

Geometry-Adaptive Explainer for Faithful Dictionary-Based Interpretability under Distribution Shift



Sungjun Lim¹, Heedong Kim¹, Andrew Lee², Kyungwoo Song¹

¹Yonsei University, ²Harvard University



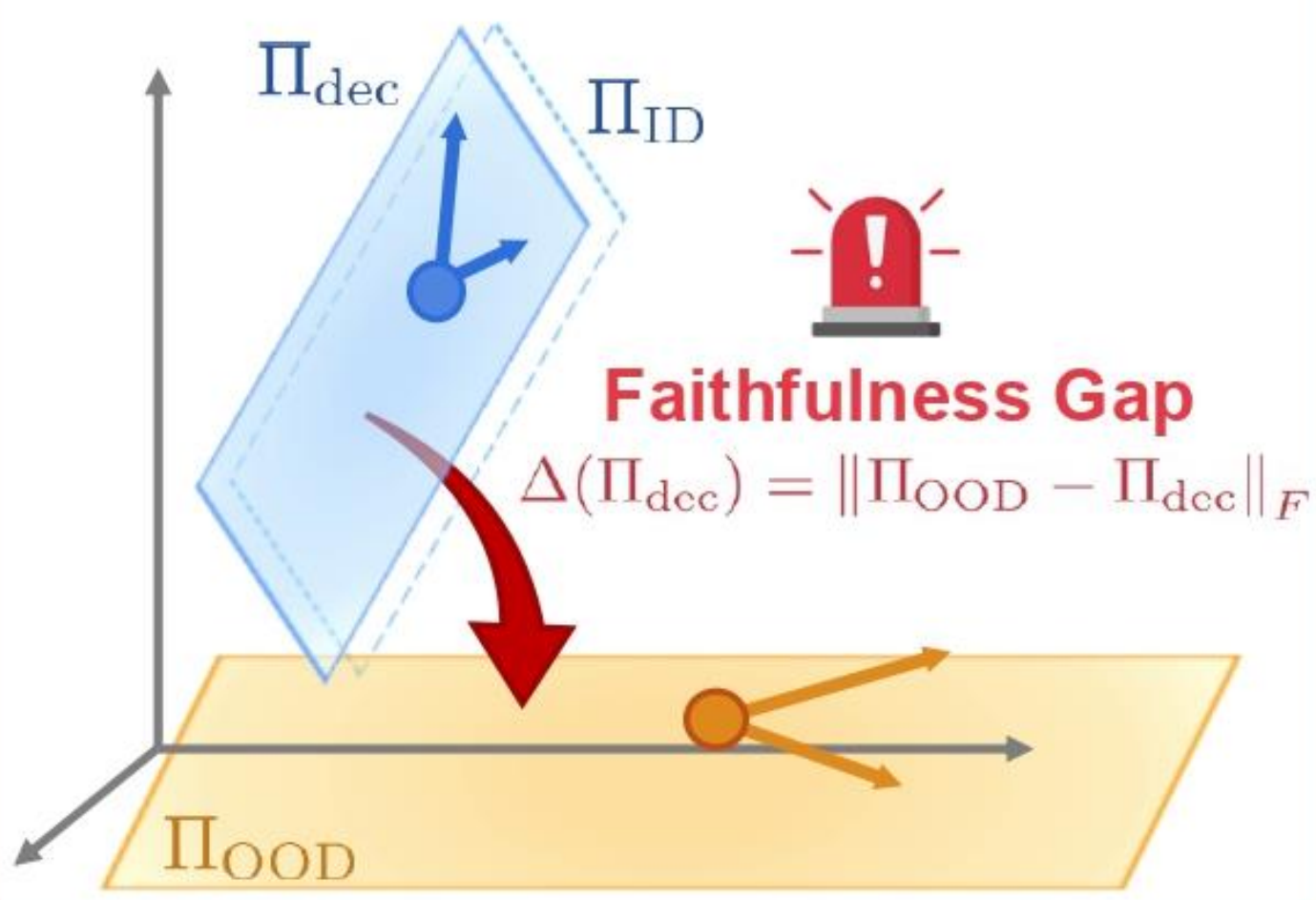
“The capital of France is Paris.”

Distribution shift

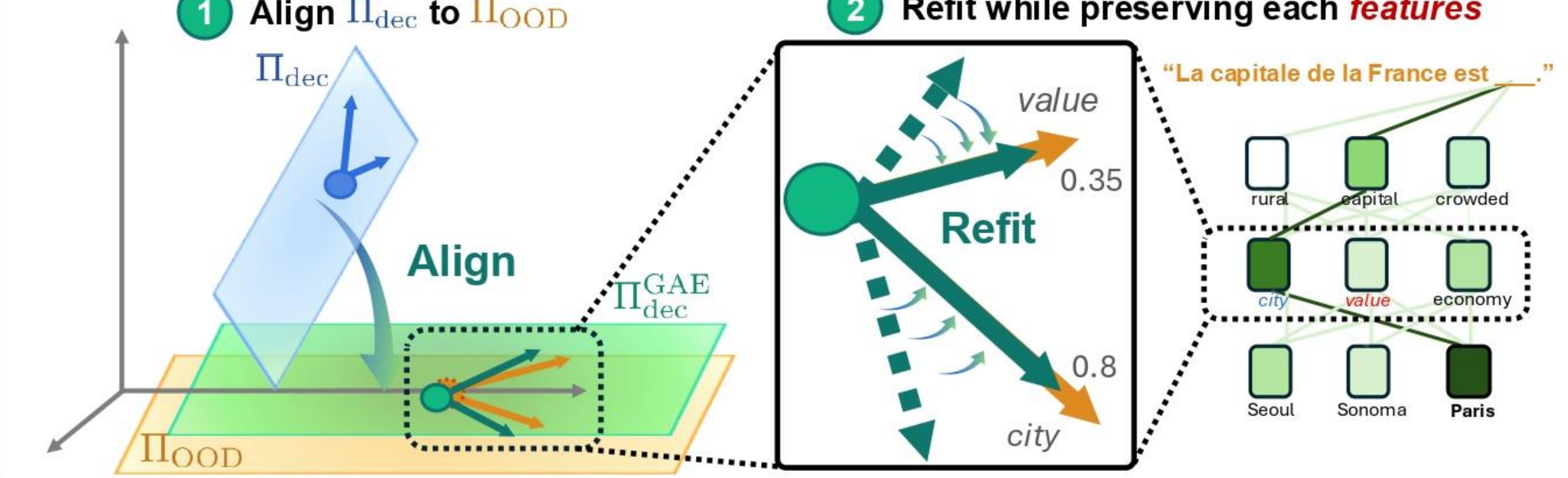


“La capitale de la France est Paris.”

! OOD Shift opens a **Faithfulness Gap**



✓ GAE: Realign while preserving features



■ Π_{dec} ID-trained explainer subspace ■ Π_{ID} ID-active subspace ■ Π_{OOD} OOD-active subspace ■ Π_{dec}^{GAE} GAE-adapted explainer subspace

OOD shift moves model’s active subspace, breaking fixed dictionary explainers. GAE closes the gap by realigning and refitting them *without retraining*.

1 Problem: OOD shift opens a faithfulness gap

Why Do Explainers Fail Under OOD Shift?

- ID-trained explainer’s dictionary (subspace):
- Under OOD shift, OOD-active subspace moves:

$$\Pi_{dec} \approx \Pi_{ID}$$

$$\Pi_{ID} \rightarrow \Pi_{OOD}$$

- This creates a **faithfulness gap**:

$$\Delta(\Pi_{dec}) = \|\Pi_{OOD} - \Pi_{dec}\|_F$$

Why does this gap matter?

- Explainer’s OOD loss is decomposed to irreducible and alignment error (explainer-dependent).

$$\mathcal{L}_{OOD}(\Pi_{dec}) = \mathcal{L}_{OOD}(\Pi_{OOD}) + \mathcal{E}_{align}(\Delta)$$

$$\mathcal{E}_{align}(\Delta) \asymp \Delta(\Pi_{dec})^2$$

- Reducing the faithfulness gap targets the fixable part of OOD faithfulness loss.**

3 Method: Geometry-Adaptive Explainer

Step 1: Align Subspace

- Estimate OOD-active subspace:

$$\hat{\Pi}_{OOD} \leftarrow \text{TopEig}(\frac{1}{N} \sum_i h_i h_i^\top)$$

- Rotate ID-aligned decoder by Procrustes alignment:

$$\tilde{W}_{dec} \leftarrow \text{ProcAlign}(W_{dec}^{ID}, \hat{\Pi}_{OOD})$$

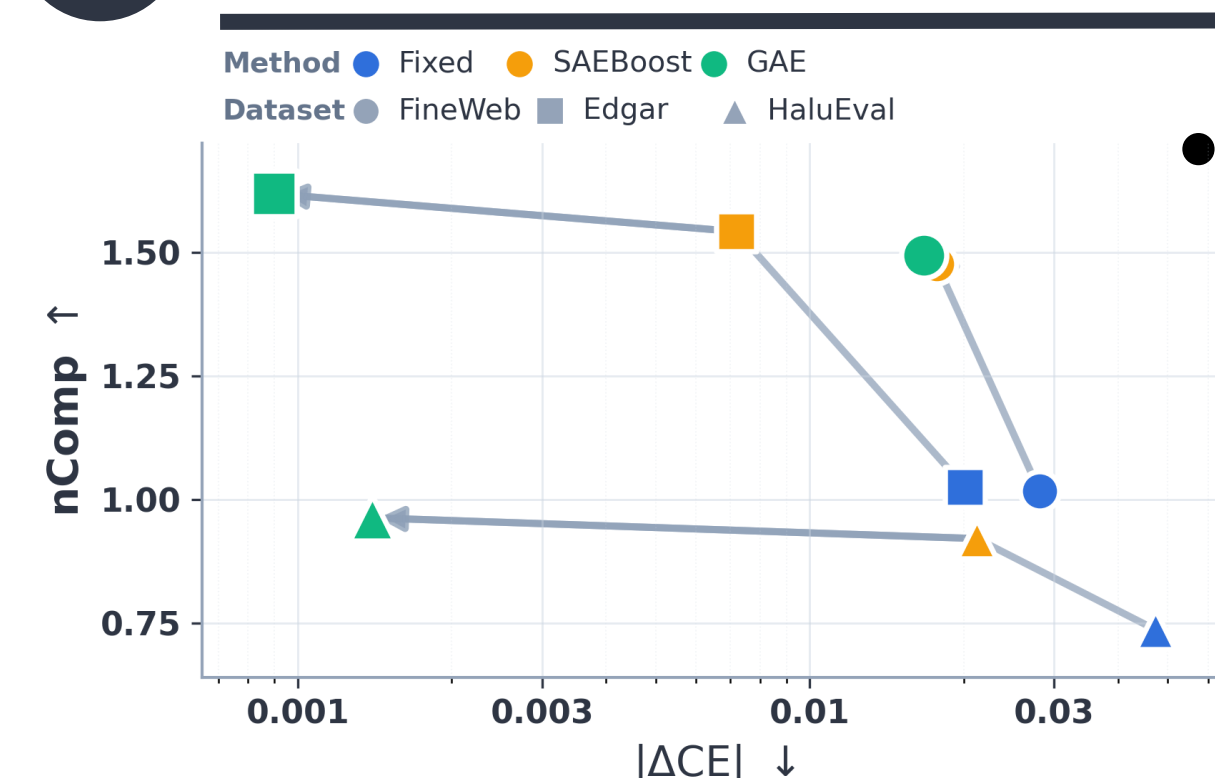
- Close the faithfulness gap while staying close to the original dictionary.**

Step 2: Refit Decoder

- Alignment fixes the subspace; Step 2 **preserves feature identities** while refitting the decoder.
- Freeze encoder: $z_i = \text{Enc}_{\text{frozen}}(h_i)$
- Closed-form constrained ridge:

$$W_{dec}^{GAE} \leftarrow \text{RidgeFit}(\{(h_i, z_i)\}_{i=1}^N; \hat{\Pi}_{OOD}, \tilde{W}_{dec})$$

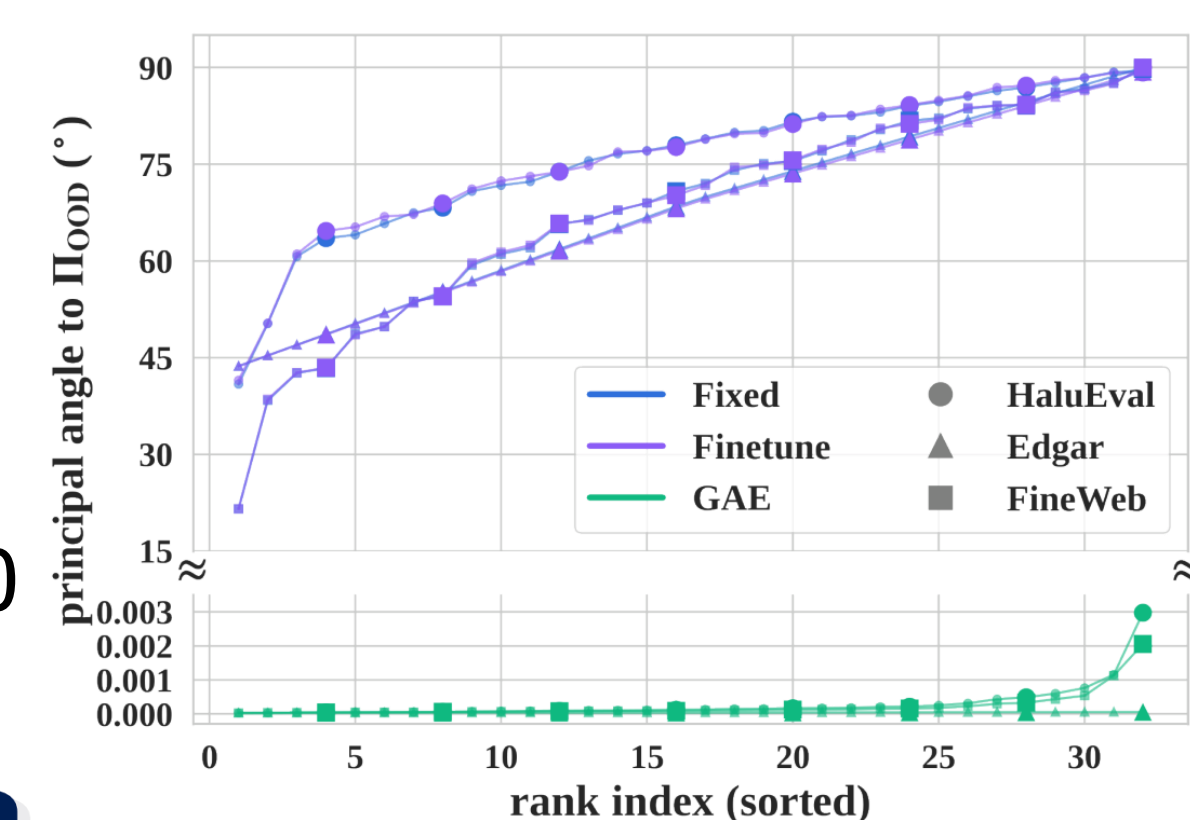
4 Results



- GAE raises faithfulness (nComp) and lowers reconstruction error (|Delta CE|) across all shifts.**

- GAE closes subspace gap:**

- Fixed/Finetune stay at 40-90 deg,
- while GAE aligns near 0 deg.



Method	Adapt. signal	Grad.	Wall-clock
Finetune	5M OOD tokens	○	2-12 min
Retrain SAEBoost FaithfulSAE	100M OOD tokens	○	39 min-4 hrs
GAE	2K OOD activations	✗	0.5-2.9 sec

- GAE adapts
 - from **2K activations**
 - without gradients**
 - in **0.5-2.9 sec**.

5 Limitations & Discussions

- Top-r SVD truncation may drop feature information in residual singular directions.
- Natural next step: adaptive rank selection, encoder adaptation, and OT/Grassmannian alignment.



Project