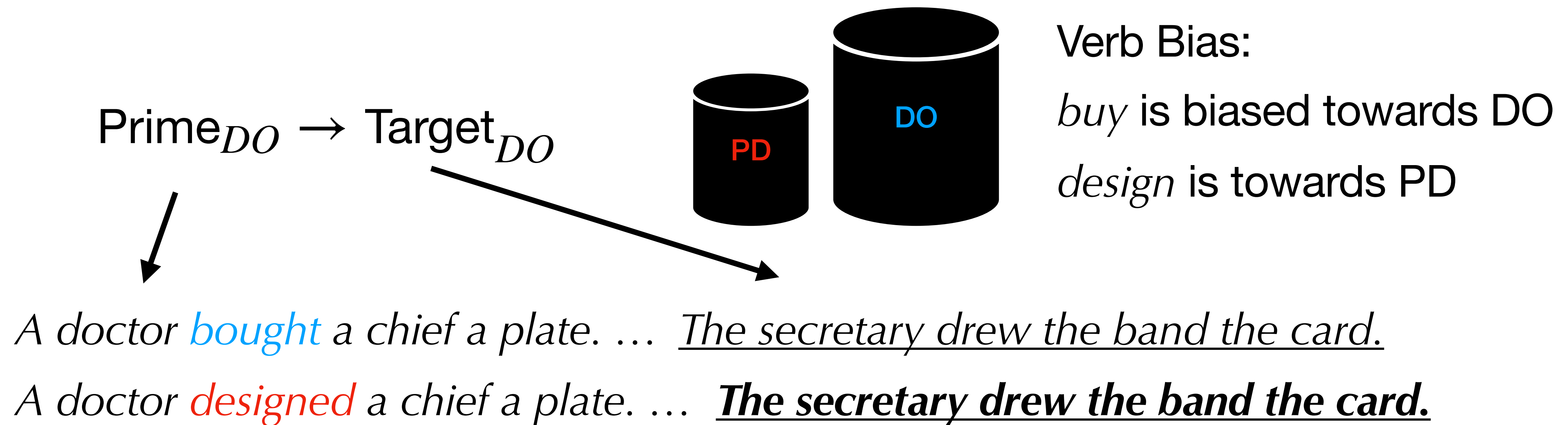
Verb Bias:

buy is biased towards DO

design is towards PD

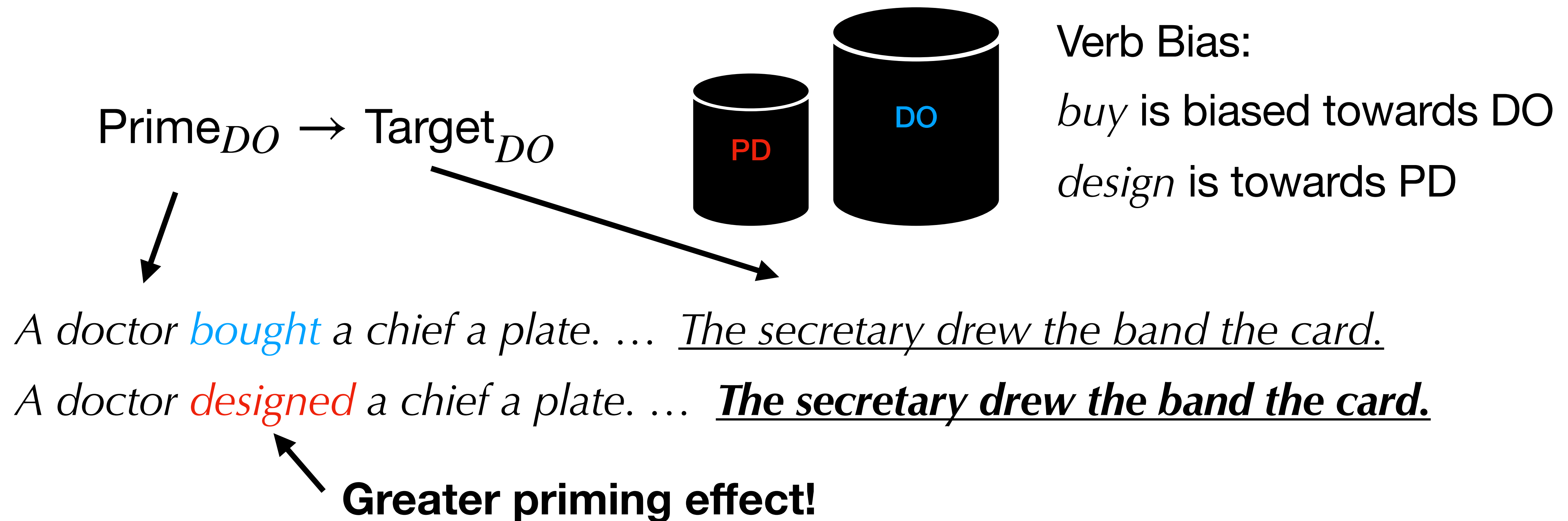
The Inverse Frequency Effect (IFE)

- **Inverse Frequency Effect:** the *less preferred* (lower frequency) syntactic structure causes a *stronger* priming effect than the more preferred (higher frequency) structural alternative.



The Inverse Frequency Effect (IFE)

- **Inverse Frequency Effect:** the *less preferred* (lower frequency) syntactic structure causes a *stronger* priming effect than the more preferred (higher frequency) structural alternative.



The Inverse Frequency Effect (IFE)

Prime_{DO} → Target_{DO}

Unprimed probability = 0.2

A doctor bought a chief a plate. ... The secretary drew the band the card.

A doctor designed a chief a plate. ... The secretary drew the band the card.

The Inverse Frequency Effect (IFE)

Prime_{DO} → Target_{DO}

Unprimed probability = 0.2

A doctor bought a chief a plate. ... The secretary drew the band the card.

A doctor designed a chief a plate. ... The secretary drew the band the card.

Prime Verb

Bring

Buy

Find

Draw

Design

The Inverse Frequency Effect (IFE)

Prime_{DO} → Target_{DO}

Unprimed probability = 0.2

A doctor bought a chief a plate. ... The secretary drew the band the card.

A doctor designed a chief a plate. ... The secretary drew the band the card.

Prime Verb

Verb PD Bias

Bring

0.23

Buy

0.27

Find

0.41

Draw

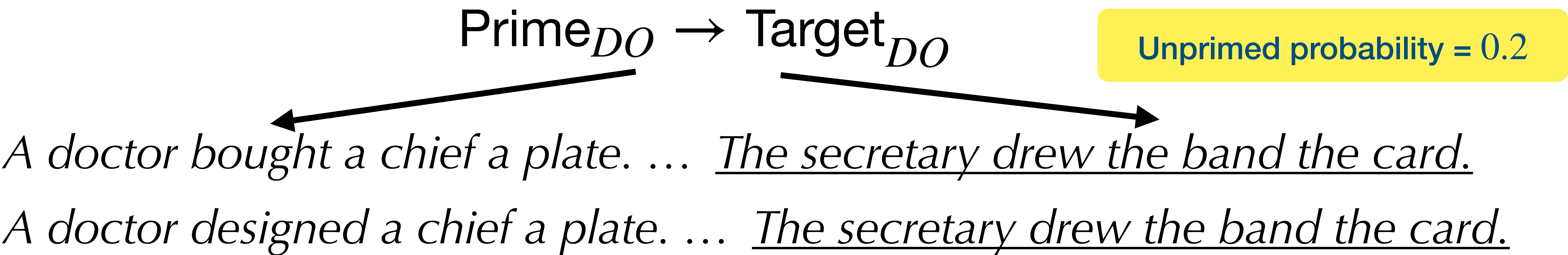
0.52

Design

0.77



The Inverse Frequency Effect (IFE)



<u>Prime Verb</u>	<u>Verb PD Bias</u>	<u>Primed probability</u> (priming magnitude)
<i>Bring</i>	0.23	0.23 (0.03)
<i>Buy</i>	0.27	0.25 (0.05)
<i>Find</i>	0.41	0.28 (0.08)
<i>Draw</i>	0.52	0.30 (0.10)
<i>Design</i>	0.77	0.33 (0.13)

The Inverse Frequency Effect (IFE)

Prime_{DO} → Target_{DO}

Unprimed probability = 0.2

A doctor bought a chief a plate. ... The secretary drew the band the card.

A doctor designed a chief a plate. ... The secretary drew the band the card.

Prime Verb

Verb PD Bias

Primed probability (priming magnitude)

Bring

0.23

0.23 (0.03)

Buy

0.27

0.25 (0.05)

Find

0.41

0.28 (0.08)

Draw

0.52

0.30 (0.10)

Design

0.77

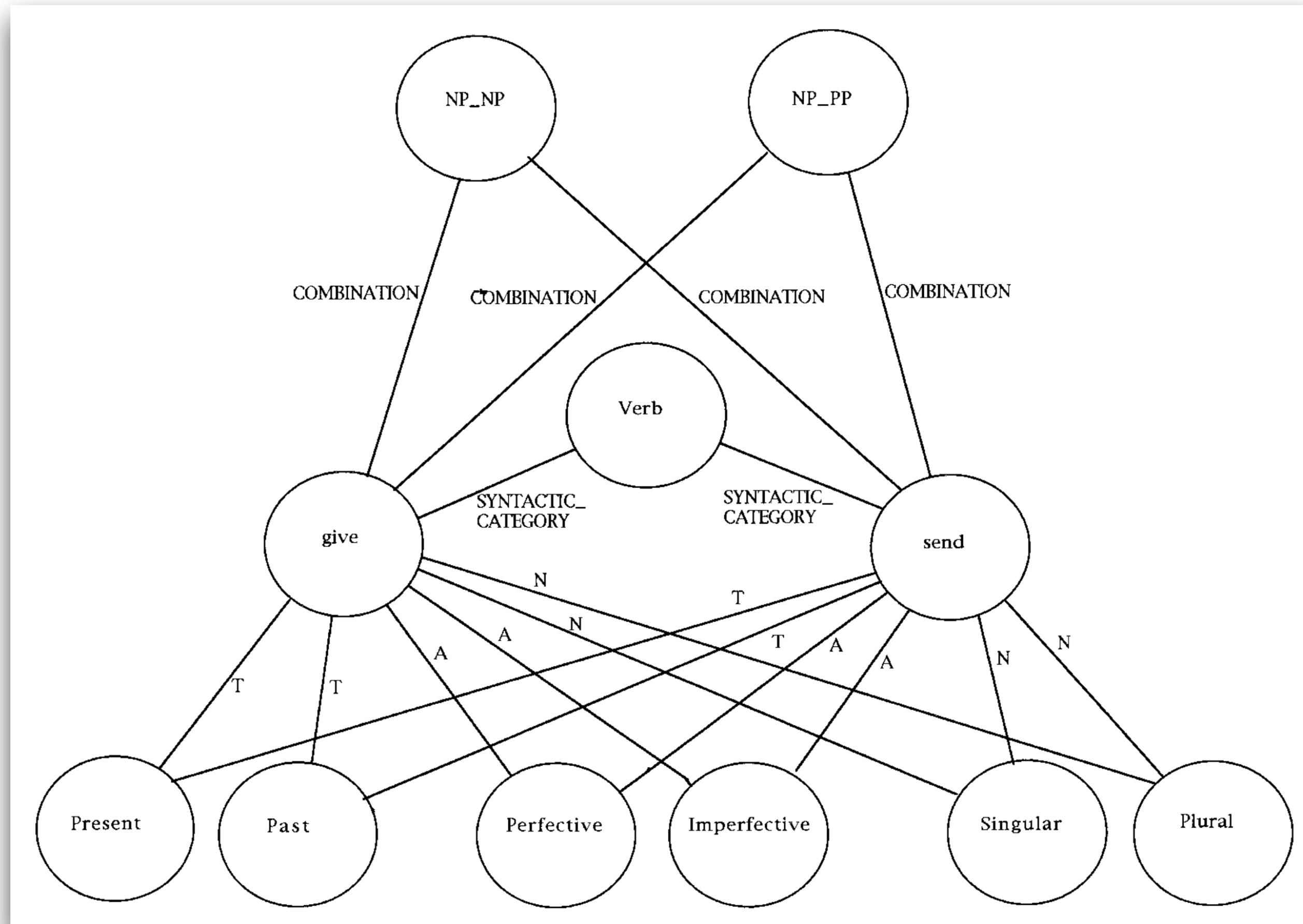
0.33 (0.13)

priming in DO

larger PD biases

greater priming effects

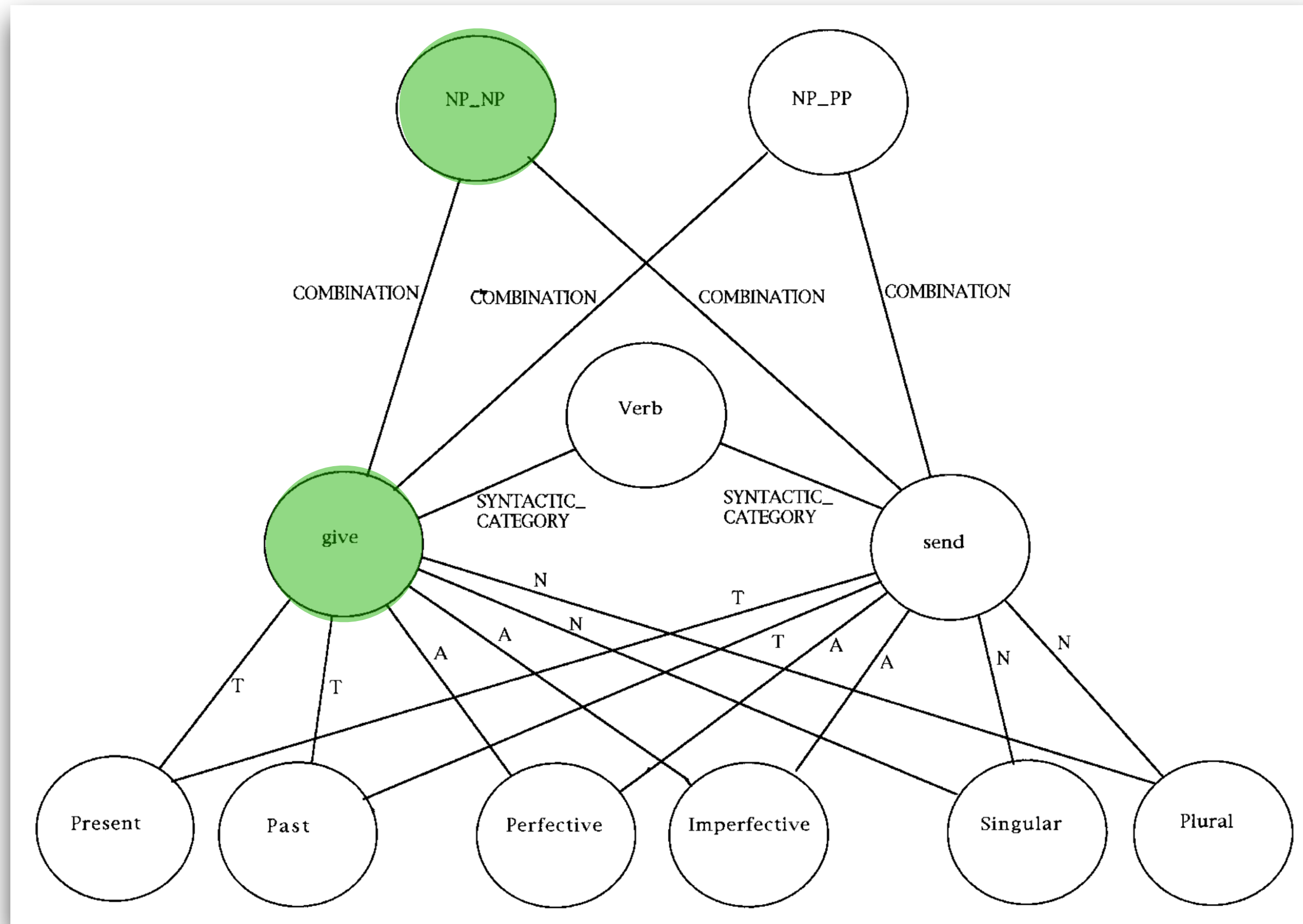
Theory 1: Transient Activation



← Structural Nodes

← Lexical Nodes

Theory 1: Transient Activation

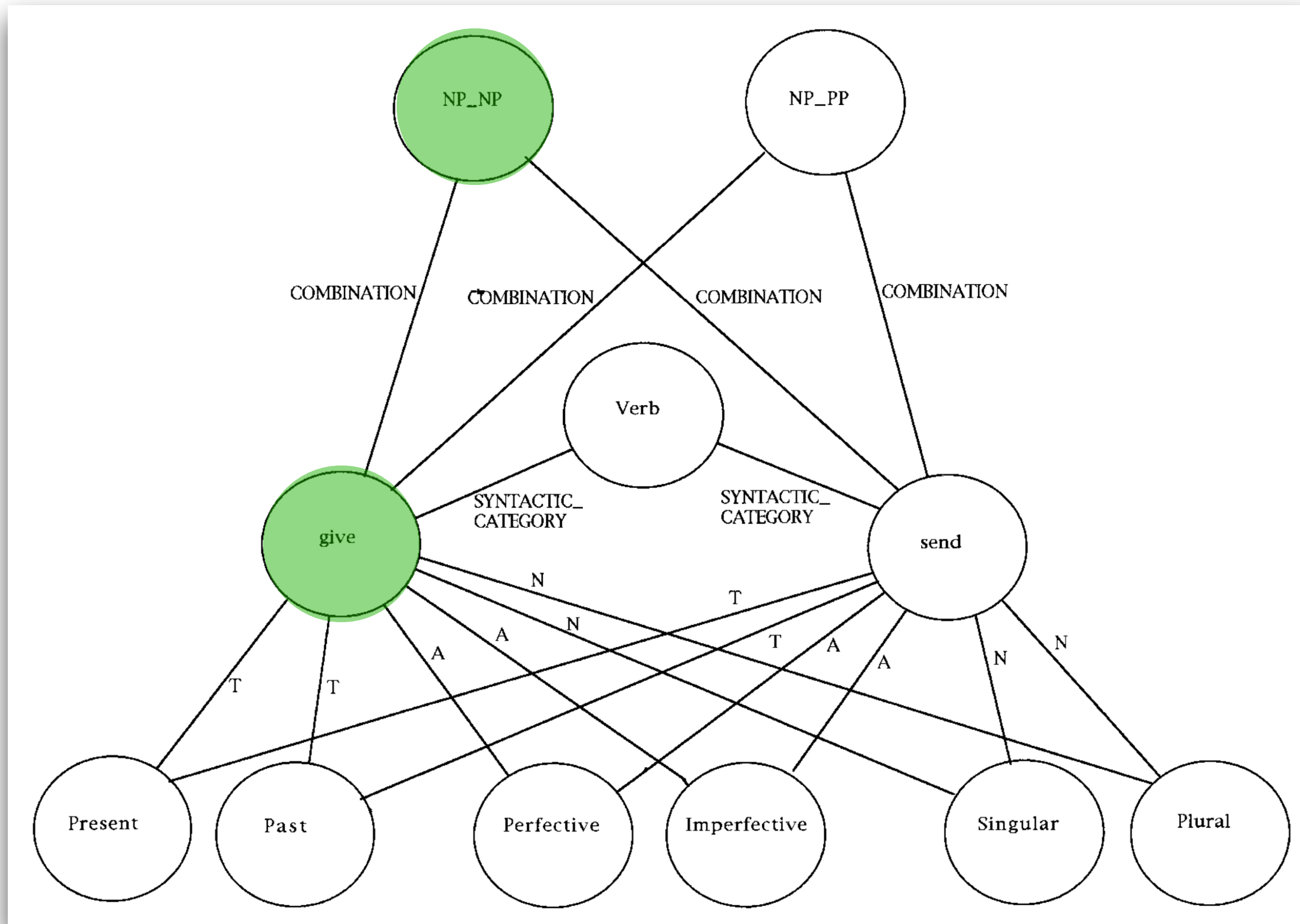


← Structural Nodes

← Lexical Nodes

Theory 1: Transient Activation

Quickly decaying activations



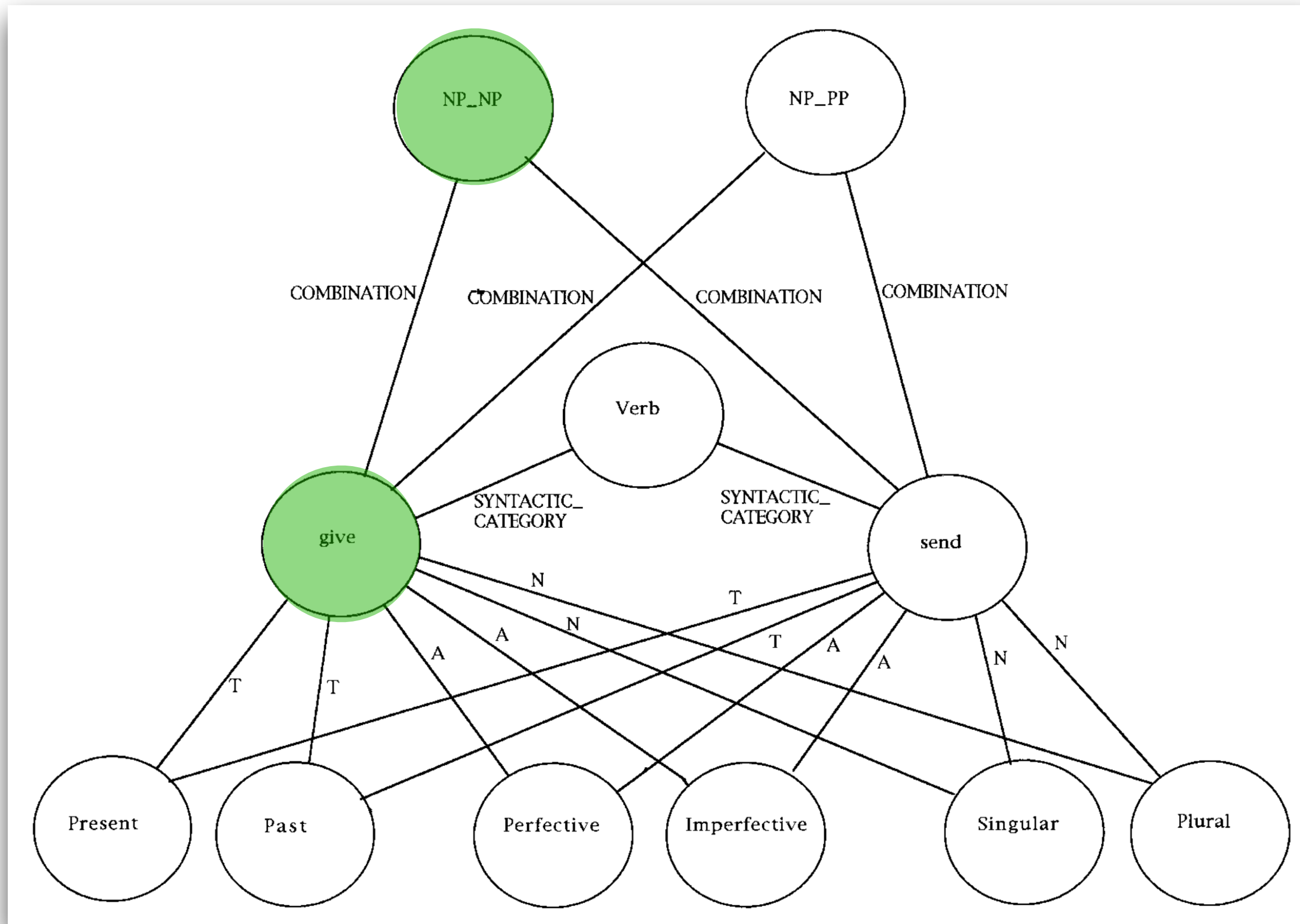
← Structural Nodes

Increasing the activation values of the nodes, which persists through next productions.

← Lexical Nodes

Theory 1: Transient Activation

Quickly decaying activations



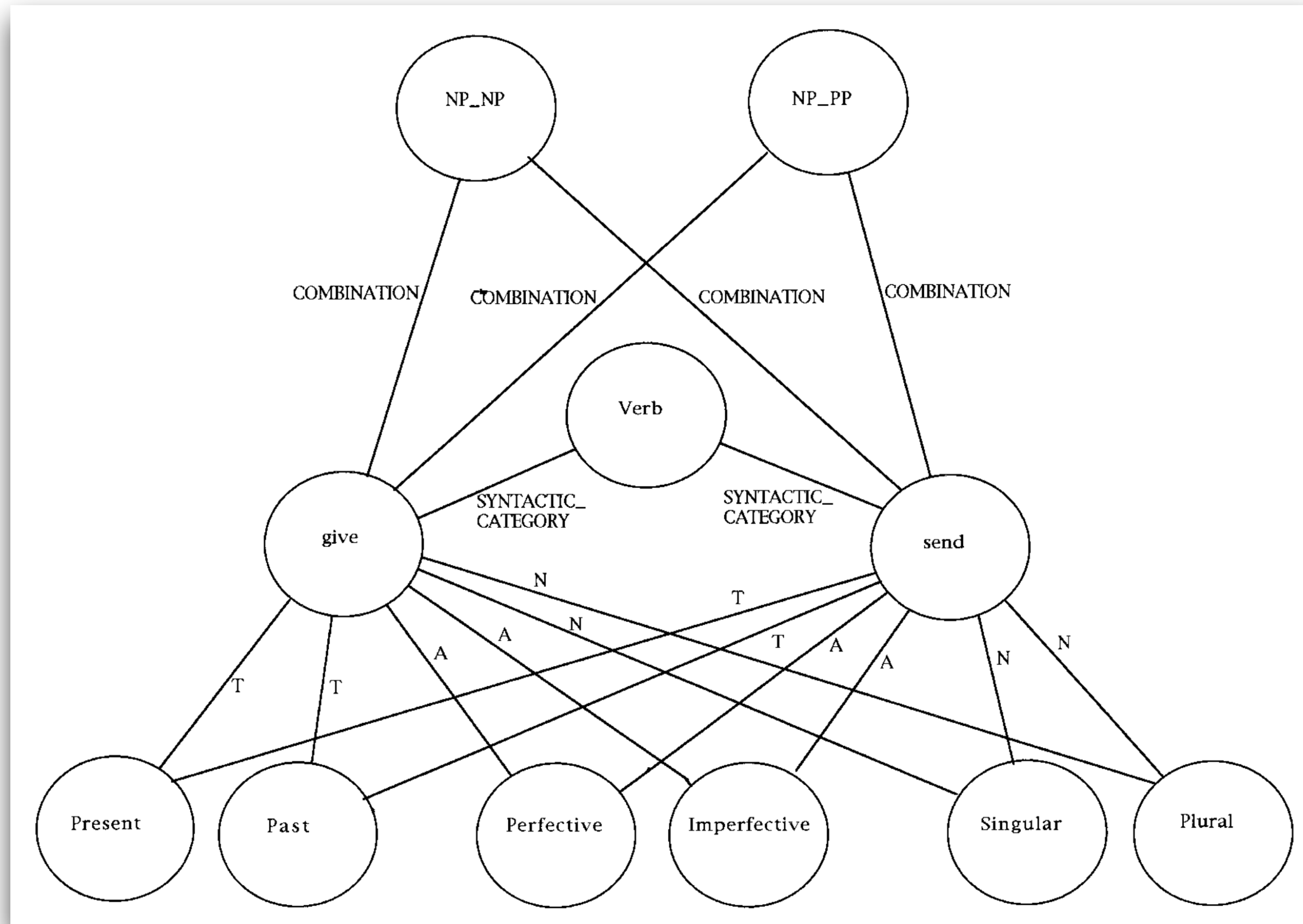
← Structural Nodes

Increasing the activation values of the nodes, which persists through next productions.

← Lexical Nodes

- An activation-based mechanism
- Explaining short-term effect
- Inverse frequency effect? ❌

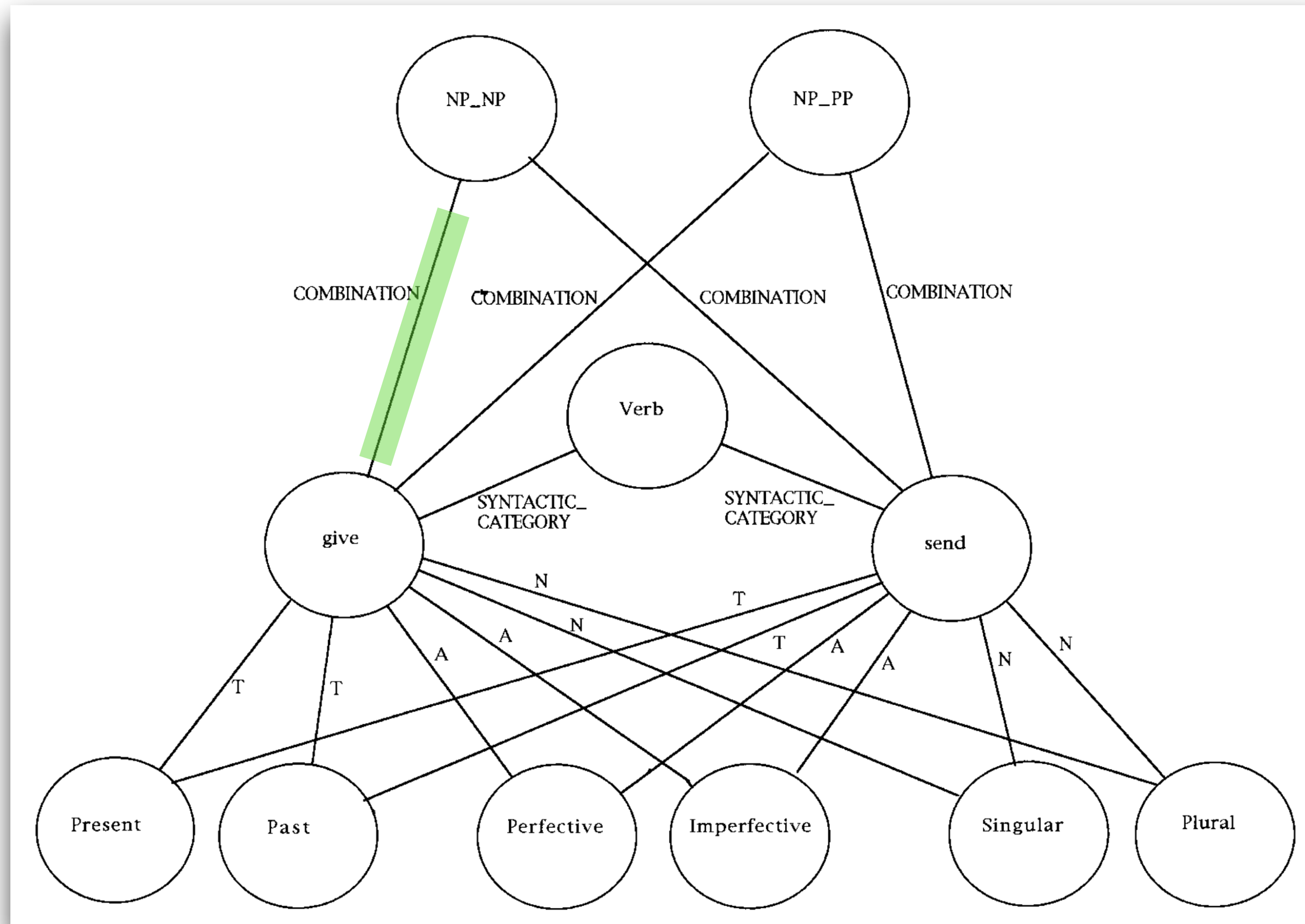
Theory 2: Implicit Learning



← Structural Nodes

← Lexical Nodes

Theory 2: Implicit Learning

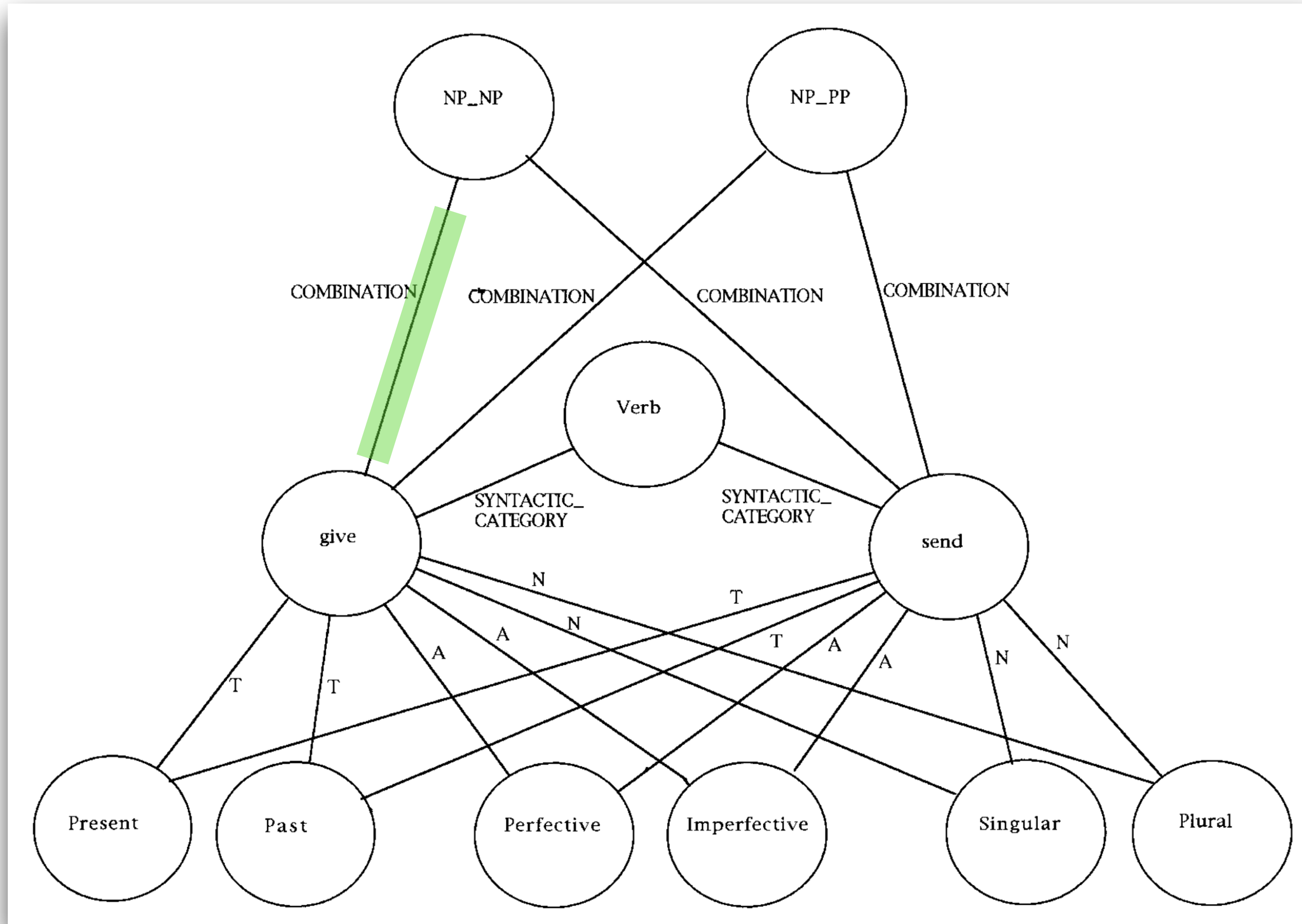


← Structural Nodes

← Lexical Nodes

Theory 2: Implicit Learning

Slowly decaying weight changes



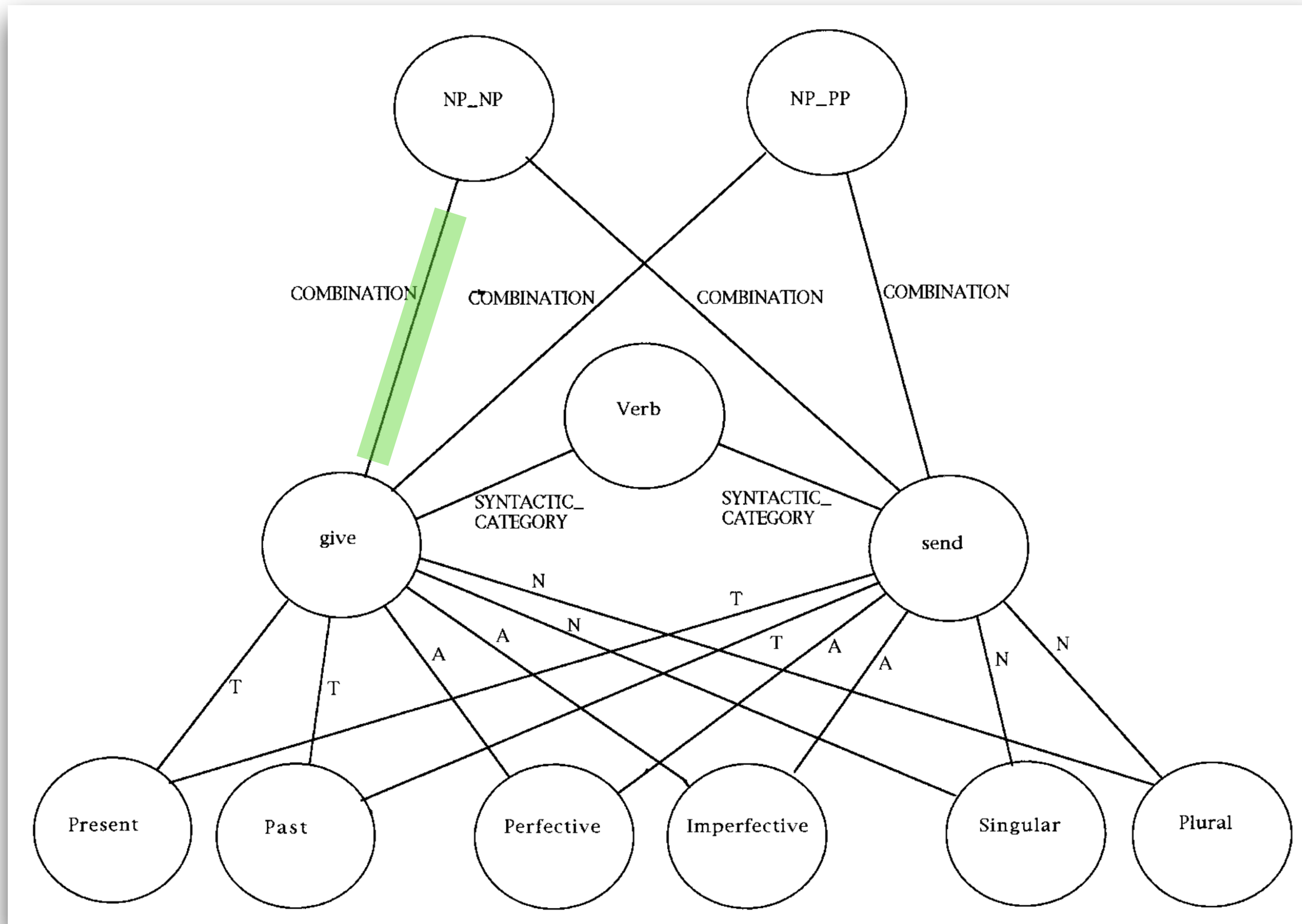
← Structural Nodes

Increasing the link weights proportional to the error signal between verb-bias prediction and the actual input prime.

← Lexical Nodes

Theory 2: Implicit Learning

Slowly decaying weight changes



← Structural Nodes

Increasing the link weights proportional to the error signal between verb-bias prediction and the actual input prime.

← Lexical Nodes

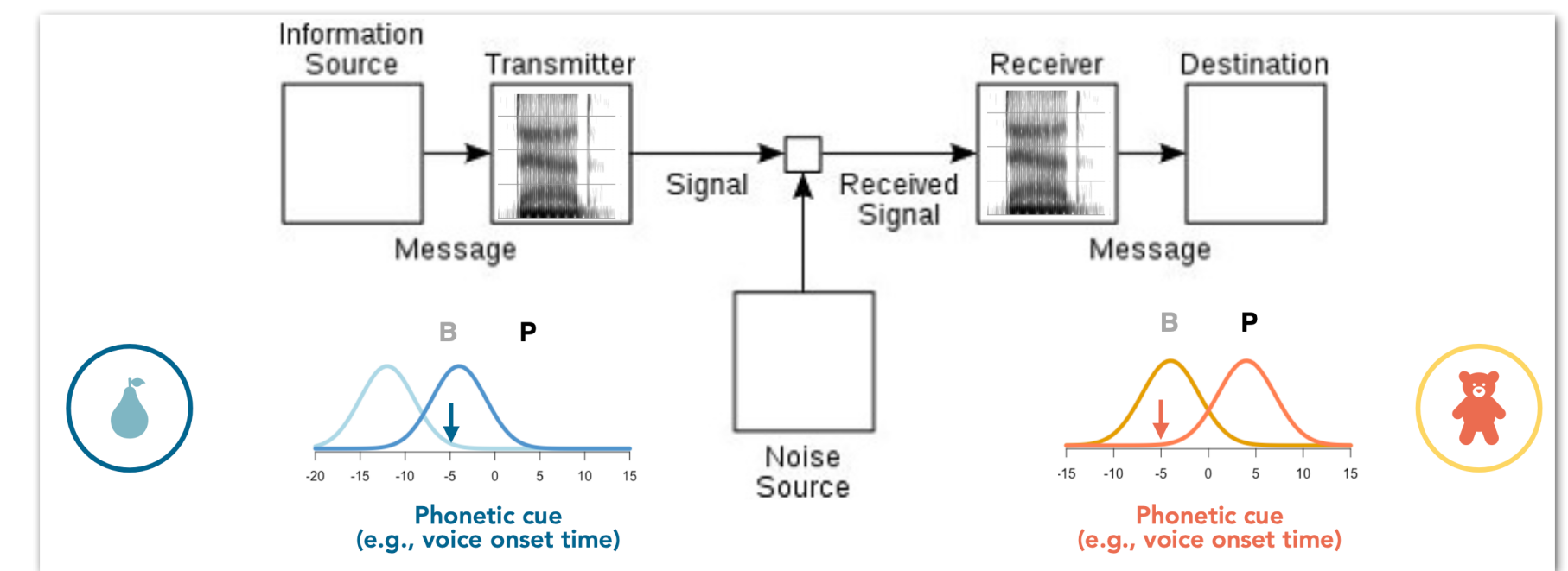
- An error-driven learning / adaptation mechanism
- Explaining long-term effect
- Inverse Frequency effect? ✓

Side Notes: Priming as Linguistic Adaptation

[e.g., Xie & Kurumada 2023; Lu et al. 2025; Feng 2022; etc.]

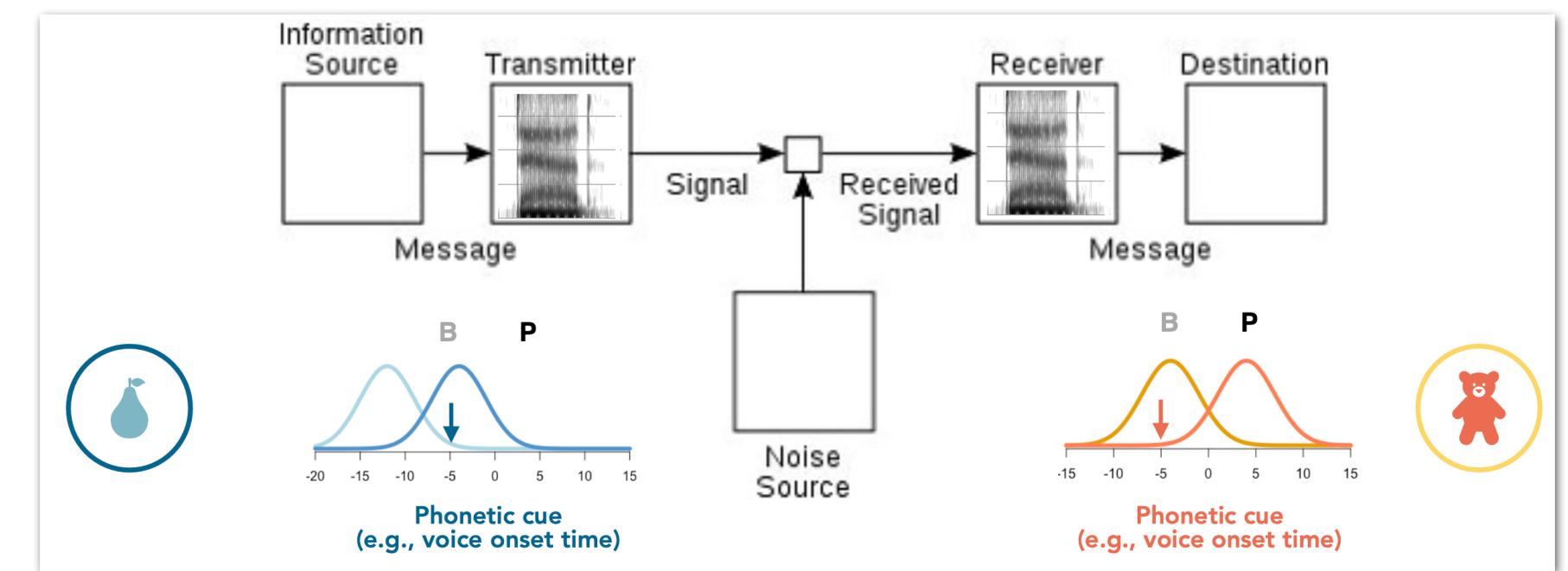
Side Notes: Priming as Linguistic Adaptation

- **Linguistic Adaptation:** adjustment of distributional knowledge / expectations to accommodate local input;
 - Speech adaptation & perception;
 - Syntactic satiation;
 - Informativeness calibration & presupposition/pragmatic routine calibration;



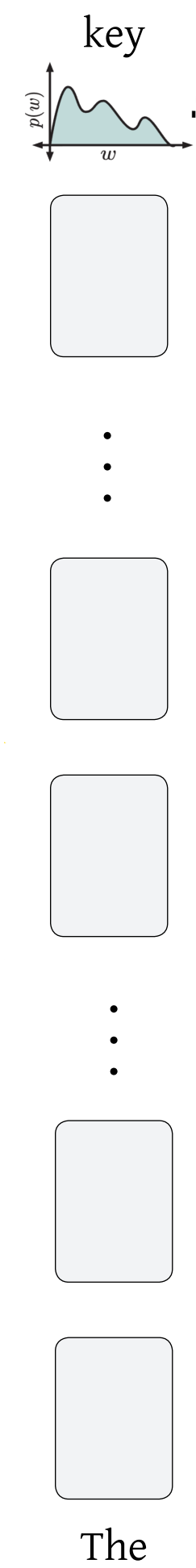
Side Notes: Priming as Linguistic Adaptation

- **Linguistic Adaptation:** adjustment of distributional knowledge / expectations to accommodate local input;
 - Speech adaptation & perception;
 - Syntactic satiation;
 - Informativeness calibration & presupposition/pragmatic routine calibration;
- **Predictive coding gives rise to linguistic adaptation**
 - Predictive coding: the system constantly predicts (linguistic) input and update the mental model based on the mismatch between predicted vs. actual input.
 - **The IFE is the *signature* that the update is error-driven**, not just residual activation: rarer structures create bigger prediction error → bigger update.



In-Context Learning in LLMs

A Primer on Intuitions behind LLMs



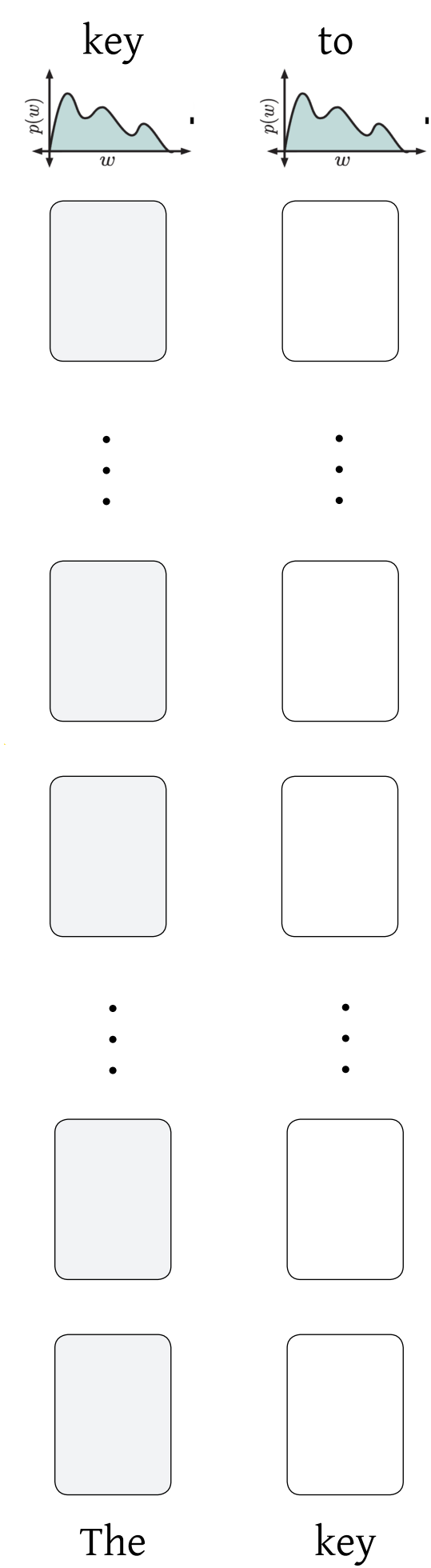
Time steps to process one sentence

A Primer on Intuitions behind LLMs



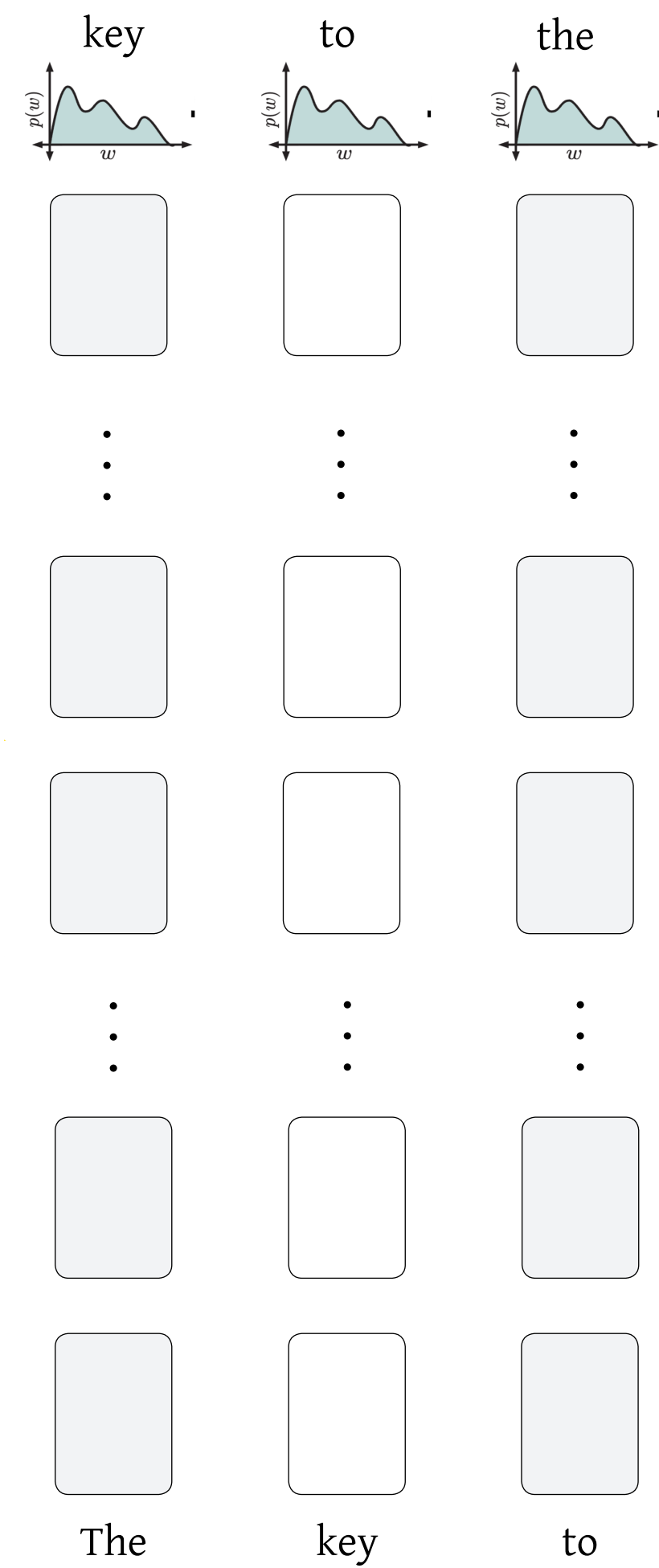
Time steps to process one sentence

A Primer on Intuitions behind LLMs



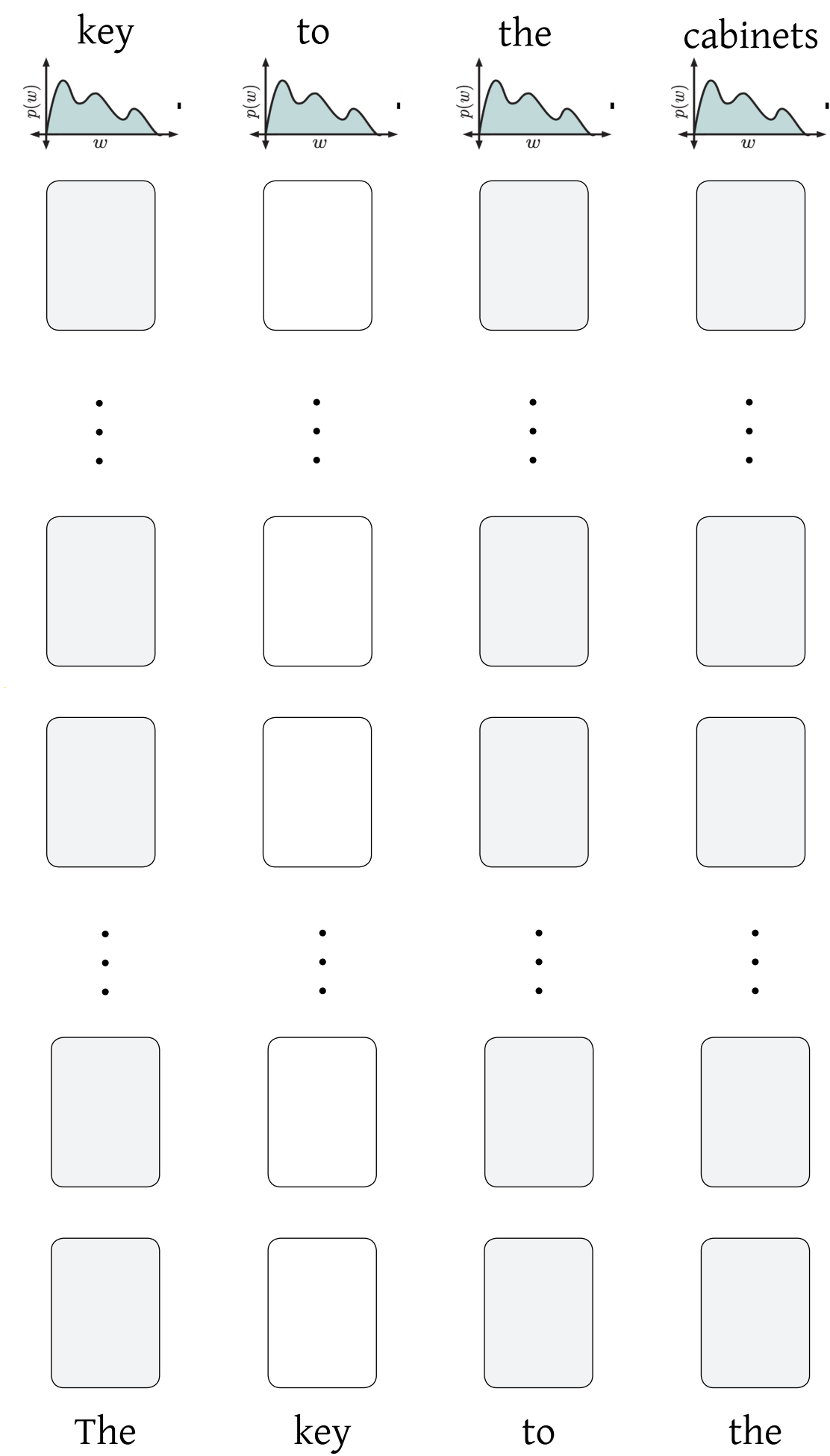
Time steps to process one sentence

A Primer on Intuitions behind LLMs



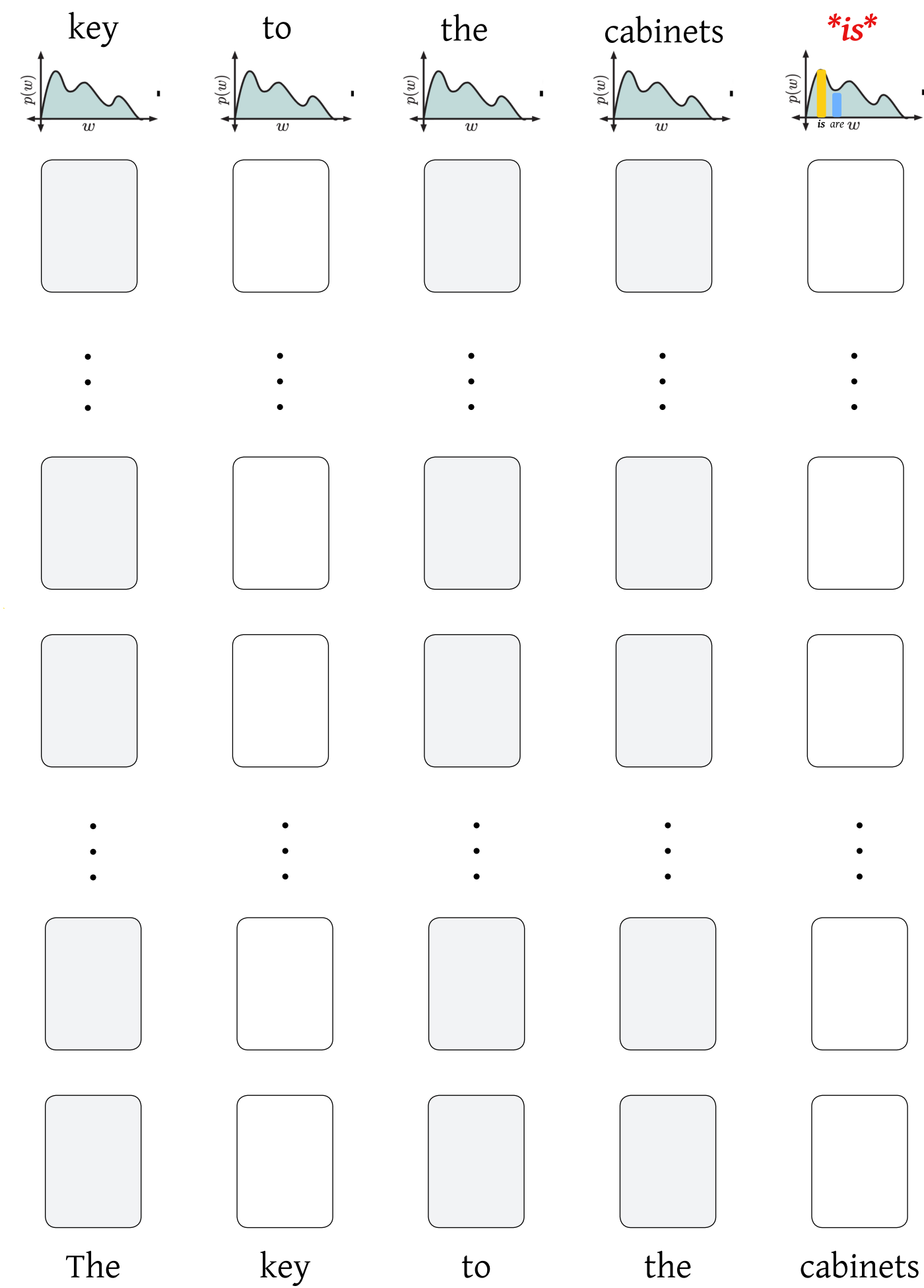
Time steps to process one sentence

A Primer on Intuitions behind LLMs



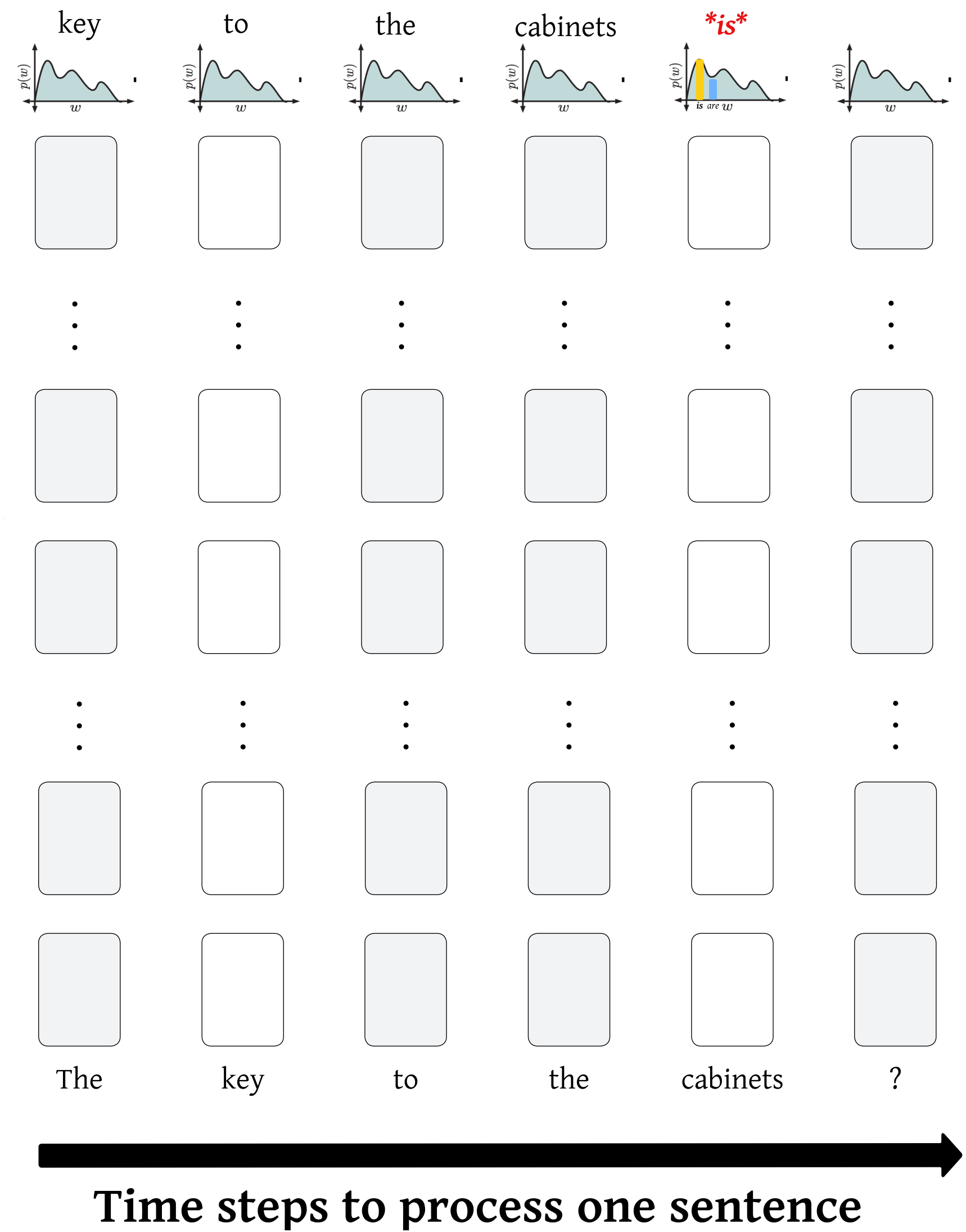
Time steps to process one sentence

A Primer on Intuitions behind LLMs

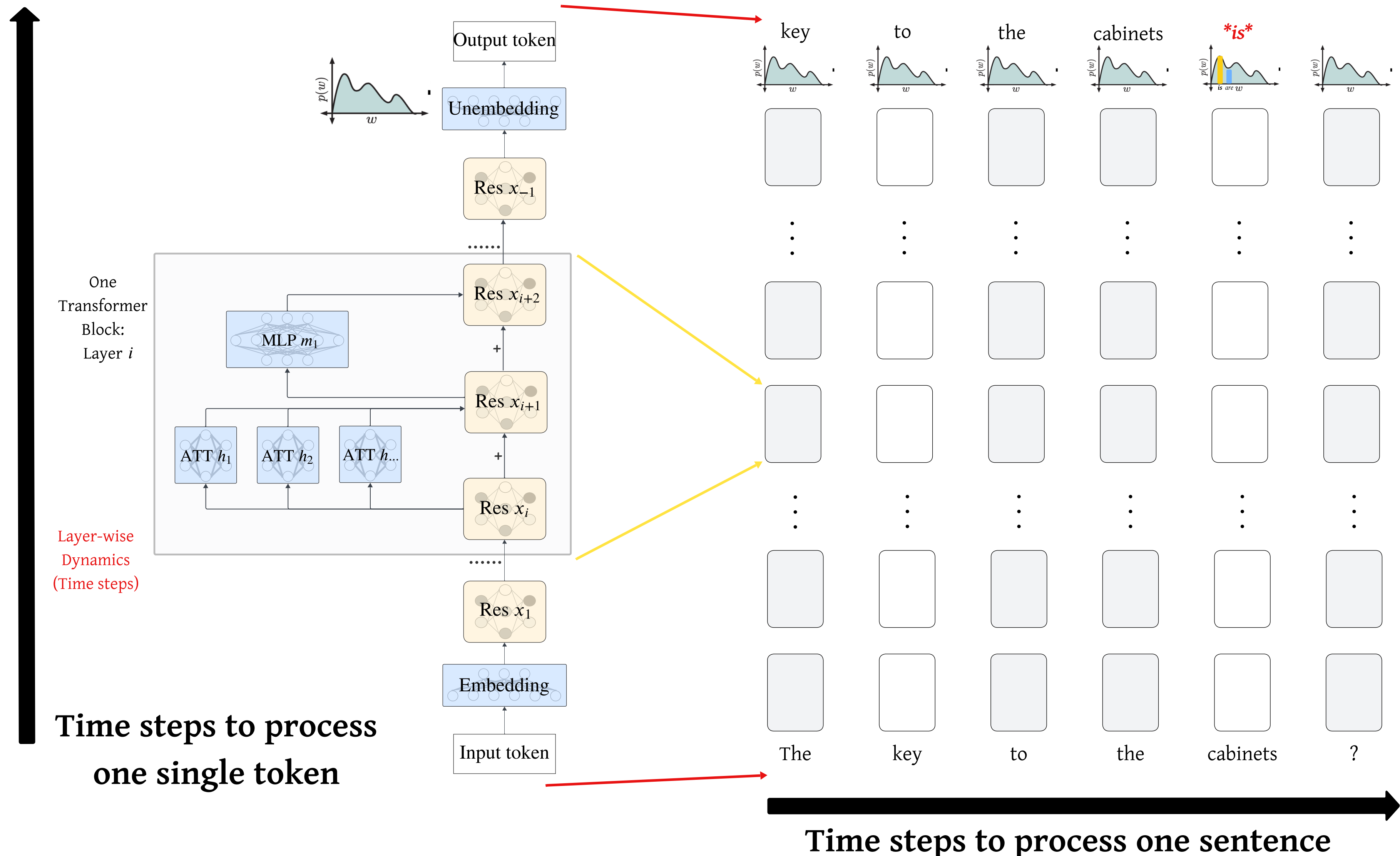


Time steps to process one sentence

A Primer on Intuitions behind LLMs



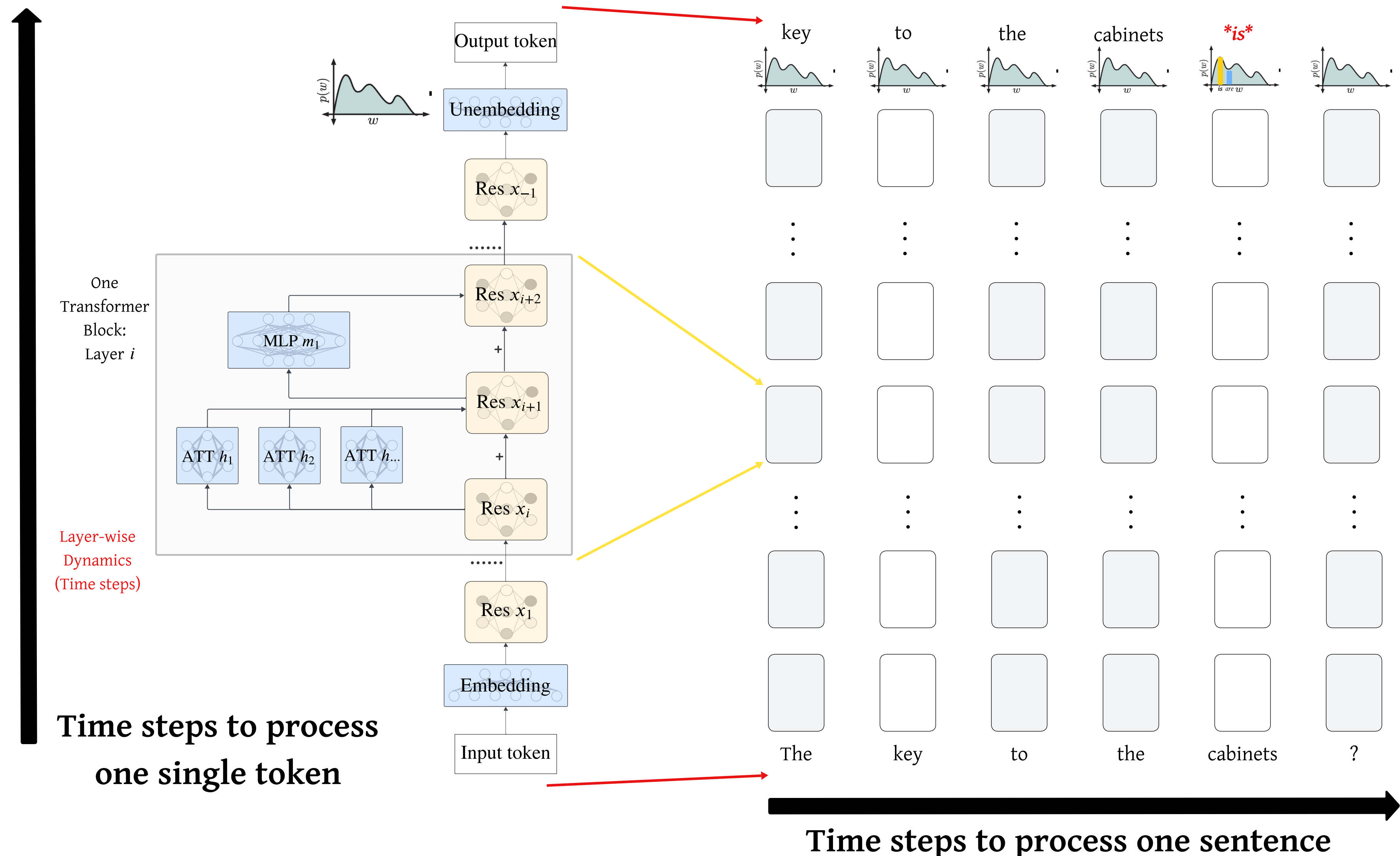
A Primer on Intuitions behind LLMs



A Primer on Intuitions behind LLMs

Modern LLMs:
machines of *next-token prediction*

- **Parameters:**
trainable weights
updated during
training & fine-
tuning → long-term
memory-ish;
- **Activations:**
temporary variables
generated during
running → working
memory;



In-Context Learning (ICL) in LLMs

In-Context Learning: several demonstration-answer pairs of a task in context improves models' performance on this task.

- No explicit description of the task;
- No weight (long-term memory) update!

In-Context Learning (ICL) in LLMs

In-Context Learning: several demonstration-answer pairs of a task in context improves models' performance on this task.

- No explicit description of the task;
- No weight (long-term memory) update!

The diagram shows a sequence of tokens in a blue box, illustrating the concept of In-Context Learning. The tokens are arranged in a grid-like structure, with the last row being shorter than the others. The tokens are: &(), --, >, (), &, 234, --, >, 342, @#\$, --, >, #\$, @, rte, --, >.

In-Context Learning (ICL) in LLMs

In-Context Learning: several demonstration-answer pairs of a task in context improves models' performance on this task.

- No explicit description of the task;
- No weight (long-term memory) update!

$$1 * 1 = 2$$

$$3 * 3 = 6$$

$$4 * 2 = 6$$

$$7 * 1 = ??$$

In-Context Learning (ICL) in LLMs

In-Context Learning: several demonstration-answer pairs of a task in context improves models' performance on this task.

- No explicit description of the task;
- No weight (long-term memory) update!

Language Models are Few-Shot Learners				
Tom B. Brown*	Benjamin Mann*	Nick Ryder*	Melanie Subbiah*	
Jared Kaplan†	Prafulla Dhariwal	Arvind Neelakantan	Pranav Shyam	Girish Sastry
Amanda Askell	Sandhini Agarwal	Ariel Herbert-Voss	Gretchen Krueger	Tom Henighan
Rewon Child	Aditya Ramesh	Daniel M. Ziegler	Jeffrey Wu	Clemens Winter
Christopher Hesse	Mark Chen	Eric Sigler	Mateusz Litwin	Scott Gray
Benjamin Chess		Jack Clark	Christopher Berner	
Sam McCandlish	Alec Radford	Ilya Sutskever	Dario Amodei	
OpenAI				

1 * 1 = 2

3 * 3 = 6

4 * 2 = 6

7 * 1 = ??

ICL vs. In-*Weights* Learning

[Brown et al. 2020]

[Brown et al. 2020]

ICL vs. *In-Weights* Learning

One-shot

In addition to the task description, the model sees a single example of the task. No gradient updates are performed.

```
1 Translate English to French: ← task description
2 sea otter => loutre de mer ← example
3 cheese => ..... ← prompt
```

In-context Learning:

- No gradient updates;
- Rapid: from a few examples;
- One-shot / Few-shot learning;

Few-shot

In addition to the task description, the model sees a few examples of the task. No gradient updates are performed.

```
1 Translate English to French: ← task description
2 sea otter => loutre de mer ← examples
3 peppermint => menthe poivrée ←
4 plush girafe => girafe peluche ←
5 cheese => ..... ← prompt
```

[Brown et al. 2020]

ICL vs. In-Weights Learning

One-shot

In addition to the task description, the model sees a single example of the task. No gradient updates are performed.

```
1 Translate English to French: ← task description
2 sea otter => loutre de mer ← example
3 cheese => ..... ← prompt
```

Few-shot

In addition to the task description, the model sees a few examples of the task. No gradient updates are performed.

```
1 Translate English to French: ← task description
2 sea otter => loutre de mer ← examples
3 peppermint => menthe poivrée ←
4 plush girafe => girafe peluche ←
5 cheese => ..... ← prompt
```

[Brown et al. 2020]

In-context Learning:

- No gradient updates;
- Rapid: from a few examples;
- One-shot / Few-shot learning;

In-weights Learning (Fine-tuning):

- Gradient-based;
- Slow: need many examples;
- Standard supervised learning;

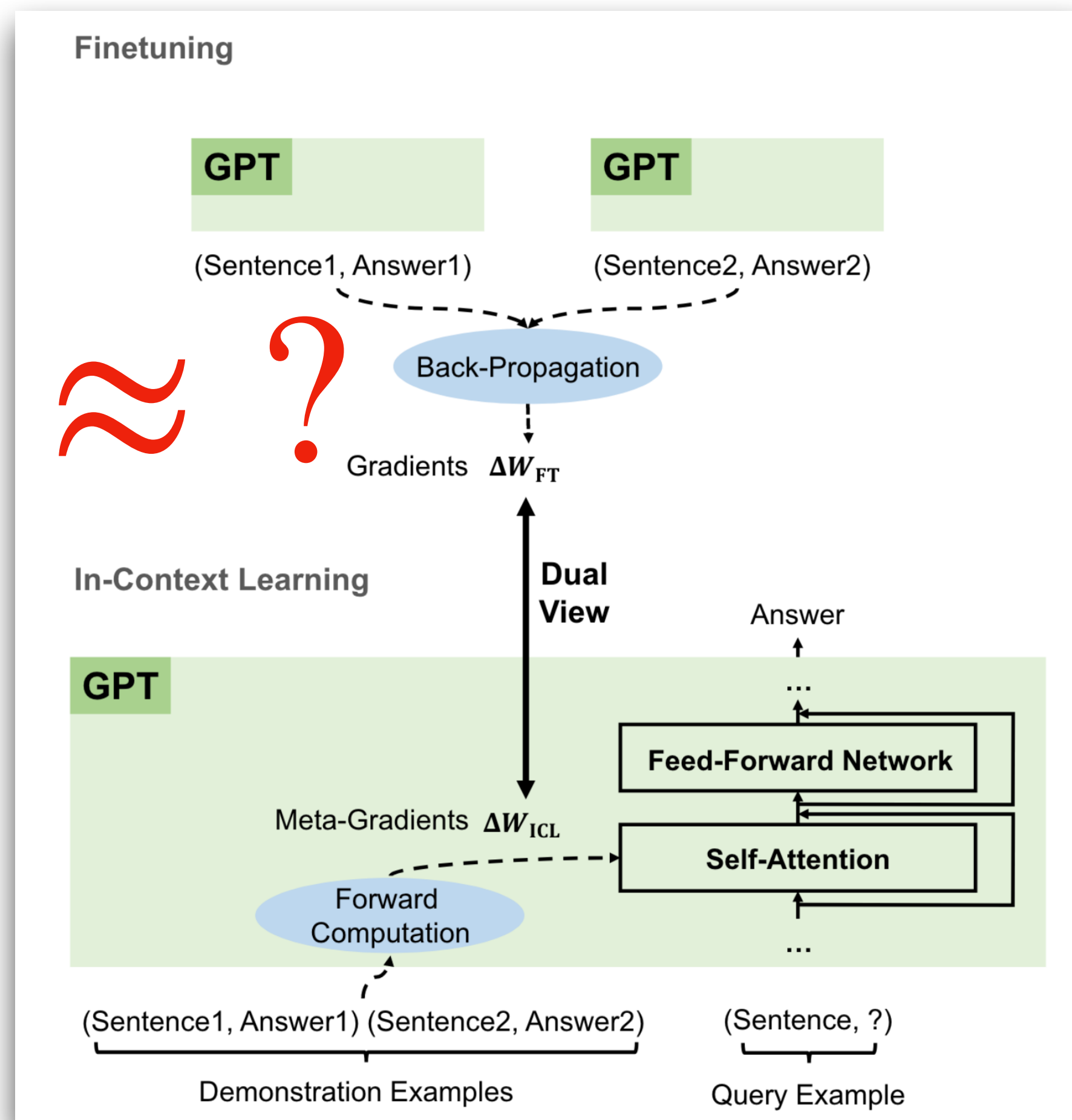
Fine-tuning

The model is trained via repeated gradient updates using a large corpus of example tasks.



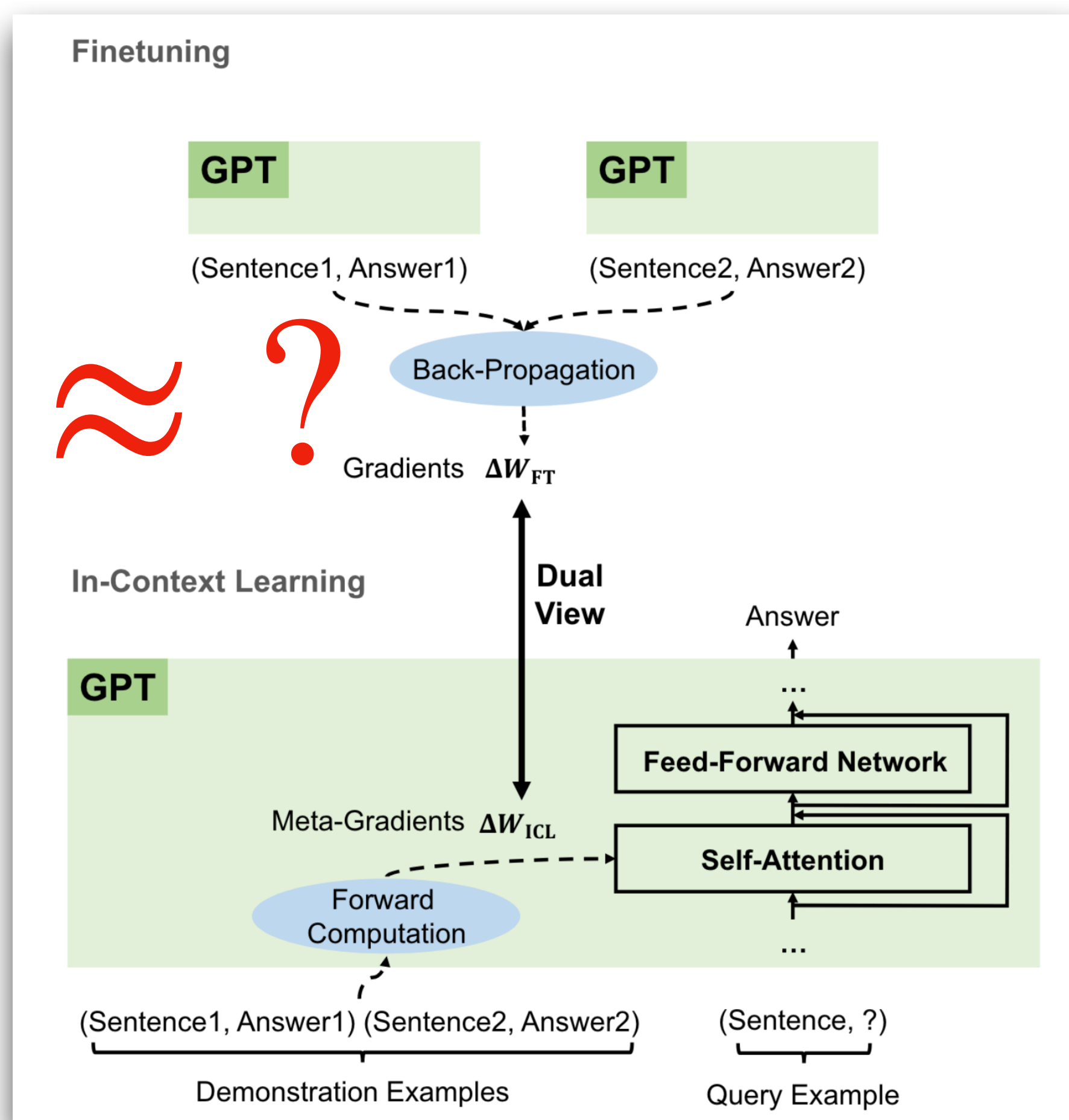
[Brown et al. 2020]

Research Question: is there error-driven learning mechanisms in ICL?



[e.g., Xie et al. 2022, von Oswald et al. 2022, Dai et al. 2023]

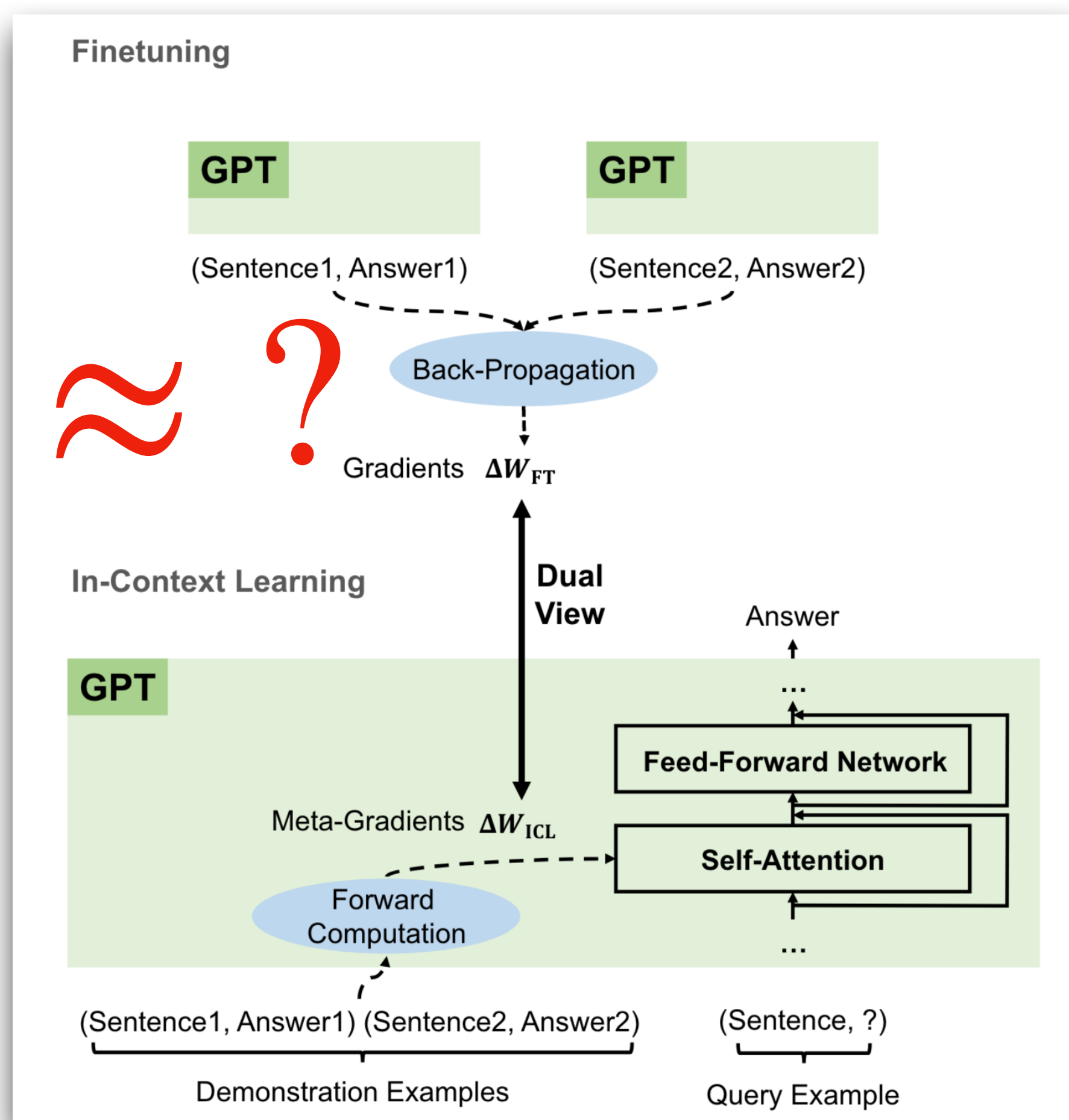
Research Question: is there error-driven learning mechanisms in ICL?



Previous claims: ICL may be implicitly doing learning

- Implicit Bayesian update; gradient descent; meta-optimization / implicit fine-tuning....

Research Question: is there error-driven learning mechanisms in ICL?



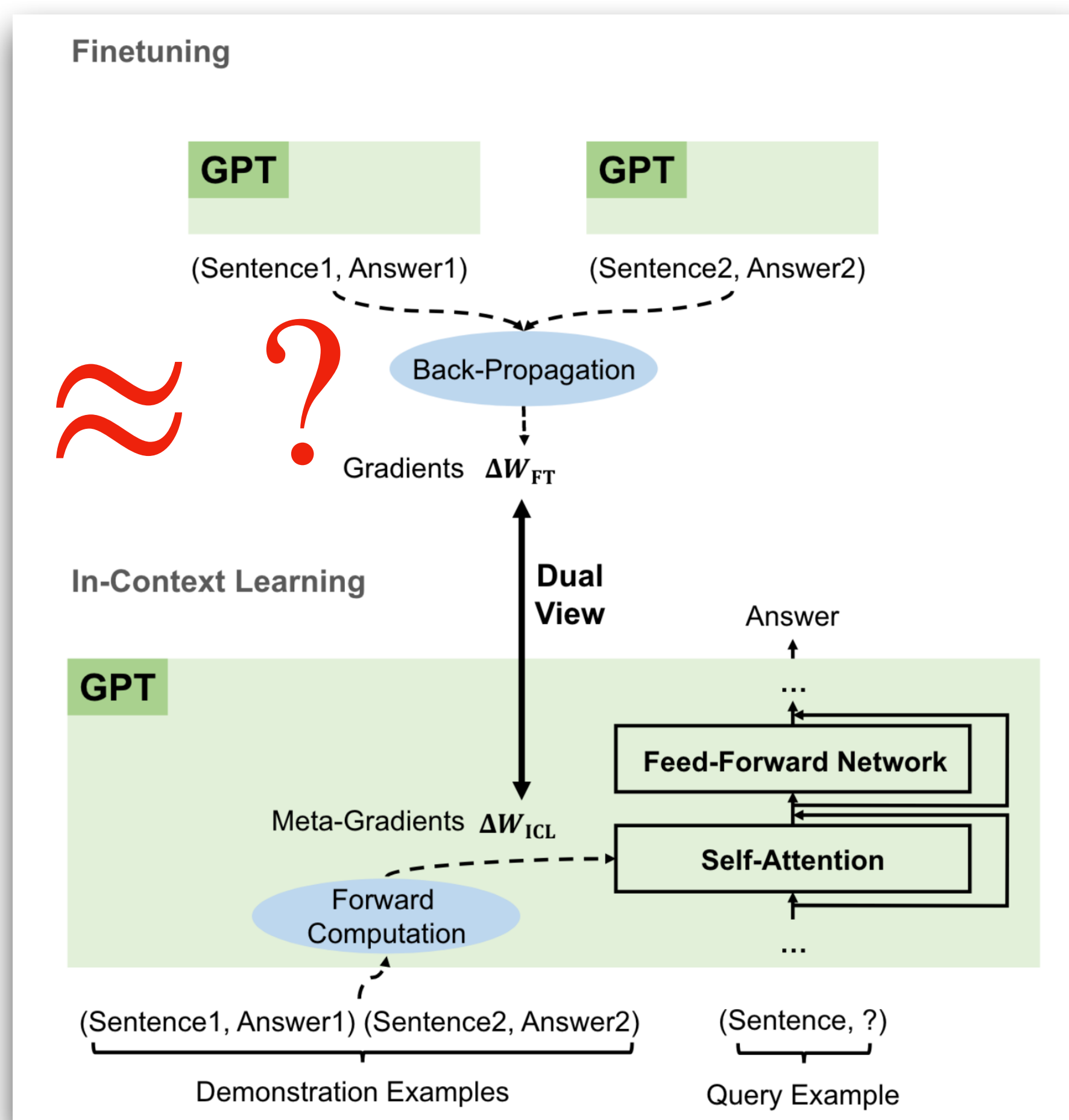
Previous claims: ICL may be implicitly doing learning

- Implicit Bayesian update; gradient descent; meta-optimization / implicit fine-tuning....

Previous studies in AI/CS:

- On toy models and toy tasks;
- In-principle claims: math derivations with assumptions on architecture;

Research Question: is there error-driven learning mechanisms in ICL?



Previous claims: ICL may be implicitly doing learning

- Implicit Bayesian update; gradient descent; meta-optimization / implicit fine-tuning....

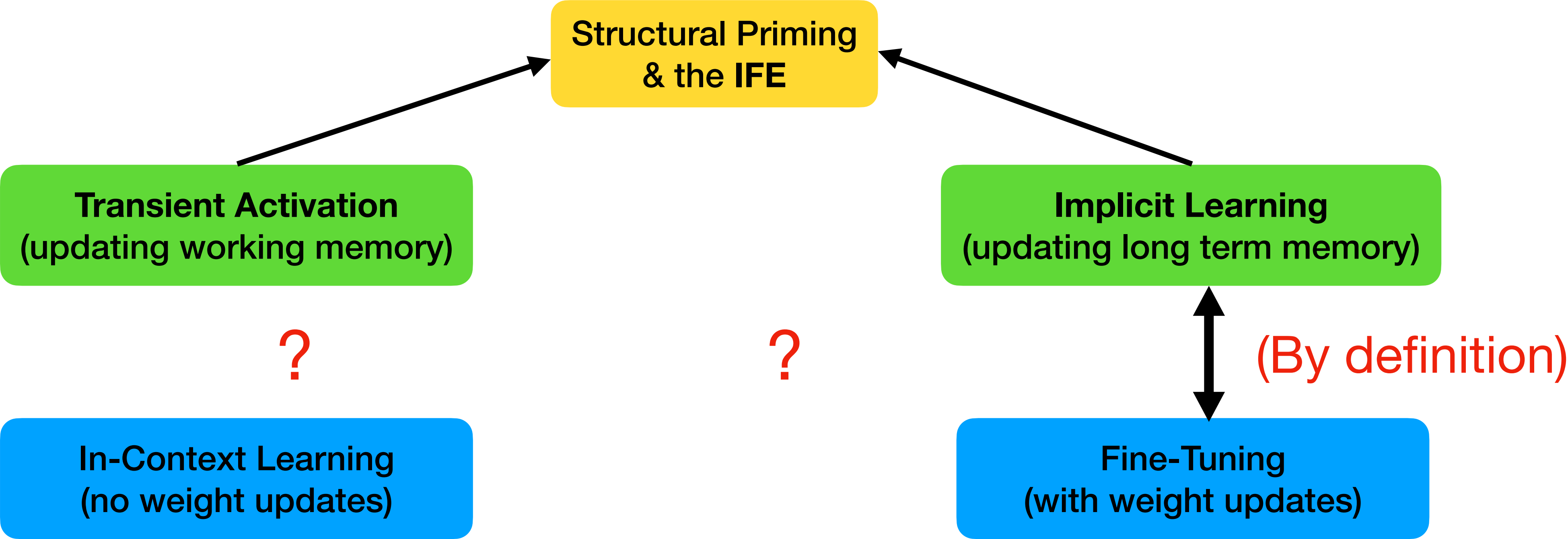
Previous studies in AI/CS:

- On toy models and toy tasks;
- In-principle claims: math derivations with assumptions on architecture;

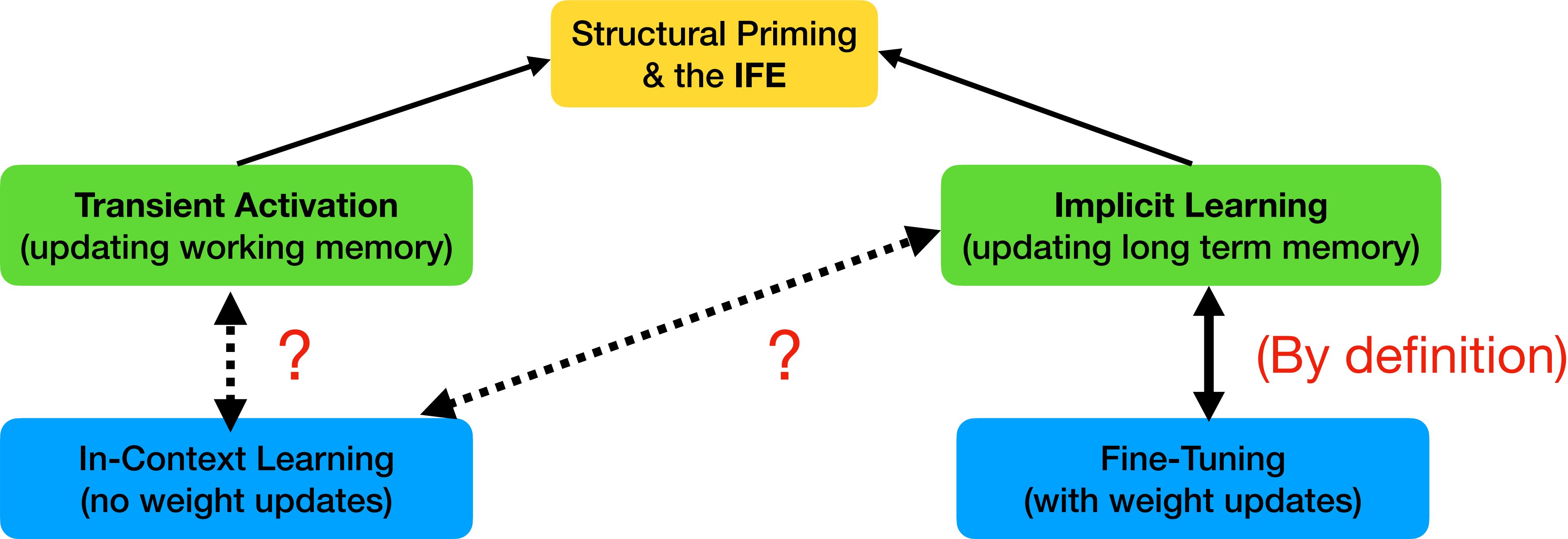
Current work:

- On pretrained, actual LLMs;
- Psycholinguistically motivated, more testing a more fine-grained effect.

Current Picture

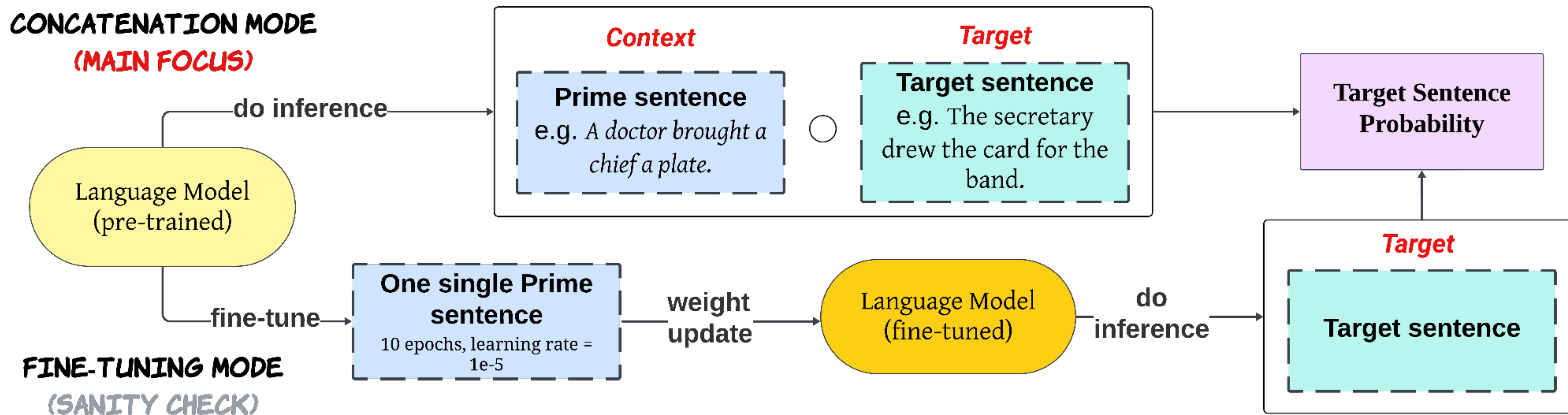


Current Picture

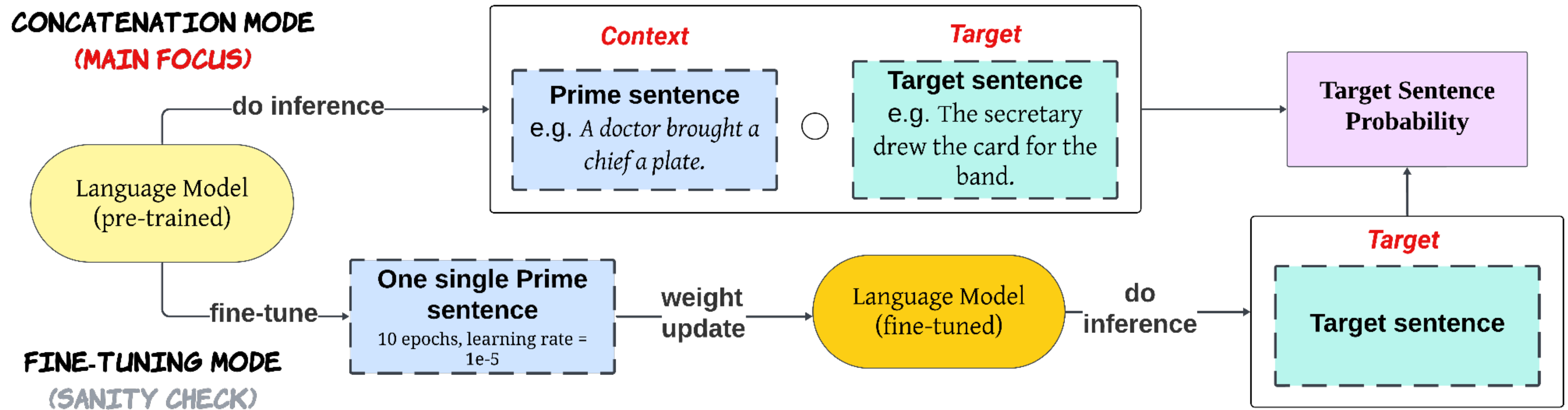


Experiment 1: Behavioral Analysis

Overview of the Current Approach



Overview of the Current Approach



Assumption from Priming Theories: only some error-driven learning mechanism could lead to the IFE. Then, we predict that:

- Fine-tuning Mode: (with weight update) IFE ✓
- Concatenation Mode: (no weight update) IFE ?

Corpus

- 22 ditransitive verbs;
- 50 target sentences per verb;
- For each target sentence, pair it with a prime sentence with each prime verb;



22 x 50 (target sentences)

x

21 (prime sentences)

=

23100 <prime, target> pairs

Corpus

- 22 ditransitive verbs;
- 50 target sentences per verb;
- For each target sentence, pair it with a prime sentence with each prime verb;

22 x 50 (target sentences)

x

21 (prime sentences)

=

23100 <prime, target> pairs



Each <prime, target> pair $\Rightarrow$ 4 structural combinations $\Rightarrow$ 92400 trials.

Prime_{DO} $\rightarrow$ Target_{DO}, Prime_{DO} $\rightarrow$ Target_{PD}, Prime_{PD} $\rightarrow$ Target_{DO}, Prime_{PD} $\rightarrow$ Target_{PD}

Corpus

- 22 ditransitive verbs;
- 50 target sentences per verb;
- For each target sentence, pair it with a prime sentence with each prime verb;

22 x 50 (target sentences)

x

21 (prime sentences)

=

23100 <prime, target> pairs



Each <prime, target> pair $\Rightarrow$ 4 structural combinations $\Rightarrow$ 92400 trials.

Prime_{DO} $\rightarrow$ Target_{DO}, Prime_{DO} $\rightarrow$ Target_{PD}, Prime_{PD} $\rightarrow$ Target_{DO}, Prime_{PD} $\rightarrow$ Target_{PD}

Prime: A doctor brought a plate to a chief.

Target: The secretary drew the card for the band.

Quantifying Verb Biases and the IFE

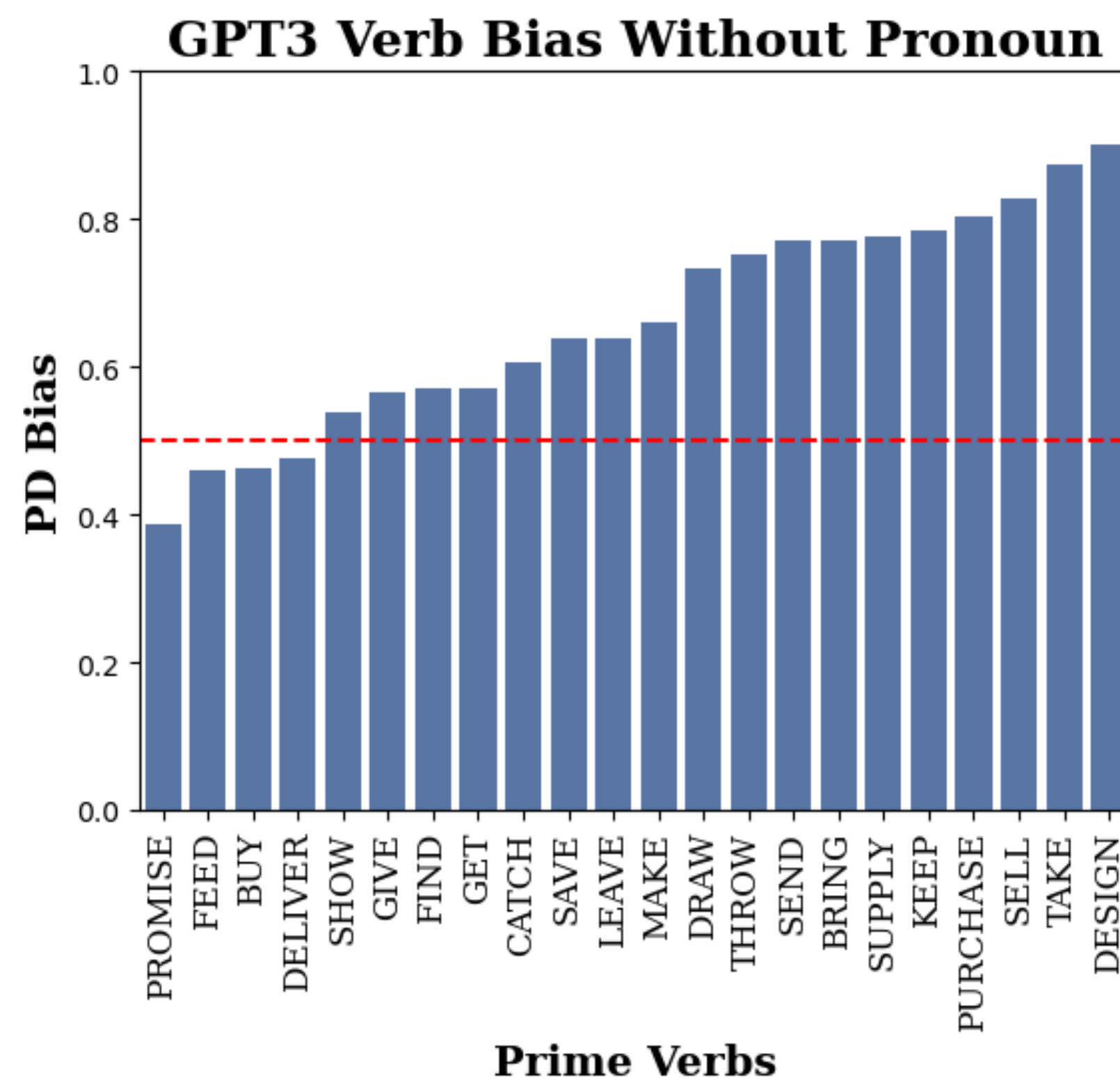
Verb Bias of verb V on structure X :

$$\textit{bias}(V, \text{PD}) = \frac{1}{|\mathcal{S}_V|} \sum_{t_{\text{PD}} \in \mathcal{S}_V} \frac{\mathcal{P}(t_{\text{PD}})}{\mathcal{P}(t_{\text{PD}}) + \mathcal{P}(t_{\text{DO}})}$$

Quantifying Verb Biases and the IFE

Verb Bias of verb V on structure X :

$$\text{bias}(V, \text{PD}) = \frac{1}{|\mathcal{S}_V|} \sum_{t_{\text{PD}} \in \mathcal{S}_V} \frac{\mathcal{P}(t_{\text{PD}})}{\mathcal{P}(t_{\text{PD}}) + \mathcal{P}(t_{\text{DO}})}$$



Quantifying Verb Biases and the IFE

Verb Bias of verb V on structure X :

$$\textit{bias}(V, \text{PD}) = \frac{1}{|\mathcal{S}_V|} \sum_{t_{\text{PD}} \in \mathcal{S}_V} \frac{\mathcal{P}(t_{\text{PD}})}{\mathcal{P}(t_{\text{PD}}) + \mathcal{P}(t_{\text{DO}})}$$

Quantifying Verb Biases and the IFE

Verb Bias of verb V on structure X :

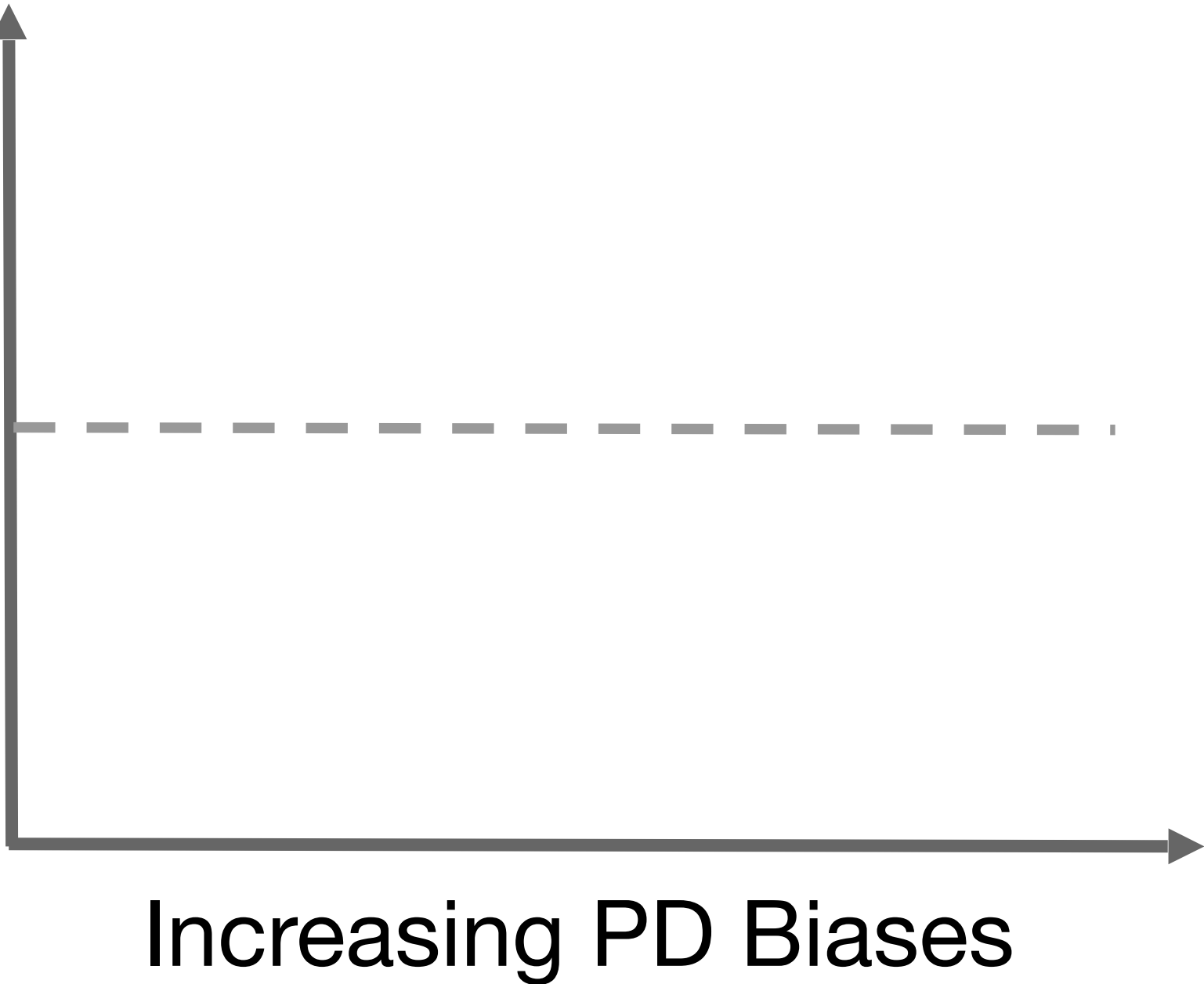
$$\textit{bias}(V, \text{PD}) = \frac{1}{|\mathcal{S}_V|} \sum_{t_{\text{PD}} \in \mathcal{S}_V} \frac{\mathcal{P}(t_{\text{PD}})}{\mathcal{P}(t_{\text{PD}}) + \mathcal{P}(t_{\text{DO}})}$$

IFE: the priming effect for verb V in DO form on PD targets

$$\textit{PrimeBias}(\text{PD}|\text{DO}, V) = \frac{1}{|T_{\text{PD}}| \cdot |P_{\text{DO}}^V|} \sum_{t_{\text{PD}} \in T_{\text{PD}}} \sum_{p_{\text{DO}}^V \in P_{\text{DO}}^V} \frac{\mathcal{P}(t_{\text{PD}}|p_{\text{DO}}^V)}{\mathcal{P}(t_{\text{DO}}|p_{\text{DO}}^V) + \mathcal{P}(t_{\text{PD}}|p_{\text{DO}}^V)}$$

Predictions on the IFE

Prime_{PD} → Target_{PD}

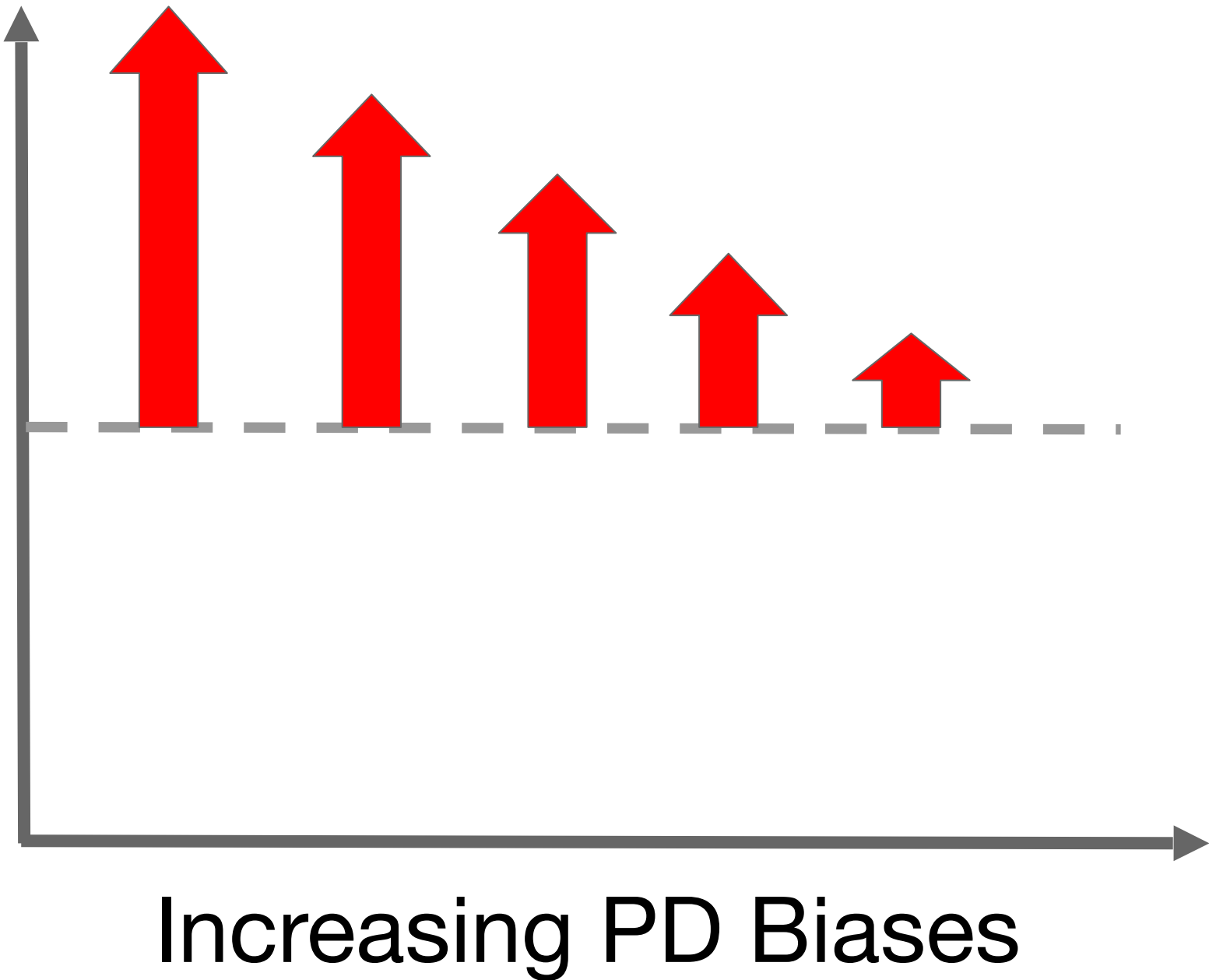


<u>Prime Verb</u>	<u>Verb PD Bias</u>	<u>Primed probability</u> (priming magnitude)
<i>Bring</i>	0.23	0.23 (0.03)
<i>Buy</i>	0.27	0.25 (0.05)
<i>Find</i>	0.41	0.28 (0.08)
<i>Draw</i>	0.52	0.30 (0.10)
<i>Design</i>	0.77	0.33 (0.13)

priming in DO *larger PD biases* *greater priming effects*

Predictions on the IFE

Prime_{PD} → Target_{PD}

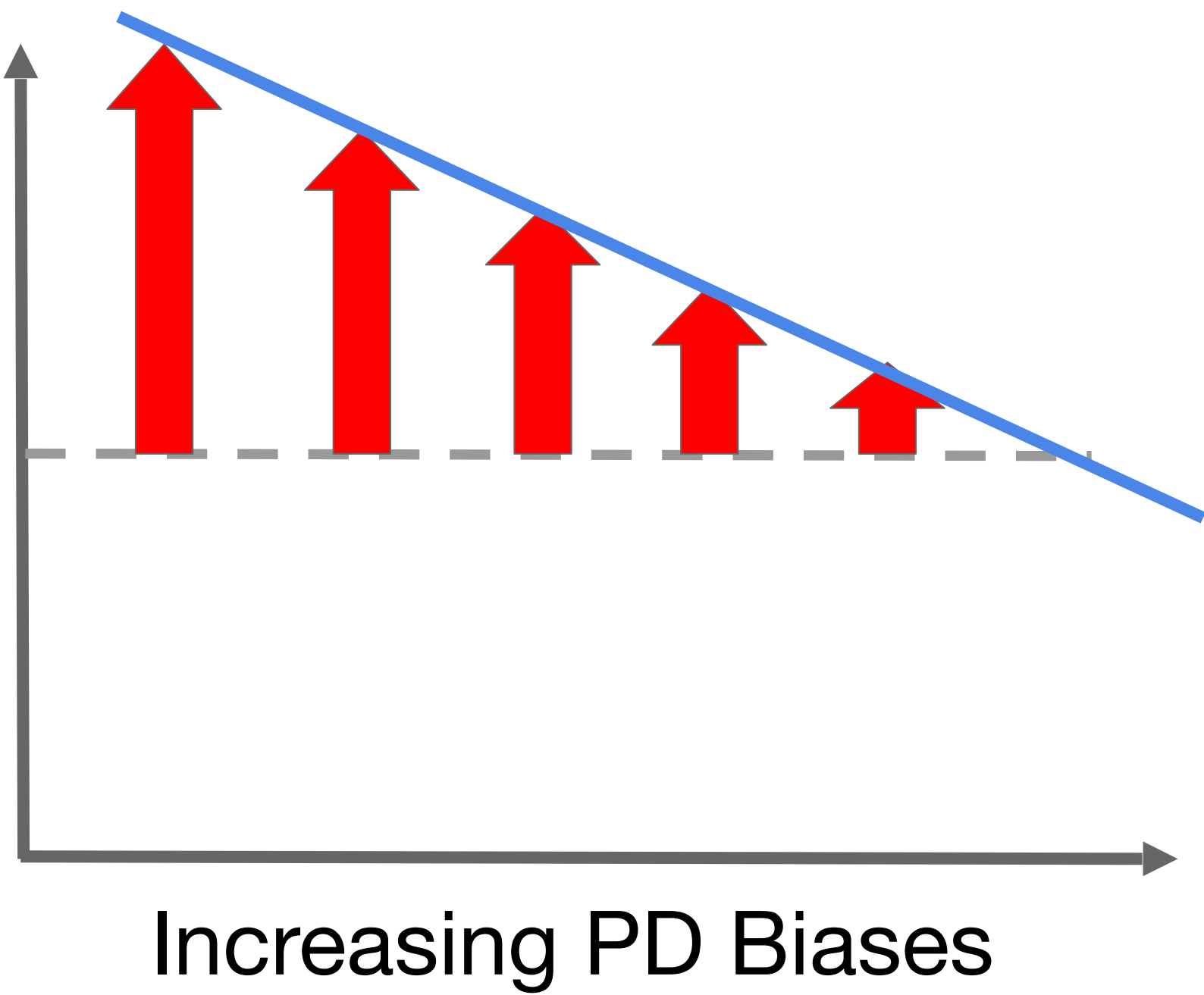


Prime Verb	Verb PD Bias	Primed probability (priming magnitude)
<i>Bring</i>	0.23	0.23 (0.03)
<i>Buy</i>	0.27	0.25 (0.05)
<i>Find</i>	0.41	0.28 (0.08)
<i>Draw</i>	0.52	0.30 (0.10)
<i>Design</i>	0.77	0.33 (0.13)

priming in DO *larger PD biases* *greater priming effects*

Predictions on the IFE

Prime_{PD} → Target_{PD}



Prime Verb	Verb PD Bias	Primed probability (priming magnitude)
<i>Bring</i>	0.23	0.23 (0.03)
<i>Buy</i>	0.27	0.25 (0.05)
<i>Find</i>	0.41	0.28 (0.08)
<i>Draw</i>	0.52	0.30 (0.10)
<i>Design</i>	0.77	0.33 (0.13)

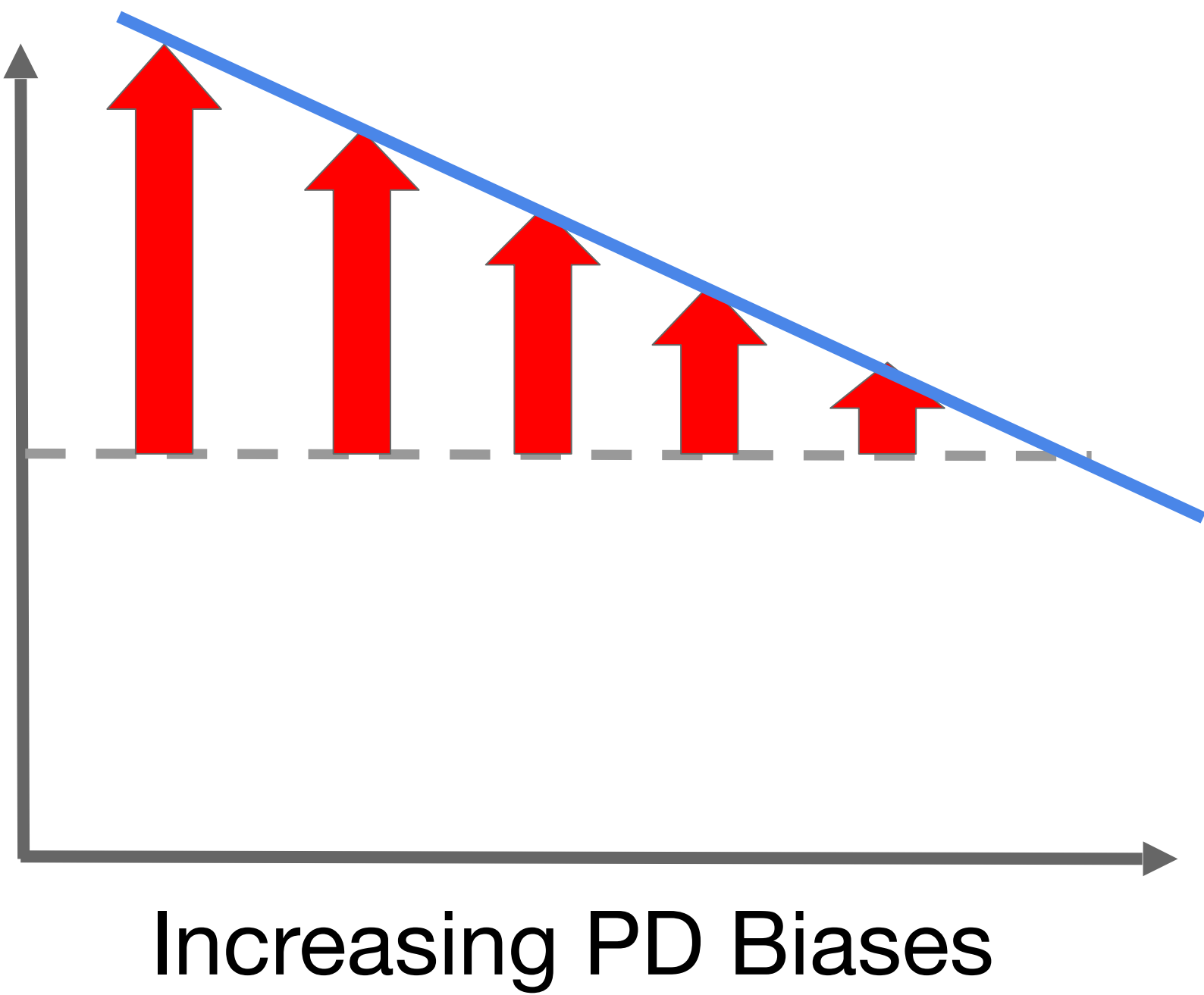
priming in DO *larger PD biases* *greater priming effects*

Predictions on the IFE

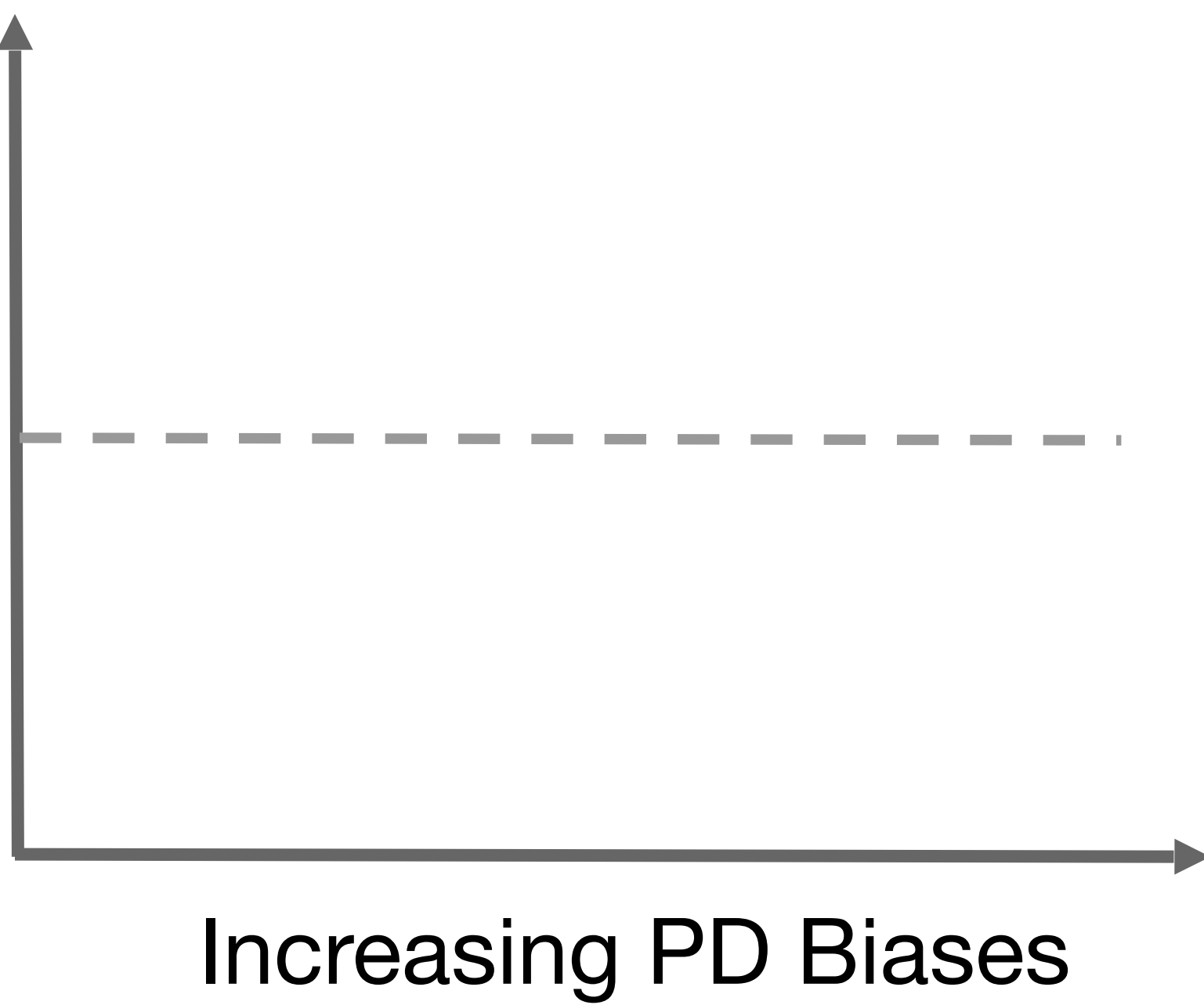
Prime Verb	Verb PD Bias	Primed probability (priming magnitude)
Bring	0.23	0.23 (0.03)
Buy	0.27	0.25 (0.05)
Find	0.41	0.28 (0.08)
Draw	0.52	0.30 (0.10)
Design	0.77	0.33 (0.13)

priming in DO *larger PD biases* *greater priming effects*

Prime_{PD} → Target_{PD}

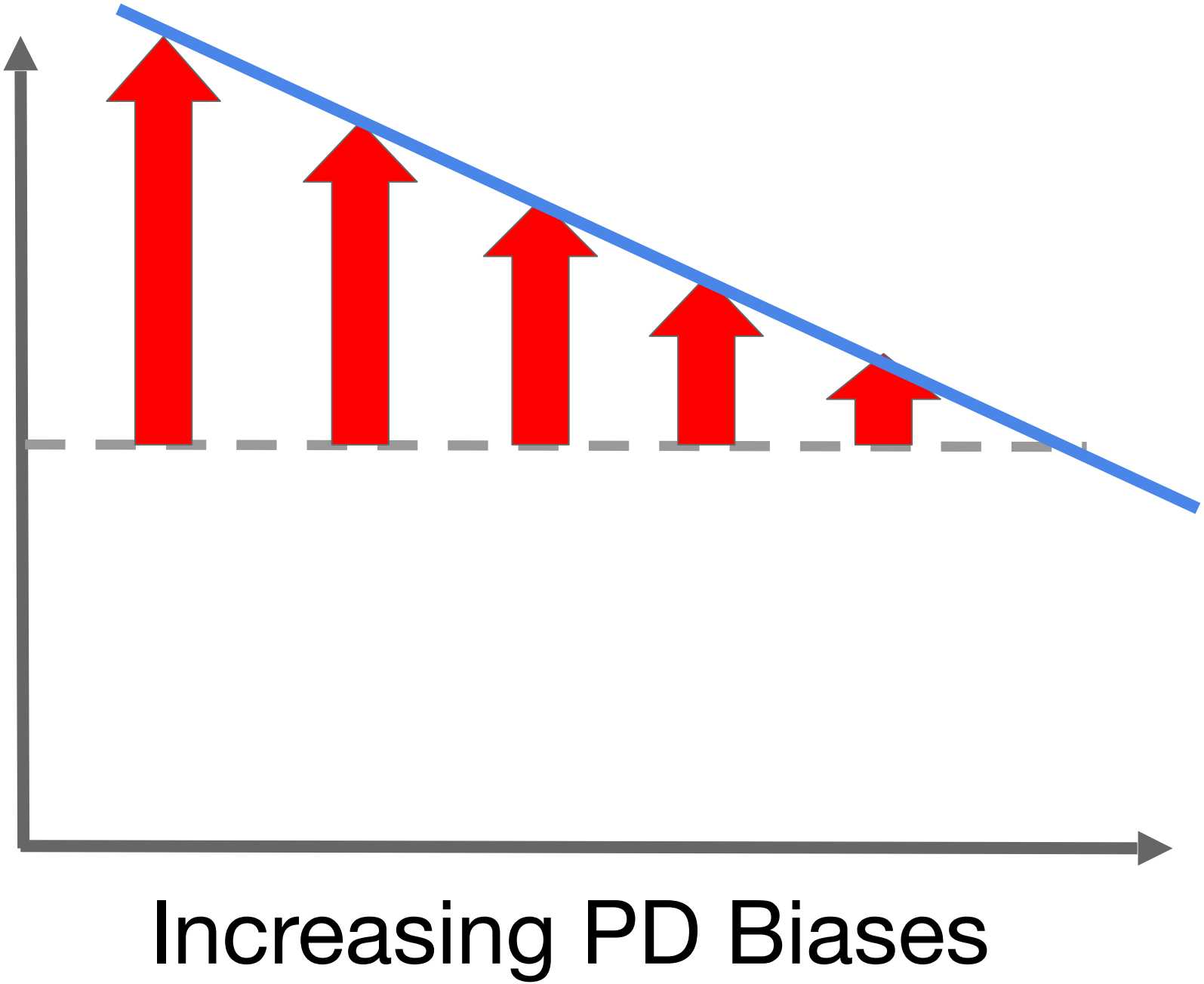


Prime_{DO} → Target_{PD}

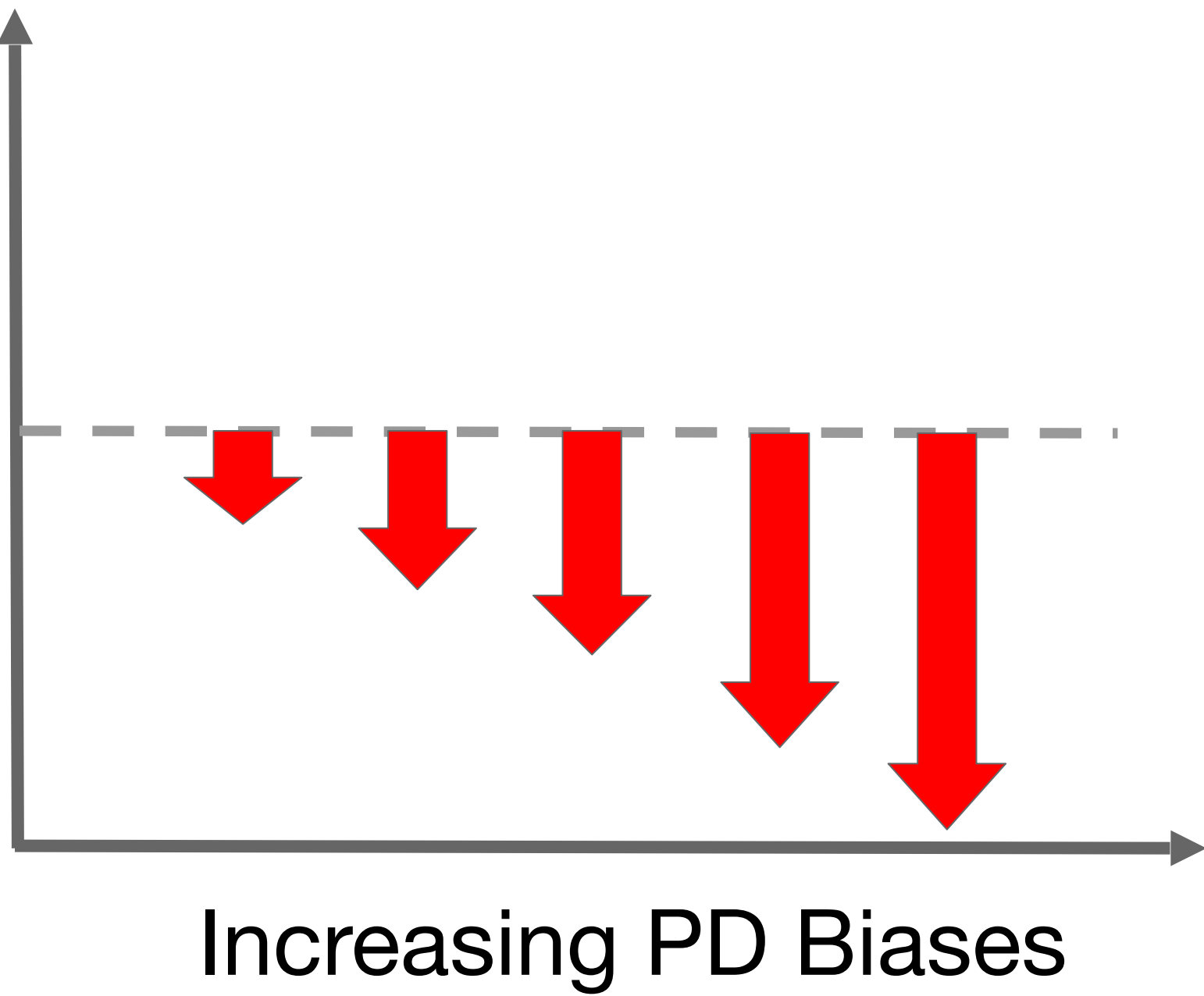


Predictions on the IFE

Prime_{PD} → Target_{PD}



Prime_{DO} → Target_{PD}



Prime Verb	Verb PD Bias	Primed probability (priming magnitude)
Bring	0.23	0.23 (0.03)
Buy	0.27	0.25 (0.05)
Find	0.41	0.28 (0.08)
Draw	0.52	0.30 (0.10)
Design	0.77	0.33 (0.13)

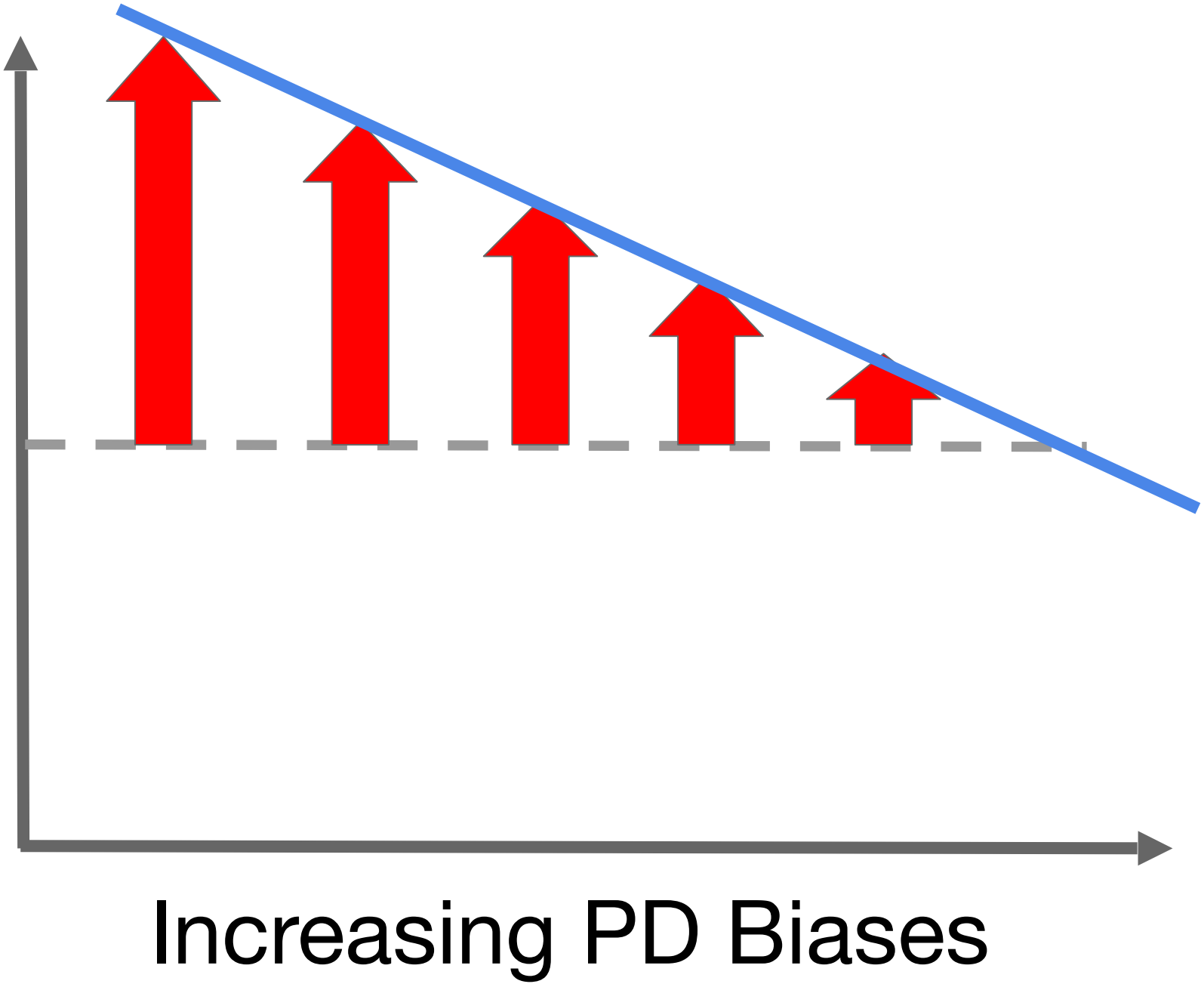
priming in DO *larger PD biases* *greater priming effects*

Predictions on the IFE

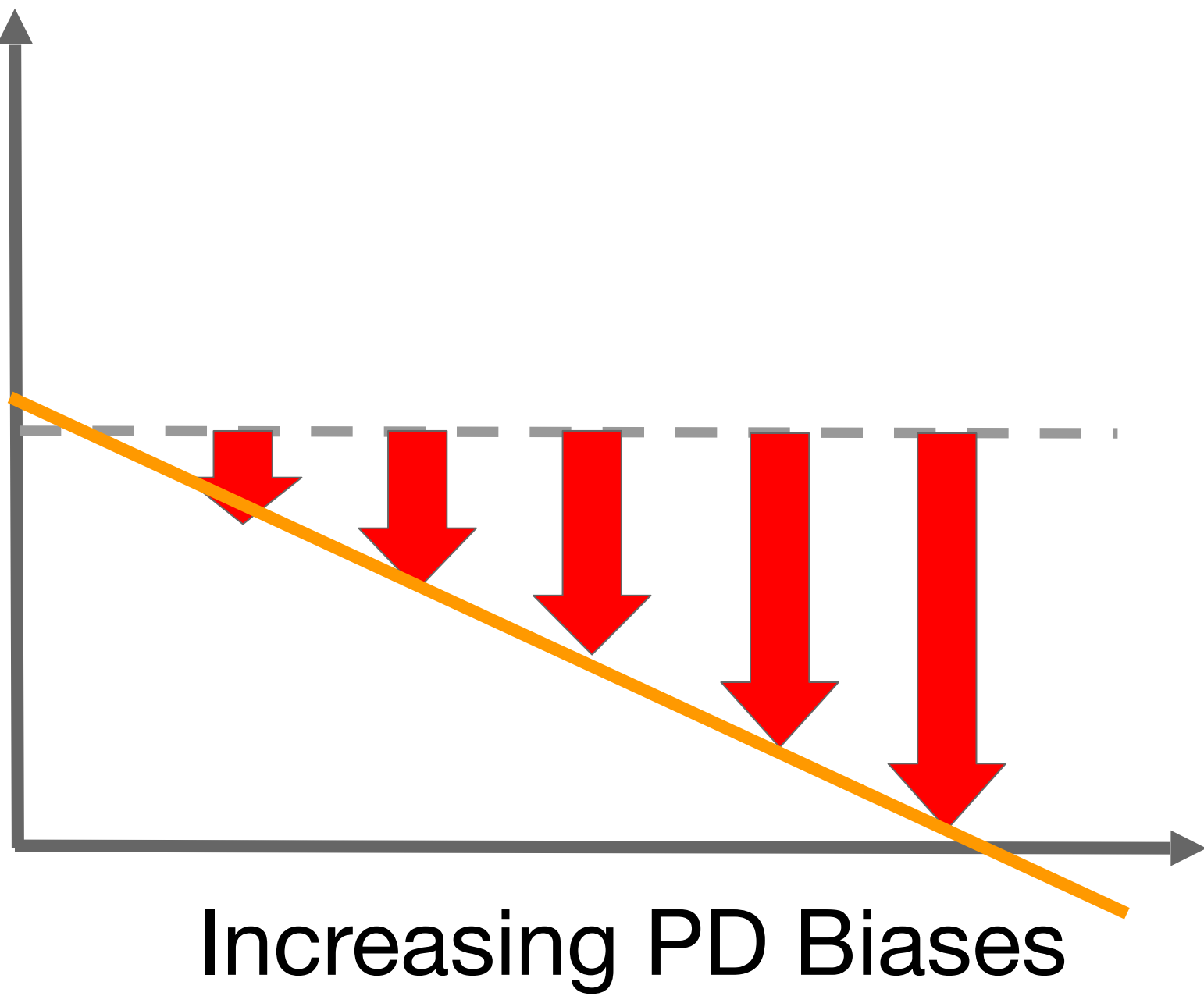
Prime Verb	Verb PD Bias	Primed probability (priming magnitude)
Bring	0.23	0.23 (0.03)
Buy	0.27	0.25 (0.05)
Find	0.41	0.28 (0.08)
Draw	0.52	0.30 (0.10)
Design	0.77	0.33 (0.13)

priming in DO *larger PD biases* *greater priming effects*

Prime_{PD} → Target_{PD}



Prime_{DO} → Target_{PD}

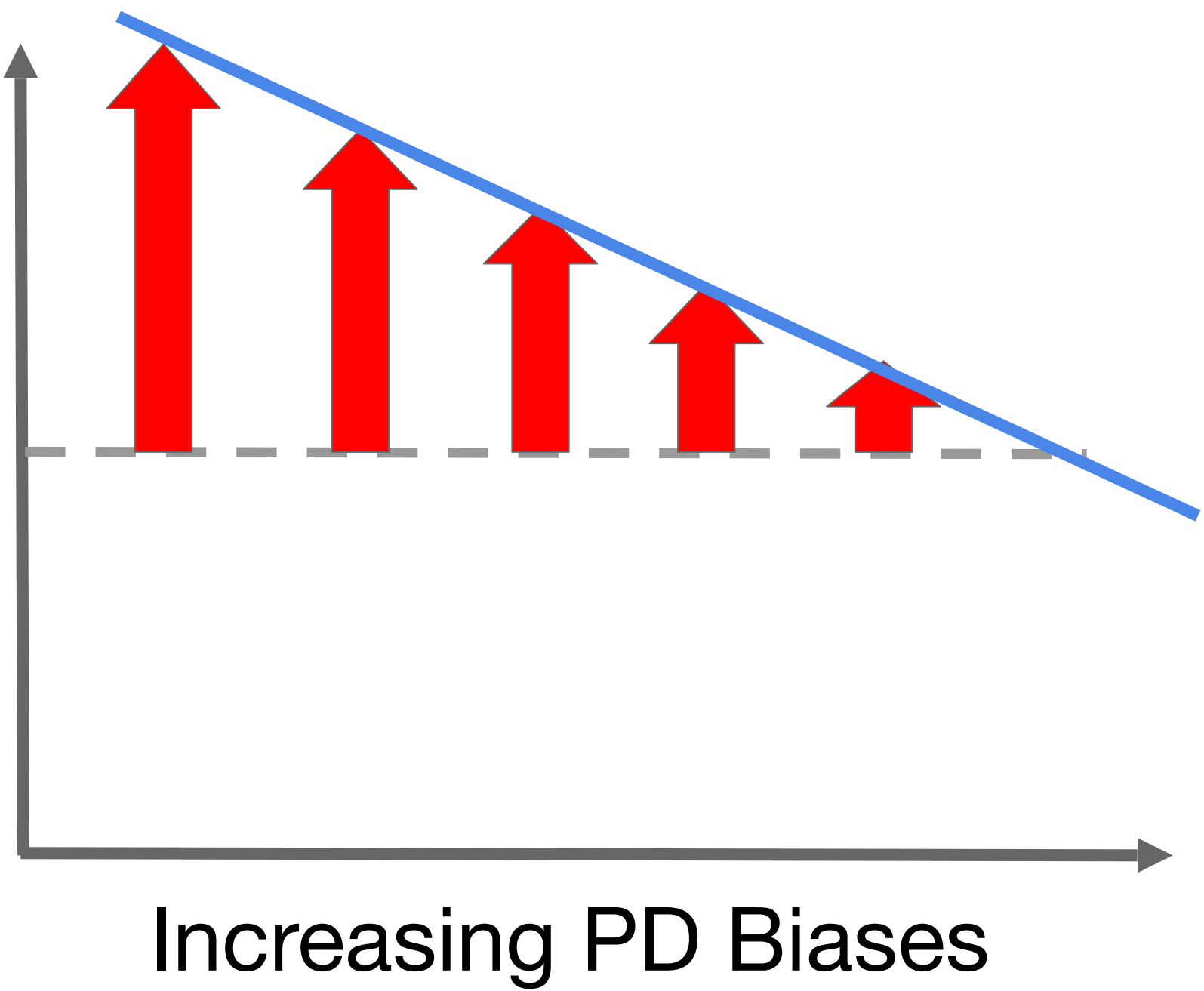


Predictions on the IFE

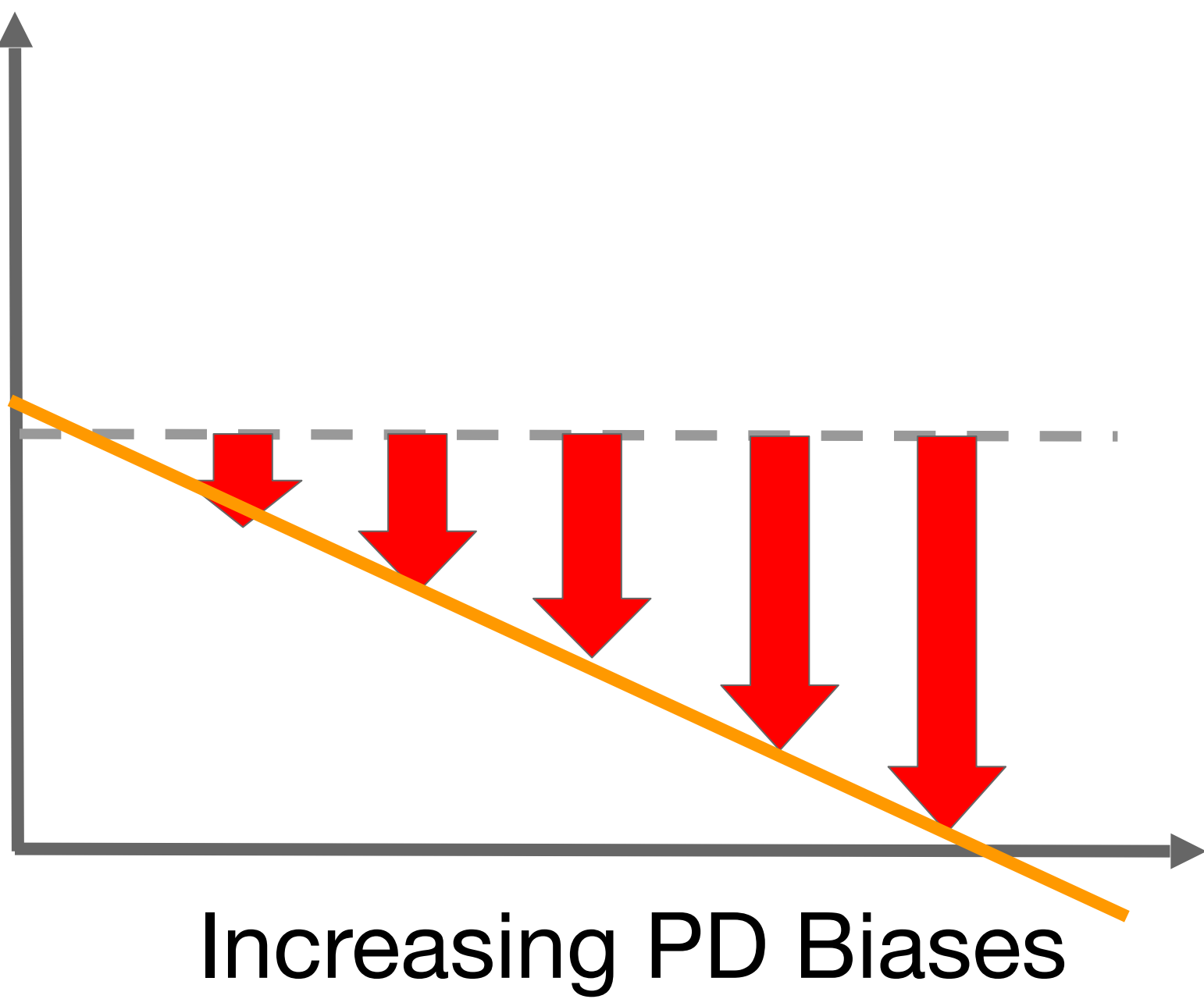
Prime Verb	Verb PD Bias	Primed probability (priming magnitude)
Bring	0.23	0.23 (0.03)
Buy	0.27	0.25 (0.05)
Find	0.41	0.28 (0.08)
Draw	0.52	0.30 (0.10)
Design	0.77	0.33 (0.13)

priming in DO *larger PD biases* *greater priming effects*

Prime_{PD} → Target_{PD}



Prime_{DO} → Target_{PD}



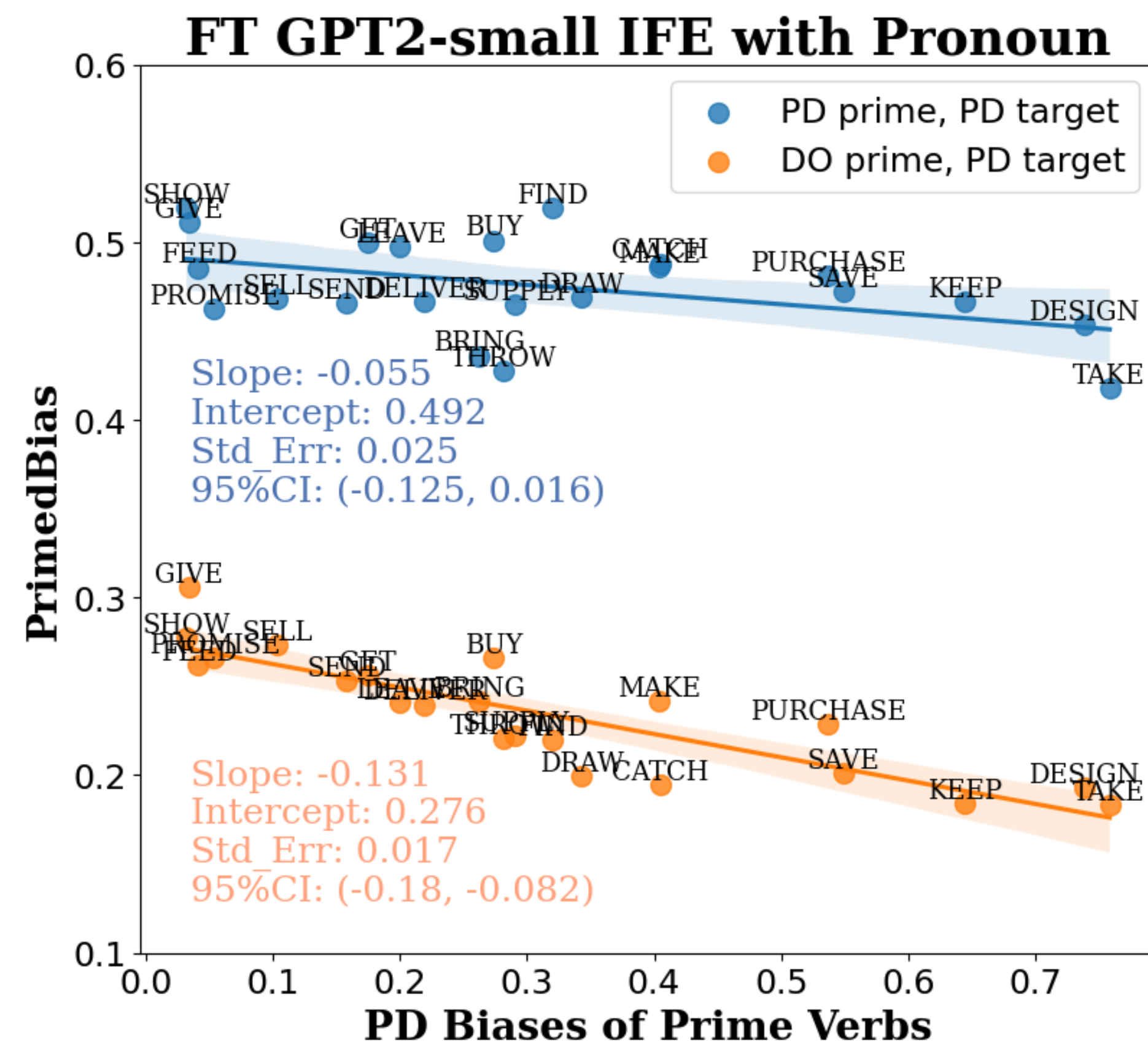
- **IFE:** double negative slopes;
- **Standard Priming:** PD-PD has higher intercept than DO-PD;

Priming in LLMs

Fine-tuning Mode: fine-tuning the model with the prime sentence and use the updated model to run the target sentence — **with weight update**;

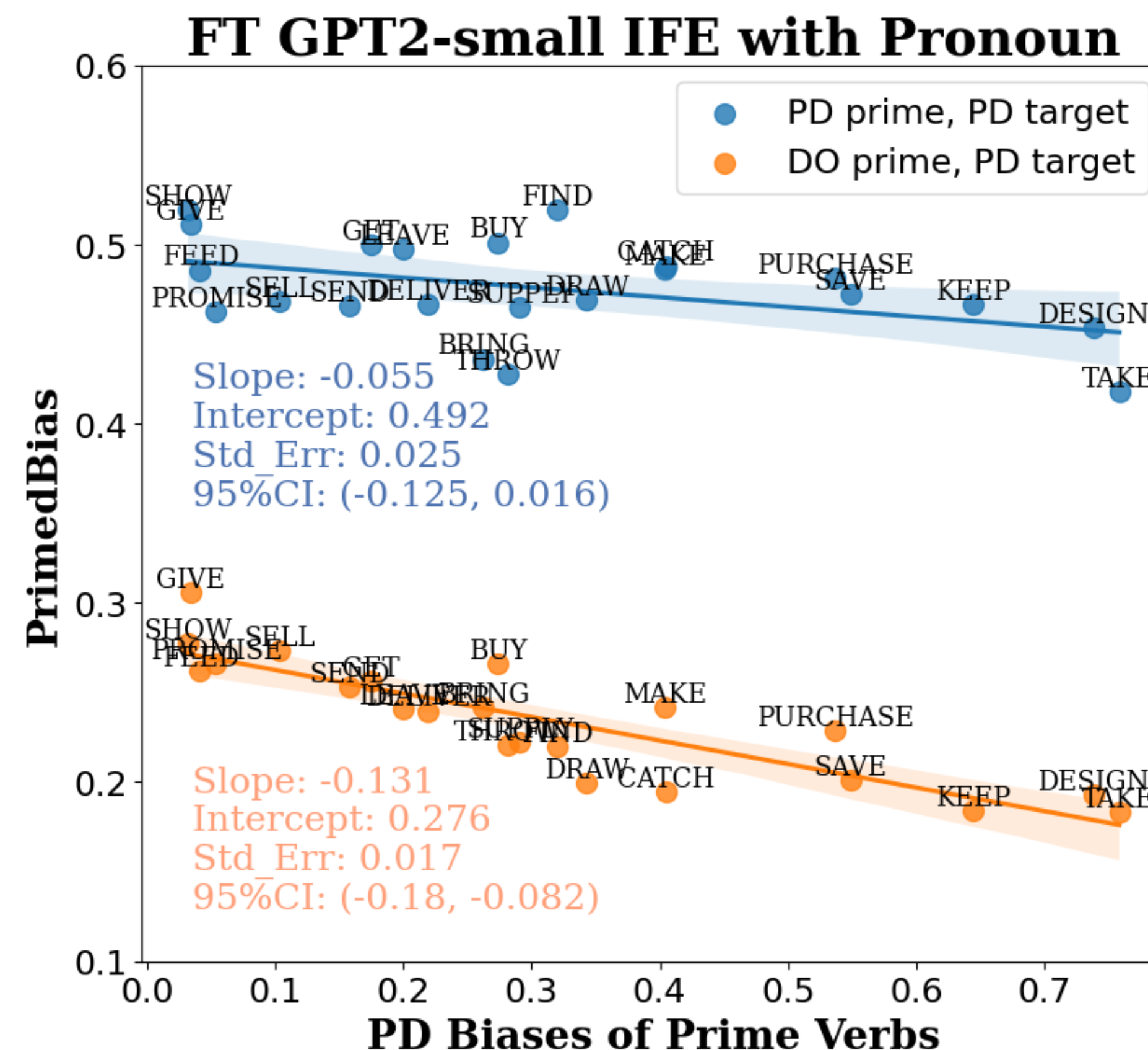
Priming in LLMs

Fine-tuning Mode: fine-tuning the model with the prime sentence and use the updated model to run the target sentence — **with weight update**;



Priming in LLMs

Fine-tuning Mode: fine-tuning the model with the prime sentence and use the updated model to run the target sentence — **with weight update**;



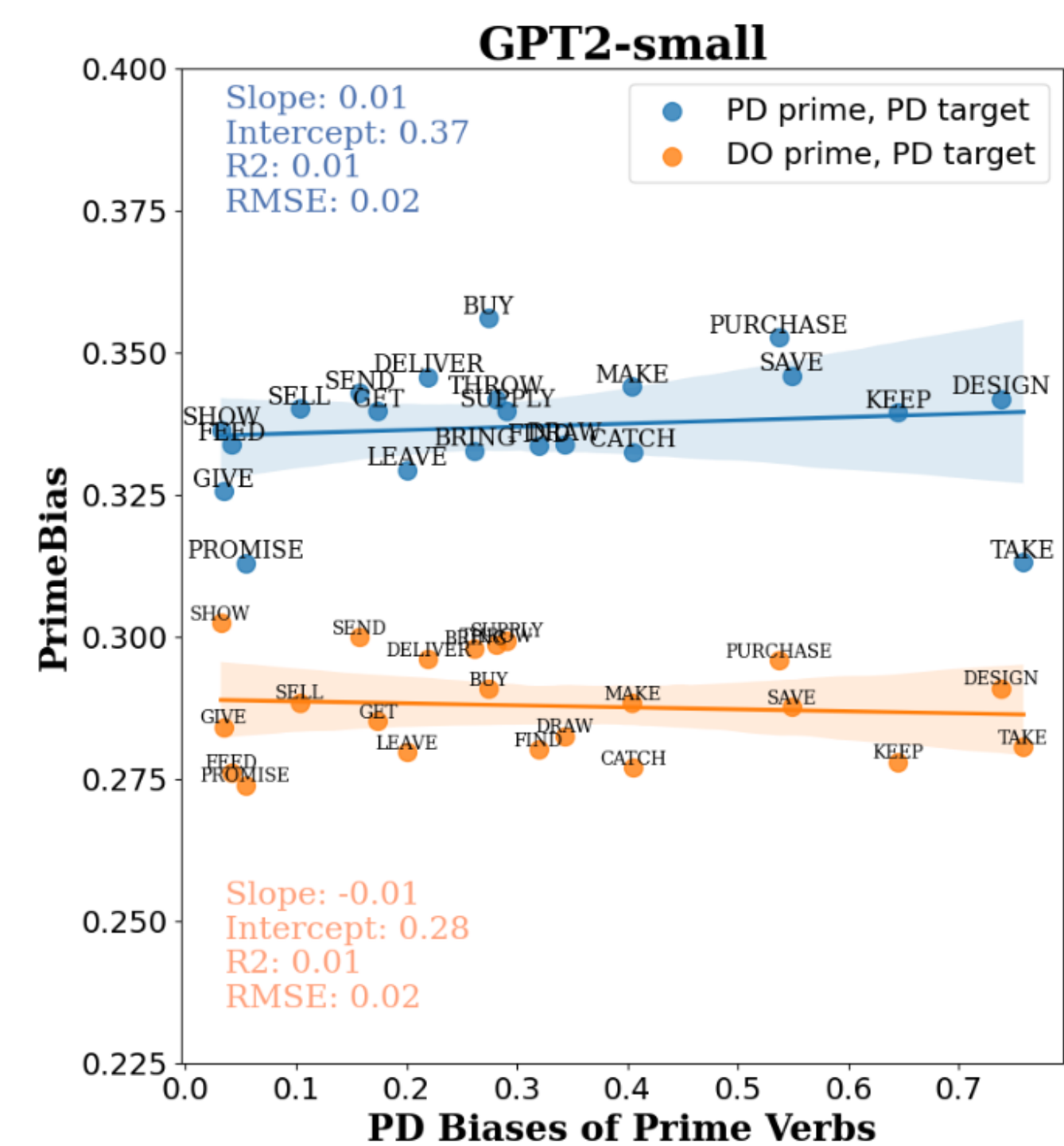
Even *GPT2-small* shows significant inverse frequency effects!

Priming in LLMs

Concatenation Mode: concatenating the prime and target sentences and run the model — **without weight update**;

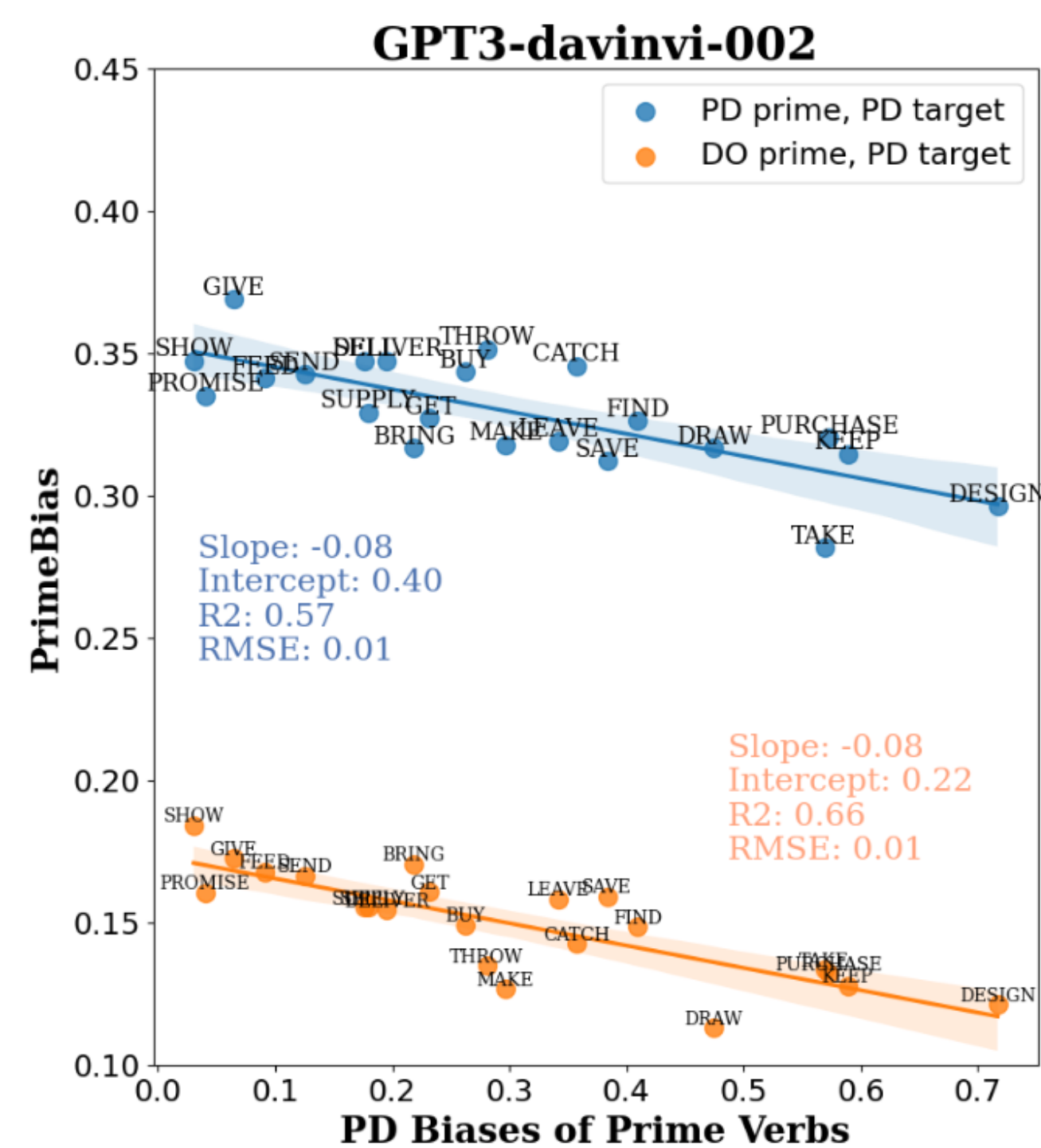
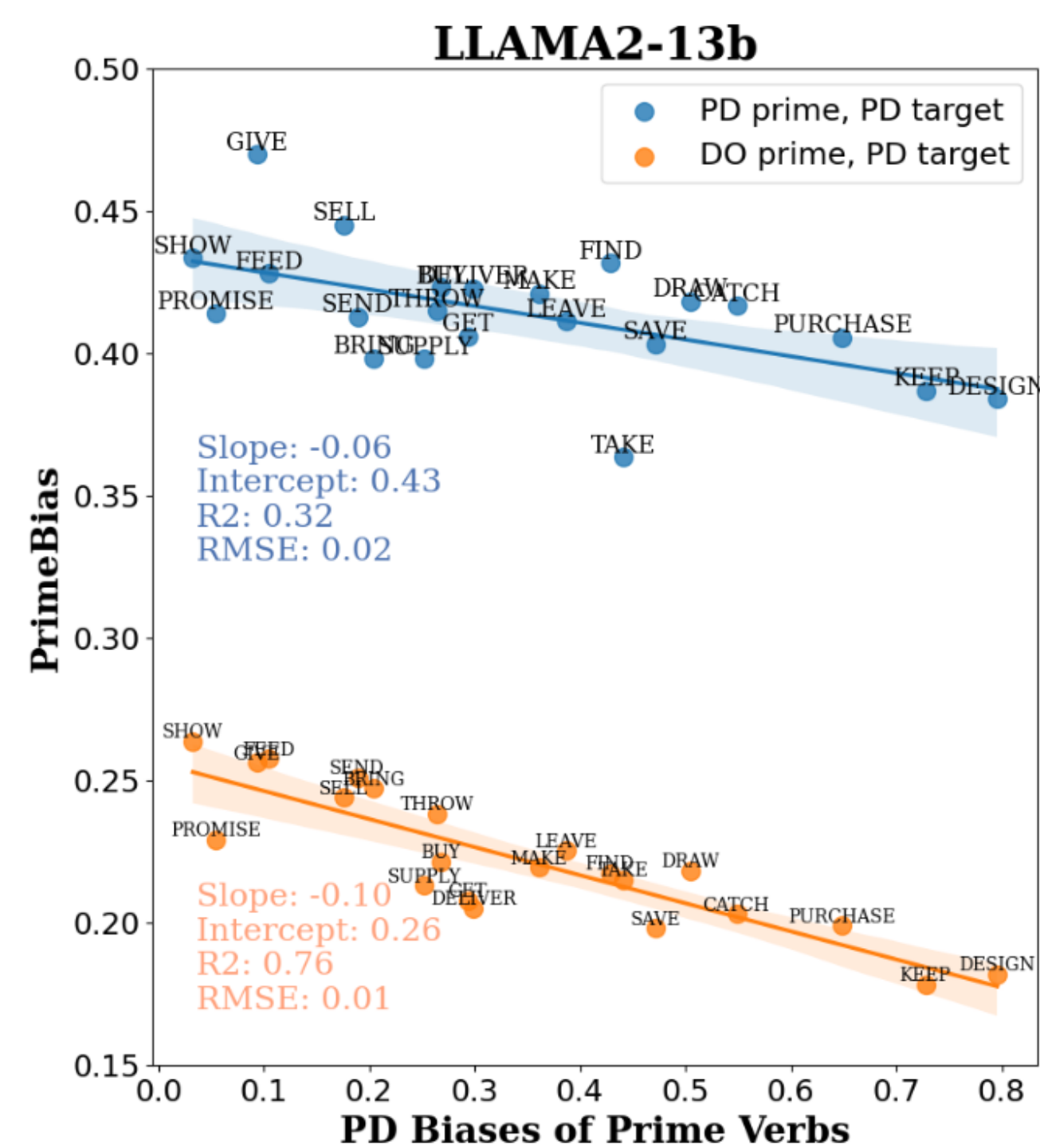
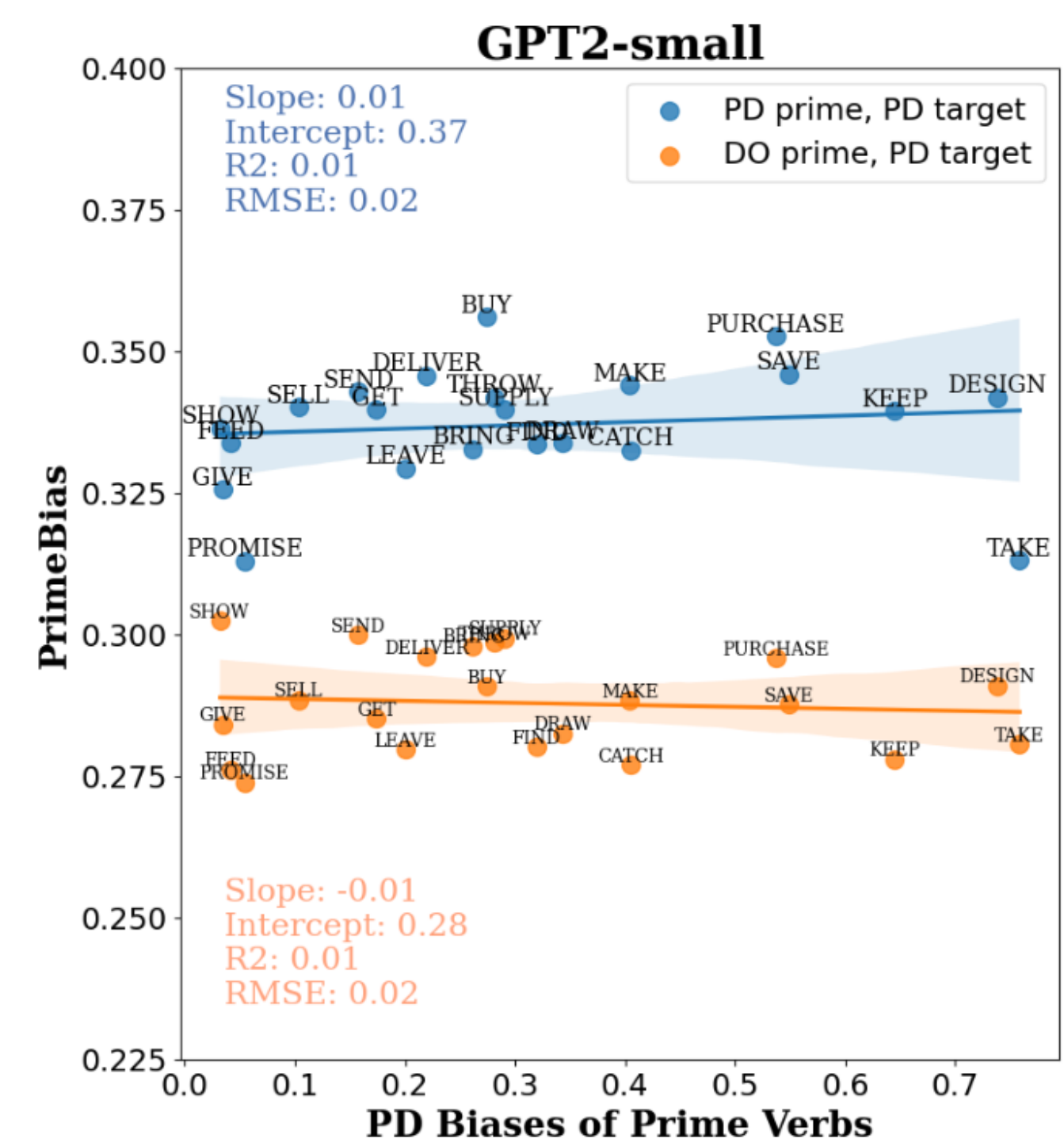
Priming in LLMs

Concatenation Mode: concatenating the prime and target sentences and run the model — **without weight update**;



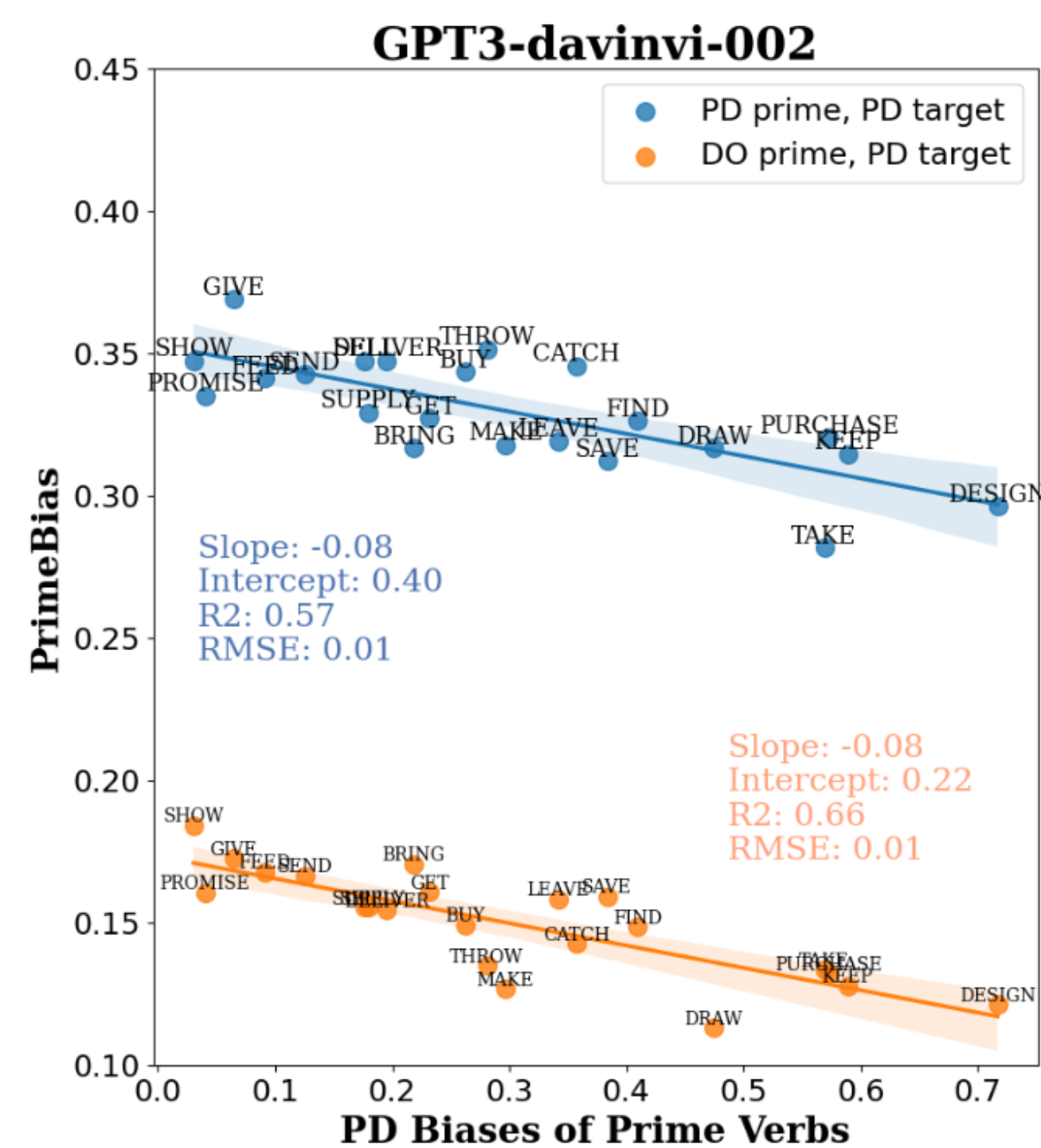
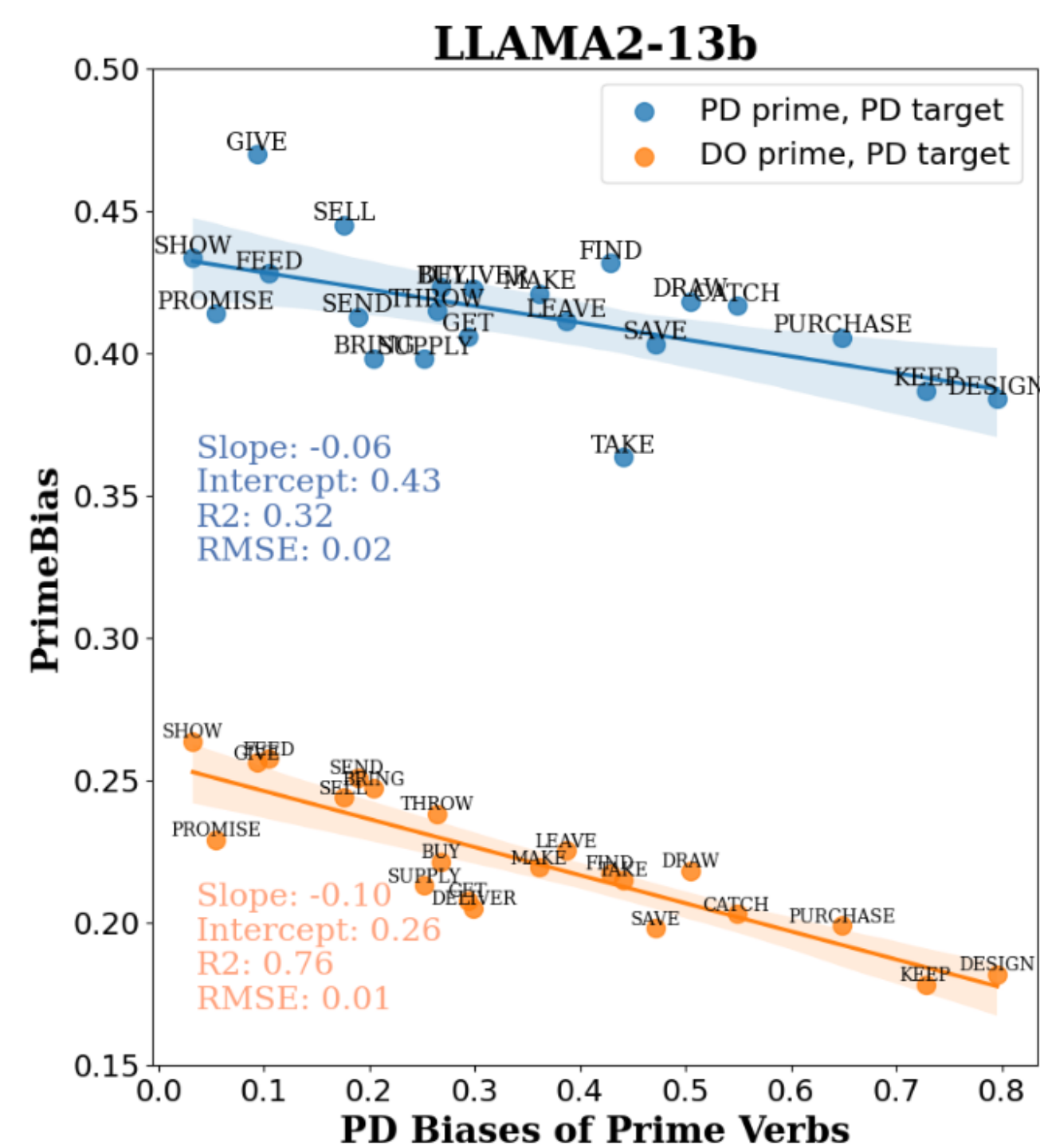
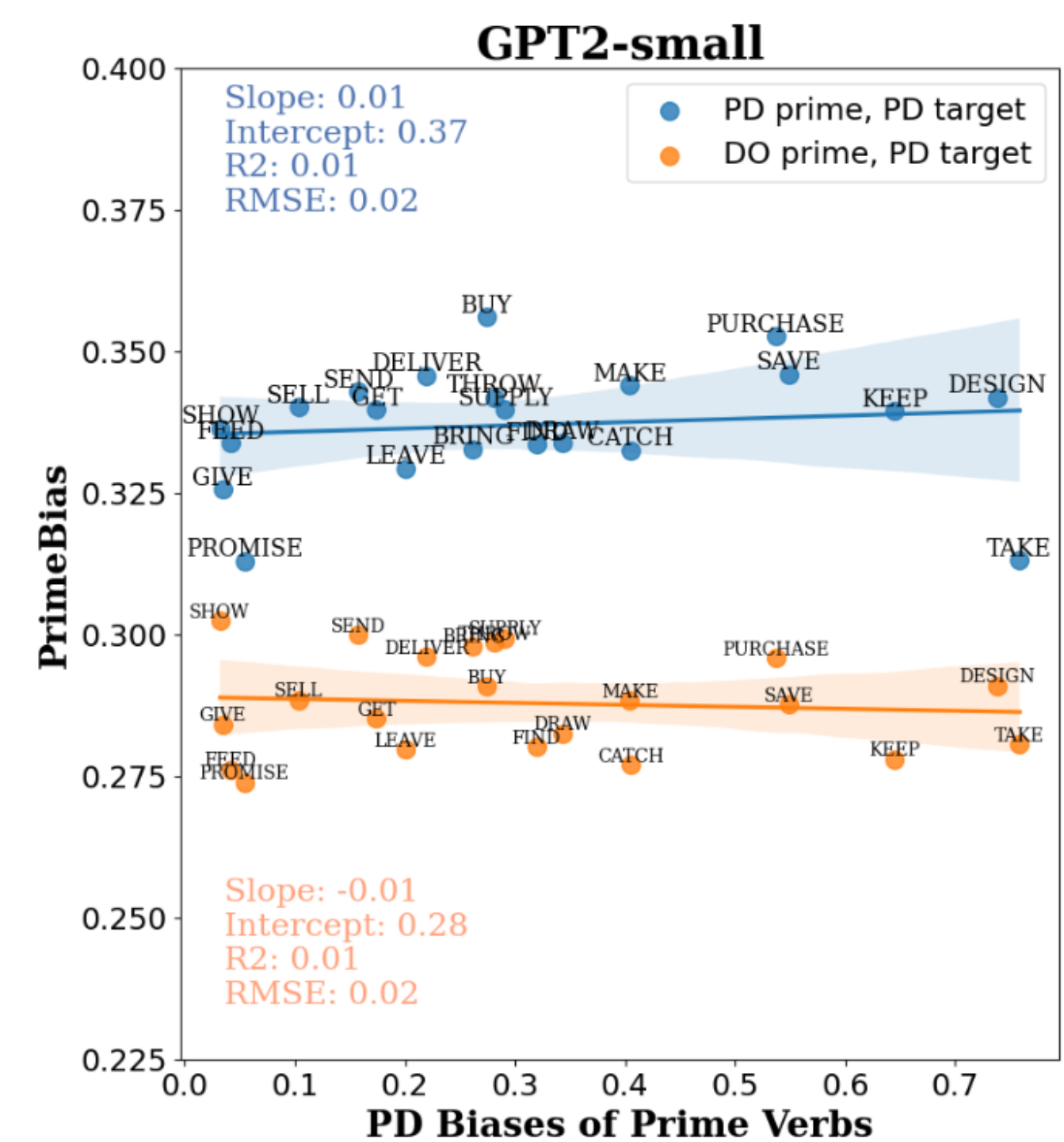
Priming in LLMs

Concatenation Mode: concatenating the prime and target sentences and run the model — **without weight update**;



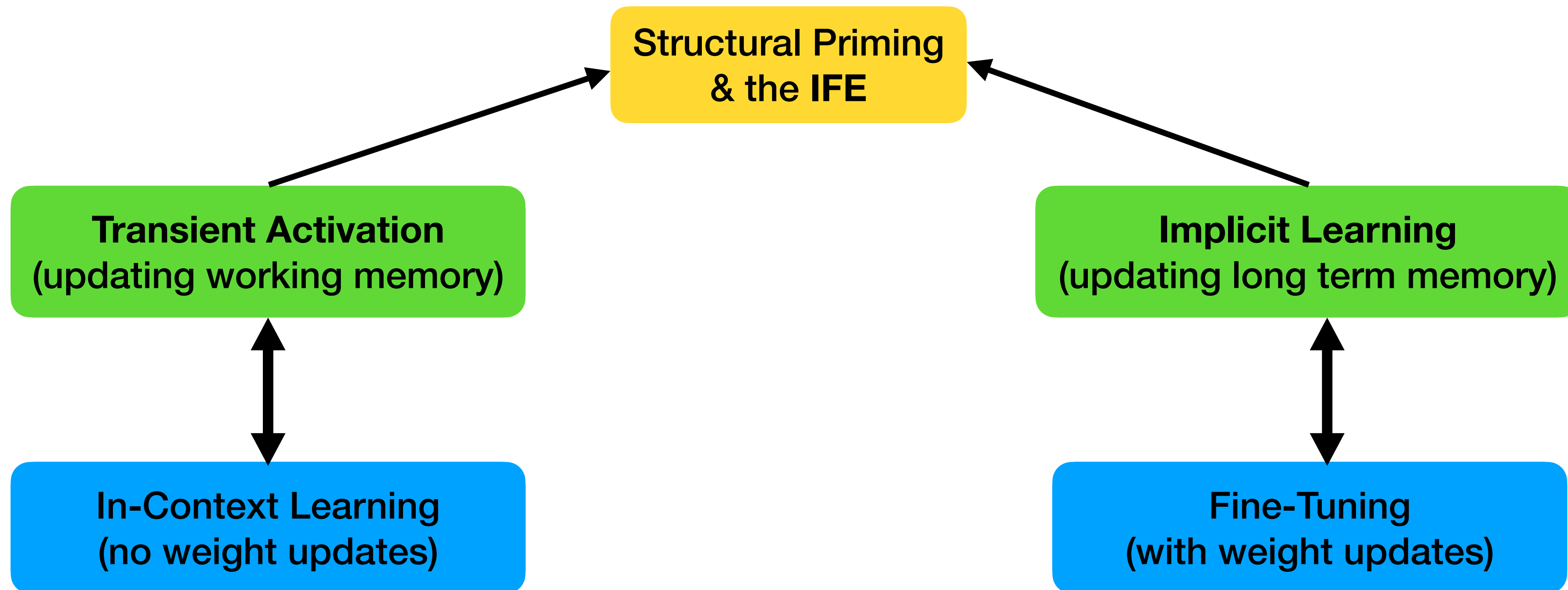
Priming in LLMs

Concatenation Mode: concatenating the prime and target sentences and run the model — **without weight update;**

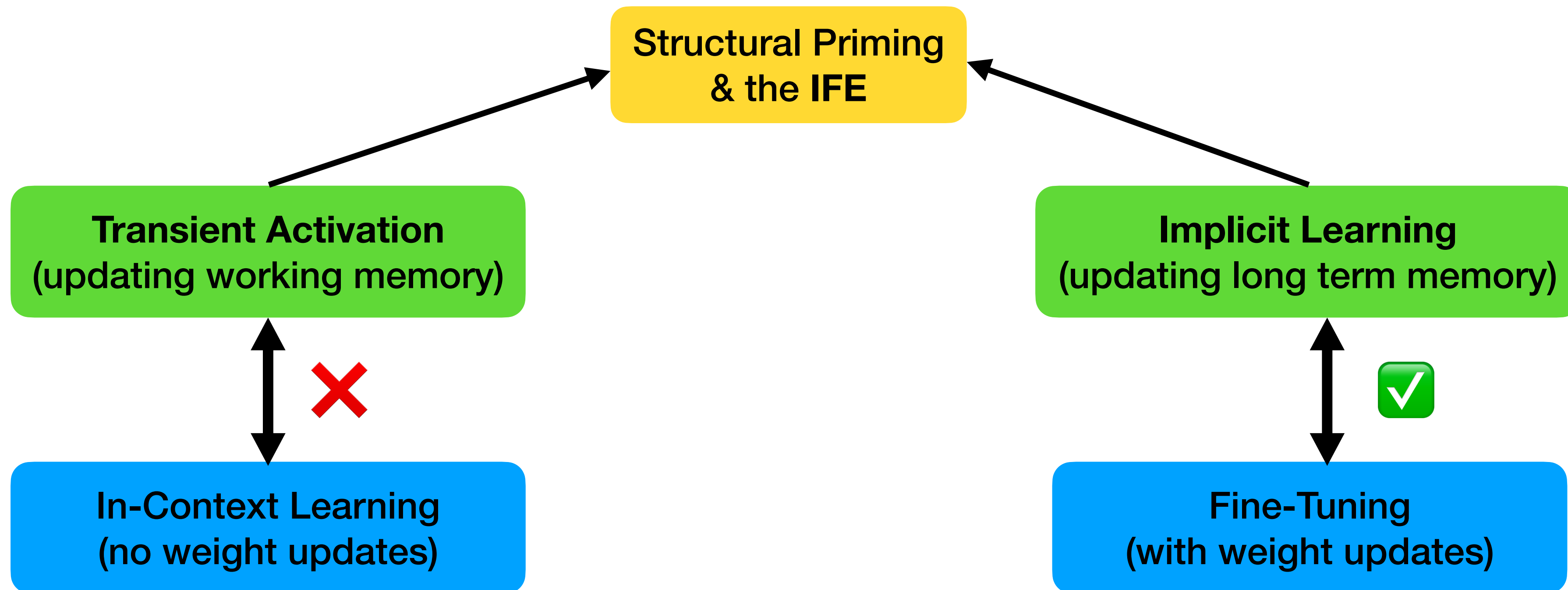


Larger models show more significant inverse frequency effects!

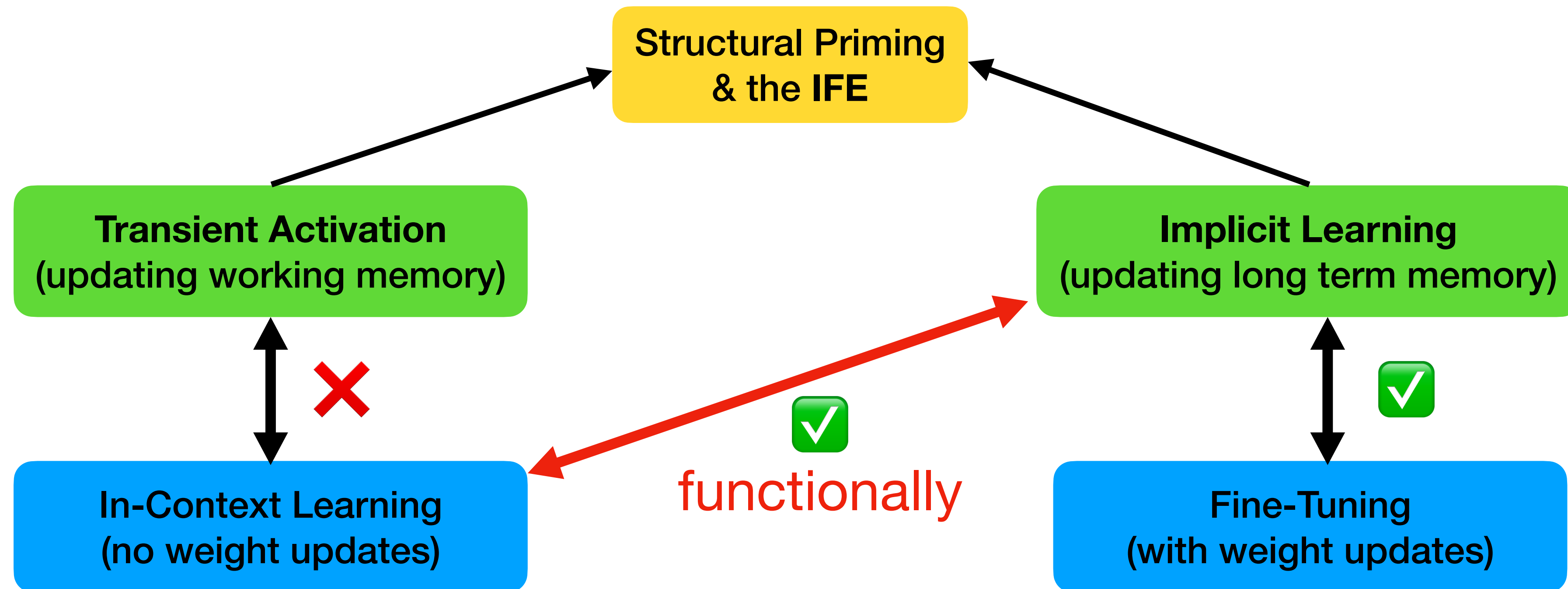
Takeaways from inverse frequency



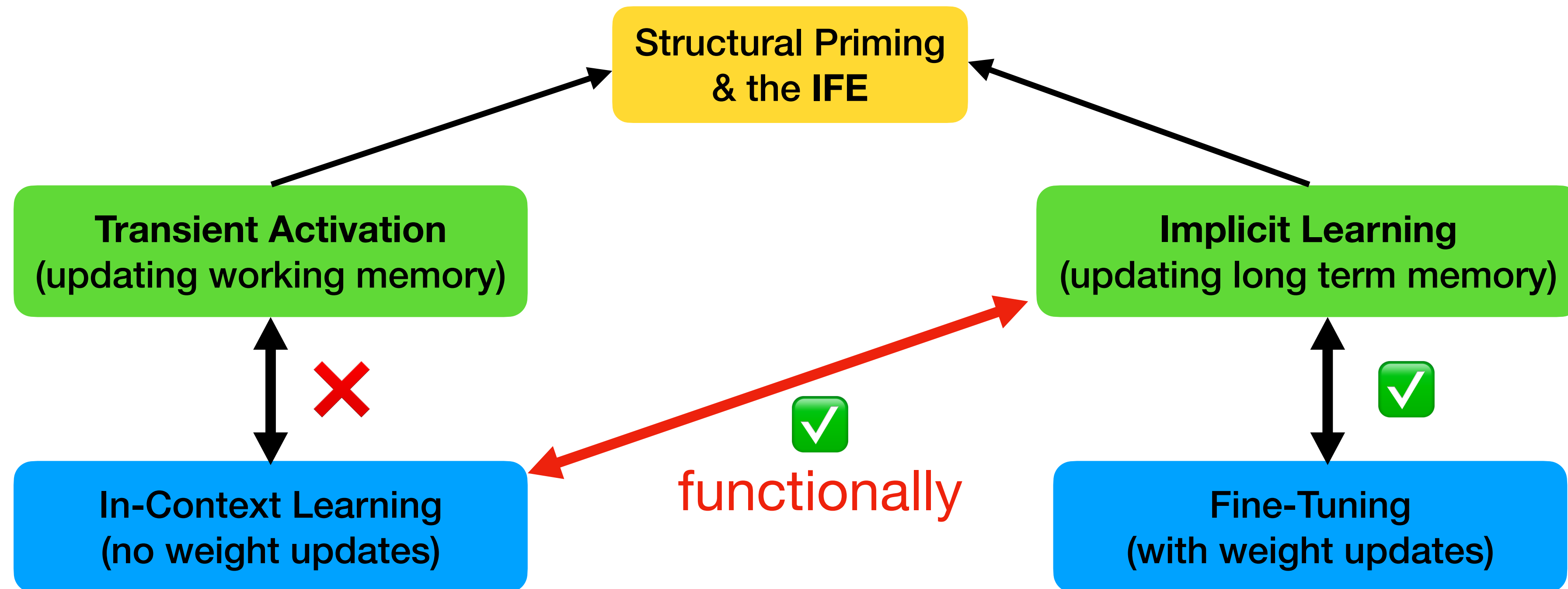
Takeaways from inverse frequency



Takeaways from inverse frequency



Takeaways from inverse frequency



We used the IFE as a diagnostic on the error-driven nature of ICL as a processing / adaptation mechanism in LLMs.

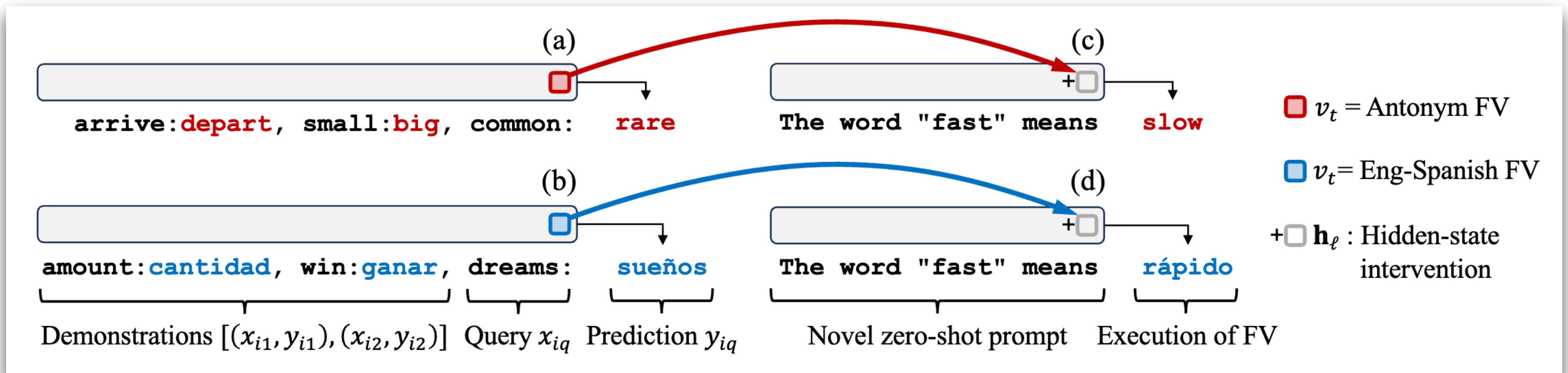
Experiment 2: Mechanistic Analysis

Research Questions:

- How to further characterize the generalized “ICL $\approx$ Priming” claim?
- How could we causally verify the involvement of verb bias?

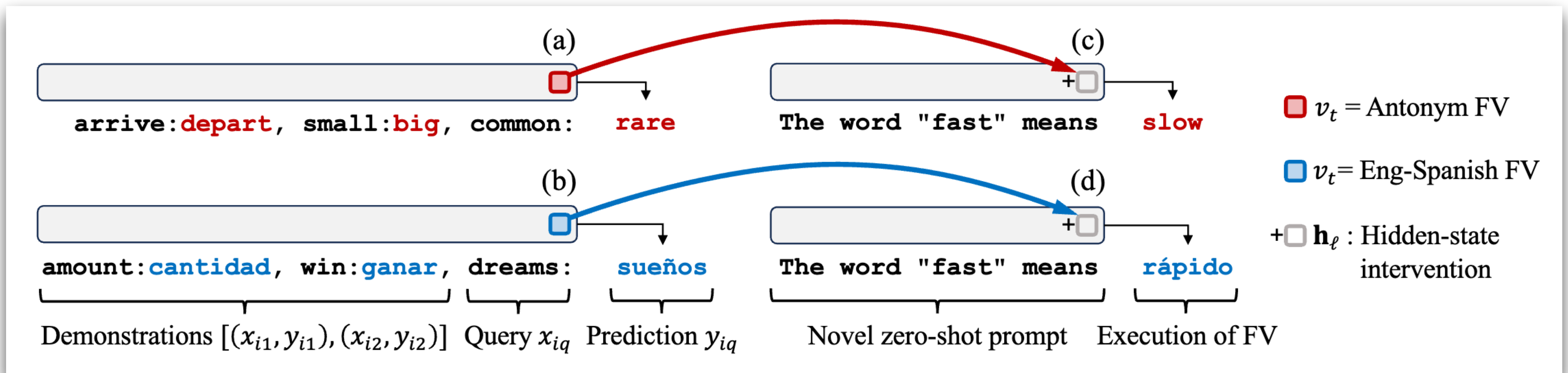
ICL leads to Function Vectors

Function vectors: compressed vector representation of the task in context;



ICL leads to Function Vectors

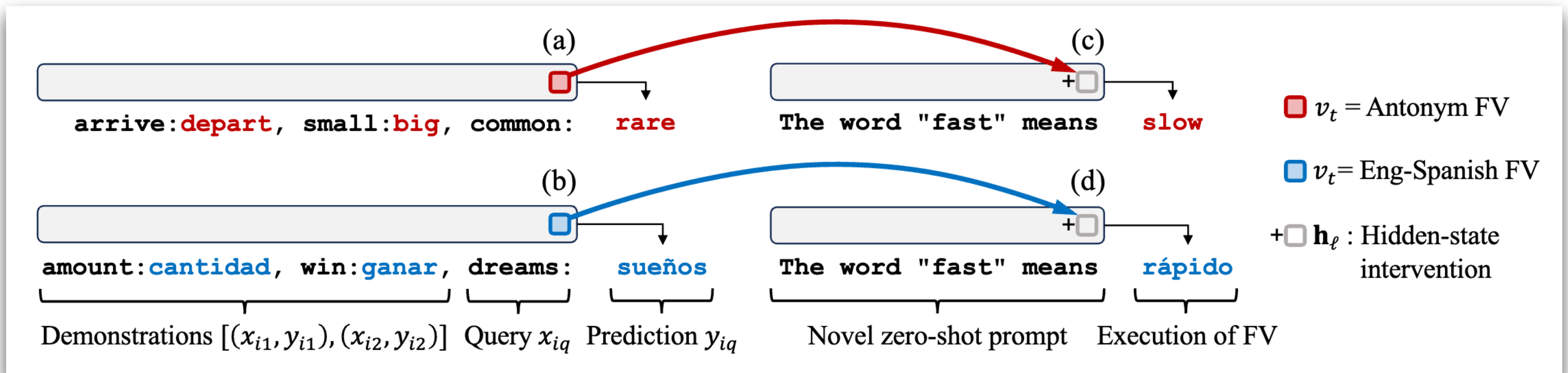
Function vectors: compressed vector representation of the task in context;



Extract mean activations
at the query position
after context

ICL leads to Function Vectors

Function vectors: compressed vector representation of the task in context;



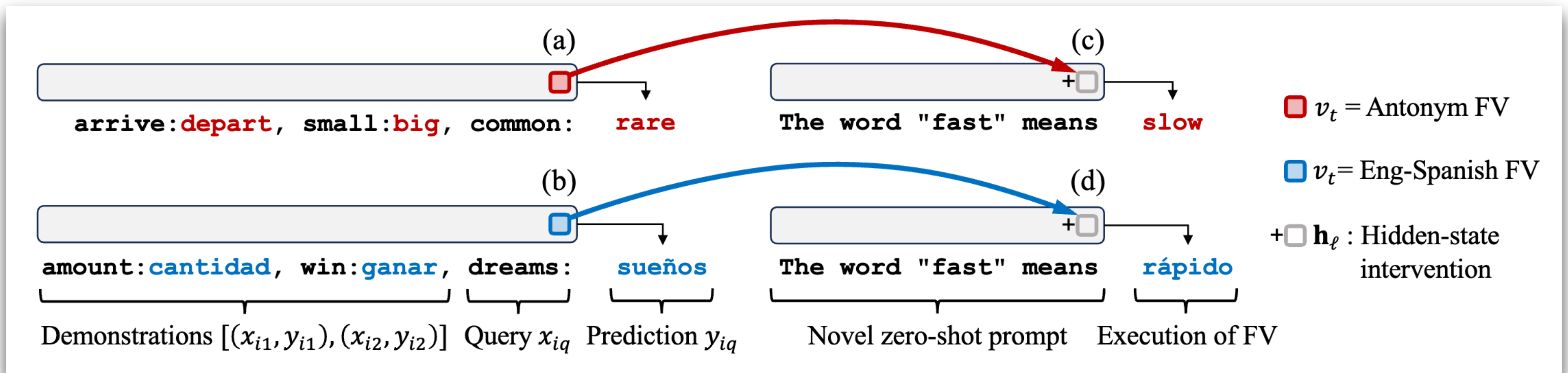
Extract mean activations
at the query position
after context



Injecting this vector
in a new inference
run (zero-shot)

ICL leads to Function Vectors

Function vectors: compressed vector representation of the task in context;



Extract mean activations
at the query position
after context

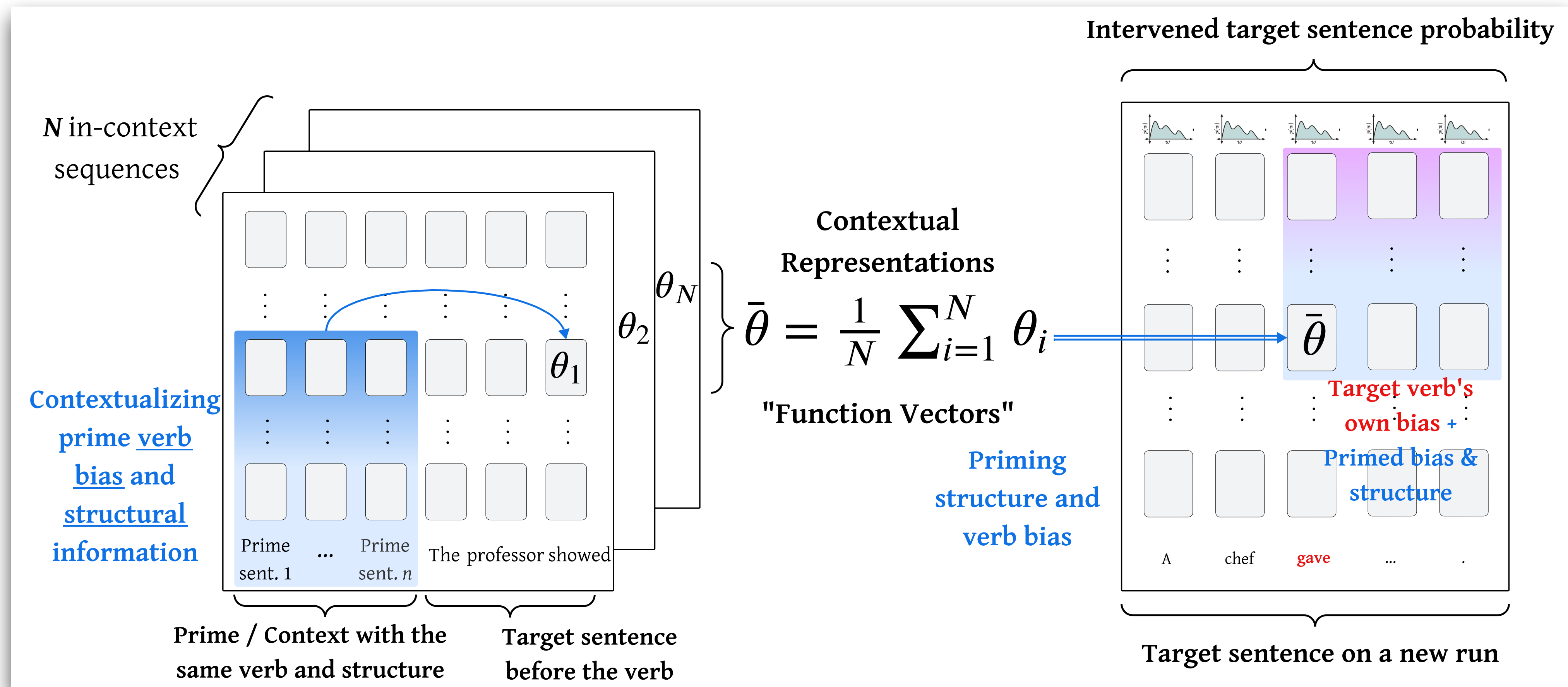


Injecting this vector
in a new inference
run (zero-shot)



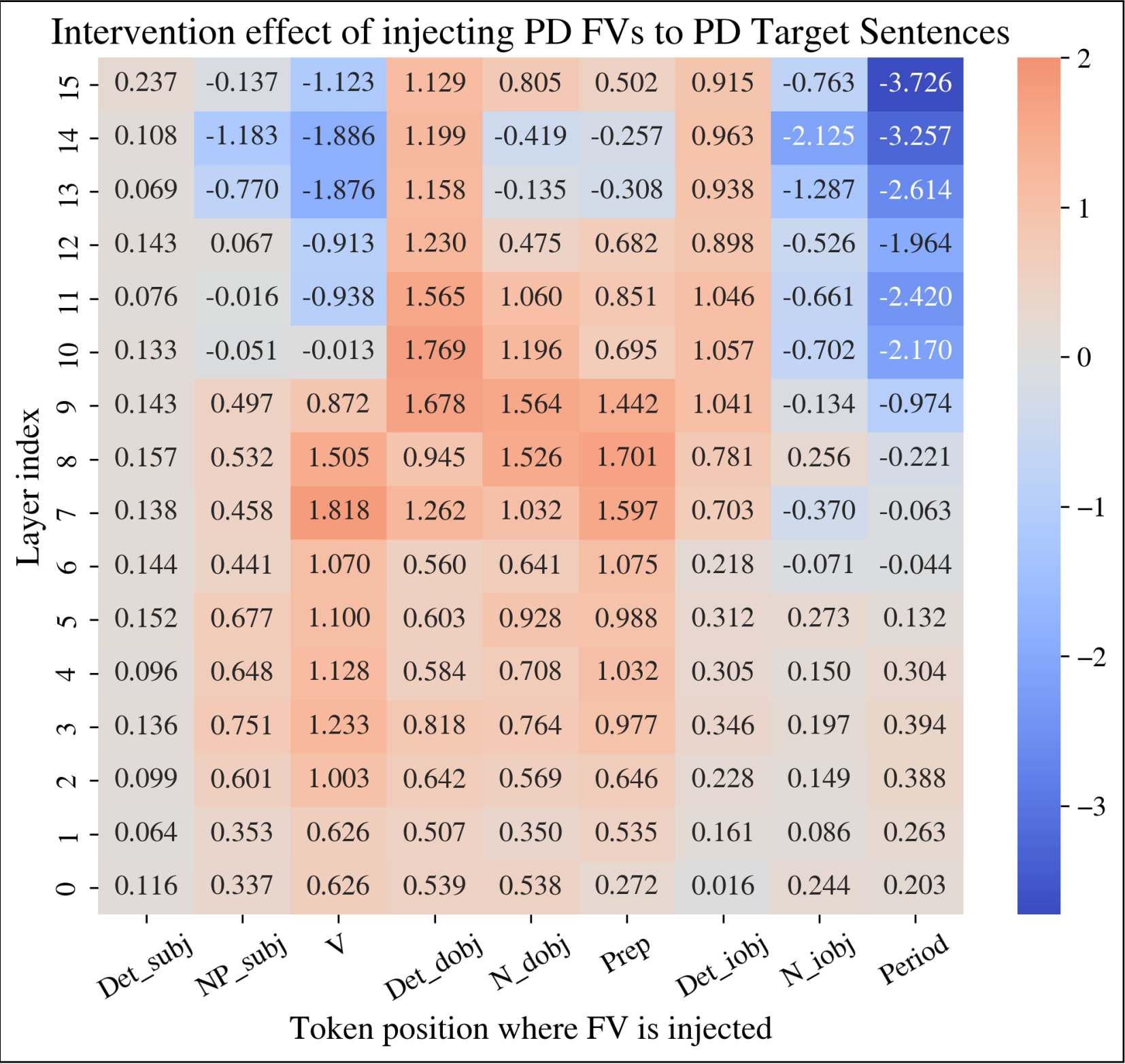
Steering model
production
towards the task

Structural Priming with Function Vectors

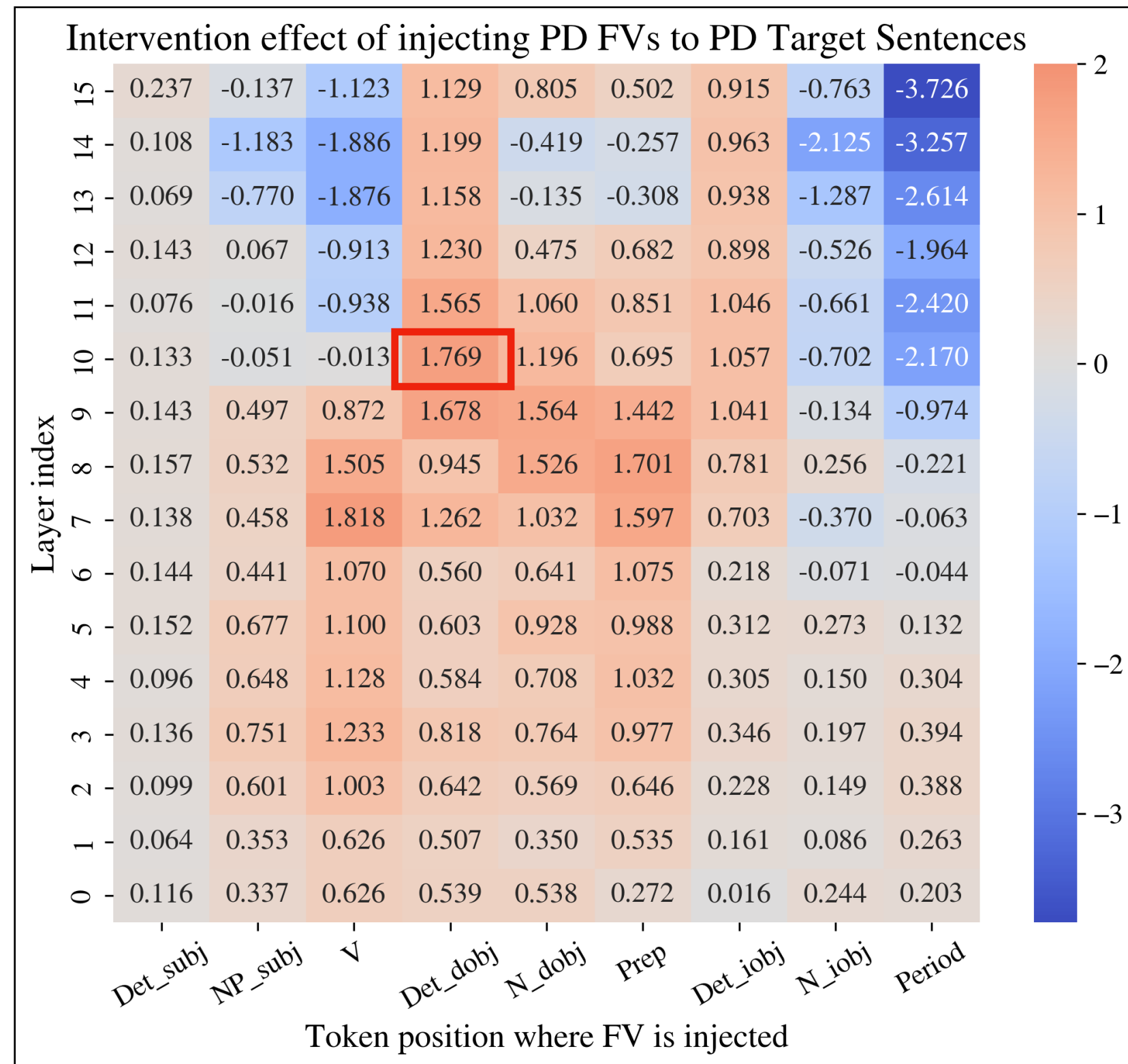


* we also did a version where function vectors are the mean of the most influential attention heads' output activations (Todd et al. 2024), leading to similar results.

Structural Priming with Function Vectors

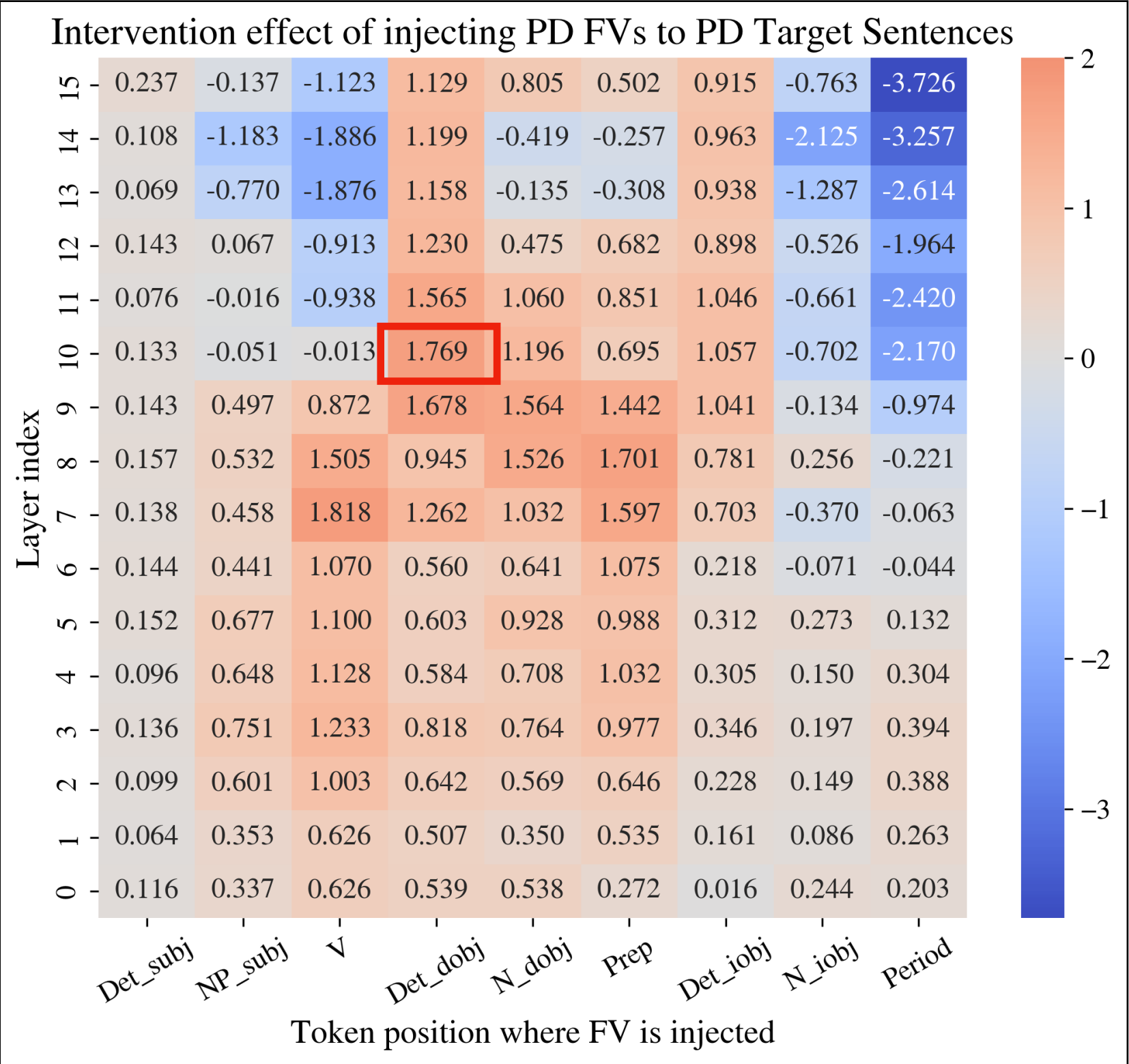


Structural Priming with Function Vectors

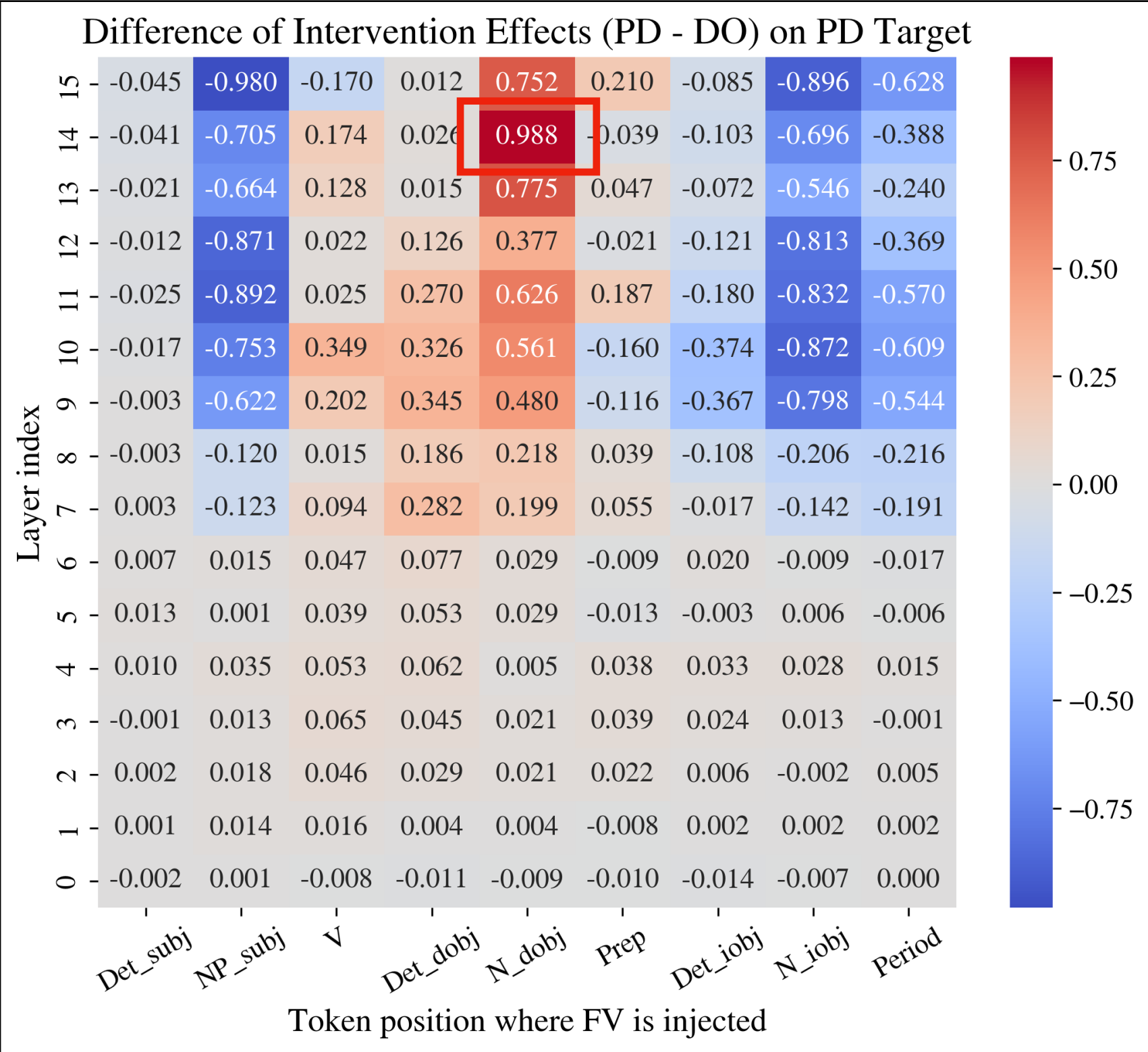


- Observing the priming effect: FVs increased target sentence probability;

Structural Priming with Function Vectors

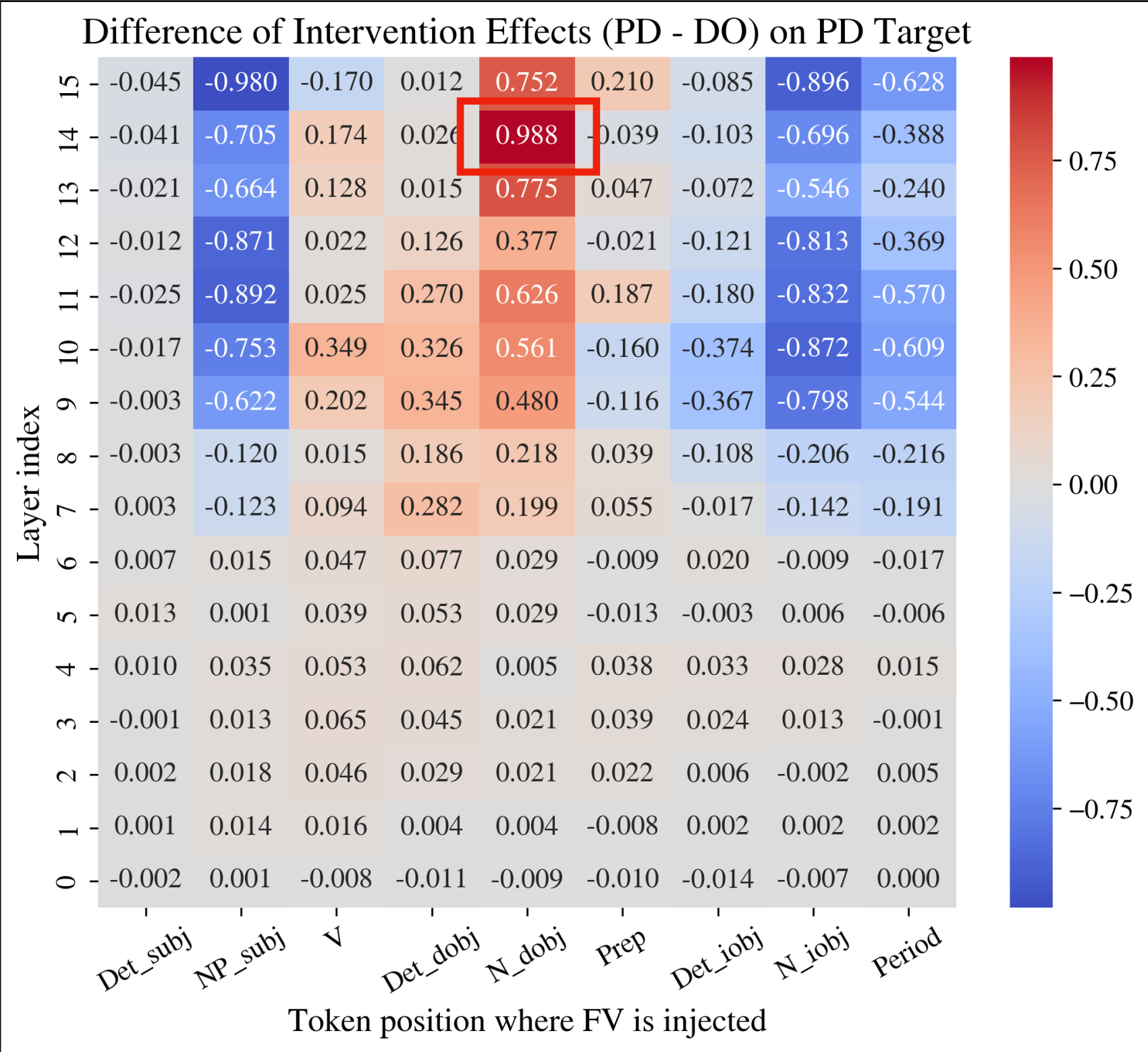
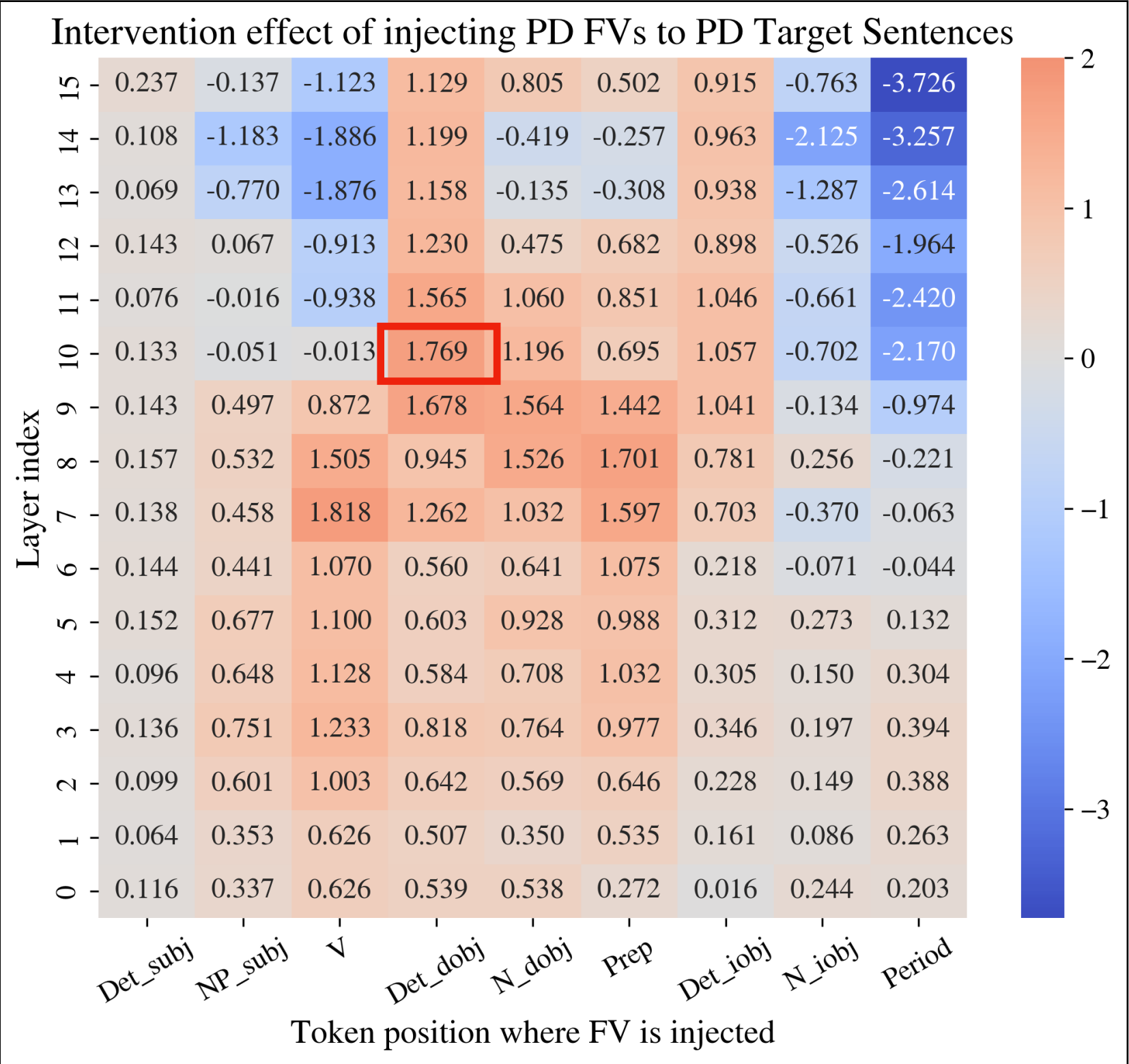


- Observing the priming effect: FVs increased target sentence probability;



- PD FVs cause a larger intervention effect than DO FVs on PD targets.

Structural Priming with Function Vectors

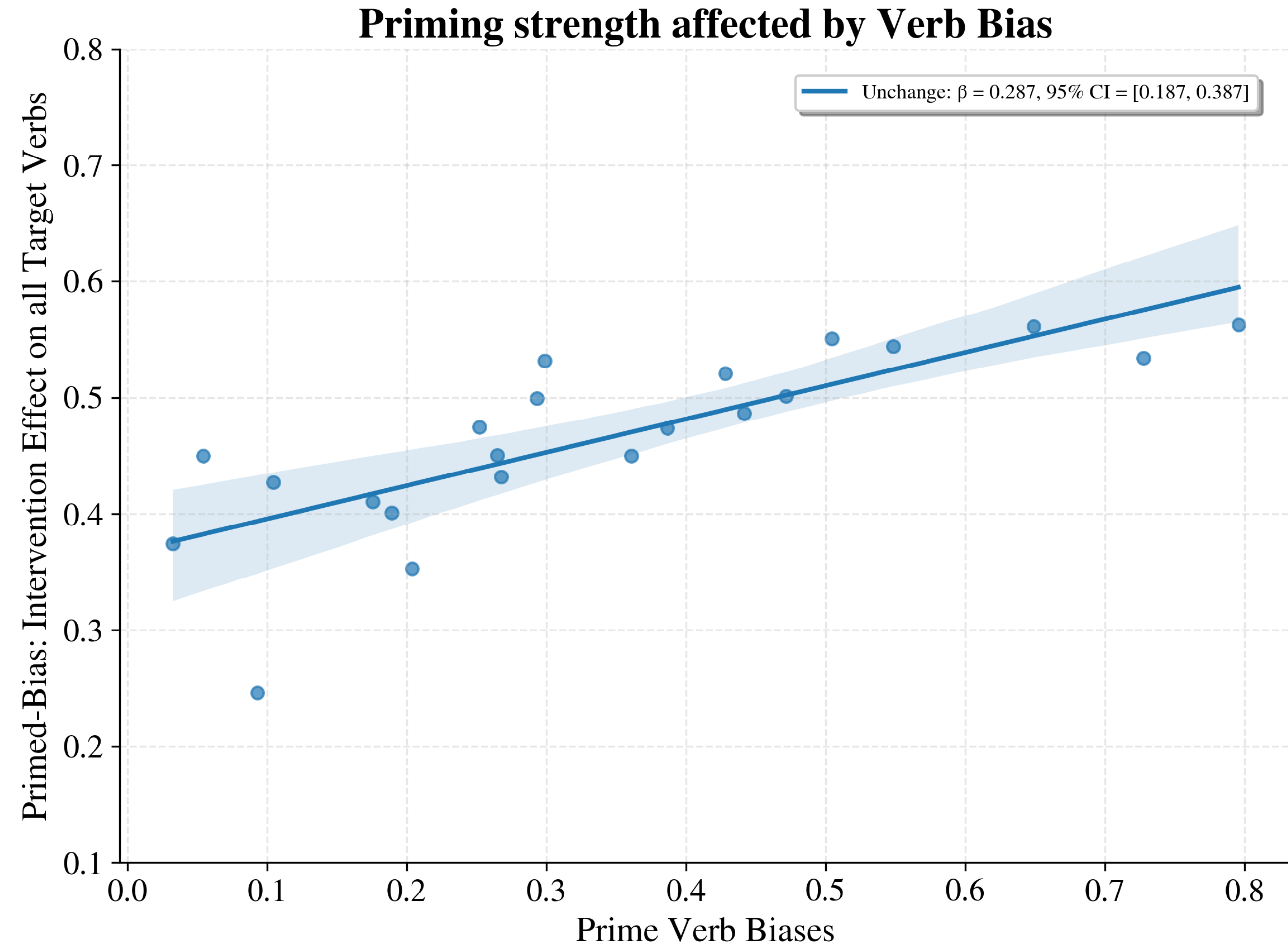


- Observing the priming effect: FVs increased target sentence probability;

- PD FVs cause a larger intervention effect than DO FVs on PD targets.

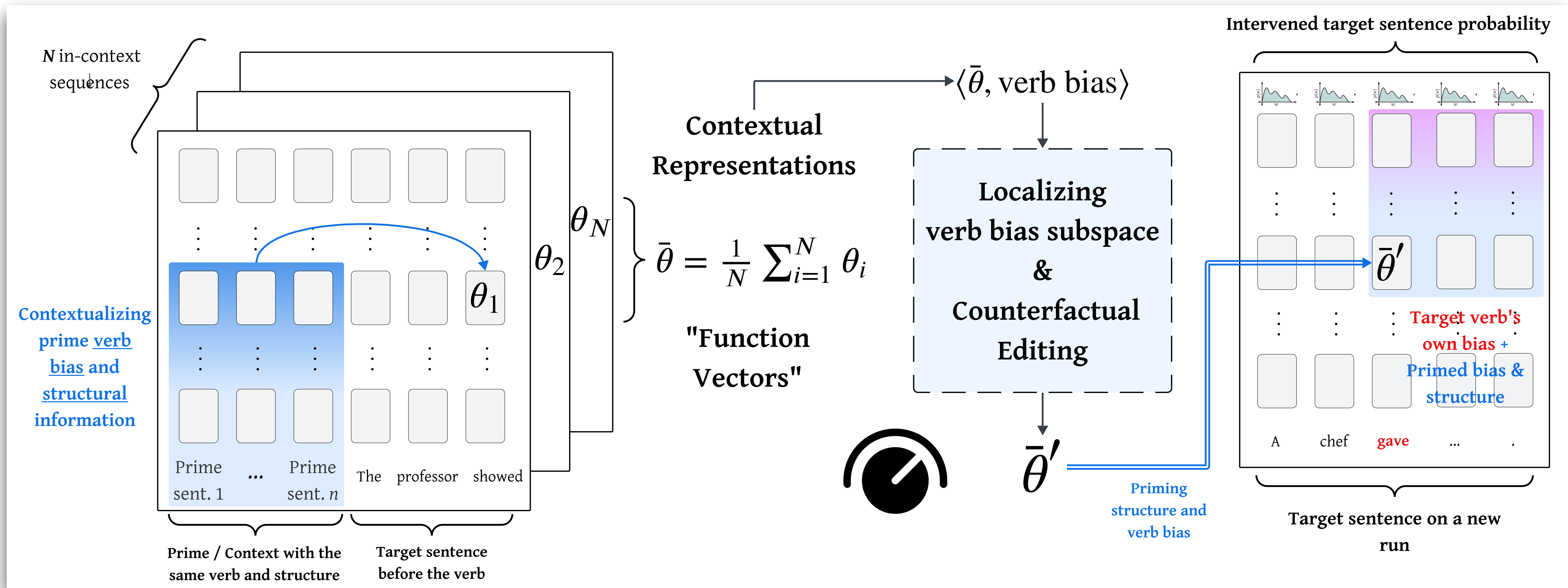
Yes, we do observe standard priming effect.

Structural Priming with Function Vectors

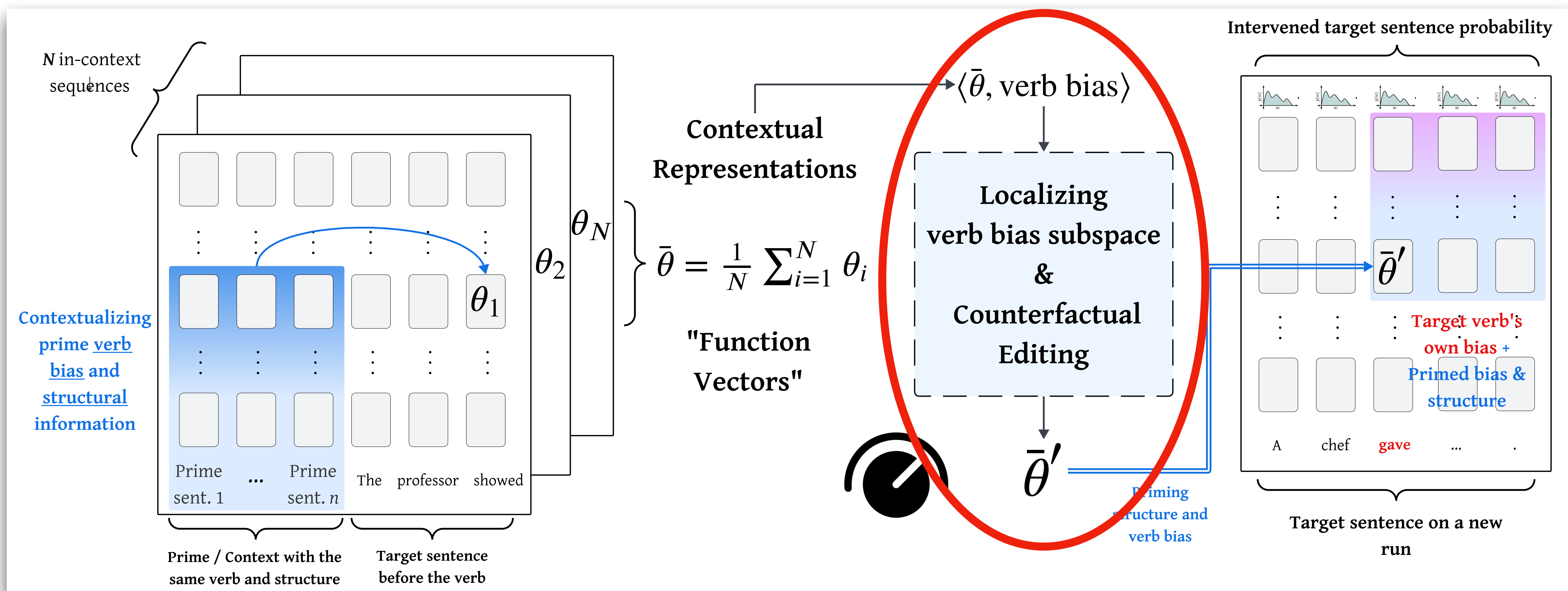


However, no IFE is observed...!

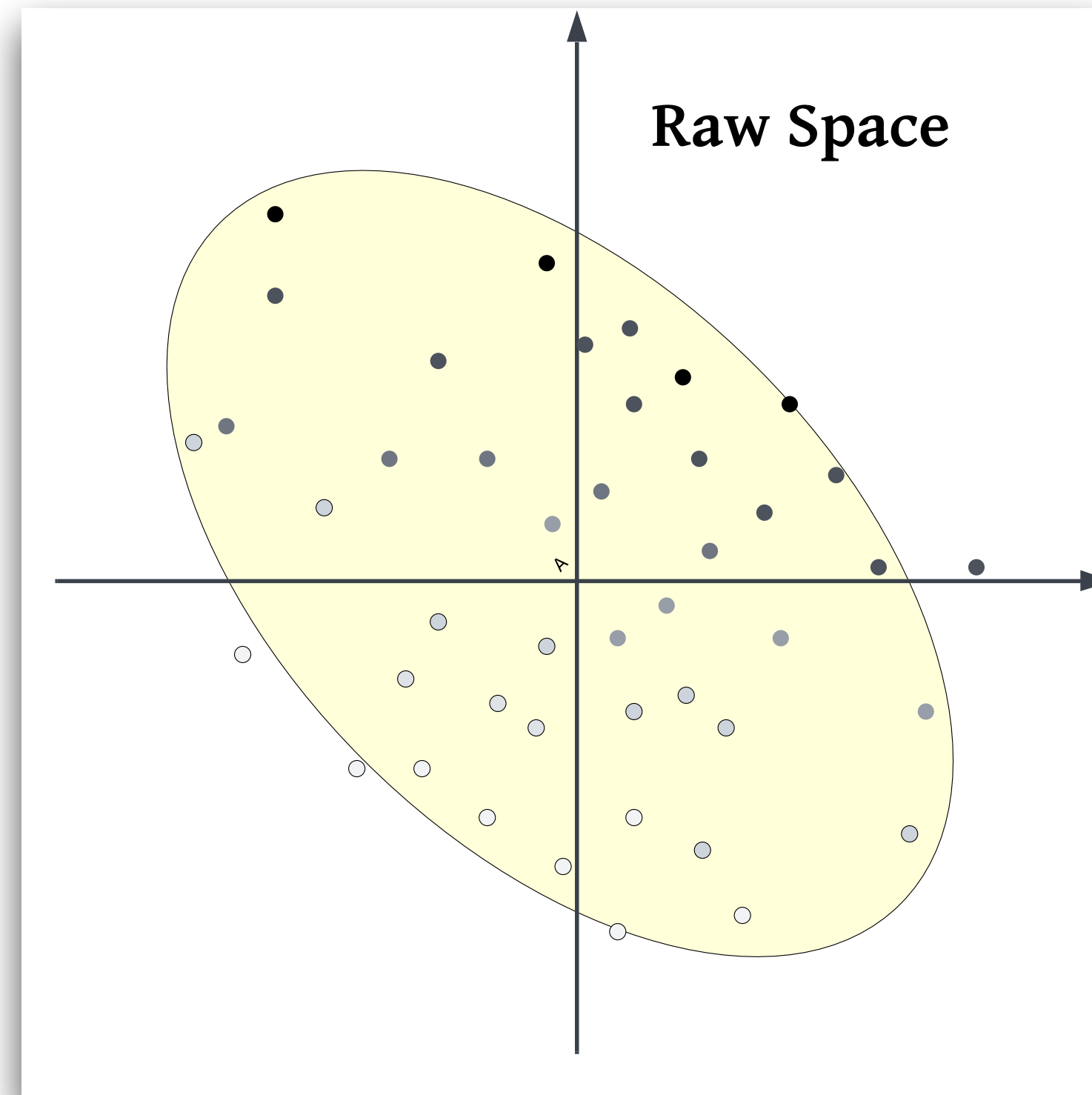
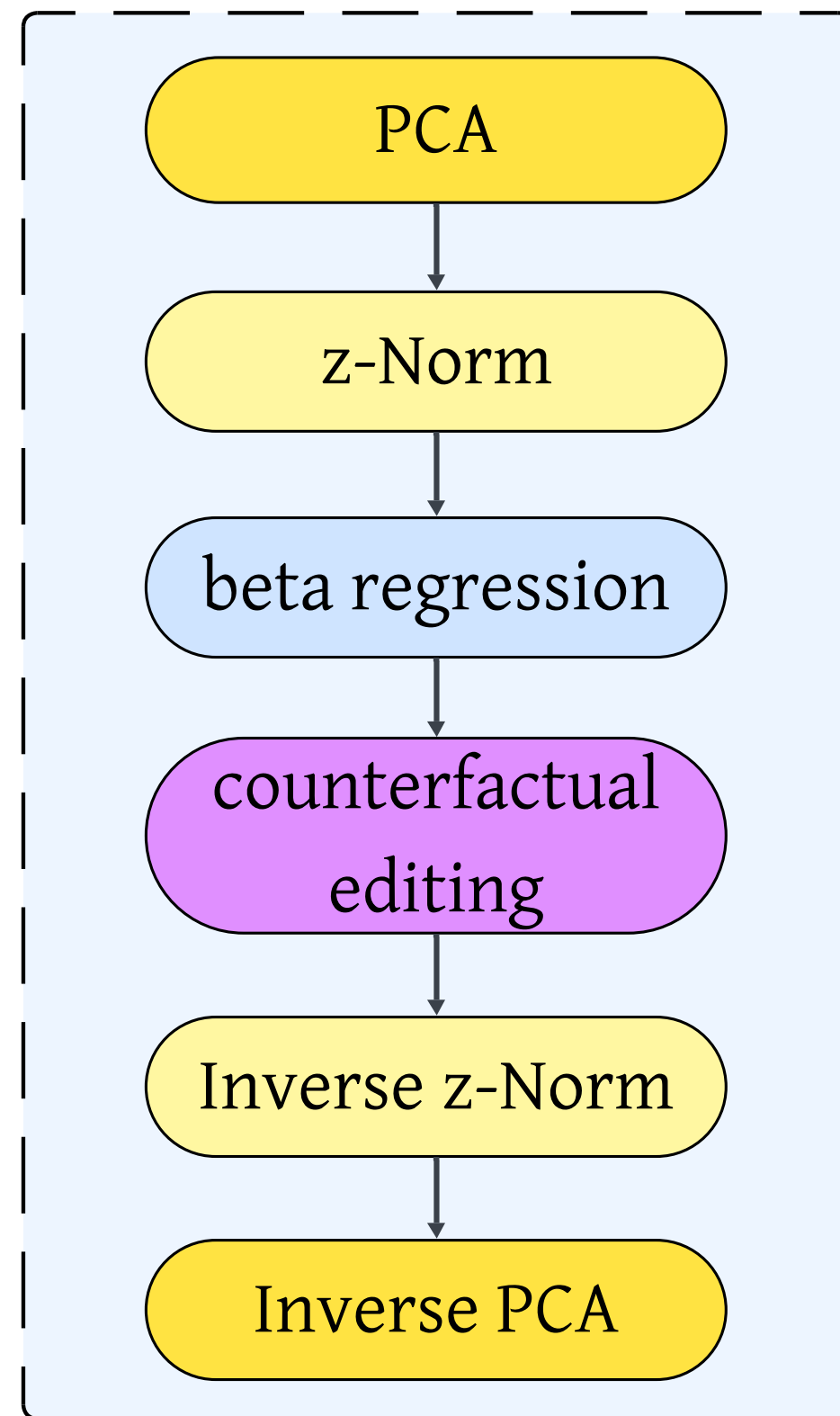
Causal Role of Verb Bias in Function Vectors?



Causal Role of Verb Bias in Function Vectors?

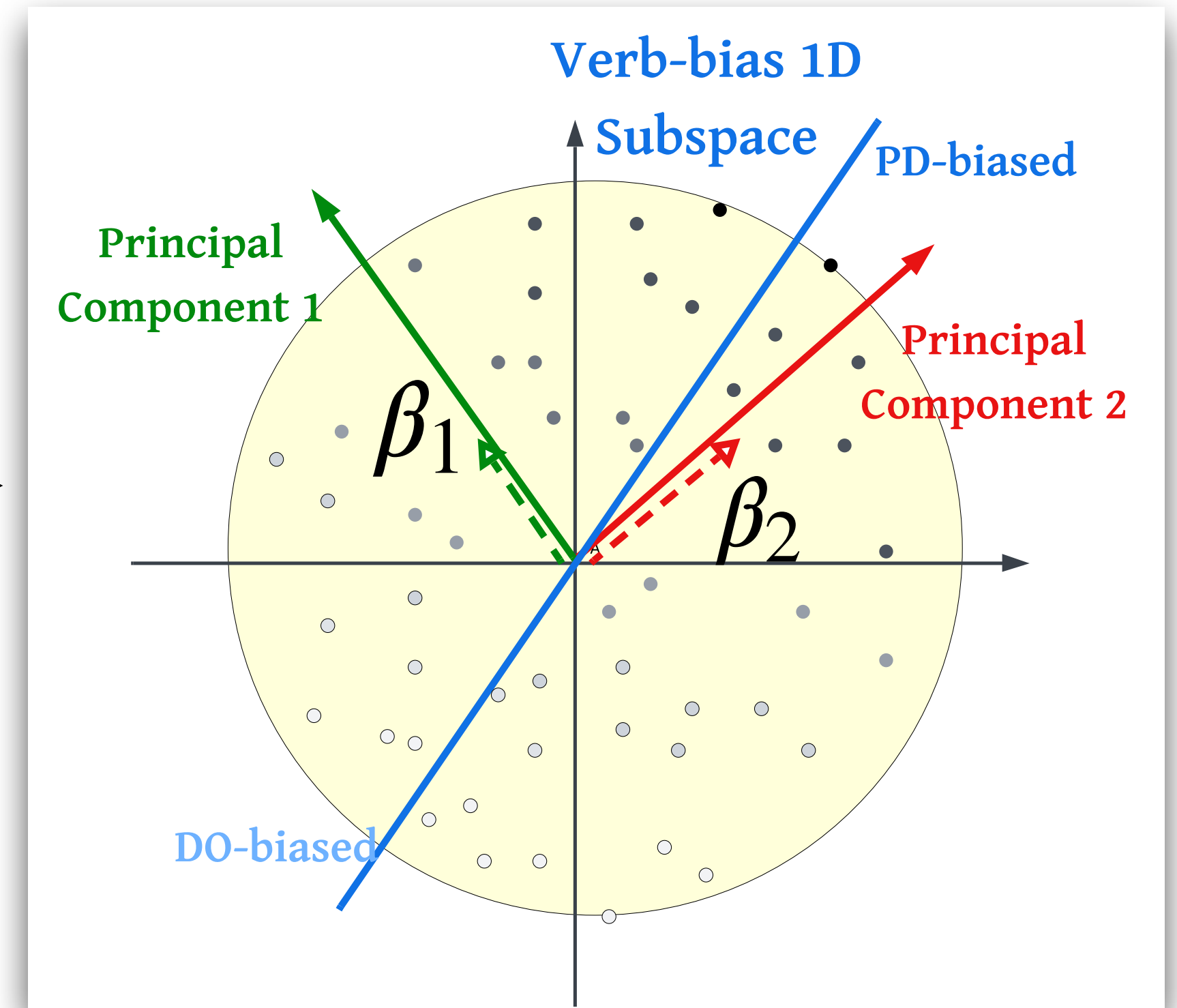
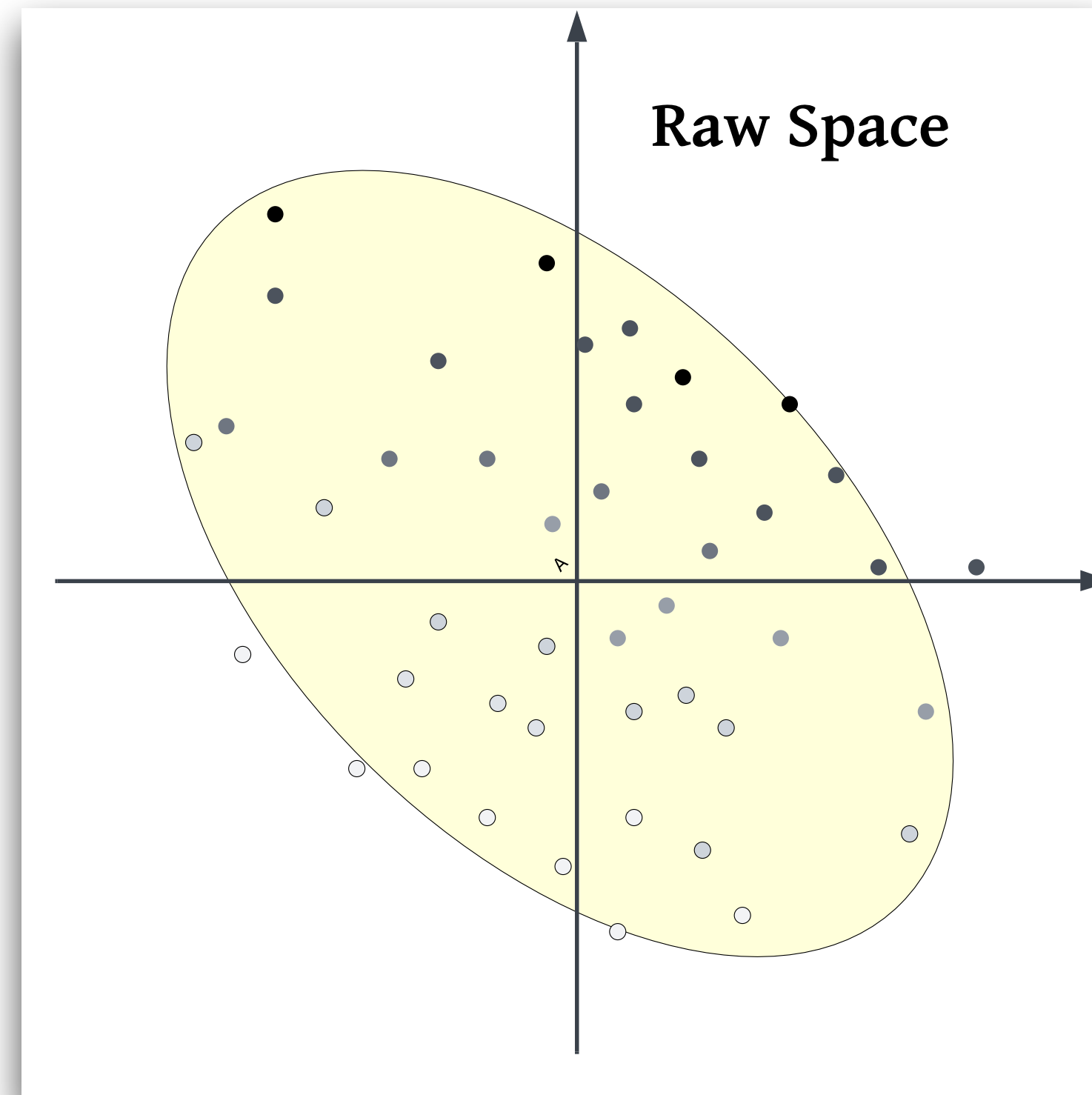
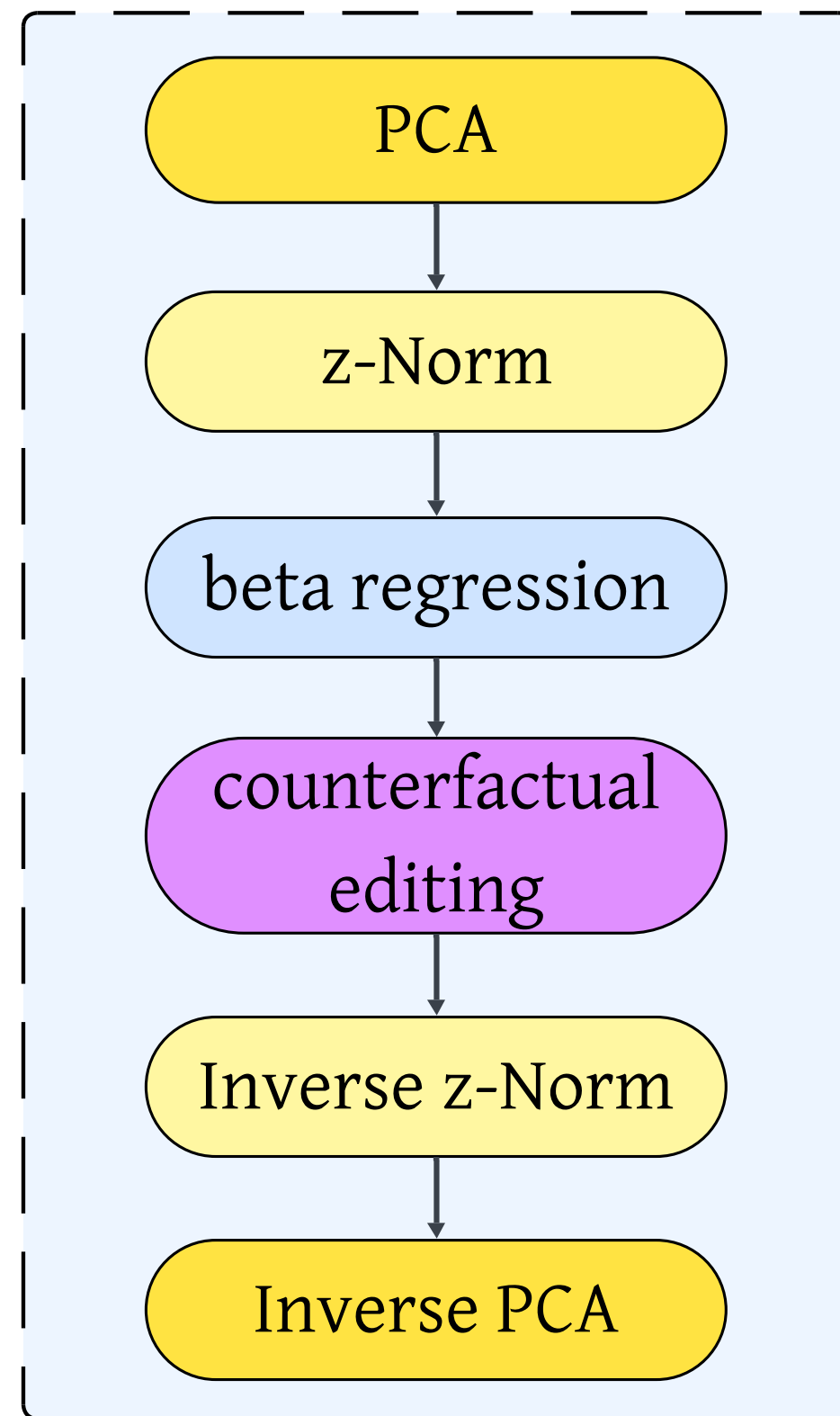


A Method for Continuous Variable Intervention



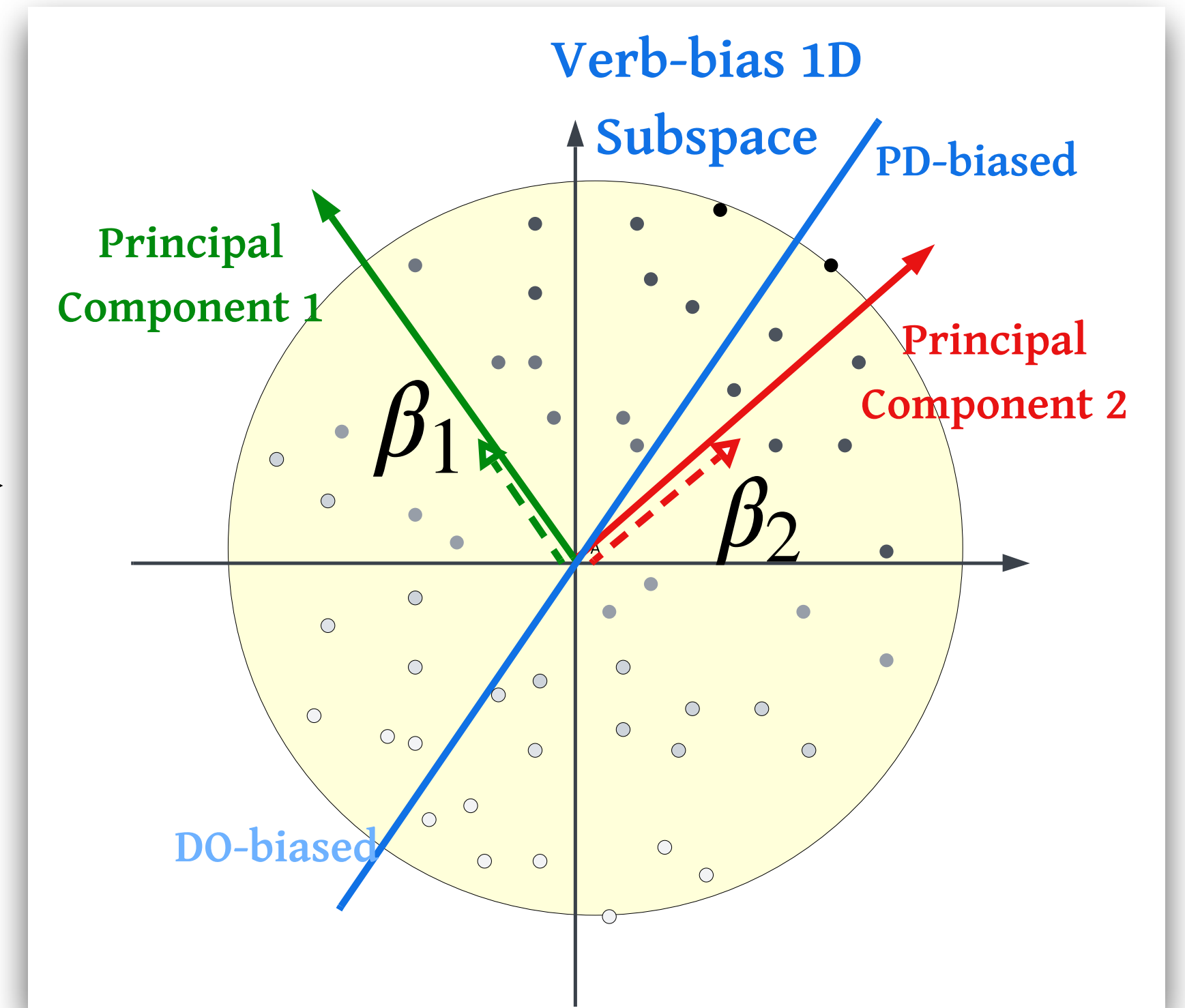
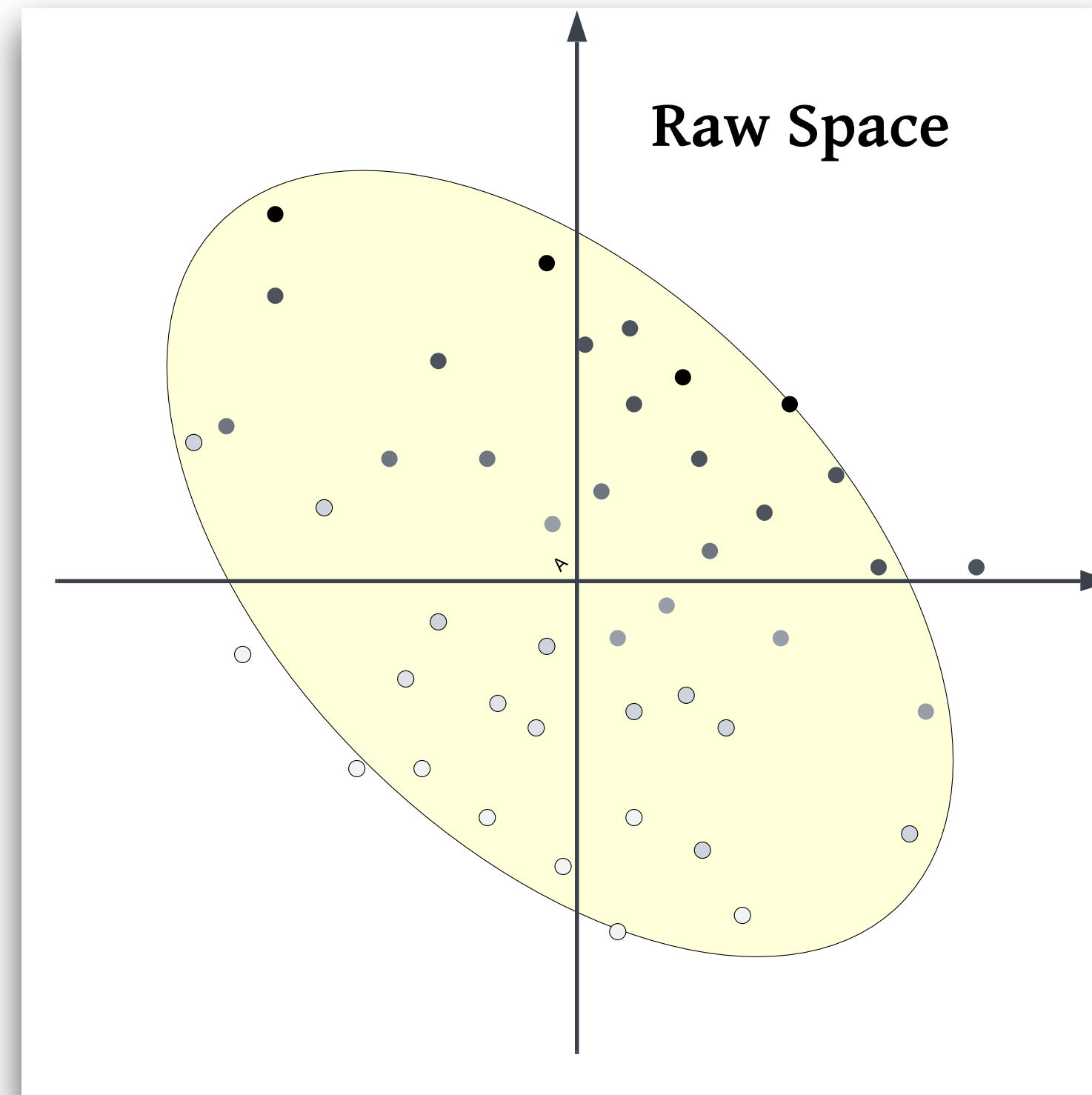
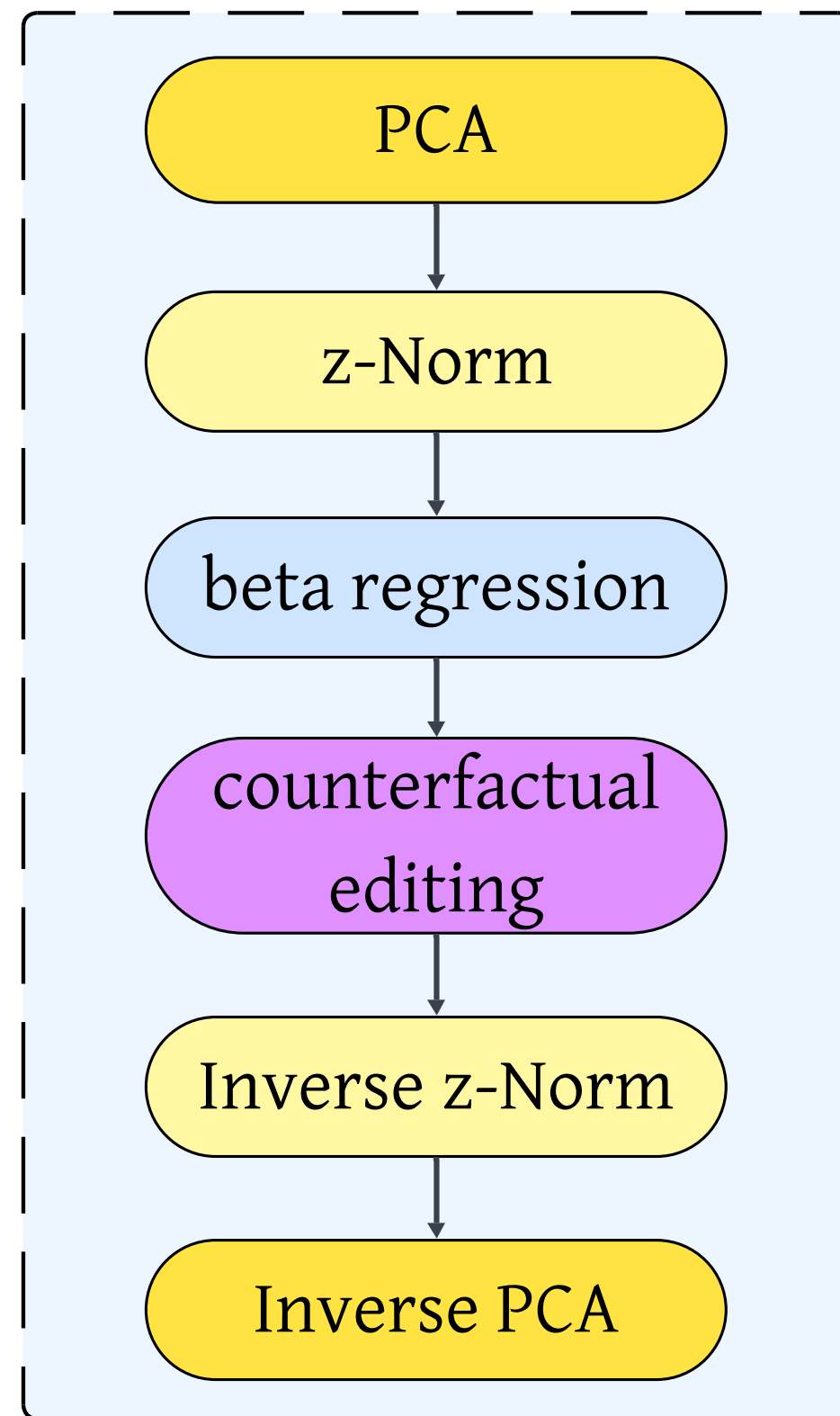
[For binary/discrete variable intervention, see Ravfogel et al. 2021, Hao & Linzen 2023, Ozaki et al. 2025]

A Method for Continuous Variable Intervention



[For binary/discrete variable intervention, see Ravfogel et al. 2021, Hao & Linzen 2023, Ozaki et al. 2025]

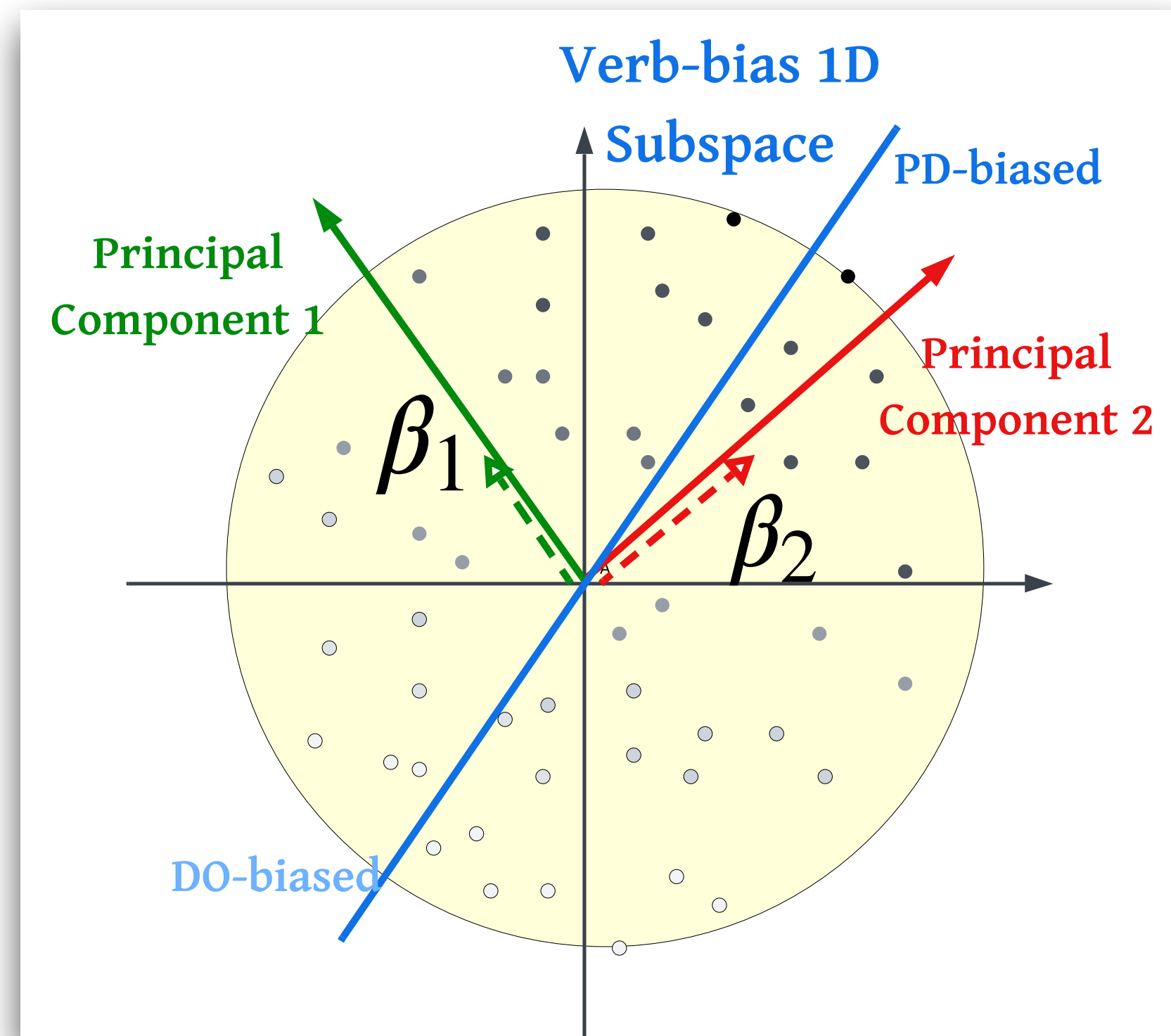
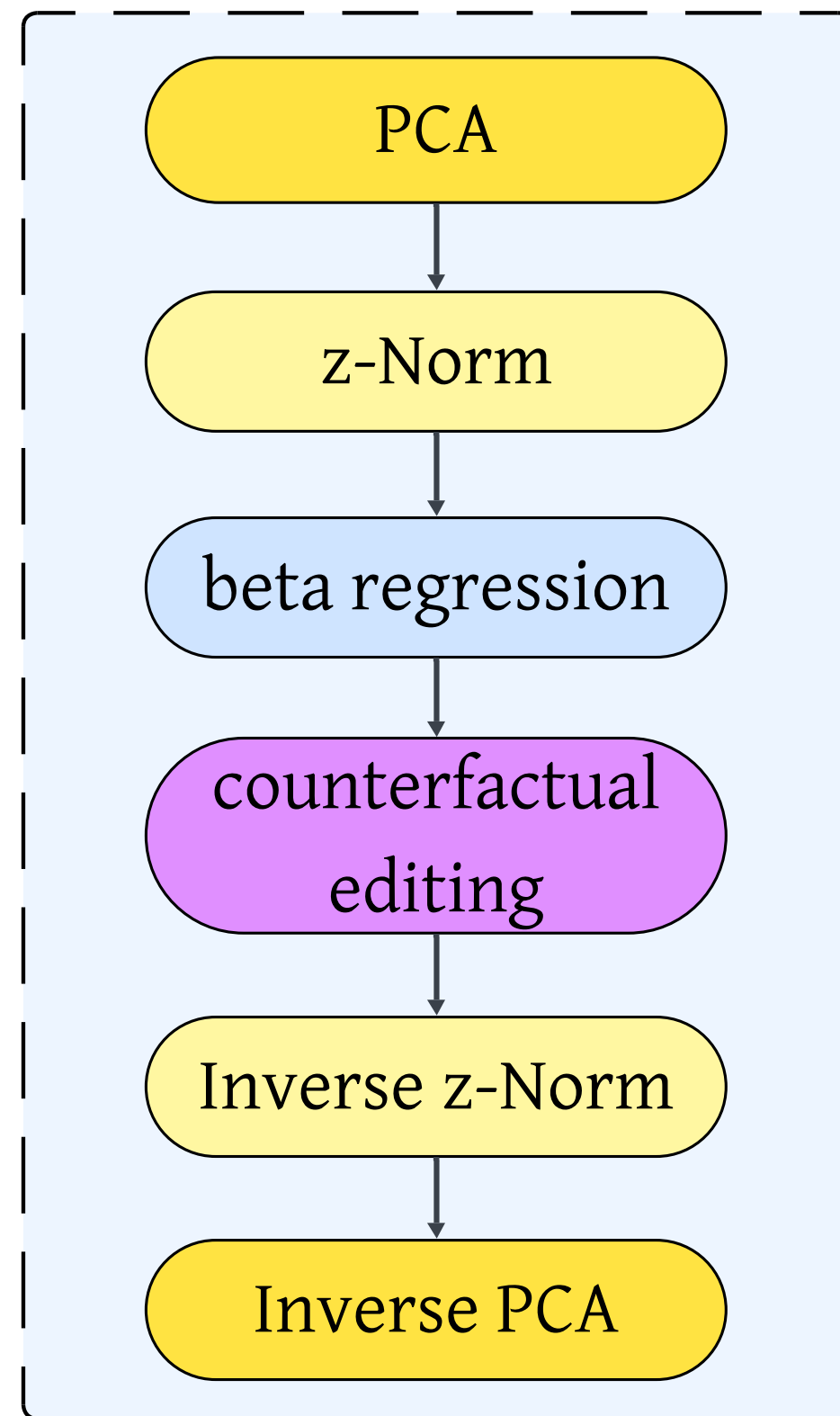
A Method for Continuous Variable Intervention



Localizing a 1-dimensional space that encodes verb bias information

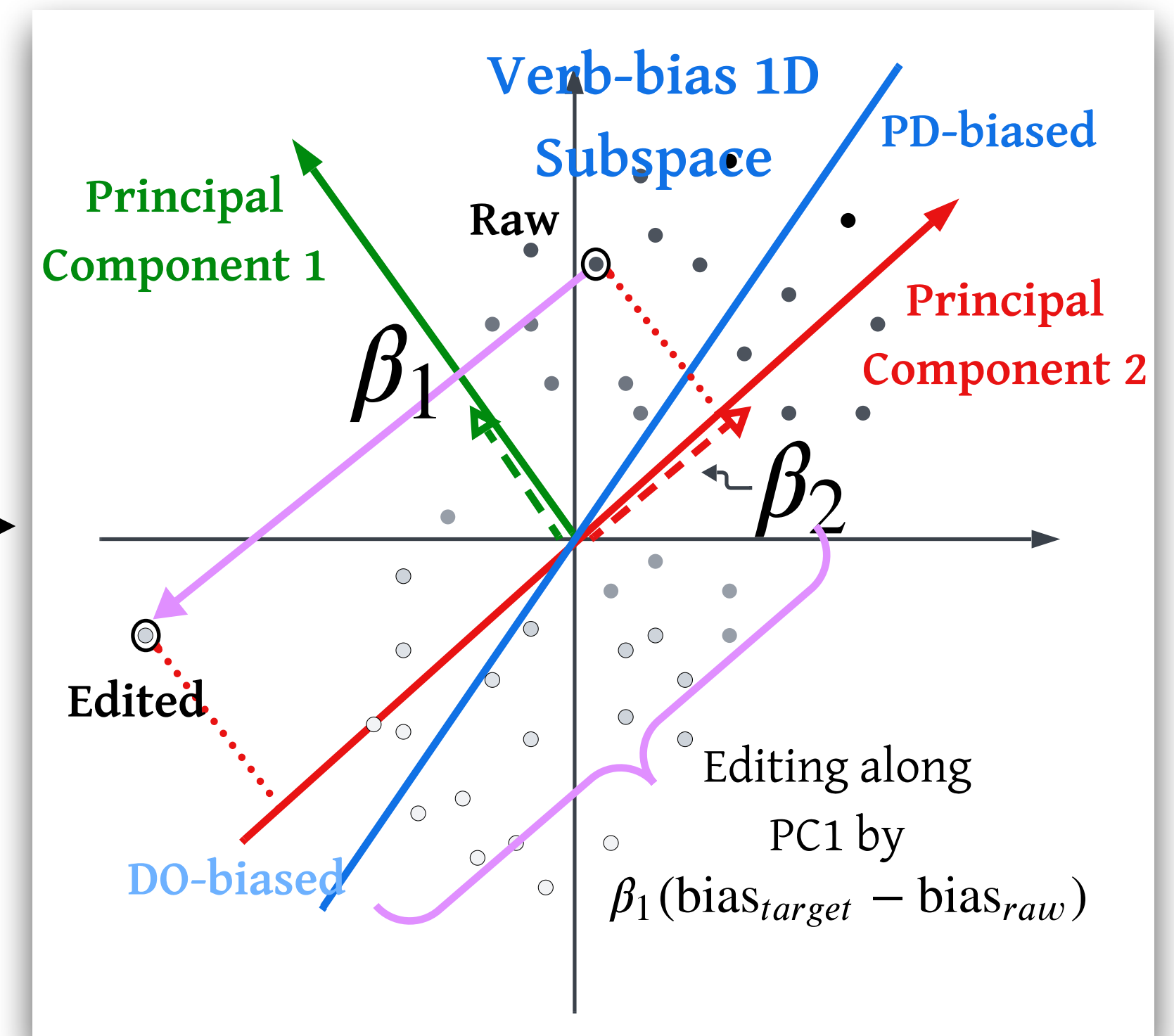
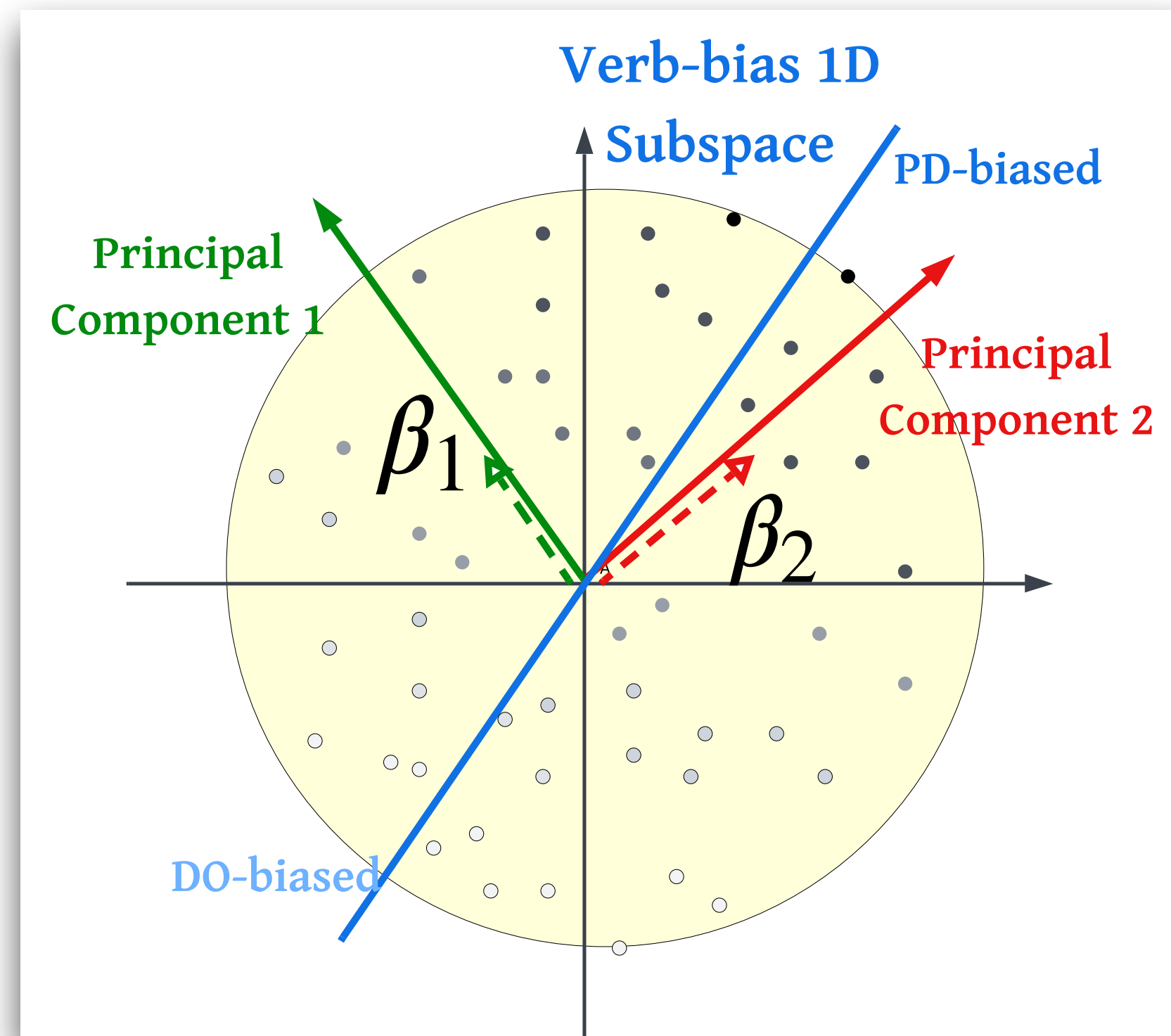
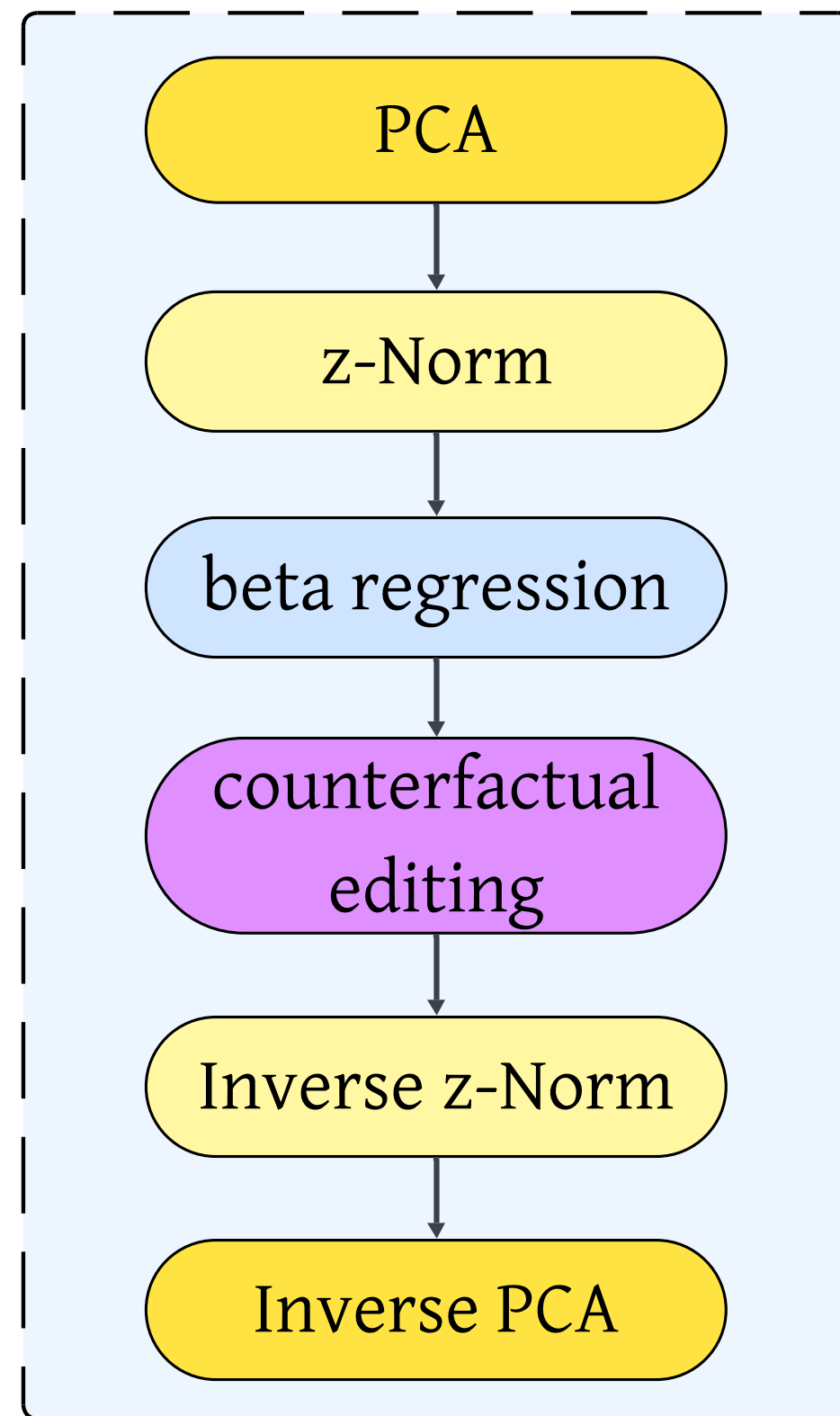
[For binary/discrete variable intervention, see Ravfogel et al. 2021, Hao & Linzen 2023, Ozaki et al. 2025]

A Method for Continuous Variable Intervention



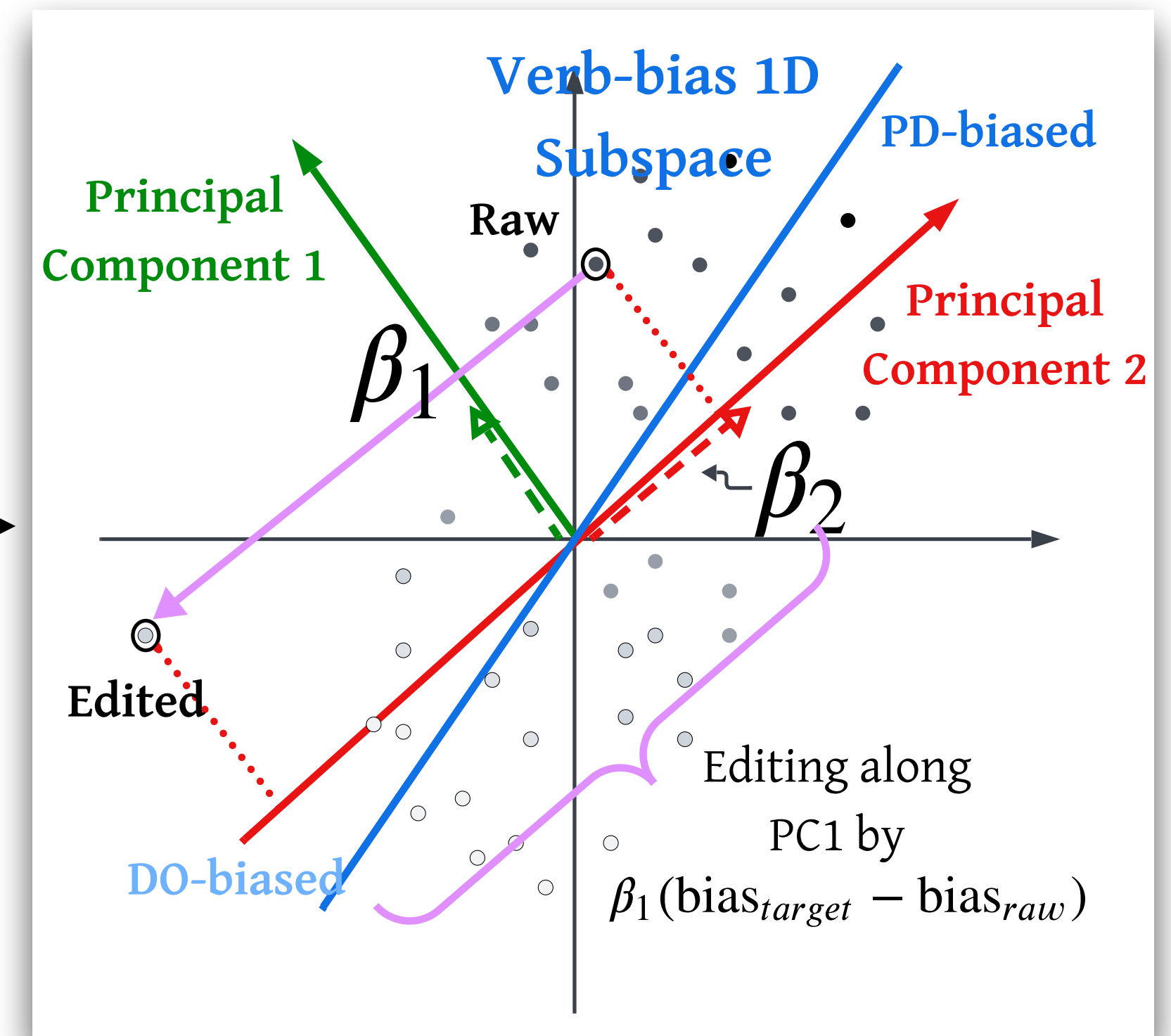
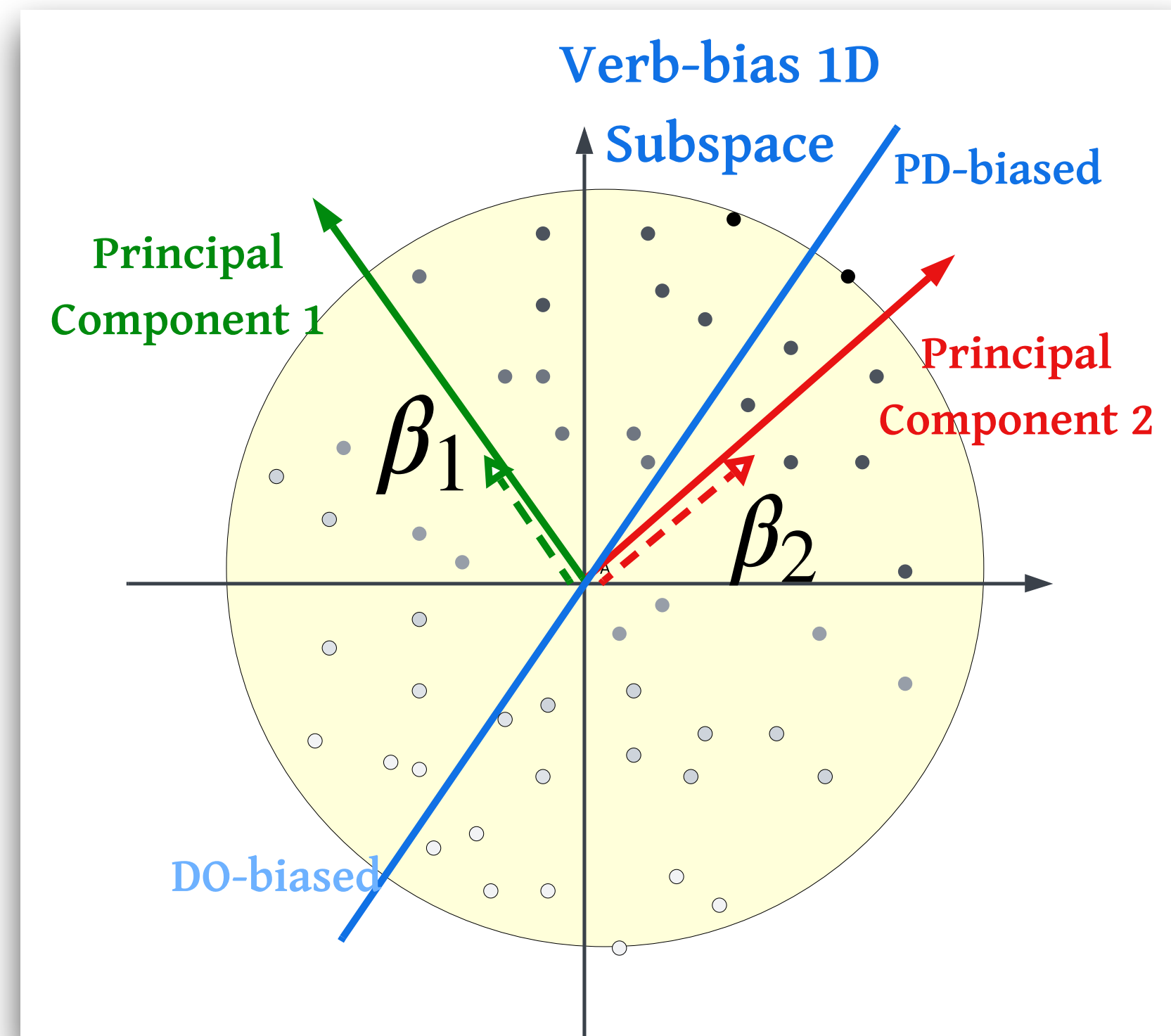
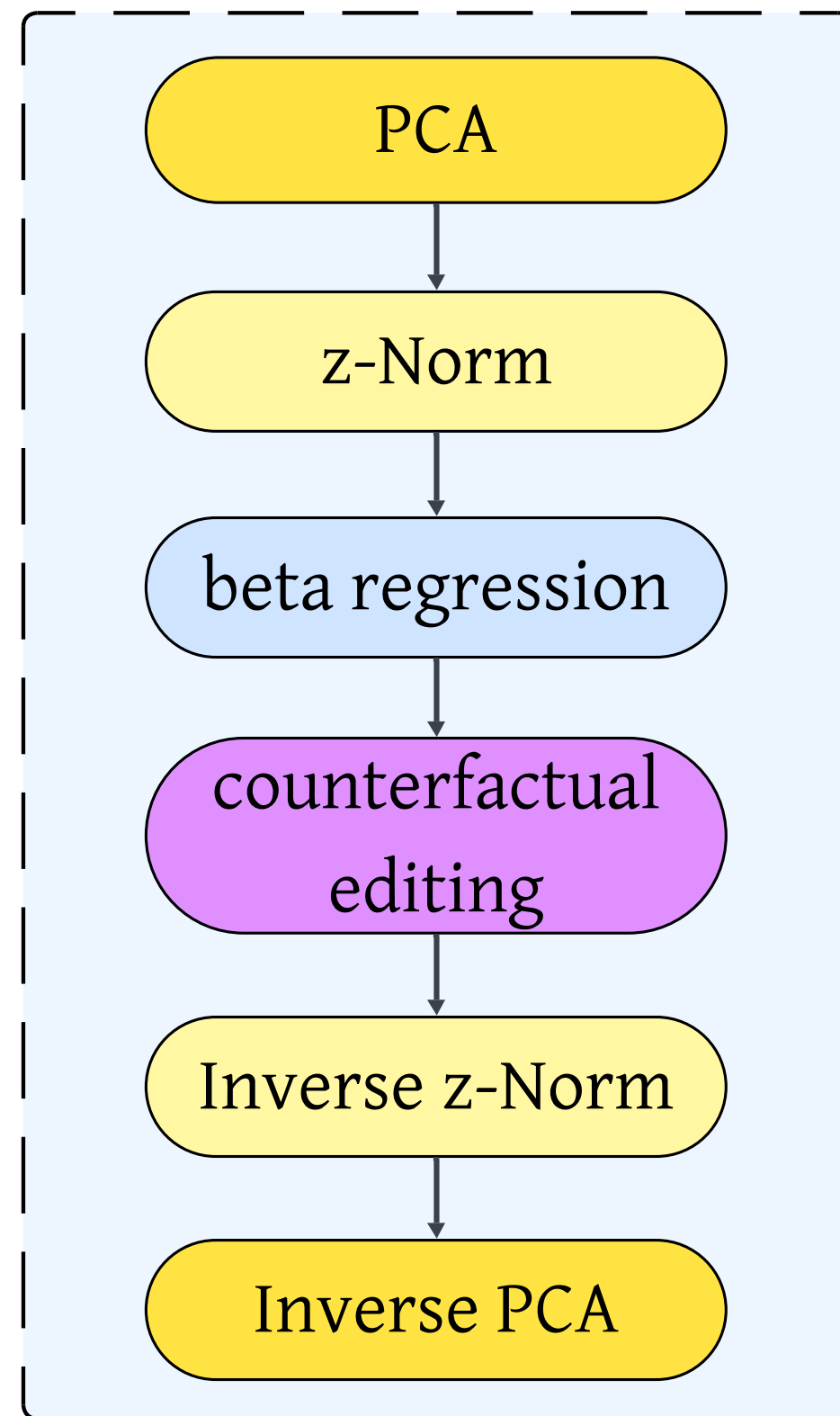
[For a close-form solution for concept erasing, see Belrose et al. 2023, LEACE]

A Method for Continuous Variable Intervention



[For a close-form solution for concept erasing, see Belrose et al. 2023, LEACE]

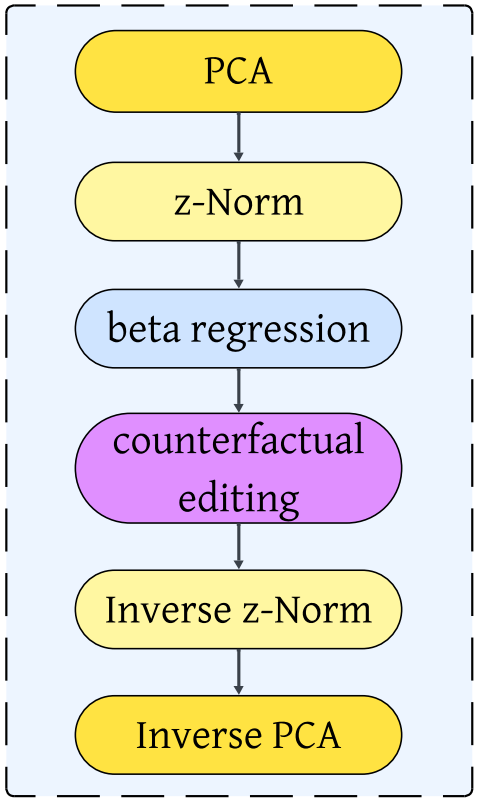
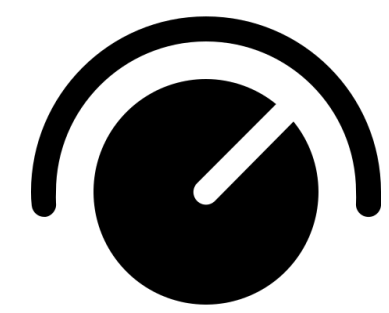
A Method for Continuous Variable Intervention



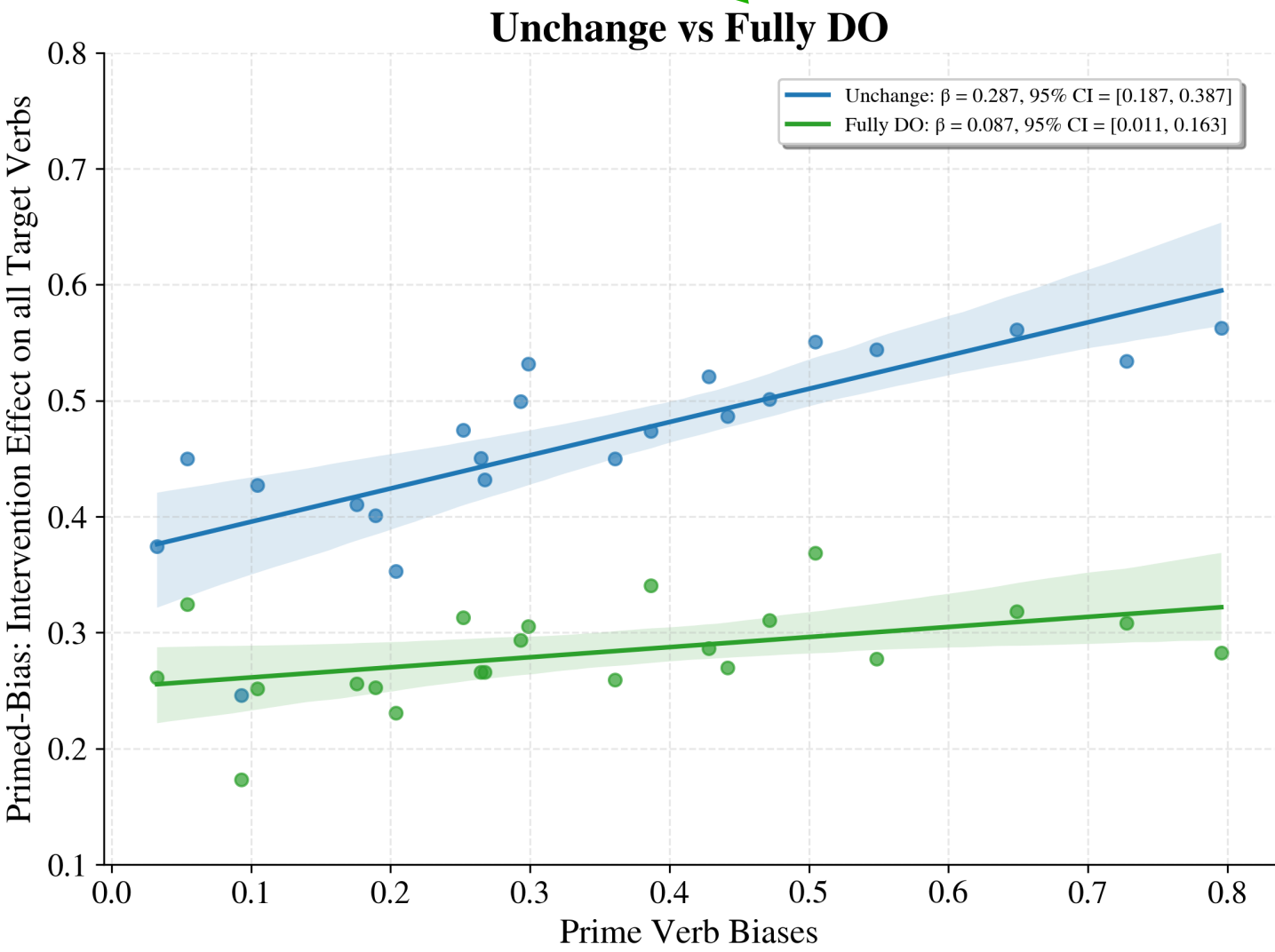
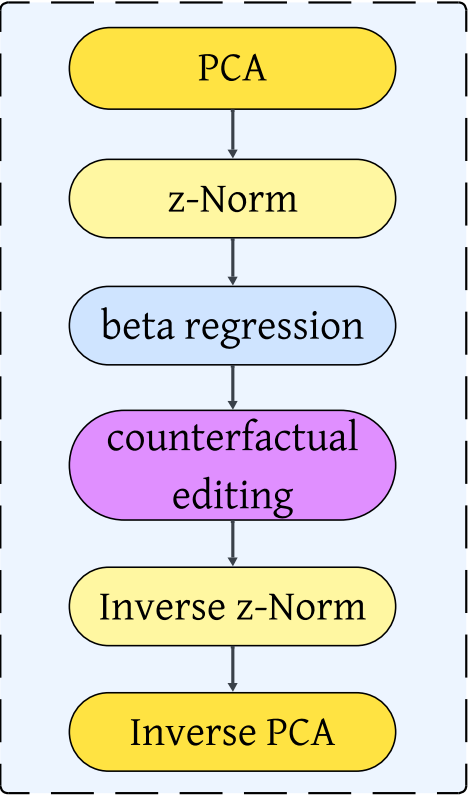
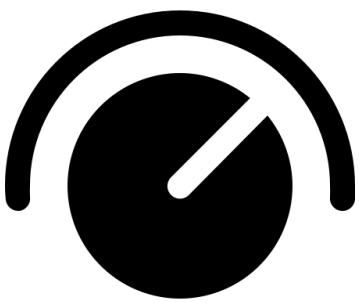
Causally manipulate the verb bias values without changing other information.

[For a close-form solution for concept erasing, see Belrose et al. 2023, LEACE]

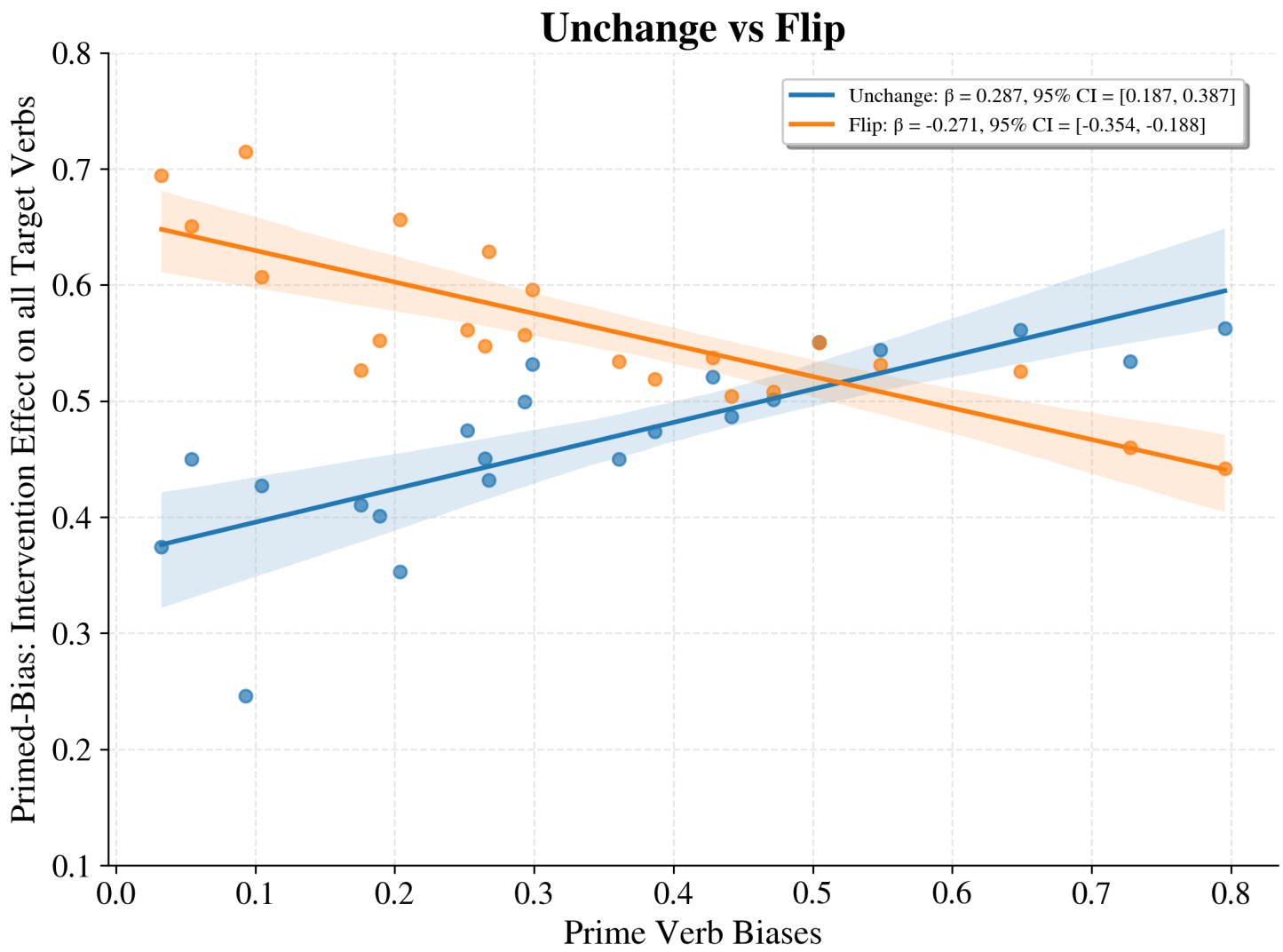
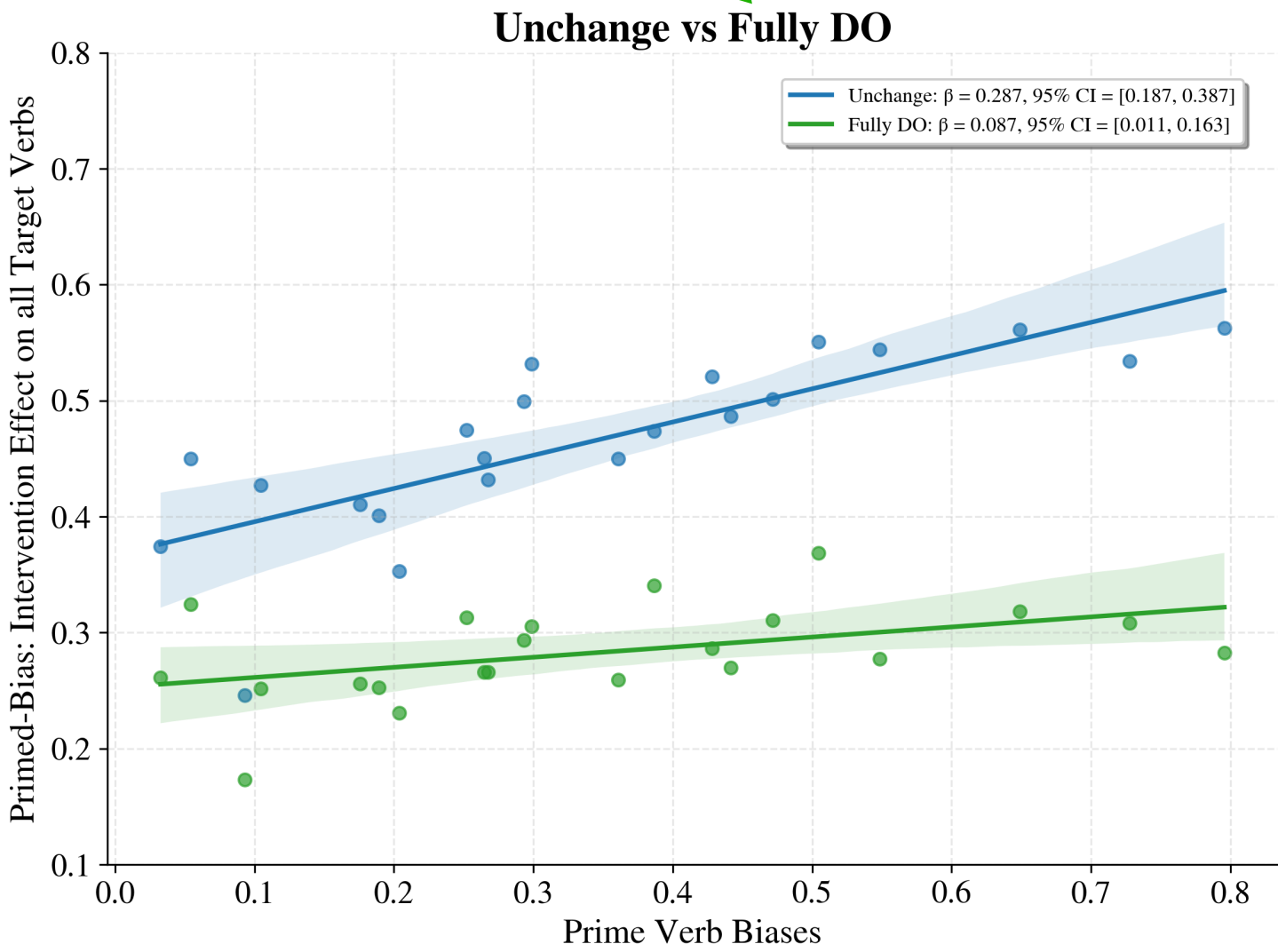
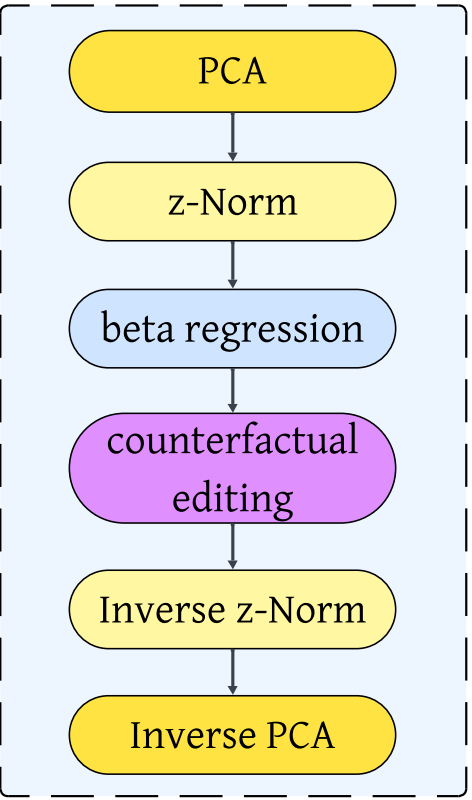
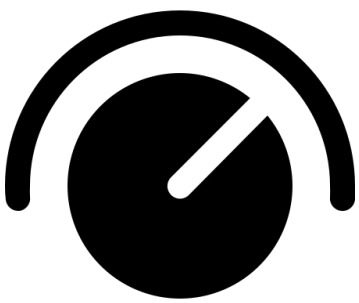
Causal Verification



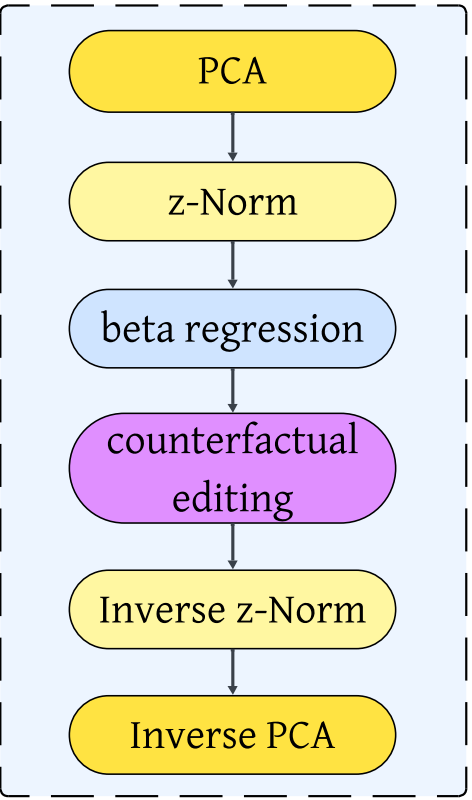
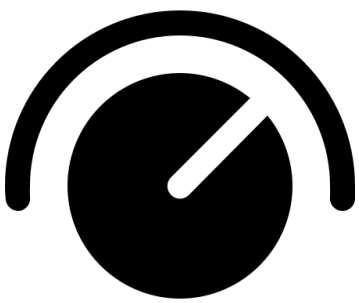
Causal Verification



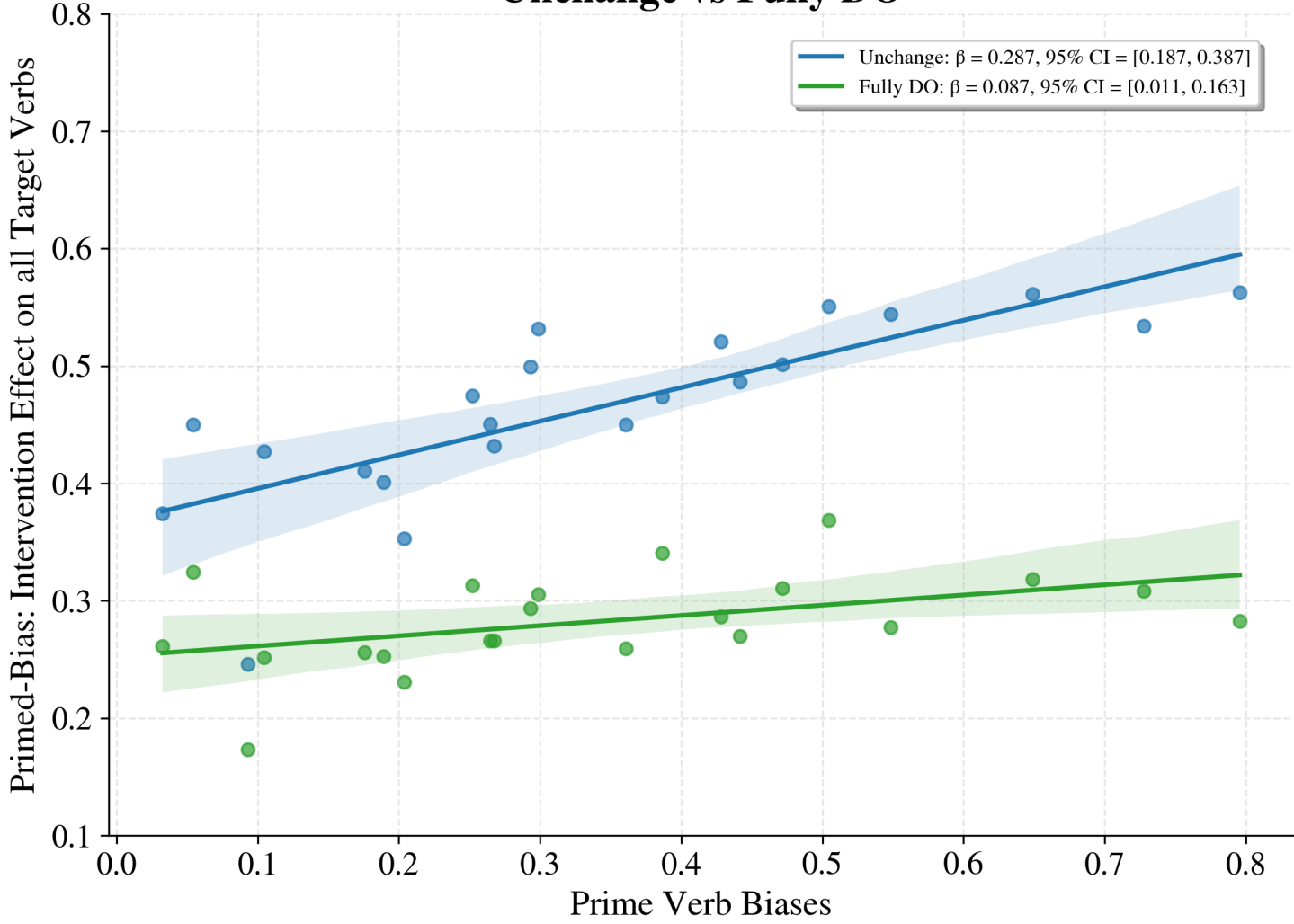
Causal Verification



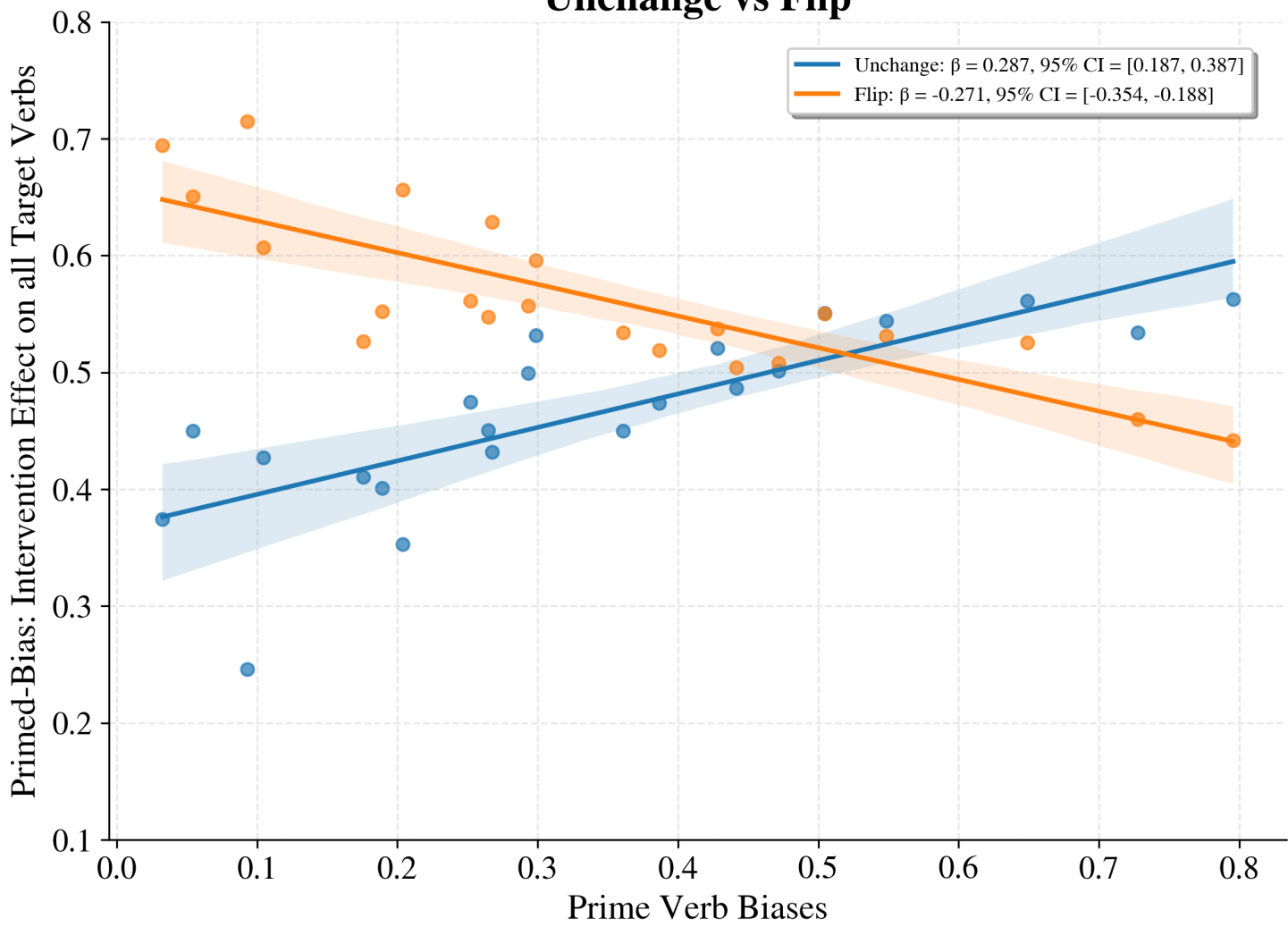
Causal Verification



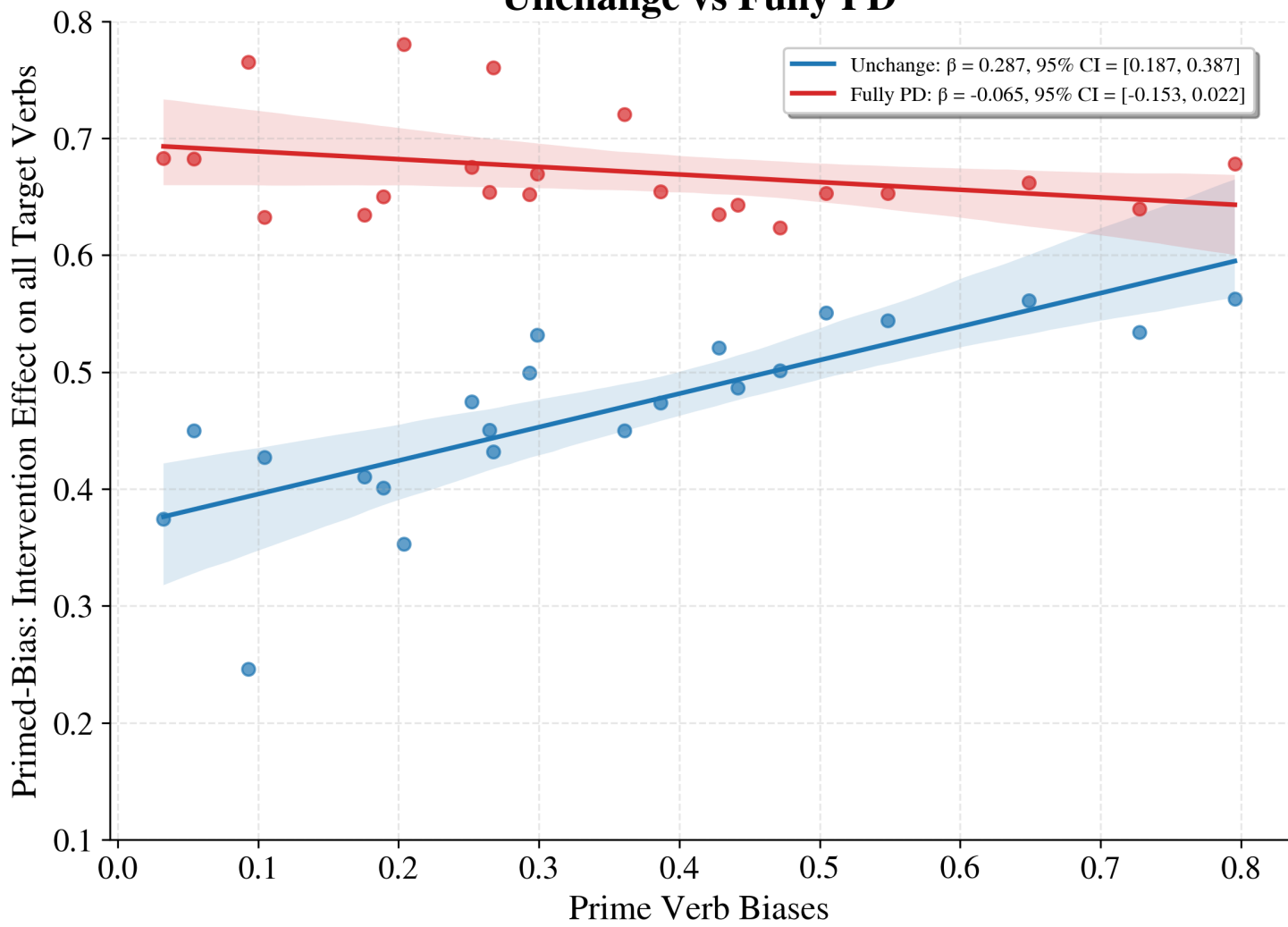
Unchange vs Fully DO



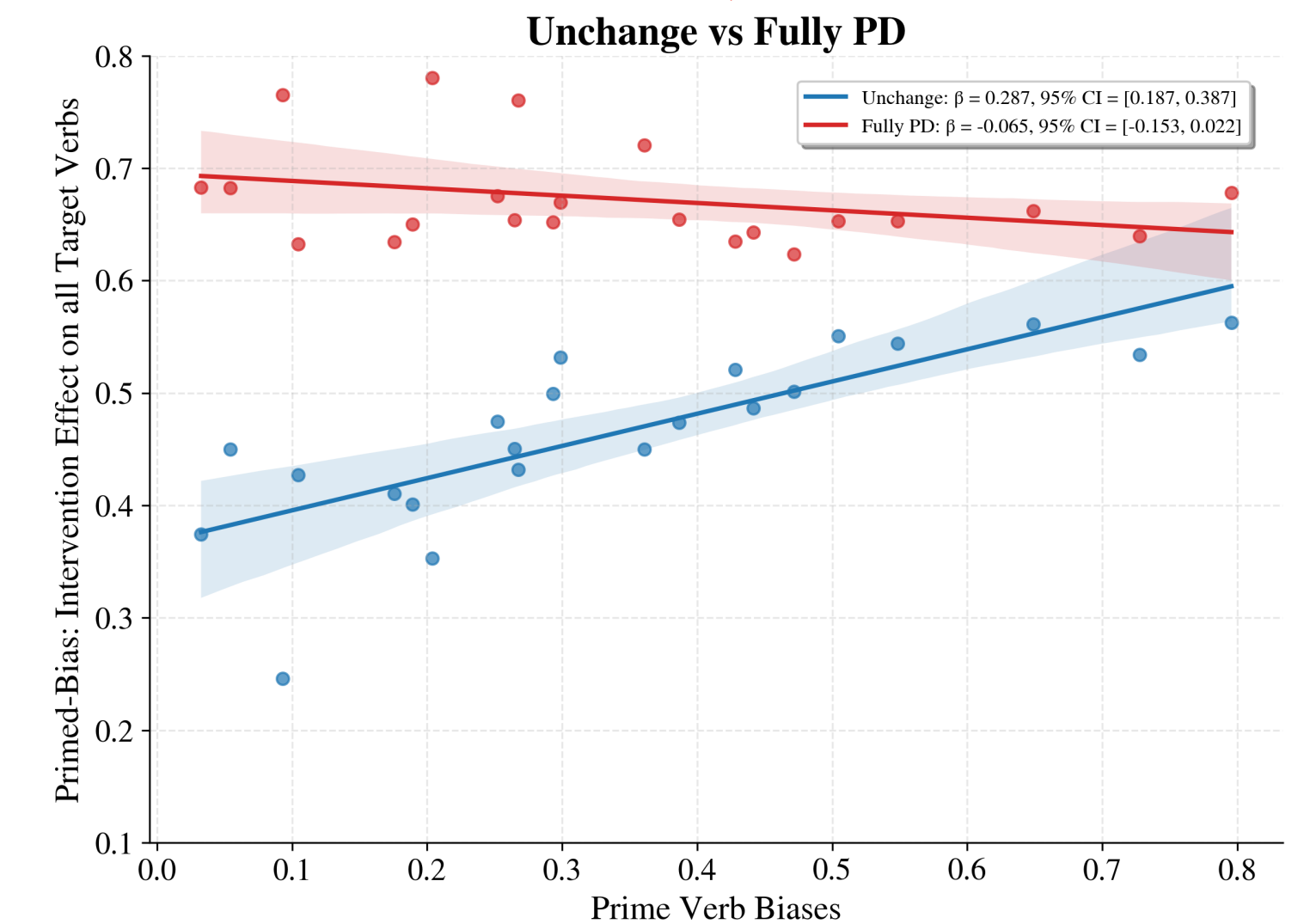
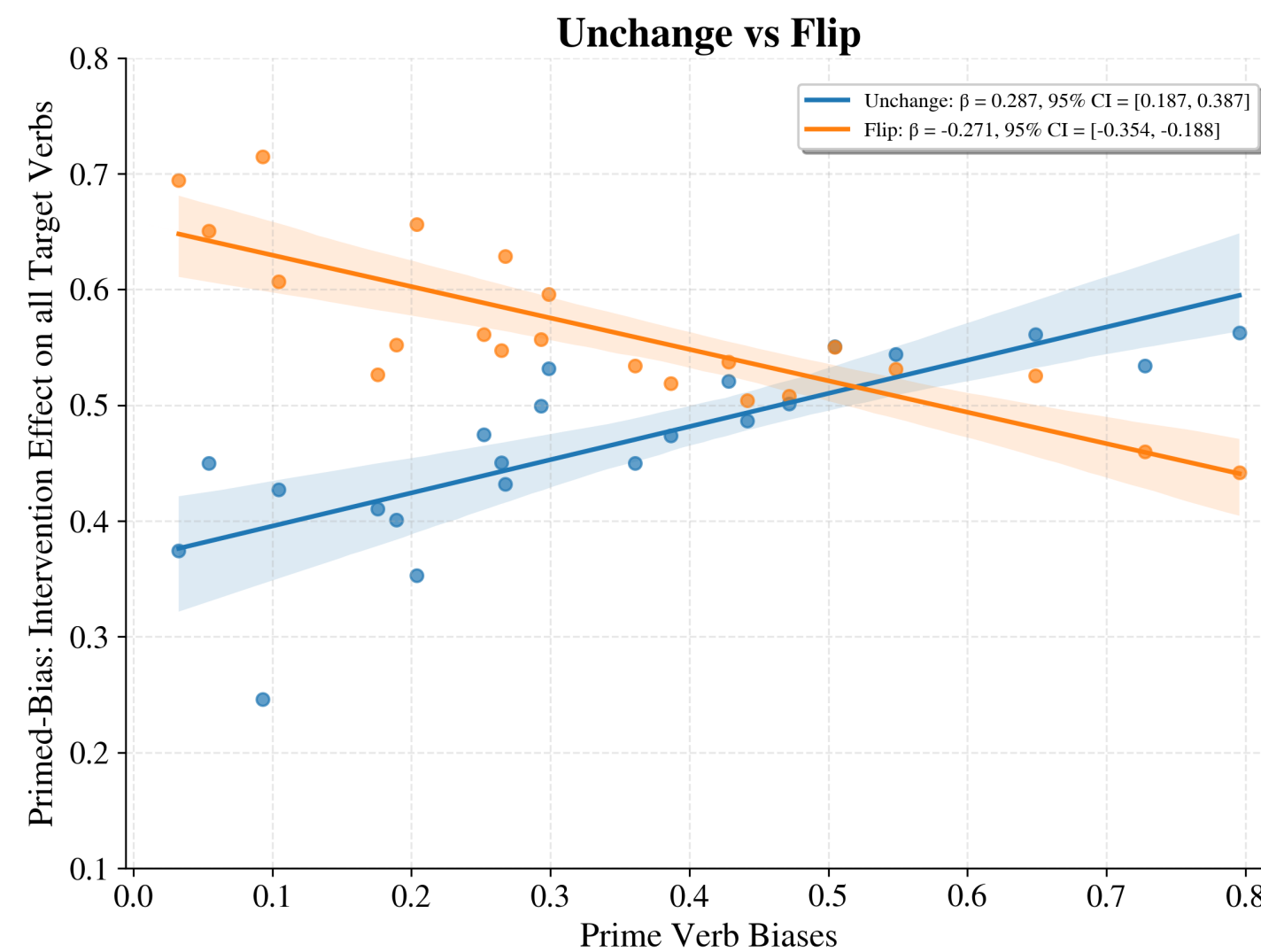
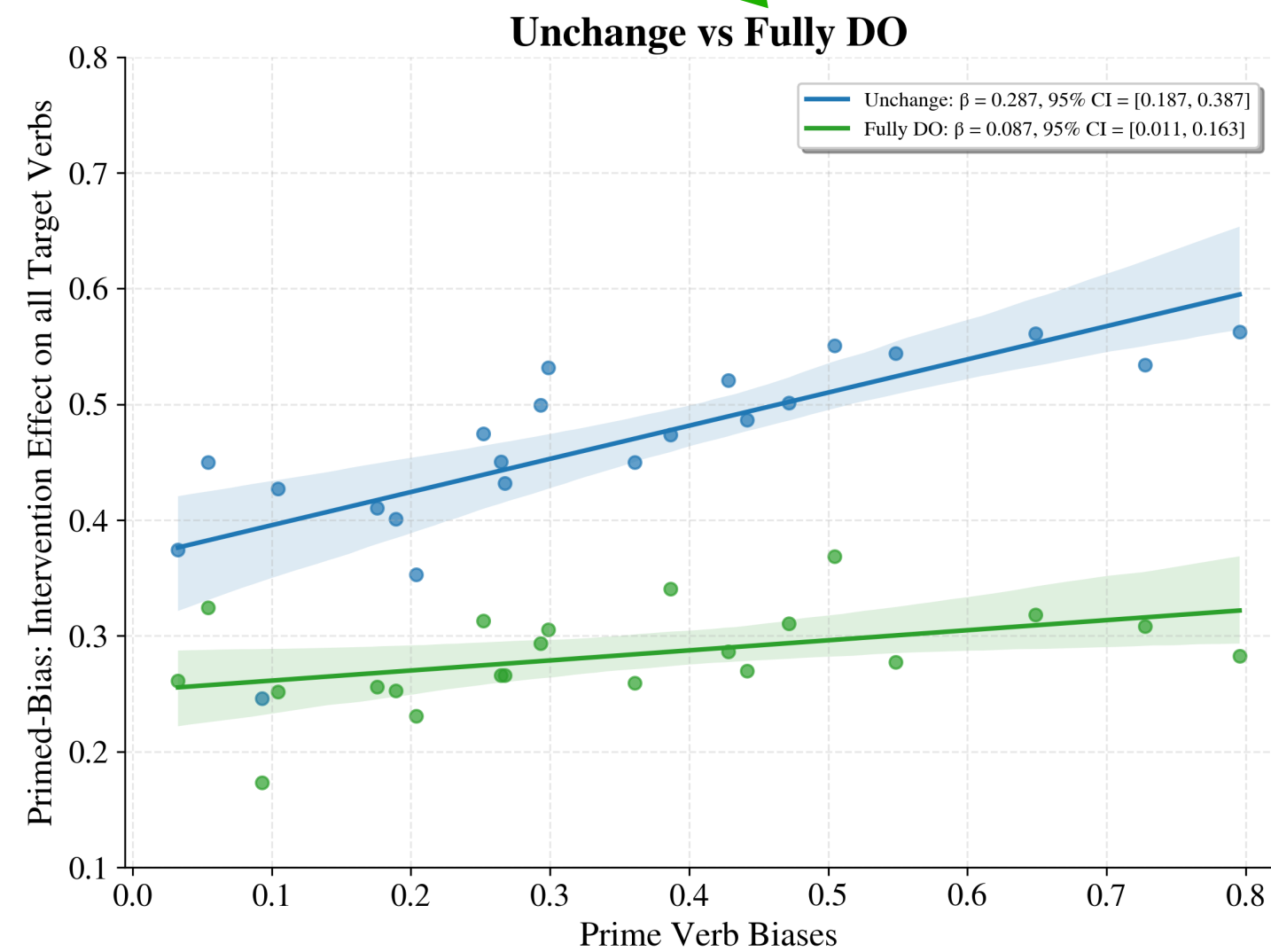
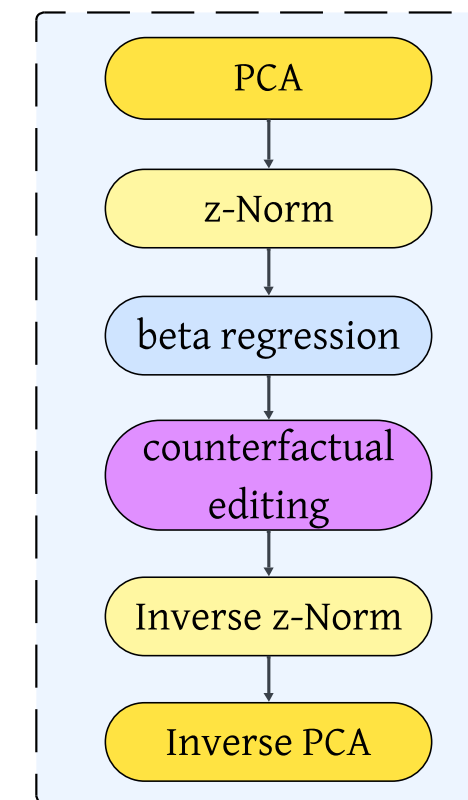
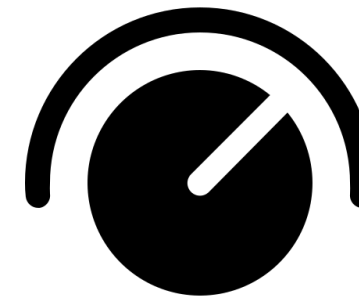
Unchange vs Flip



Unchange vs Fully PD



Causal Verification



**We showed that manipulating the verb biases causes predicted changes in the priming strength;
⇒ The verb bias representation is causally involved in LLM's structural prediction!**

Takeaways from Function Vector Studies

Takeaways from Function Vector Studies

- We use function vectors to further assess the “priming as ICL” claim:
 - **Yes, function vector do lead to standard priming effect**

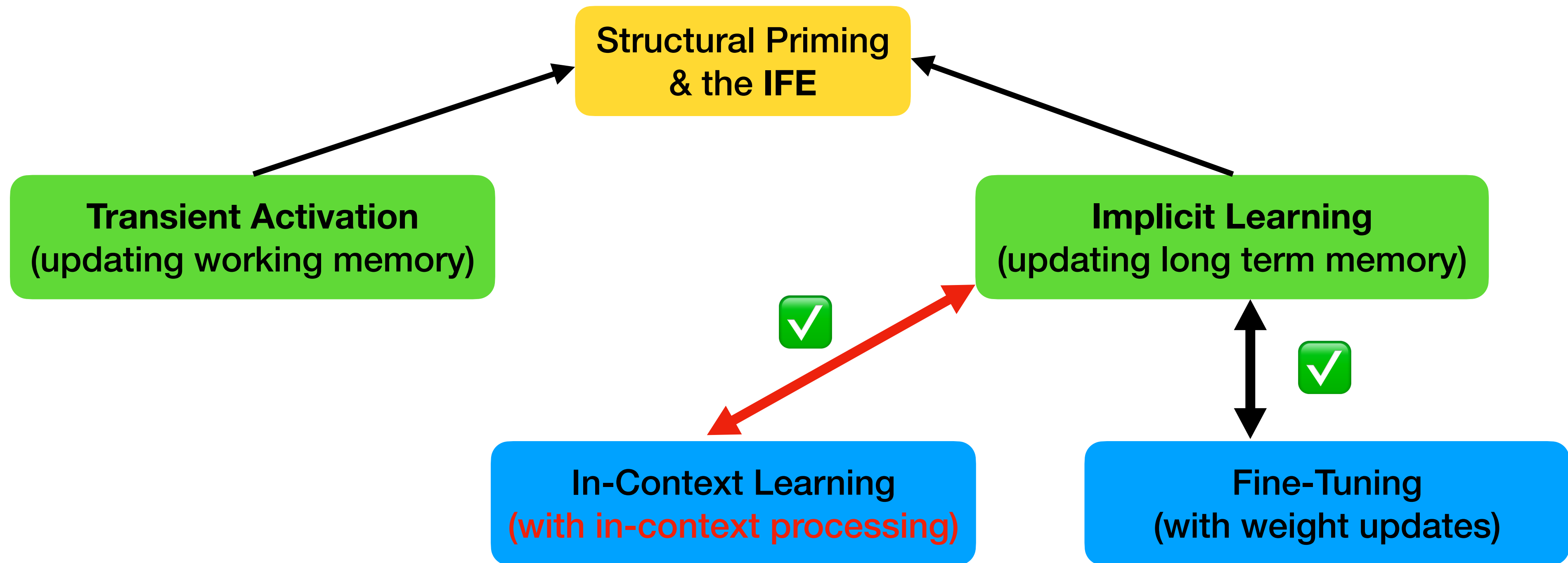
Takeaways from Function Vector Studies

- We use function vectors to further assess the “priming as ICL” claim:
 - **Yes, function vector do lead to standard priming effect**
- What abstract properties are encoded in the function vectors for priming?
 - **We found a low-dimensional subspace that encode verb bias;**
 - **Editing verb bias causally produces predicted effect on verb-specific priming magnitude;**

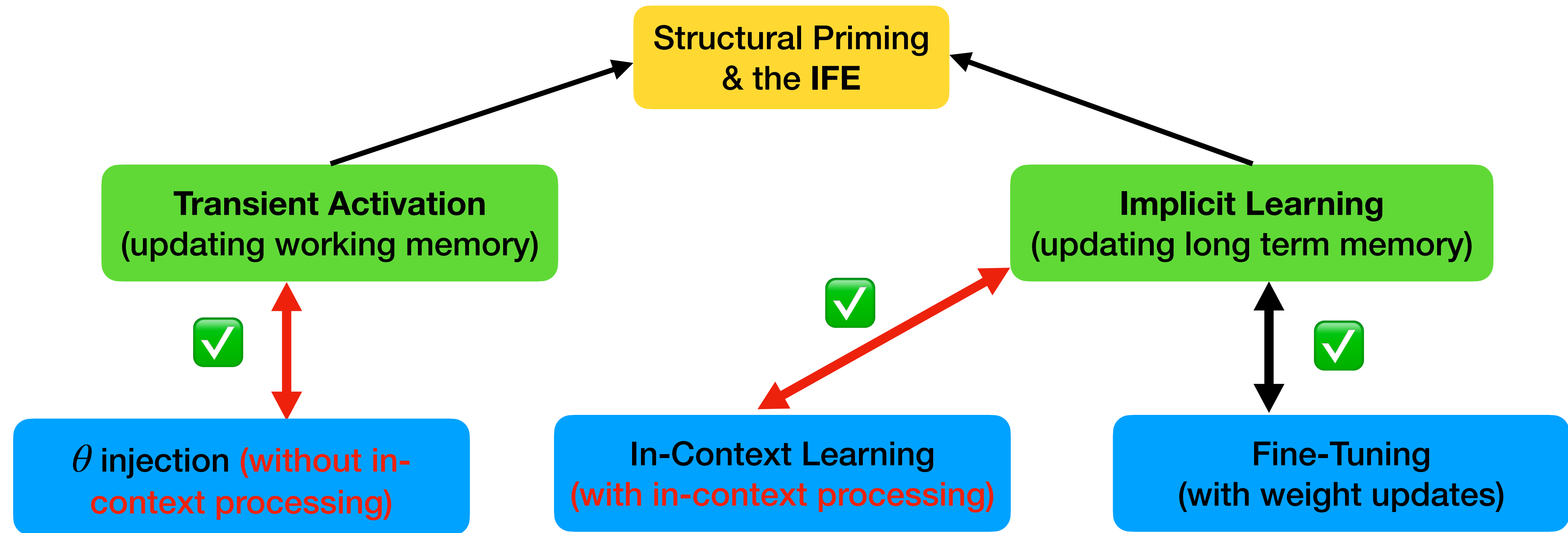
Takeaways from Function Vector Studies

- We use function vectors to further assess the “priming as ICL” claim:
 - **Yes, function vector do lead to standard priming effect**
- What abstract properties are encoded in the function vectors for priming?
 - **We found a low-dimensional subspace that encode verb bias;**
 - **Editing verb bias causally produces predicted effect on verb-specific priming magnitude;**
- However... function vector cannot be the whole story of ICL because:
 - **We see no IFE with the function vector induced priming effect;**
 - **What is the crucial ingredient about IFE that function vectors are missing?**

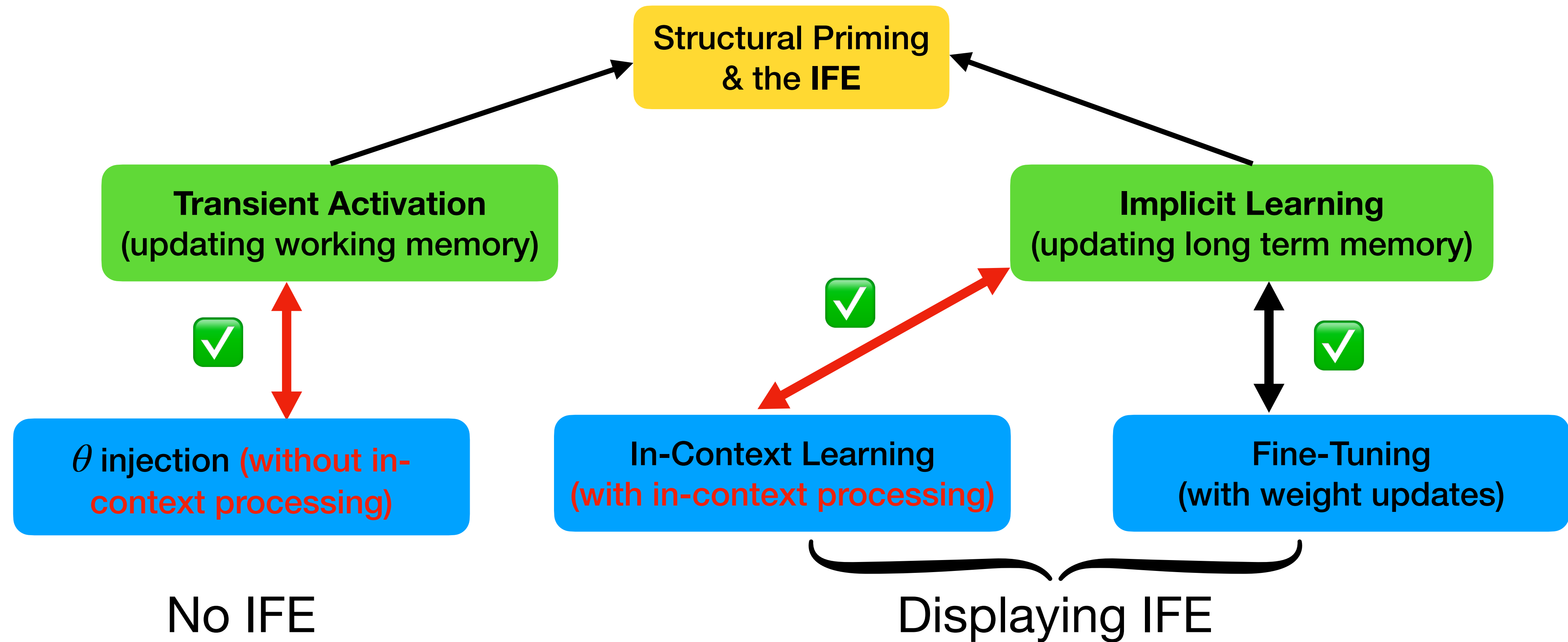
Revisiting the big picture so far



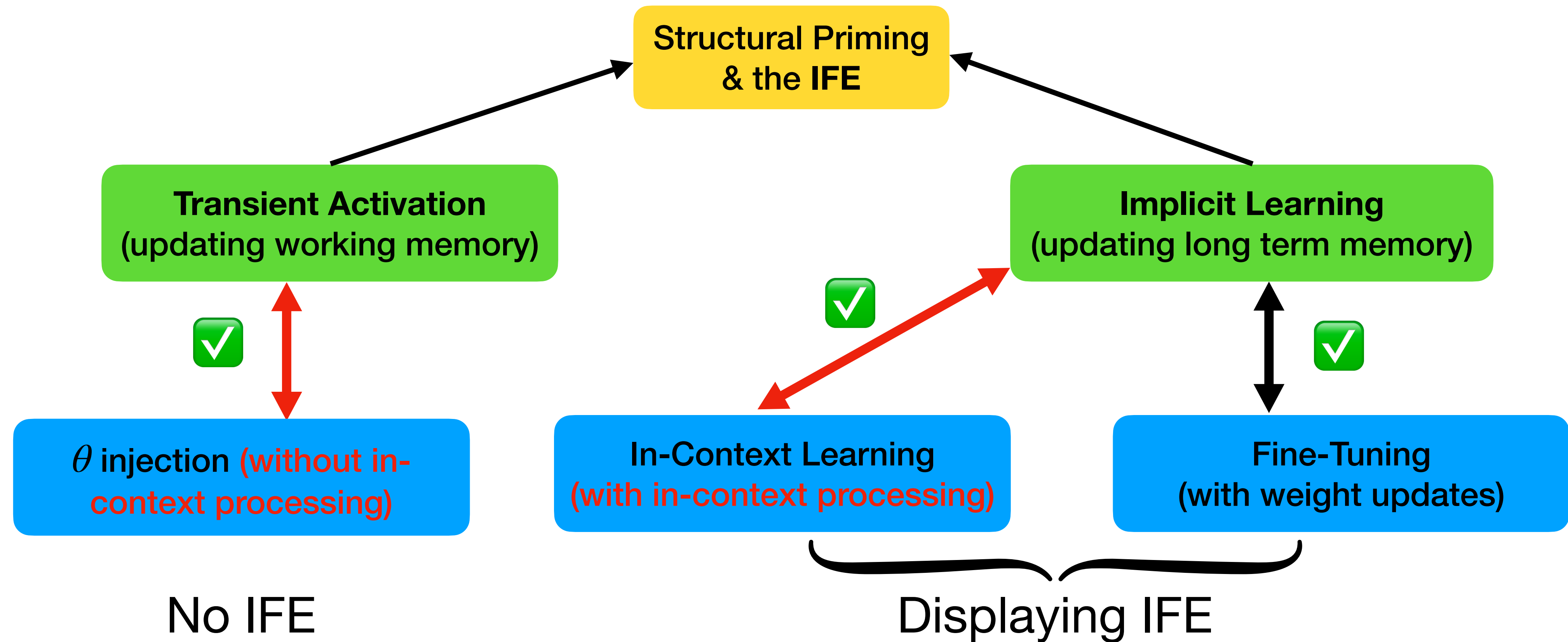
Revisiting the big picture so far



Revisiting the big picture so far



Revisiting the big picture so far



Function vectors do not capture the crucial error signal for LLMs to adapt to context, making it resembling transient activation instead of implicit learning!

High-level Takeaways

What do we learn about ICL from human priming theories?

High-level Takeaways

What do we learn about ICL from human priming theories?

- Previous / popular definition: ICL takes the form of <demo, answer> pairs;

High-level Takeaways

What do we learn about ICL from human priming theories?

- Previous / popular definition: ICL takes the form of <demo, answer> pairs;
- The IFE provides a finer-grained characterization of ICL:
 - Instead of just identifying the “static” task (e.g., antonym, translation), LLMs also dynamically adapt to the context in an error-driven way.

High-level Takeaways

What do we learn about ICL from human priming theories?

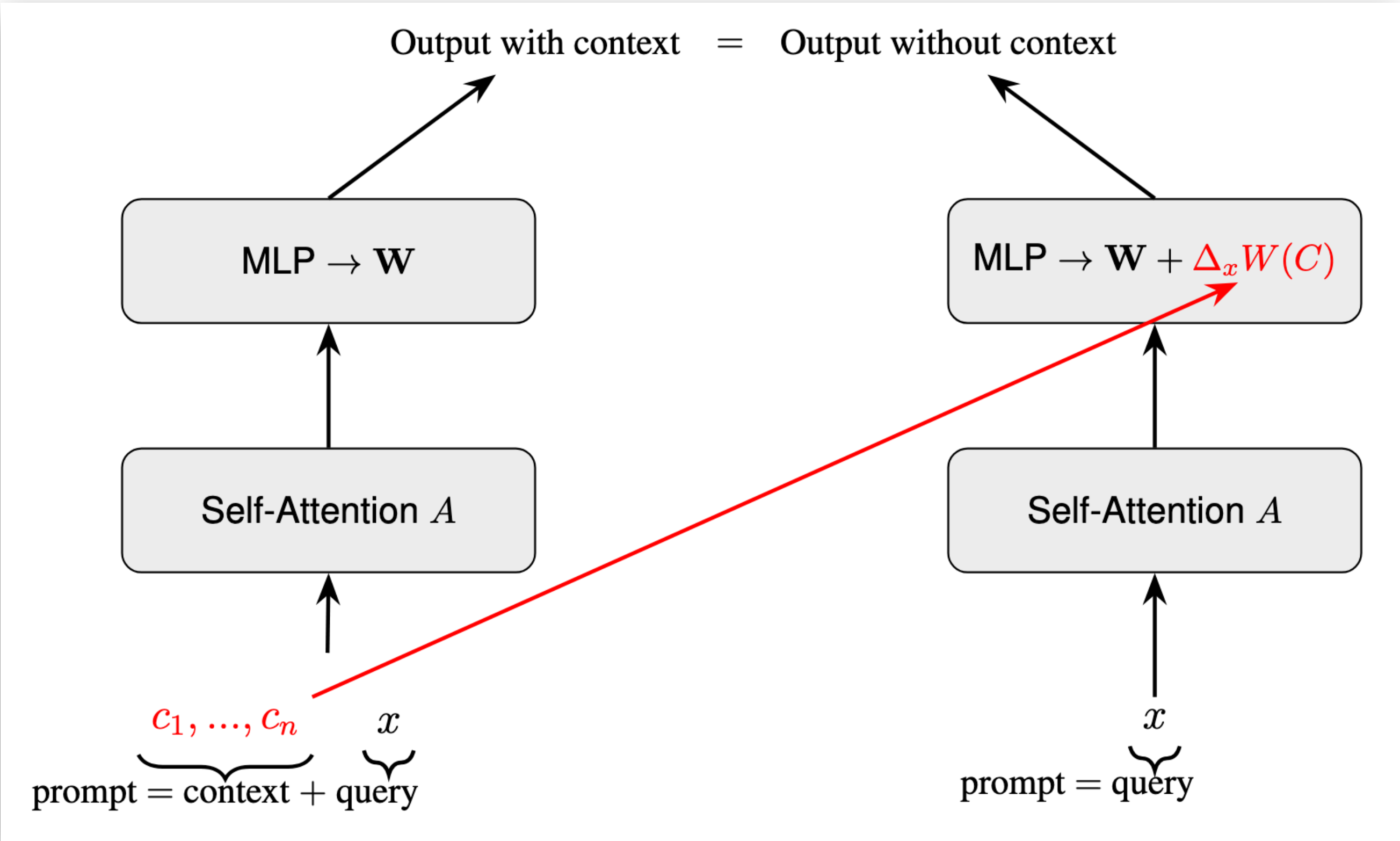
- Previous / popular definition: ICL takes the form of <demo, answer> pairs;
- The IFE provides a finer-grained characterization of ICL:
 - Instead of just identifying the “static” task (e.g., antonym, translation), LLMs also dynamically adapt to the context in an error-driven way.
- This gives rise to a generalized notion of ICL:
 - When we get rid of the “input-output” task form, then ICL becomes the more general ability of “adapting to context” — being context-sensitive in an error-driven way — e.g., Bayesian mind instantiated as predictive coding!

High-level Takeaways

What do we learn about ICL from human priming theories?

- Previous / popular definition: ICL takes the form of <demo, answer> pairs;
- The IFE provides a finer-grained characterization of ICL:
 - Instead of just identifying the “static” task (e.g., antonym, translation), LLMs also dynamically adapt to the context in an error-driven way.
- This gives rise to a generalized notion of ICL:
 - When we get rid of the “input-output” task form, then ICL becomes the more general ability of “adapting to context” — being context-sensitive in an error-driven way — e.g., Bayesian mind instantiated as predictive coding!
- Function vectors only captures the “static” part of ICL — this compression of context is indeed lossy.

Open Questions



Learning without training: The implicit dynamics of in-context learning

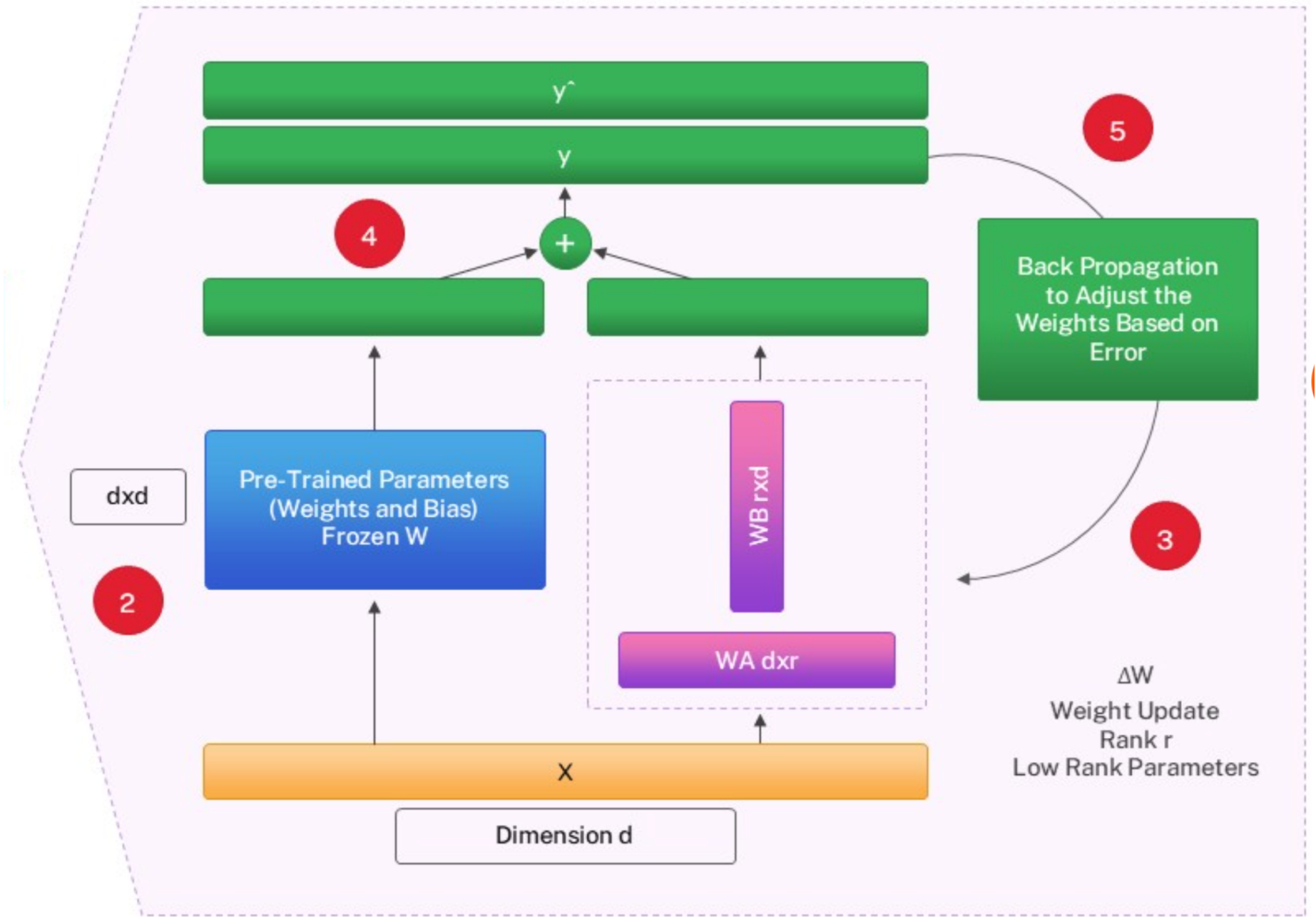
Benoit Dherin*
Google Research
dherin@google.com

Michael Munn*
Google Research
munm@google.com

Hanna Mazzawi*
Google Research
mazzawi@google.com

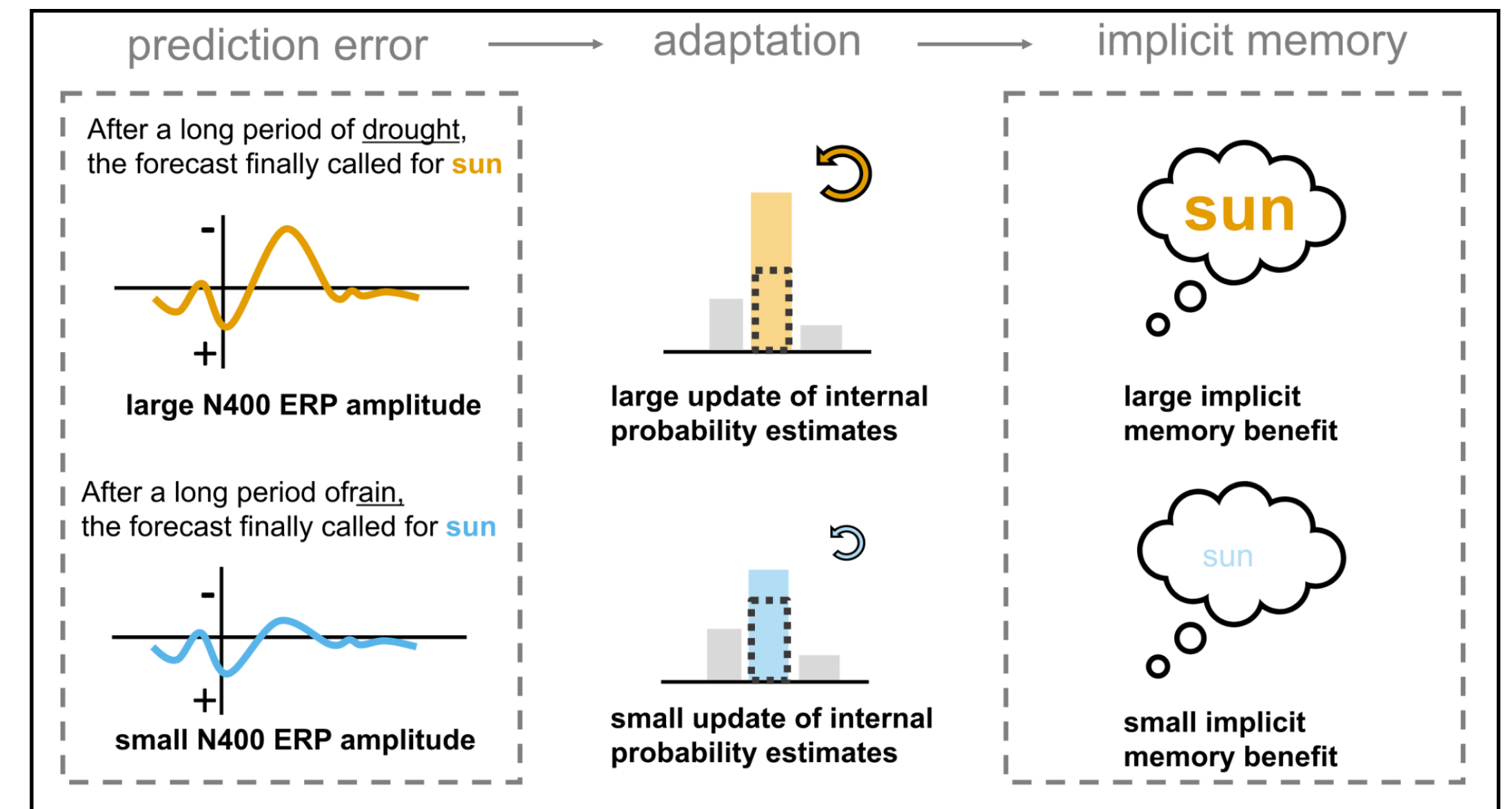
Michael Wunder
Google Research
mwunder@google.com

Javier Gonzalvo
Google Research
xavigonzalvo@google.com



Latest progress: contextual influence as implicit low-rank update on MLP weights?

Open Questions: N400 as error signal?



Open Questions: N400 as error signal?

SPECIAL ISSUE:
Cognitive Computational Neuroscience of Language

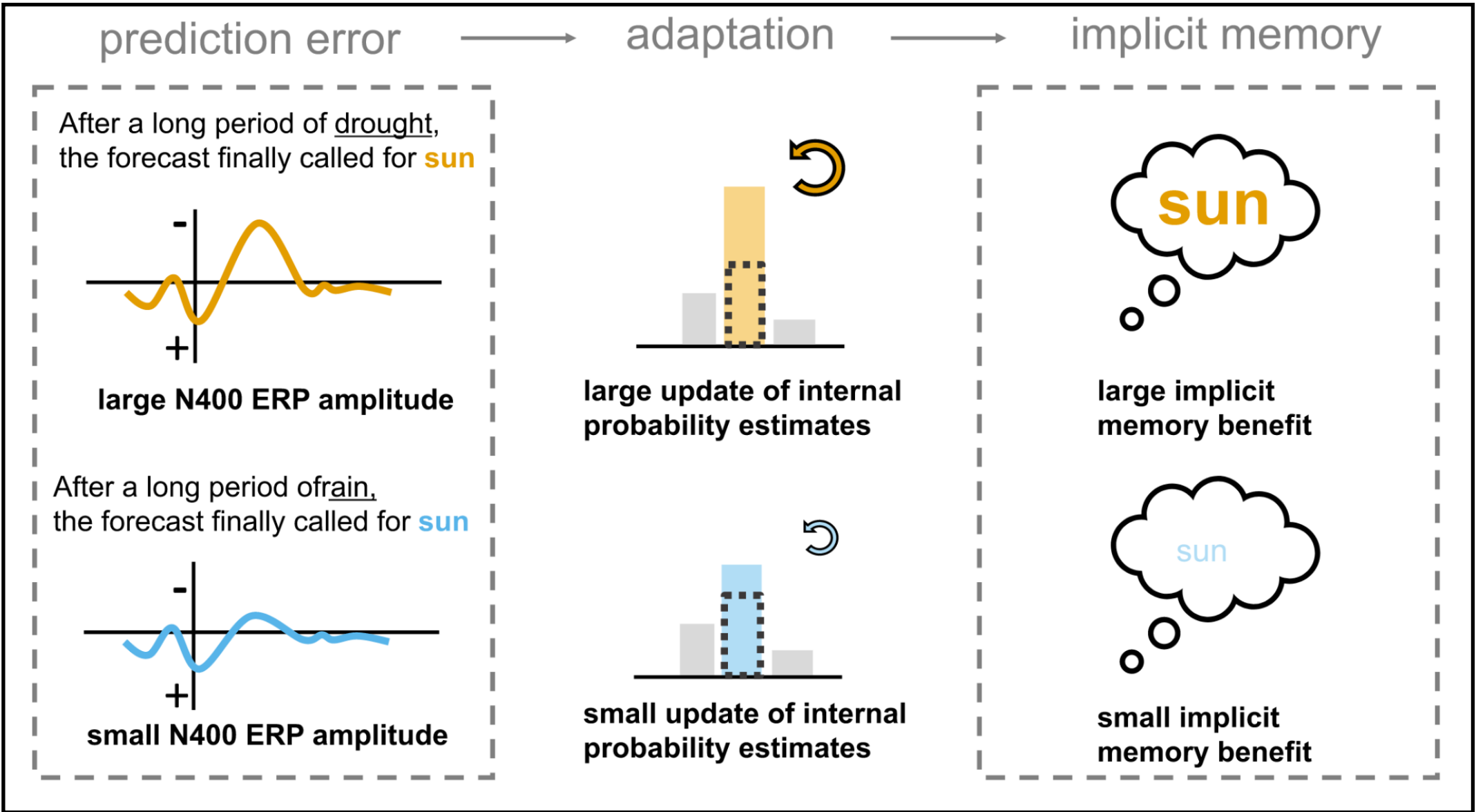
Strong Prediction: Language Model Surprisal Explains Multiple N400 Effects

James A. Michaelov¹, Megan D. Bardolph¹, Cyma K. Van Petten²,
Benjamin K. Bergen¹, and Seana Coulson¹

¹Department of Cognitive Science, University of California, San Diego, La Jolla, CA, USA

²Department of Psychology, Binghamton University, State University of New York, Binghamton, NY, USA

Keywords: distributional semantics, ERPs, N400, neural language models, predictive coding



Thanks for Listening!

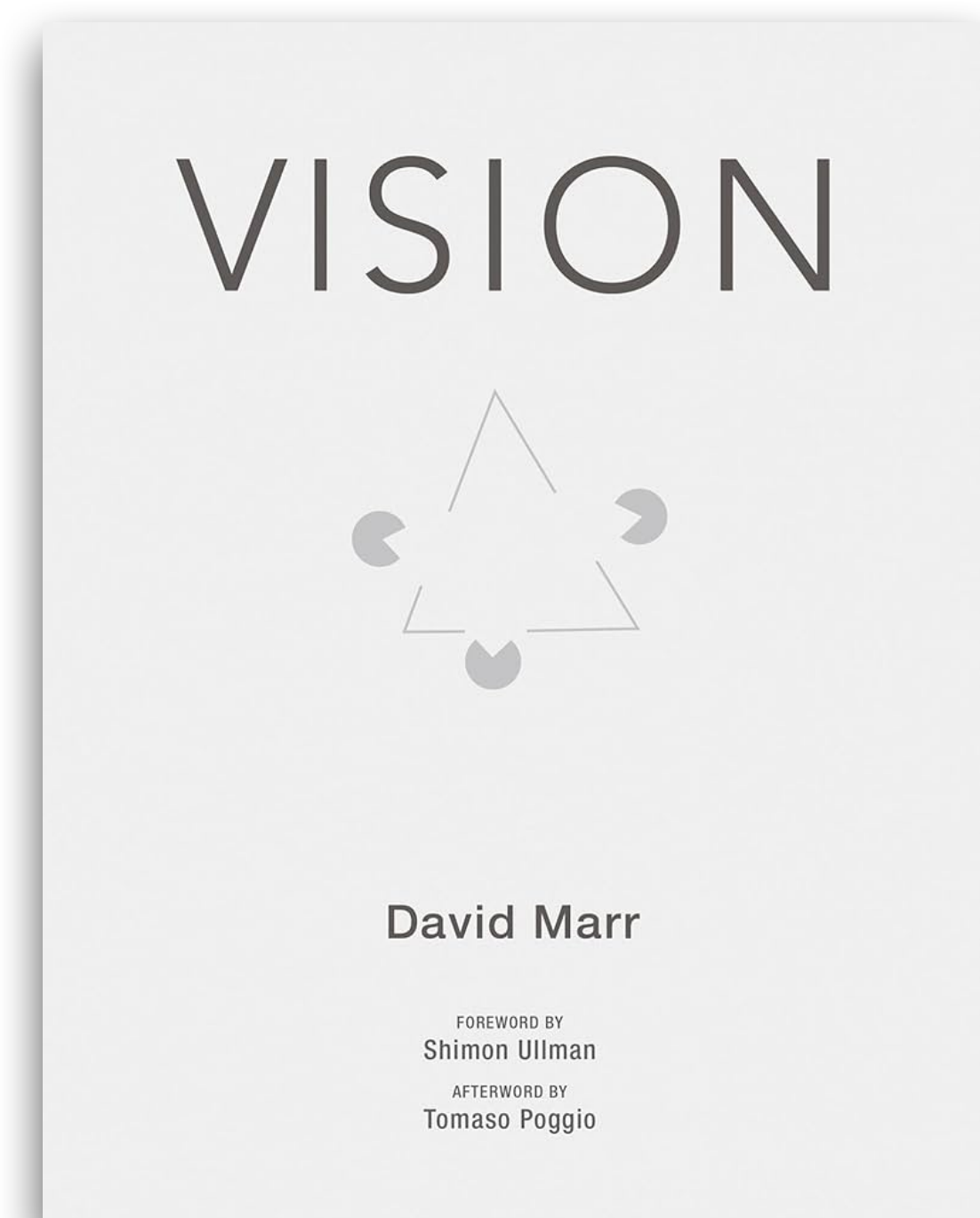


Q&A Discussion

Appendices

Revisiting Marr's Levels of Analysis

Multiple facets for studying information processing systems



[Marr, 1982]

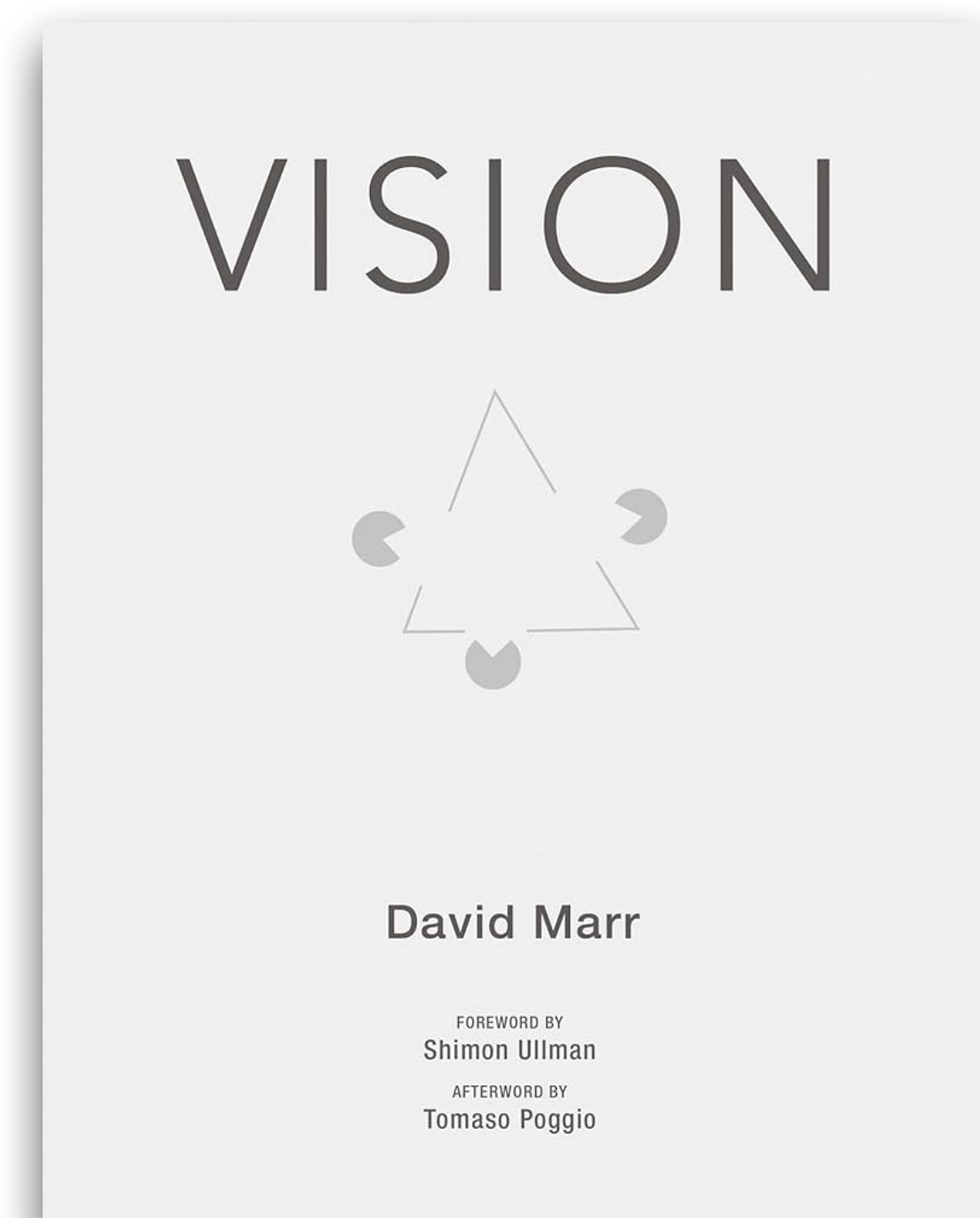
Revisiting Marr's Levels of Analysis

Multiple facets for studying information processing systems



- **Computational level** → [Theoretical linguistics]
 - [What, why] *Competence, mental grammar, abstract linguistic knowledge*

[Marr, 1982]

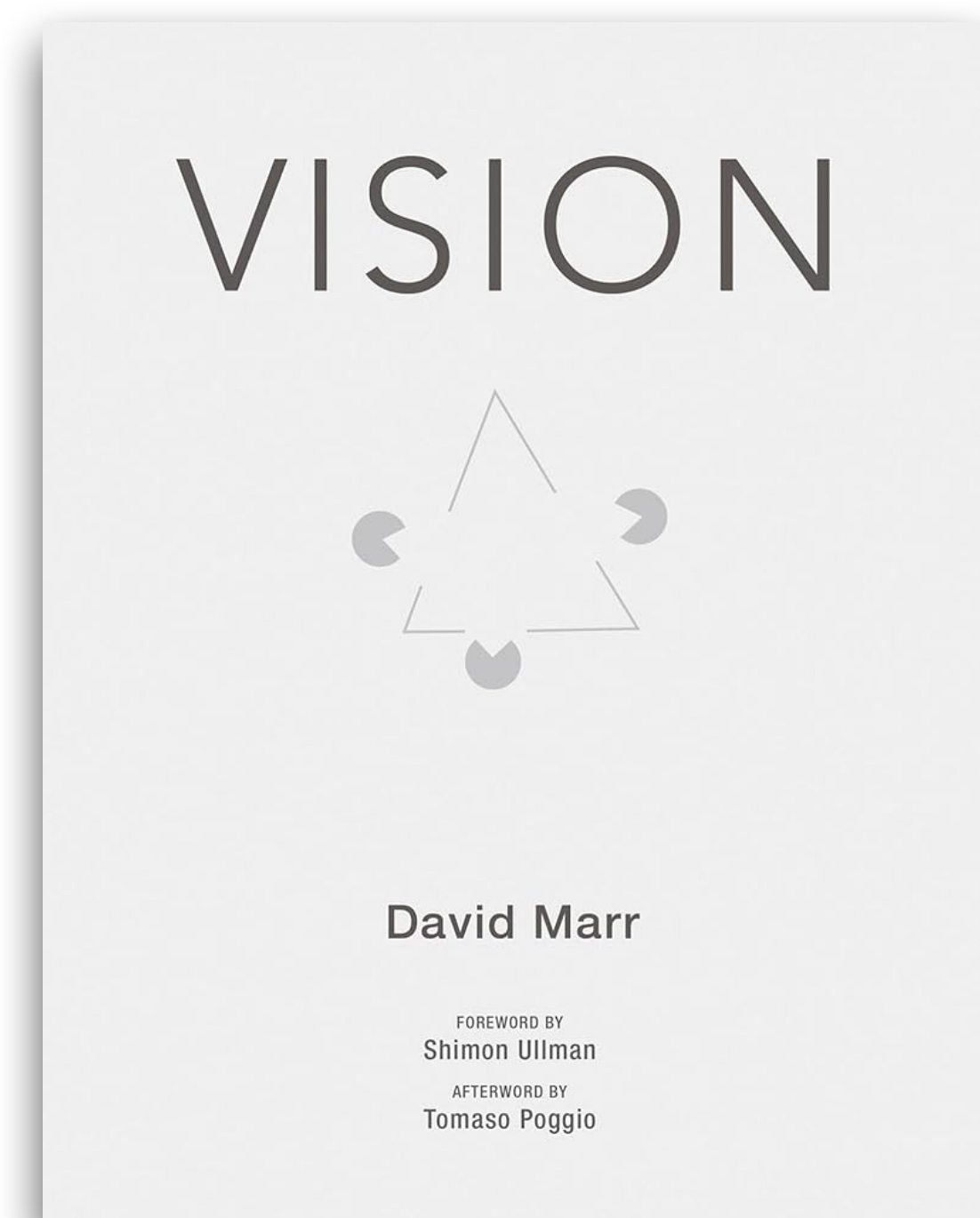


Revisiting Marr's Levels of Analysis

Multiple facets for studying information processing systems



- **Computational level** → [Theoretical linguistics]
 - [What, why] *Competence, mental grammar, abstract linguistic knowledge*
- **Algorithmic level** → [Psycholinguistics]
 - [How] *Performance, real-time processing*



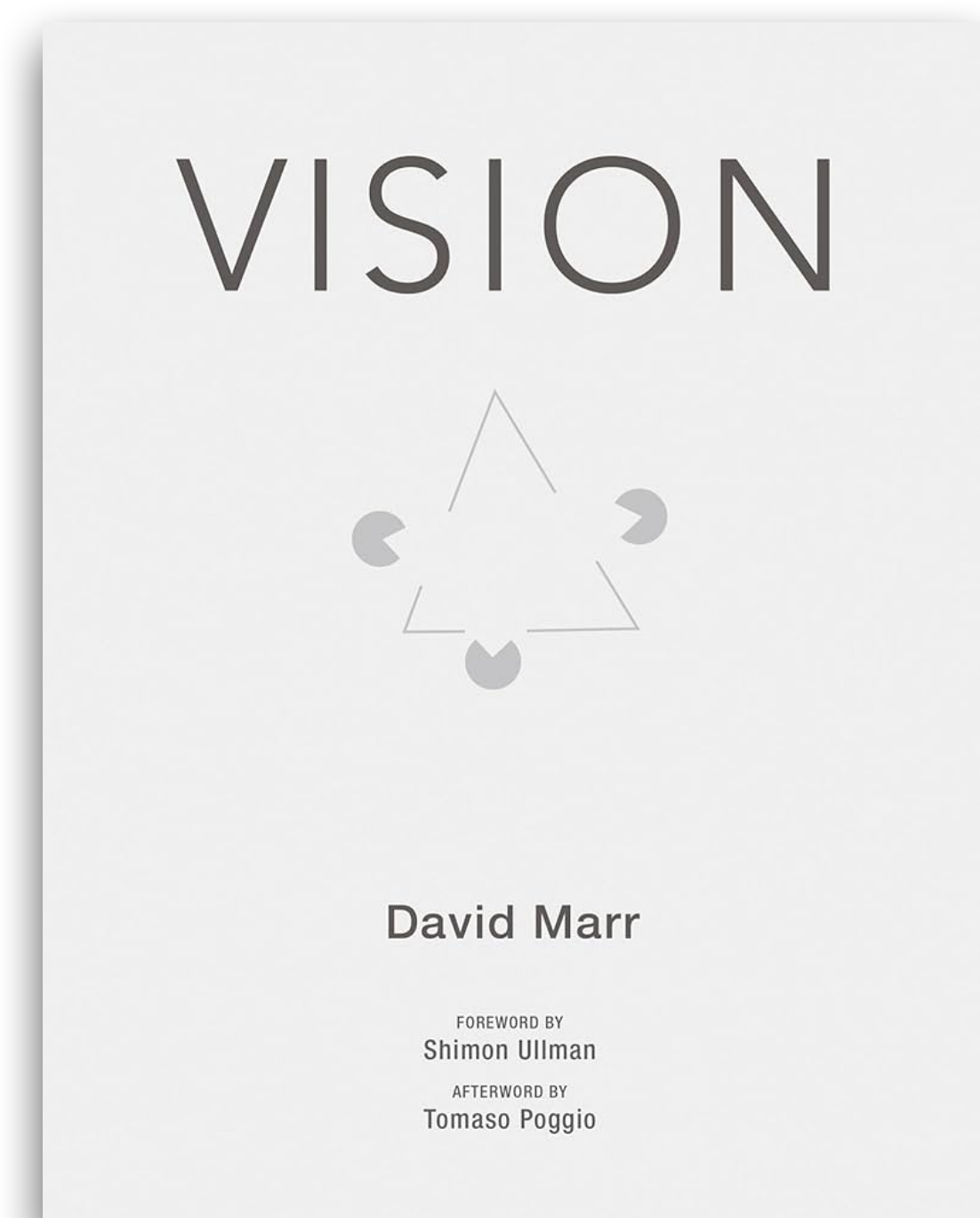
[Marr, 1982]

Revisiting Marr's Levels of Analysis

Multiple facets for studying information processing systems



- **Computational level → [Theoretical linguistics]**
 - [What, why] *Competence, mental grammar, abstract linguistic knowledge*
- **Algorithmic level → [Psycholinguistics]**
 - [How] *Performance, real-time processing*
- **Implementational level → [Neurolinguistics]**
 - [How - lower level] *Performance, neural instantiations of language*



[Marr, 1982]

Current Proposal: The Full Picture

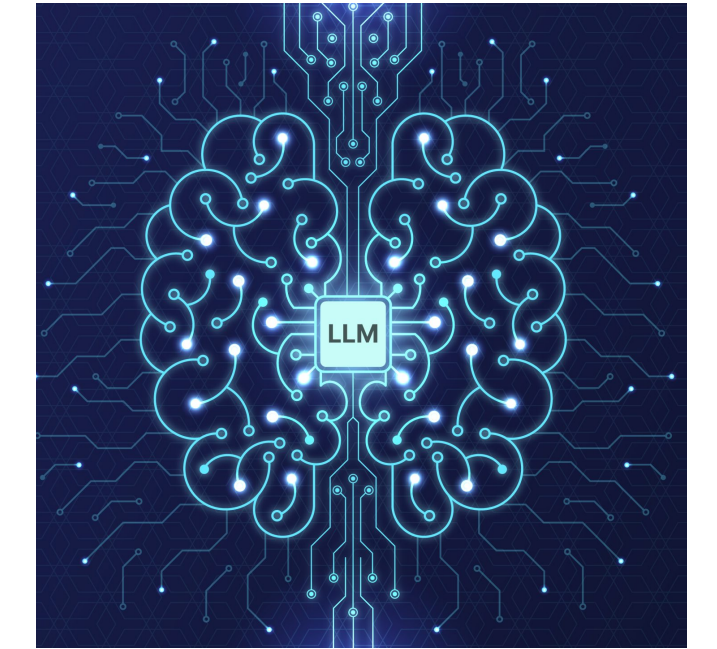
Current Proposal: The Full Picture

Computational
Level

Current Proposal: The Full Picture

**Computational
Level**

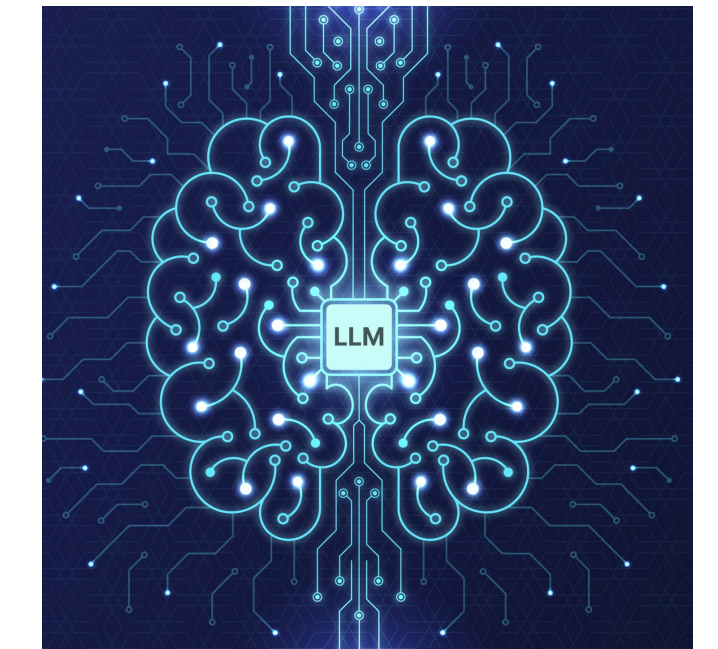
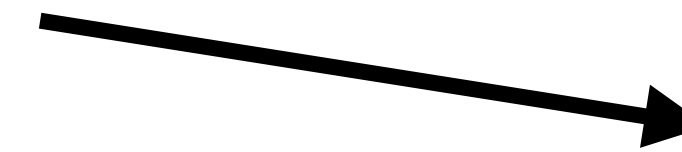
Language in the broader design of intelligent systems



Current Proposal: The Full Picture

**Computational
Level**

Language in the broader design of intelligent systems



**Algorithmic
Level ONE**

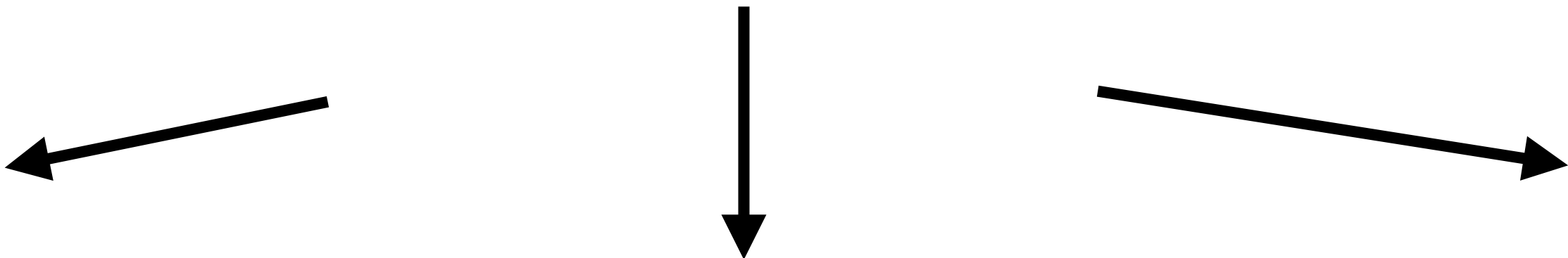
Current Proposal: The Full Picture

Computational
Level

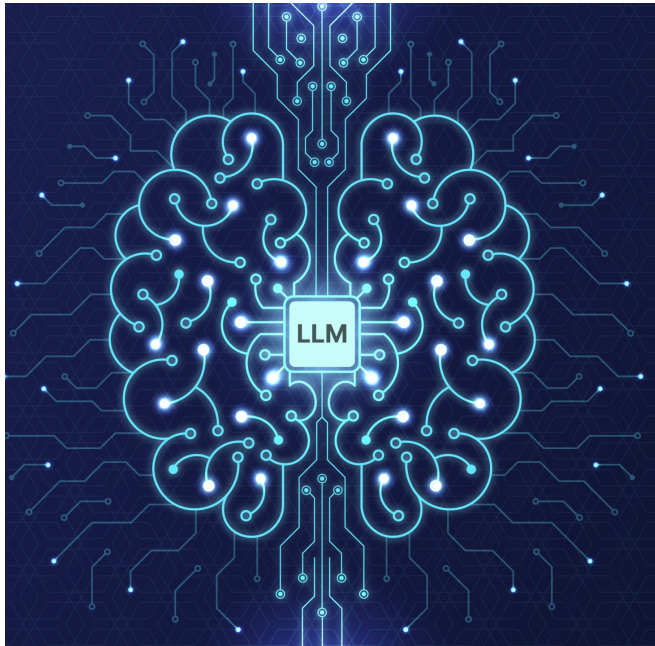
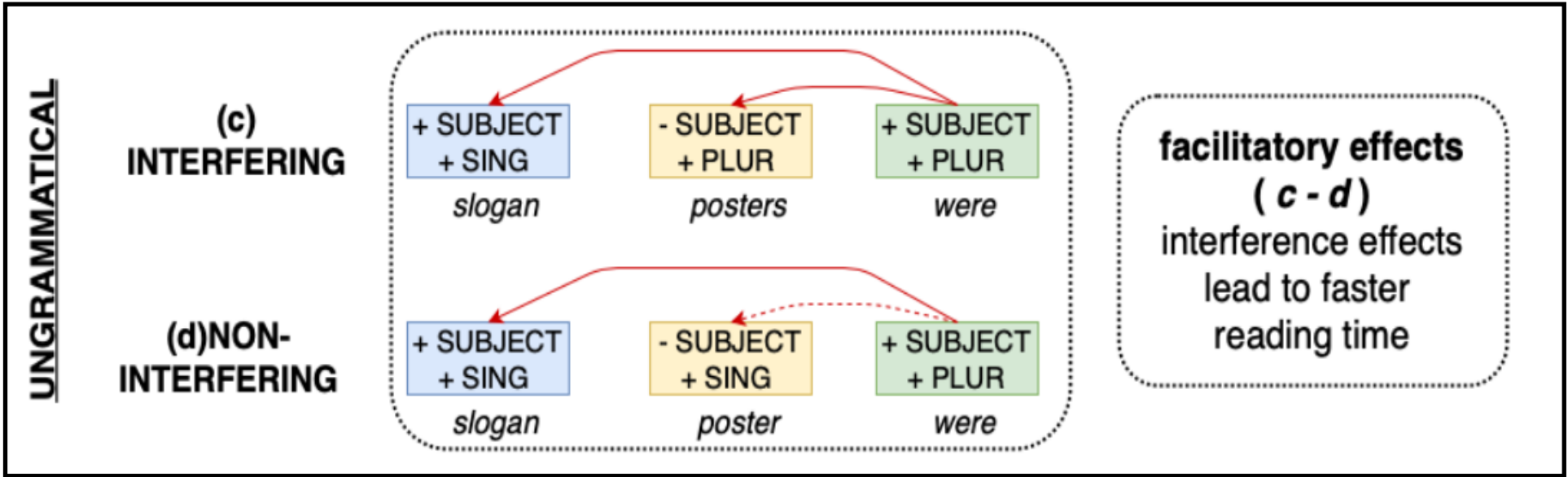
Language in the broader design of intelligent systems



Algorithmic
Level ONE



Psycholinguistic theories (?)



An “algorithm” of sentence processing

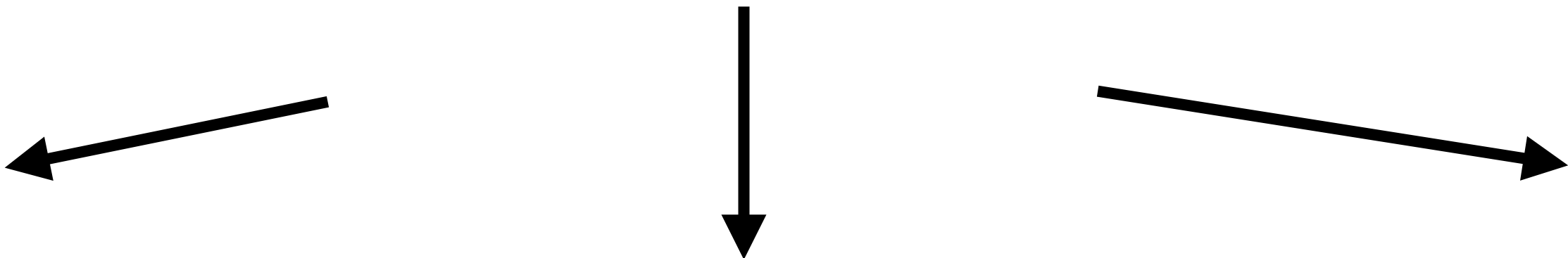
Current Proposal: The Full Picture

Language in the broader design of intelligent systems

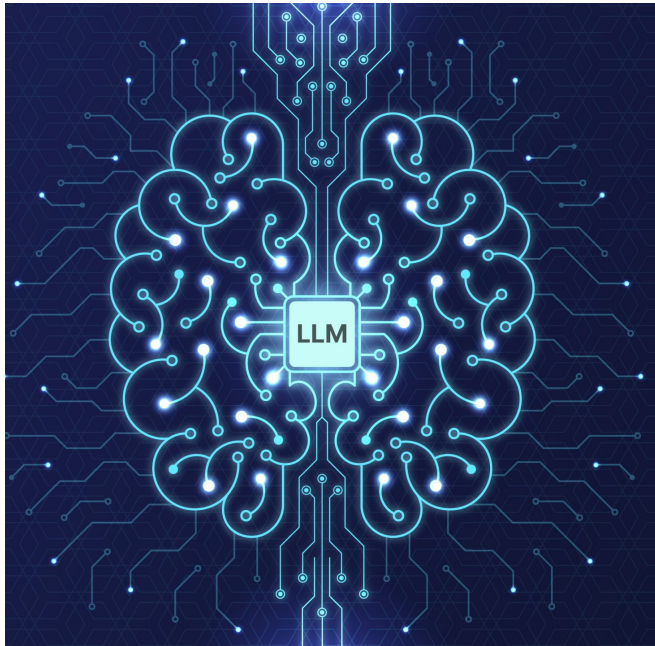
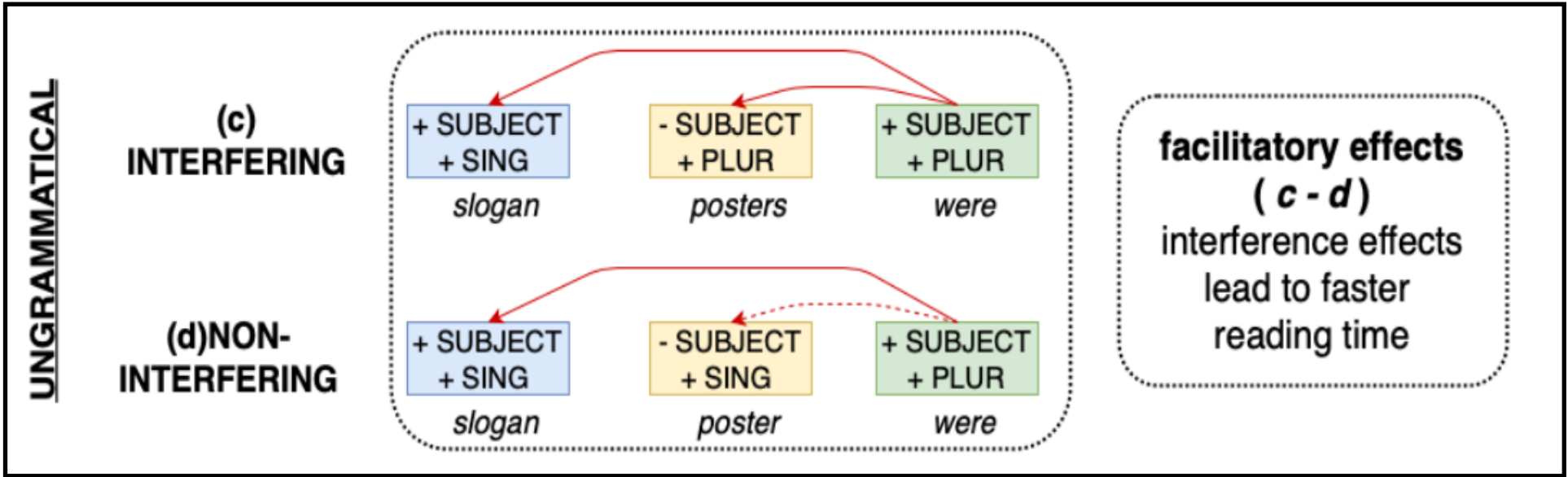
Computational
Level



Algorithmic
Level ONE



Psycholinguistic theories (?)



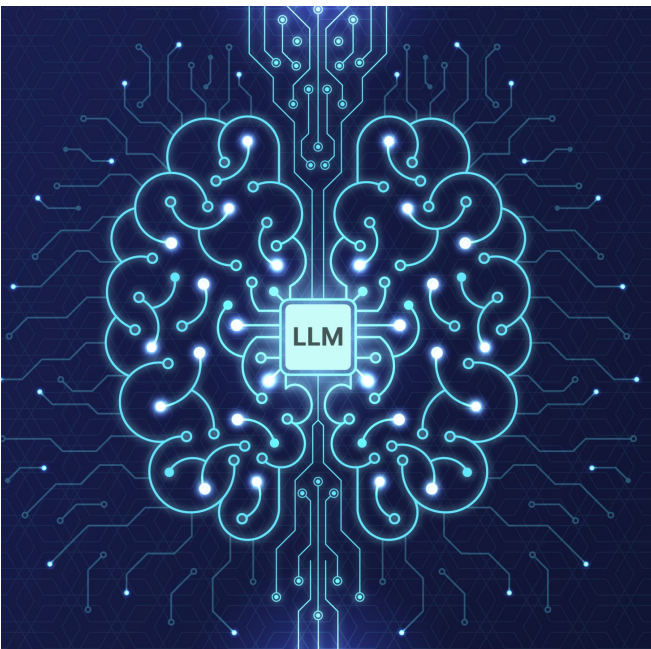
Algorithmic
Level TWO

An “algorithm” of sentence processing

Current Proposal: The Full Picture

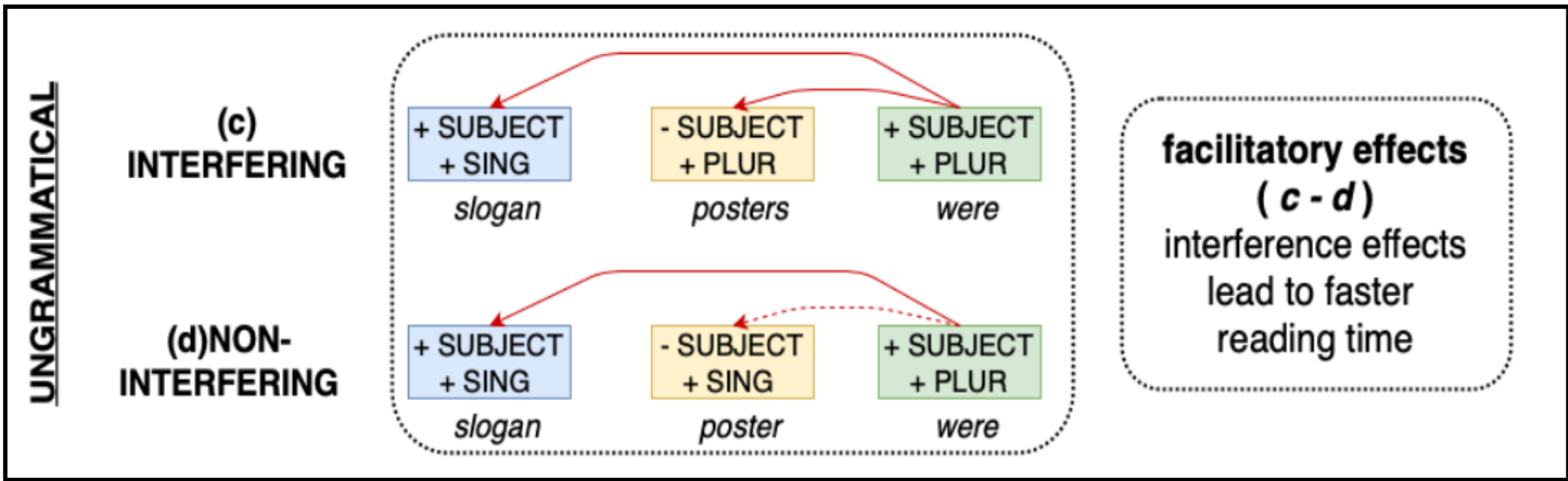
Computational
Level

Language in the broader design of intelligent systems



Algorithmic
Level ONE

Psycholinguistic theories (?)

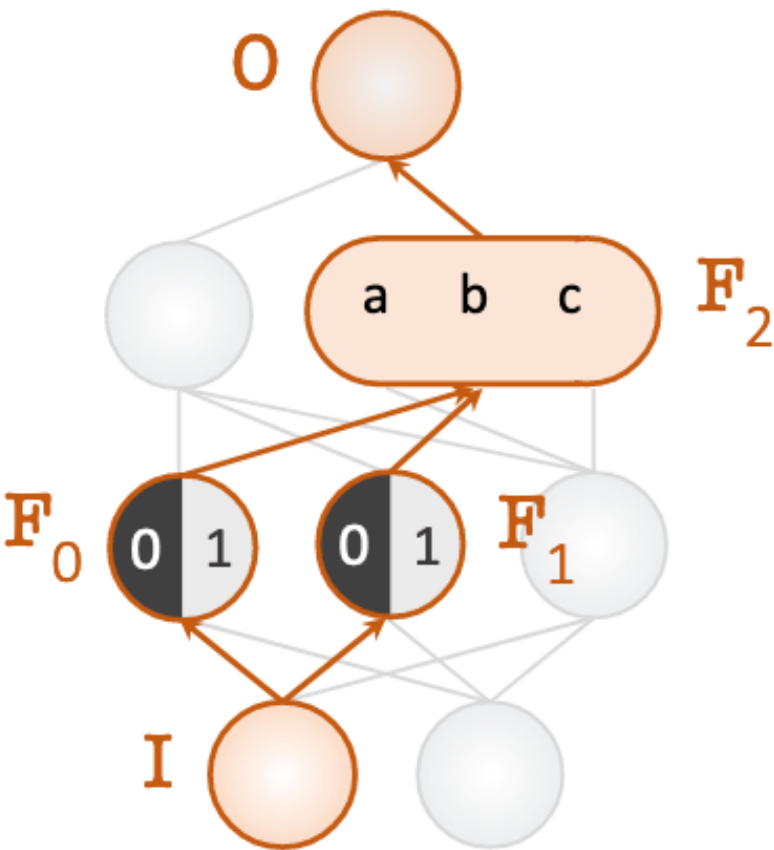
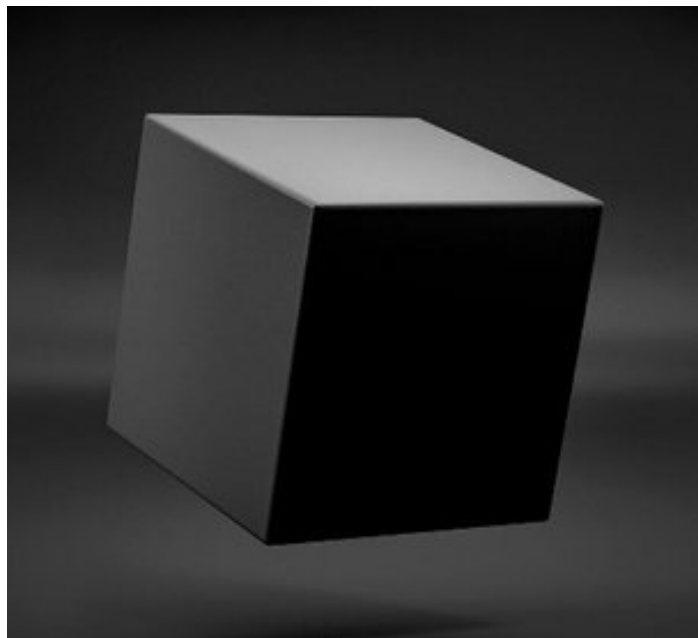


Functionally
approximating

Functionally
approximating

Algorithmic
Level TWO

An “algorithm” of sentence processing



LMs are shaped by what they are trained on!

Data Distributional Properties Drive Emergent In-Context Learning in Transformers

Stephanie C.Y. Chan
DeepMind

Adam Santoro
DeepMind

Andrew K. Lampinen
DeepMind

Jane X. Wang
DeepMind

Aaditya K. Singh
University College London

Pierre H. Richemond
DeepMind

James L. McClelland
DeepMind, Stanford University

Felix Hill
DeepMind

Burstiness as a distributional property!

Compositional Structures!

Rethinking the Role of Demonstrations: What Makes In-Context Learning Work?

Sewon Min^{1,2} Xinxu Lyu¹ Ari Holtzman¹ Mikel Artetxe²
Mike Lewis² Hannaneh Hajishirzi^{1,3} Luke Zettlemoyer^{1,2}

Parallel Structures in Pre-training Data Yield In-Context Learning

Yanda Chen¹ Chen Zhao^{2,3} Zhou Yu¹ Kathleen McKeown¹ He He²

Parallelism!

⇒ The IFE as a diagnostic of the error-driven learning
mechanism in ICL!