



Ph.D. Dissertation Defense

Paraphrase Sensitivity in Medical Vision-Language Models

What medical vision-language models get wrong when a clinician rewords the question, why, and what to do about it

Presenter

Binesh Sadanandan

Ph.D. Candidate

bsada1@unh.newhaven.edu



Thesis Website

<https://bineshkumar.me/phd-thesis/>

Advisor

Vahid Behzadan, Ph.D.

Associate Professor of CS

Advisor and Committee Chair

vbehzadan@newhaven.edu

This work is intended solely for research purposes. We present a systematic audit of medical vision-language models that reveals how phrasing variations and deployment gating can mask text-driven selection, posing serious risks in clinical settings.

The purpose of exposing this vulnerability is not to promote harm, but to enable the research community and model developers to address these critical safety issues. By identifying how these failures occur, we can contribute to the development of more robust and reliable medical AI systems.

Our research shows that current selective prediction mechanisms may be insufficient for clinical deployment. We believe transparent reporting of these findings leads to greater benefit for the field and allows clinical institutions to make informed decisions about AI deployment.

All experiments were conducted on publicly available datasets (MIMIC-CXR, PadChest, VinDR-CXR-VQA, NIH Chest X-ray) following established ethical guidelines for AI safety research. No patient-identifying information was accessed beyond the de-identified public releases.

Responsible disclosure: all findings are, under review or accepted with full code and evaluation artifacts.

Show a Medical Vision Language Model (VLM) the exact same chest X-ray, ask the exact same clinical question two different ways, and the answer can reverse

1

PSF-Med: audited benchmark

26,850 clinical questions, 92,856 equivalence-filtered pairs, six models, three countries, clinician-audited and public

2

Consistency is not grounding

Image-reliance controls separate text-driven answers from image-grounded ones

3

PSF has a locatable model specific mechanism and an identified cause

The answer commits at layer 16 - the training-phrasing distribution causes the sensitivity

4

A cheap targeted fix, audited against its own claim.

LoRA on 0.1% of parameters cuts flips 59% with accuracy held

5

Safety screen for Deployment

Paraphrase consistency crossed with image reliance gives four cells; the Dangerous cell (consistent + text-driven) dominates most models.

6

A deployment signal that costs one forward pass.

Predictive entropy predicts flips across two architecture families; grounding still needs a second pass

PER YEAR

2B+

chest radiographs acquired worldwide, the most common diagnostic image in medicine.

US WORKFORCE

5% vs 2%

annual imaging-volume growth against radiology residency growth. AAMC projects a 17,000 to 42,000 shortfall by 2033.

MISS RATE

4-30%

senior radiologist output under pressure. Miss rates for subtle findings

AI is being positioned to close the gap, and chest radiography is where it lands first

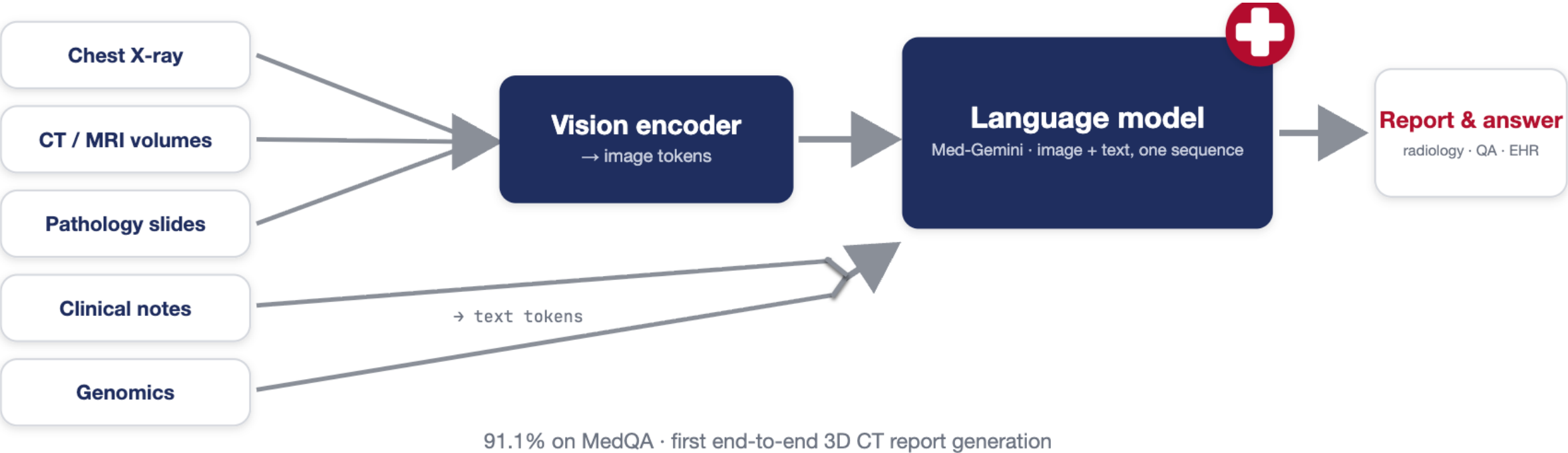
Medicare approves new technology add-on payment for inpatient radiology AI solution

Marty Stempniak | August 14, 2026 | [Radiology Business](#) | [Artificial Intelligence](#)



Sources: AAMC physician workforce projections (17,000-42,000 radiologists, pathologists and psychiatrists by 2033); Siemens Healthineers radiology workforce report; Akhter et al., Front. Big Data 2023
Radiology Business, "Medicare Approves NTAP for Inpatient Radiology AI," Aug. 2026..

From Language Modeling to the Medical Vision Language Model



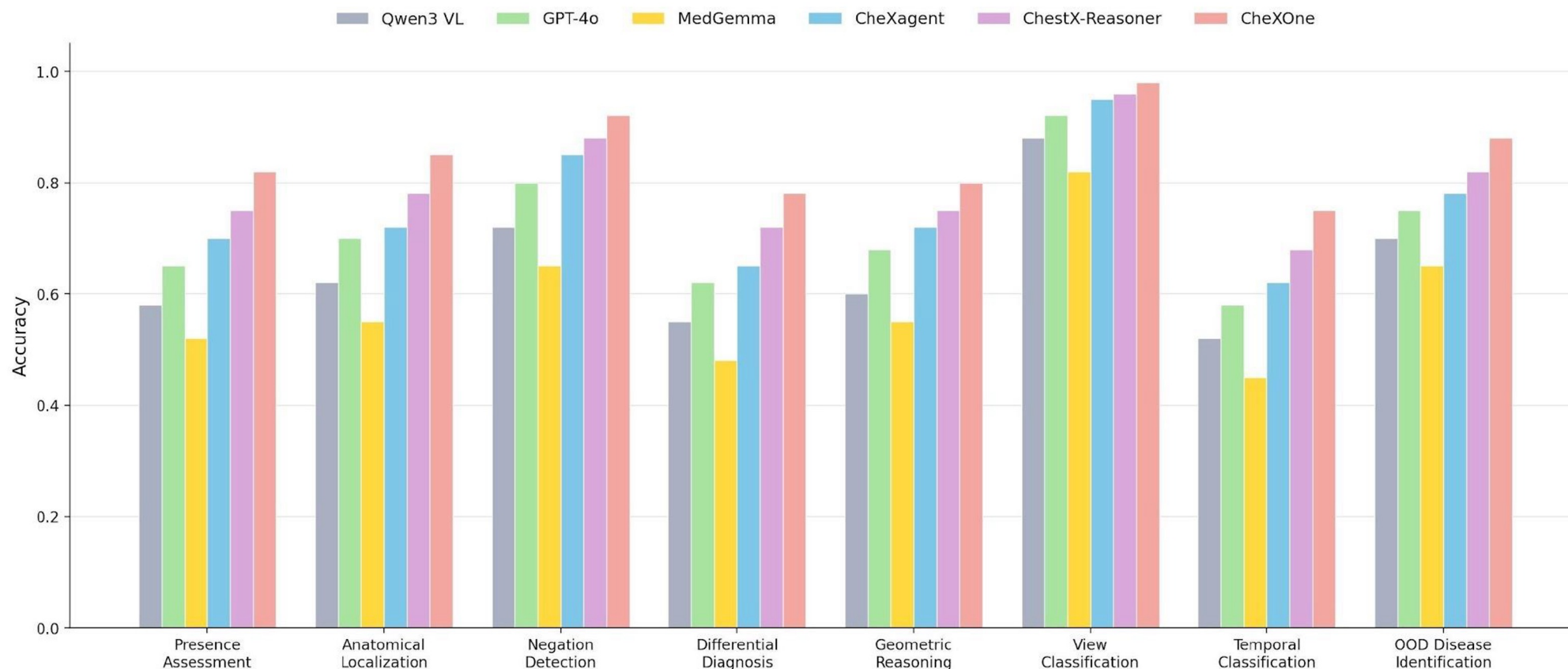
Medical models became **broadly multimodal**.

Saab et al., 2024 (Med-Gemini)



Leading Domain-Adapted CXR Models

400+
Open weight Models in Hugging
face*

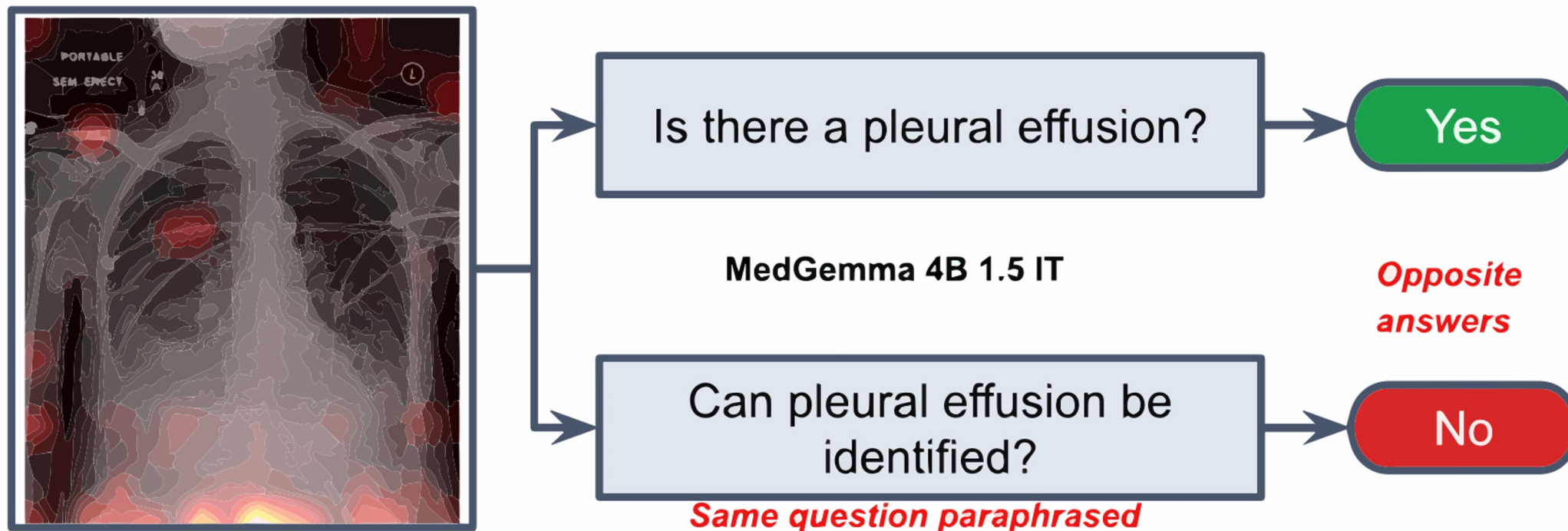


Source : Zhang, Y. et al. (2026). A Reasoning-Enabled Vision-Language Foundation Model for Chest X-ray Interpretation.

The benchmark numbers tell a promising story and *these scores come from one fixed phrasing per question*

*Source: Hugging Face Model Hub (query: "chest x-ray", retrieved August 2026, results 443 base)

Paraphrase-Sensitive Failure



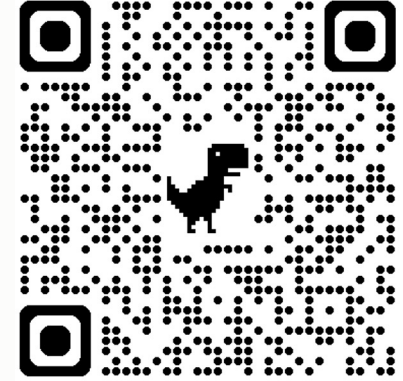
Source: Sadanandan et al., "PSF-Med," CVPRW Multimodal Foundation Models for Biomedicine (MMFM-BIOMED), 2026.

~~MedGemma, Paraphrased Question, Different Answer~~

*Source: Hugging Face Model Hub (retrieved August 23rd 2026, 3 models)

Same X-ray. Same question, asked twice.

GPT-5.6 Sol — OpenAI's flagship model



Source: Sadanandan et al., "PSF-Med," CVPRW Multimodal Foundation Models for Biomedicine (MMFM-BIOMED), 2026.

Even the frontier models show Paraphrase Sensitive Failure

Two questions are semantically equivalent for a clinician. The model answers them differently on the same image

Synonym · medical jargon

Q₁ *Is there cardiomegaly?*

Q₂ *Does the image show an enlarged cardiac silhouette?*

Negation · question polarity

Q₁ *Is there a pneumothorax?*

Q₂ *No pneumothorax, correct?*

Scope · word order

Q₁ *Any nodule in the left lower zone?*

Q₂ *In the lower-left region, is a nodule visible?*

Hedging · epistemic register

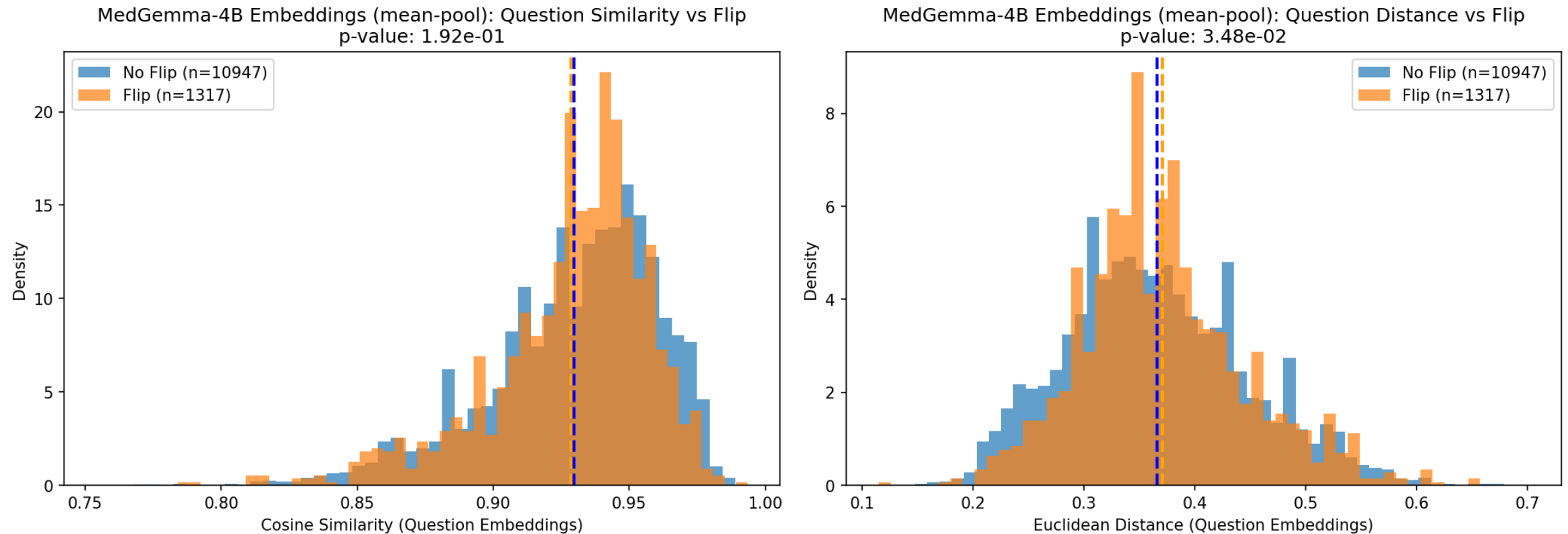
Q₁ *Is the lung clear?*

Q₂ *Would you call this lung clear?*

Source: Sadanandan et al., "PSF-Med," CVPRW Multimodal Foundation Models for Biomedicine (MMFM-BIOMED), 2026.

Semantic invariance is a hard requirement. If $Q_1 \equiv Q_2$ clinically, answers must match

How Similar are Paraphrases ?



Source: Sadanandan et al., "PSF-Med," CVPRW Multimodal Foundation Models for Biomedicine (MMFM-BIOMED), 2026.

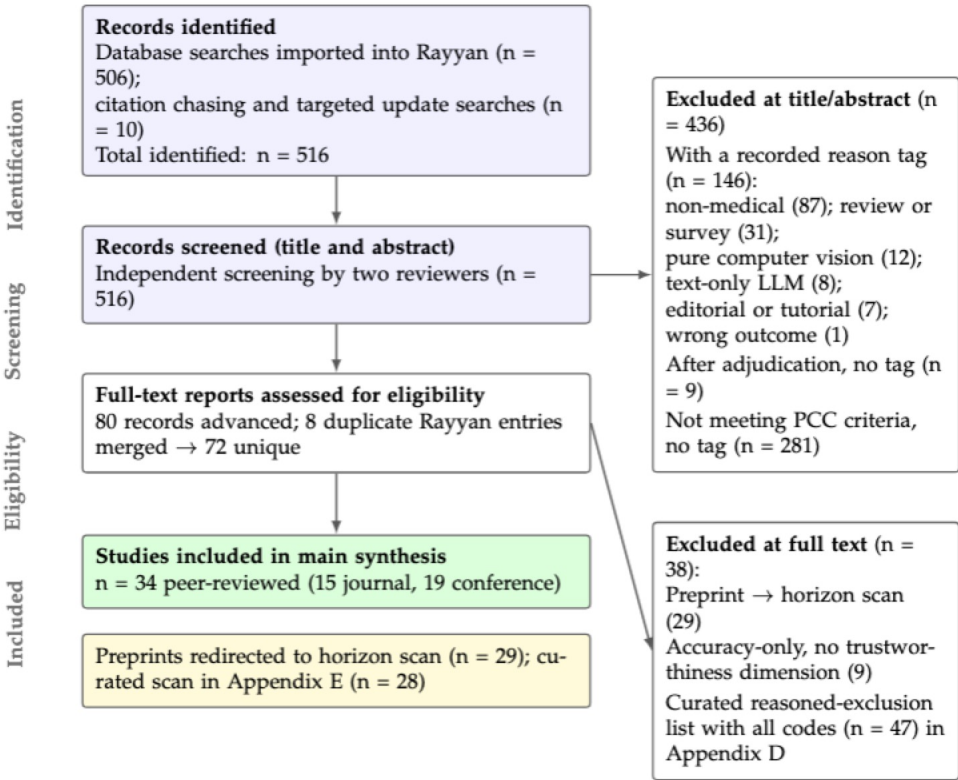
Clinicians naturally rephrase questions; a model that changes its answer under rephrasing cannot be trusted in deployment

Trustworthiness Evaluation of Medical Vision-Language Models: A Scoping Review of Robustness, Grounding, Hallucination, and Uncertainty

Binesh Sadanandan^{1,*} Amirhossein Karimi¹ Bibek Upadhyay² Vahid Behzadan¹

¹SAIL Lab, University of New Haven, West Haven, CT, USA

²Saginaw Valley State University, MI, USA *Corresponding



Method (PRISMA-ScR)

- ❖ 516 records screened → 34 trustworthiness evaluations of medical VLMs included
- ❖ 8 trustworthiness dimensions mapped: datasets, metrics, and controls cataloged per study

Source: Sadanandan et al., "Trustworthiness Evaluation of Medical VLMs," JMIR AI (Under Review), 2026.

1 of 34

applies any **language-side (paraphrase)** perturbation

1 of 34

reports a standard **calibration metric**

0

use **conformal or selective prediction**

0

evaluate **deployment safeguards** or leakage checks

0

apply **mechanistic interpretability** or test faithfulness

4 of 34

stratify results by **demographic subgroup**

11 of 34

report **external or out-of-distribution** validation

1

purpose-built trustworthiness dataset; all others repurpose accuracy benchmarks

7

studies propose **any mitigation**

7 of 34

release their data (18 of 34 release code)

n = 34 included studies

Source: Sadanandan et al., "Trustworthiness Evaluation of Medical VLMs," JMIR AI (Under Review), 2026. As of March 25th 2026

Output: the MiTEC checklist

- An 8-item Minimum Trustworthiness Evaluation Checklist for medical VLM

MiTEC: Minimum Trustworthiness Evaluation Checklist for Medical VLMs

1. **Text-only baseline.** Report model performance when the image is withheld. If the text-only baseline approaches image-conditioned performance, the model may not use visual information.
2. **Image perturbation or swap test.** Replace the target image with an unrelated image from the same modality. A model that maintains its answer has not grounded its response in the image.
3. **Prompt paraphrase consistency.** Evaluate on at least 3 clinically equivalent rephrasings of each query. Report the flip rate (proportion of queries whose answer changes).
4. **Calibration metrics.** Report at least one calibration metric (ECE, Brier score, or area under the generalized risk-coverage curve [AUGRC]) for the model's expressed confidence. If the model does not output confidence scores, apply semantic entropy or consistency-based estimation.
5. **Subgroup stratification.** Stratify results by at least one demographic variable (sex, age, or race/ethnicity where available). Report performance gaps alongside aggregate metrics.
6. **External or out-of-distribution evaluation.** Evaluate on at least one dataset not used in training or development. Report performance degradation relative to the primary evaluation set.
7. **Selective prediction curve.** Report a risk-coverage curve showing how accuracy changes as the model abstains on its least-confident predictions. Report the area under the risk-coverage curve (AURC) or AUGRC.
8. **Reproducibility artifacts.** Release code, evaluation data, and model weights (or provide a clear access mechanism). At minimum, provide a data availability statement.

Items 1–2 address grounding (is the model using the image?). Items 3–4 address reliability (is the model consistent and calibrated?). Item 5 addresses equity (is the model fair?). Item 6 addresses generalization (does the model work beyond its training distribution?). Item 7 addresses deployment readiness (can the model safely abstain?). Item 8 addresses reproducibility.

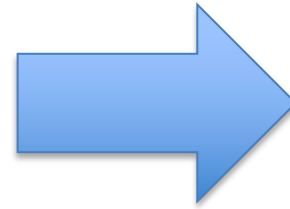
Source: Sadanandan et al., "Trustworthiness Evaluation of Medical VLMs," JMIR AI (Under Review), 2026.

Paraphrase consistency is a named test element.

N-7874

PUBLIC DOCKET

Discussion paper and request for feedback, docket FDA-2026-N-7874. Twenty-six discussion questions, open for public comment.



R.1

ROBUSTNESS ELEMENT

Consistency across semantically equivalent paraphrases, named in Appendix A as a device benchmarking element.

Variation in non-safety-critical phrasing may be acceptable, but variation in safety-critical behaviors such as escalation, refusal, or diagnostic conclusions is treated as a failure

Source: FDA CDRH, GenAI-Enabled Devices Discussion Paper, Appendix A, element R.1 (Aug 2026)

A medical vision-language model can reverse a diagnostic answer when a clinician asks the same question in different words, and standard evaluation reported by model cards miss the reversal

1 · Measure

2 · Diagnose

3 · Explain + fix

4 · Safety-evaluate

5 · Deploy

Publications at Peer Reviewed Venues

Thrust	The question it answers	Peer Reviewed & Accepted At
1 Measure	How often, and how much, do medical VLMs fail?	2 nd Multimodal Foundation Models for Biomedicine @ CVPR 2026
2 Diagnose	Is a consistent answer actually grounded in the image?	Medical Reasoning with VLMs @ CVPR 2026 40% Acceptance Rate
3 Explain + fix	What is the mechanism, and can a targeted fix reduce it?	7th Annual Conference on Health, Inference, and Learning 2026 - 30% Acceptance rate
4 Safety-evaluate	Can we screen a prediction for hidden text-reliance?	8th Interpretability of Machine Intelligence in Medical Image Computing @ MICCAI 2026 34% Acceptance rate
5 Deploy	Is there a single-pass signal to gate at deployment?	7 th Edition of Uncertainty for Safe Utilization of Machine Learning in Medical Imaging @ MICCAI 2026 32% Acceptance rate

THRUST 1 OF 5

Measure

How often, and how badly, do medical vision-language models change its answers on clinically equivalent rephrasings?

RQ1.1 magnitude · RQ1.2 transformation type · RQ1.3 model scale · RQ1.4 distribution shift

Chapter 3 · PSF-Med benchmark · Accepted, MMFM-BIOMED at CVPR 2026

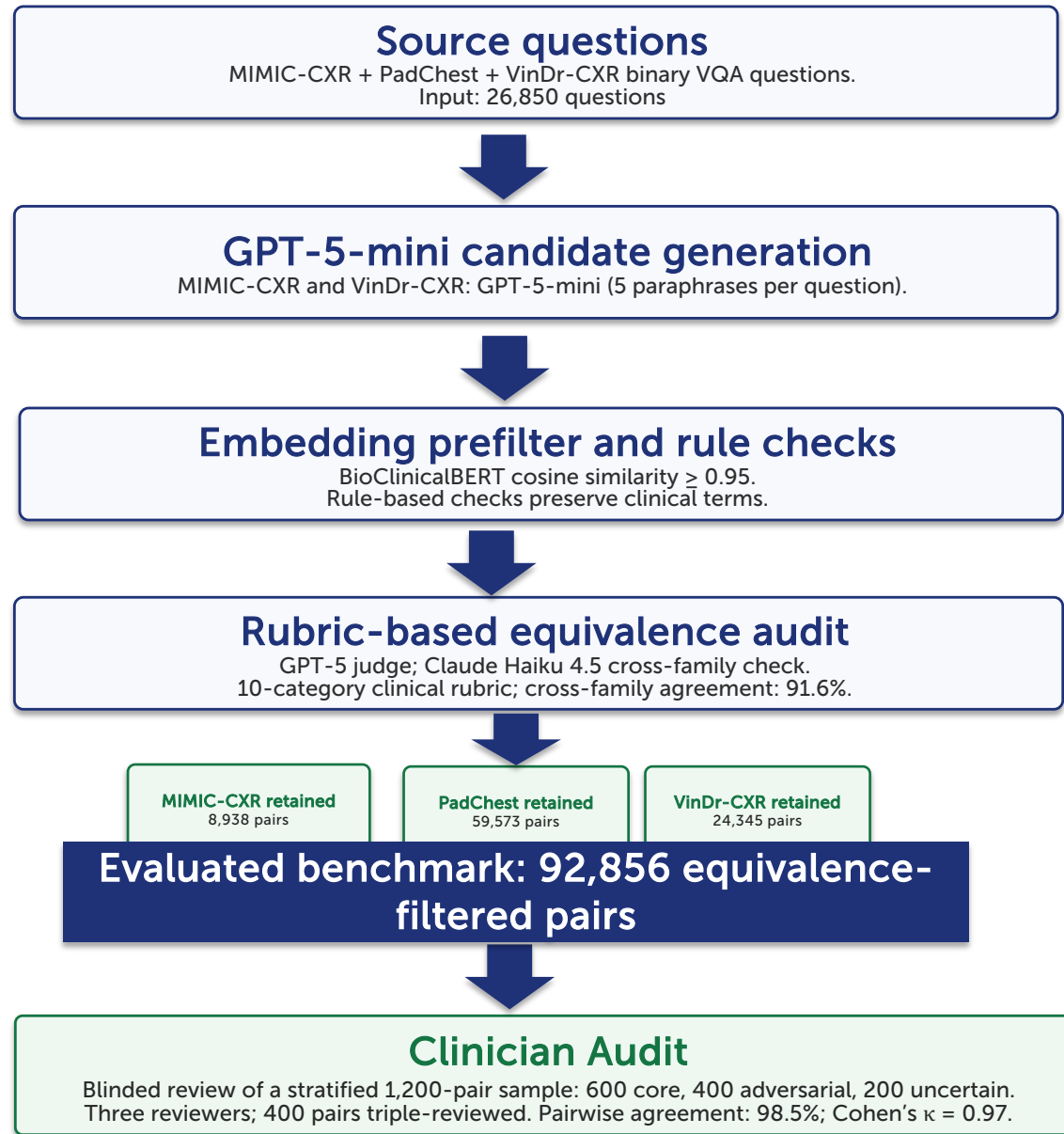
1 · Measure

2 · Diagnose

3 · Explain + fix

4 · Safety-evaluate

5 · Deploy



PSF-Med: A Clinician-Audited Benchmark for Paraphrase Sensitivity
in Medical Vision-Language Models

Binesh Sadanadan^{1*} Vahid Behzadan¹ Lekshmy Jayan² Arun Gopinatha Kurup³

¹SAIL Lab, University of New Haven, West Haven, CT, USA

²Dept. of Radio Diagnosis, Mount Zion Medical College Hospital, Adoor, Kerala, India

³Badr Al Samaa Hospitals, Nizwa, Muscat, Oman

Accepted at Multimodal Foundation Models for Biomedicine:
CVPR,2026

Datasets:

 saillab/psf-med



 like

0

Follow

 Secure and Assured Int...

35

Tasks:

 Visual Question Answering

Languages:

 English

Size:

10K<n<100K

Tags:

medical

radiology

chest-x-ray

paraphrase

vlm

benchmark

+ 1

License:

 cc-by-4.0

 Dataset card

 Files

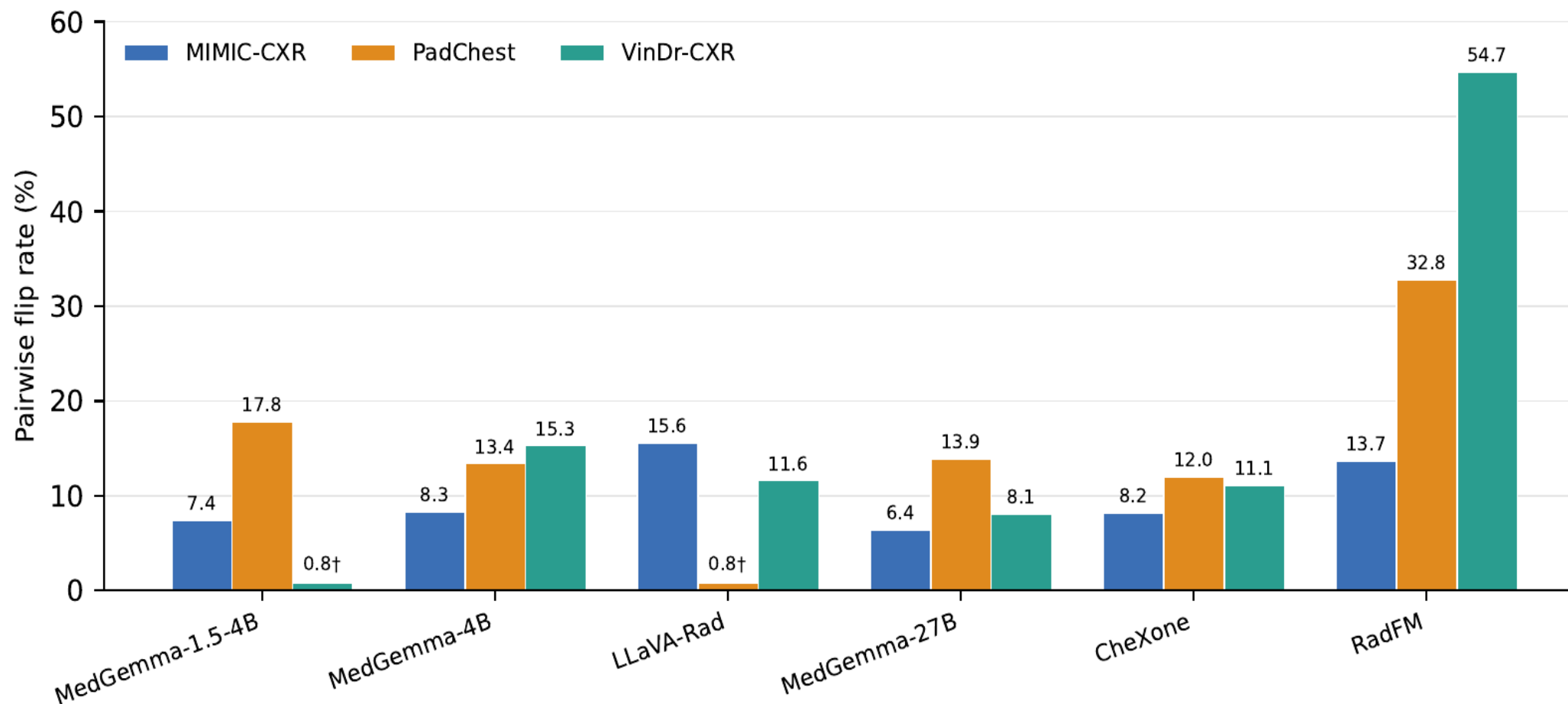
 xet

 Community

1

Source: Sadanandan et al., "PSF-Med," CVPRW Multimodal Foundation Models for Biomedicine ,2026.

Flip rates across models and populations



Pairwise flip rate on the equivalence-filtered PSF-Med binary subset.

<https://huggingface.co/spaces/saillab/psf-med-leaderboard>

Frontier models as a bounded case study

#	MODEL	MIMIC-CXR	PADCHEST	VINDR-CXR	POOLED	PAIRS
1	Kimi K3 default	6.99	3.10	8.06	5.77 [5.19–6.41]	5,613
2	Claude Opus 5 thinking off	8.47	5.41	10.18	7.79 [7.11–8.52]	5,613
3	Claude Opus 5 adaptive thinking	6.99	7.36	13.18	9.12 [8.40–9.90]	5,613
4	GPT-5.6 Sol default	8.92	3.28	17.85	9.51 [8.77–10.31]	5,613
5	Claude Haiku 4.5 default	4.55	12.01	16.18	11.28 [10.48–12.13]	5,613

Badges name the reasoning or thinking setting each arm used.

The live leaderboard shows the unrestricted per-arm cells: <https://saillab-psf-med-leaderboard.static.hf.space/index.html>

On a matched subset, turning reasoning on does not remove paraphrase sensitivity, and in two of three models it makes it significantly worse.

THRUST 1 | MEASURE | Flips are common, and the obvious suspects do not explain them

RQ1.1

Is PSF common?

YES

Between **6.4%** and **54.7%** of paraphrase pairs flip depending on model and dataset, an 8.5x spread across the board.

Experiment: full PSF-Med run, six VLMs across MIMIC, PadChest, and VinDr.

RQ1.2

Does transformation type explain it?

NO

Audited dataset, the four main types sit within two points on MIMIC, **8.7%** to **11.1%**.

Experiment: 122,778 pairs screened by rubric judge, 1,200-pair clinician check.

RQ1.3

Does scale buy consistency?

NO

The **27B** wins MIMIC (6.4%) and VinDr (8.1%), but the 4B beats it on PadChest. No size ranks consistently.

Experiment: MedGemma ladder, 4B vs 1.5-4B vs 27B on identical pairs.

RQ1.4

Does distribution shift amplify it?

PARTLY

Out-of-distribution rates run **1.3x** to **4.0x** the MIMIC baseline, but the trend is not monotonic for most models.

Experiment: same six VLMs across MIMIC (US), PadChest (Spain), VinDr (Vietnam).

The failure is real and frequent, and it does not trace to type, scale, or shift. So the cause has to live somewhere else

THRUST 2 OF 5

Diagnose

When the answer does not flip, is it actually grounded in the image, or is the model reading the question alone?

RQ2.1 text prior versus image · RQ2.2 consistency-grounding trade-off · RQ2.3 attention on flipped answers

Chapter 4 · Consistent but Dangerous · Accepted, CVPR 2026 Workshop (MedReasoner)

1 · Measure

2 · Diagnose

3 · Explain + fix

4 · Safety-evaluate

5 · Deploy

Image-Reliance Controls

Condition	What the model gets	What it isolates
Real image	The chest X-ray and the question	the standard setting
Text-only	Question alone; NO image tokens supplied	whether the image is needed at all
Gray placeholder	All pixels RGB 128: tokens present, no diagnostic content	sensitivity to image CONTENT, holding token count fixed
Opposite-label swap	A different patient's X-ray with the opposite label for the same finding	whether the answer tracks what the image shows

Metrics :

- ❖ Text-only agreement: the fraction answered identically with the image removed. Higher means more text-reliant.
- ❖ Image-agnostic rate: the fraction answered identically under ALL THREE image conditions.

Low Flip Rate Is Not Grounding

107-pair diagnostic set · all three models see the identical pairs

Metric	MedGemma-4B	MedGemma-27B	LLaVA-Rad
Flip rate (lower = more consistent)	42.1%	14.0%	6.5%
Text-only agreement (higher = less image use)	66.4%	85.0%	96.3%
Swap sensitivity (lower = ignores image content)	30.8%	19.6%	10.3%
Image-agnostic rate (higher = more image-blind)	43.0%	60.7%	85.0%

Source : Sadanandan & Behzadan, "Consistent but Dangerous," CVPRW Med-Reasoner, 2026, pp. 6961–6968.

1. Near-imperceptible stability test

JPEG recompression at quality 95 + 5% global contrast adjustment · no clinical information added · 500-question PadChest sample

ANSWER CHANGES UNDER A SMALL PIXEL EDIT

Base MedGemma-4B	Targeted LoRA	Full LoRA
1.2%	2.0%	1.4%
0.6% under contrast alone		

The visual pathway is stable.

Rewording flips are more than an order of magnitude higher than changes under pixel edits. Paraphrase sensitivity is text-side processing, not a fragile image encoder.

2. Sampling noise issue

All models use greedy decoding at temperature 0, so every answer difference is deterministic and attributable to the input change.

Attention on the annotated pathology box	Flip cases	No-flip cases	Change
Ground Truth coverage share of total attention mass falling inside the box	8.4%	14.2%	−41%
Precision share of the model's strongest attention falling inside the box	10.6%	29.0%	−63%

Two-sample t-test on Ground Truth coverage, flip versus no-flip: $p < 0.05$.

What this supports

- ❖ Flipping co-occurs with a measurable spatial grounding failure, not only a language-side one.
- ❖ Where the model looks hardest degrades most
- ❖ The encoder is stable under pixel edits, so this shift in attention is driven from the language side.

What bounds it

- ❖ Attention is low even when the model is consistent: This is weakly grounded versus less grounded.
- ❖ Correlational. It does not show that attention (lack thereof) causes the flip.

Source : Sadanandan & Behzadan, "Consistent but Dangerous," CVPRW Med-Reasoner, 2026, pp. 6961–6968.

THRUST 2 · DIAGNOSE | The steadier the answers, the less the model looks at the scan

Three grounding controls on the consistency winners · 107-pair diagnostic set · 200 annotated PadChest questions

RQ2.1

YES

Do models lean on text priors over the image?

LLaVA-Rad reproduces **96.3%** of its answers with no image at all, and **85.0%** survive both controls. Even MedGemma-4B, the most image-driven backend, moves on only **30.8%** of swaps.

Experiment: text-only and opposite-label swap controls on the 107-pair diagnostic set. Image blanked to gray, then swapped for a scan showing the opposite finding.

RQ2.2

YES

Is there a consistency-grounding trade-off?

Flip rate and text-only agreement line up in a near-perfect negative ordering, $r = -0.997$. The best-in-class **6.5%** flip rate comes with the weakest grounding of the three.

Experiment: flip rate against text-only agreement for the three backends on identical pairs.

RQ2.3

YES

Do flipped answers attend less to the right anatomy?

Attention on the annotated region runs **8.4%** on flips against **14.2%** on non-flips. The gap holds at $p < 10^{-8}$.

Experiment: attention-box overlap on 200 annotated PadChest questions (100 flip, 100 stable), MedGemma-4B final decoder layer, text-to-image attention inside the radiologist's bounding box.

Consistency can be bought with text priors, and flipped answers come with weaker grounding. So a low flip rate can hide a model that barely uses the image, which is why the next chapters pair every consistency number with a grounding control.

THRUST 3 OF 5

Explain and Fix

What is the mechanism behind the flip, and can a targeted intervention reduce it without costing accuracy?

RQ3.1 sparse-autoencoder mechanism · RQ3.2 targeted LoRA · RQ3.3 does the fix land where the diagnosis points

Chapter 5 · Mechanistically-Guided Targeted LoRA · Accepted, CHIL 2026

1 · Measure

2 · Diagnose

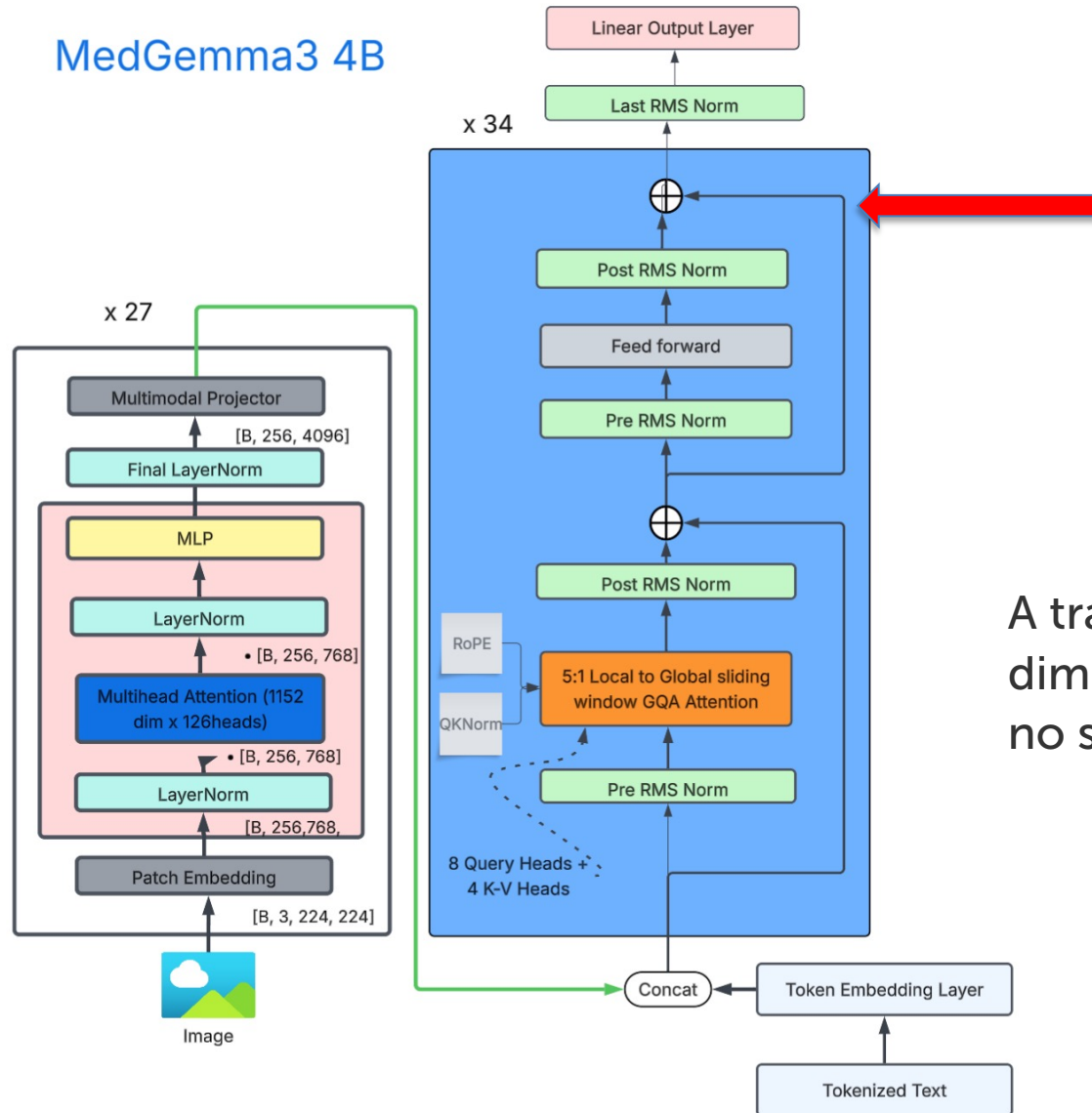
3 · Explain + fix

4 · Safety-evaluate

5 · Deploy

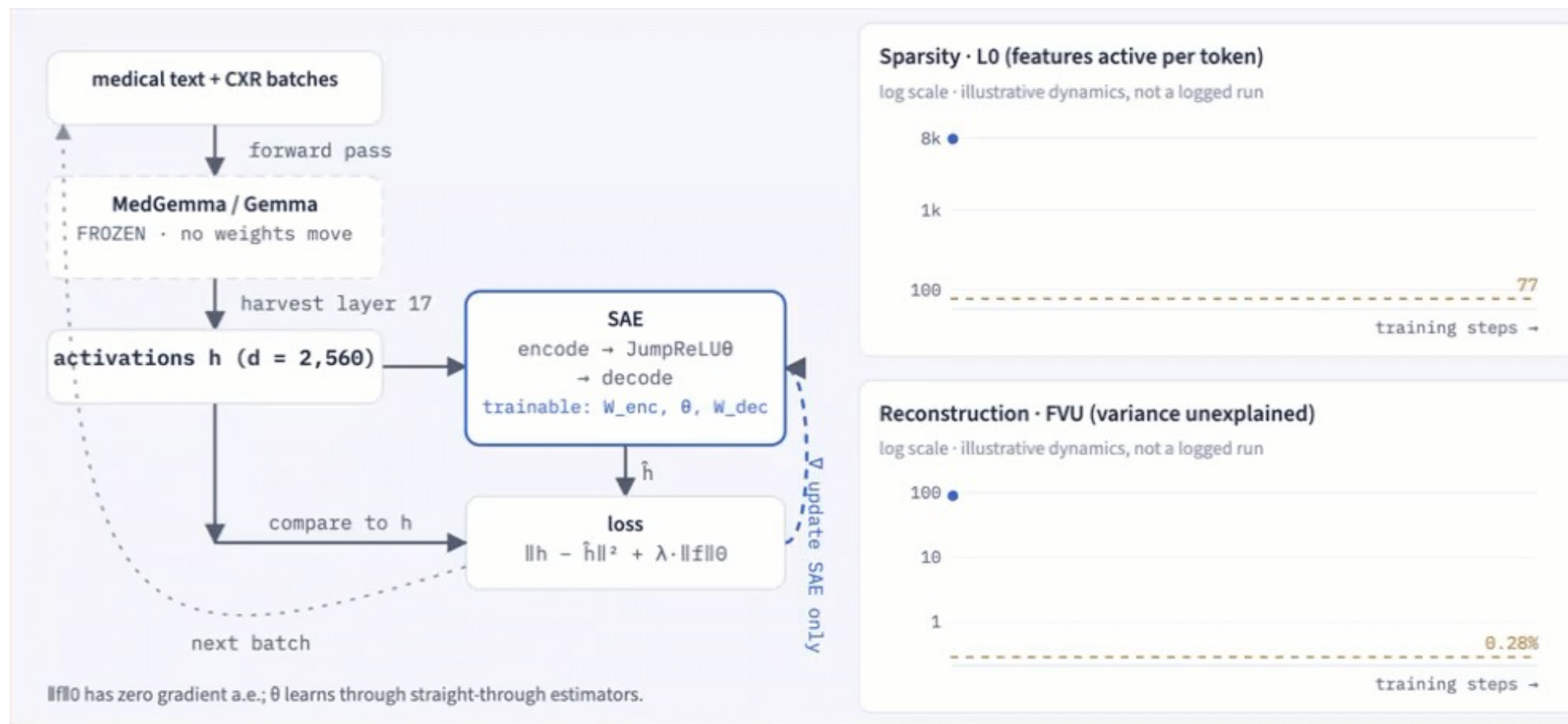
Method A - Sparse Auto-Encoders (SAE)

MedGemma3 4B



A transformer carries more concepts than it has dimensions, so concepts sit in superposition and no single neuron is one concept

Method A - Sparse Autoencoders



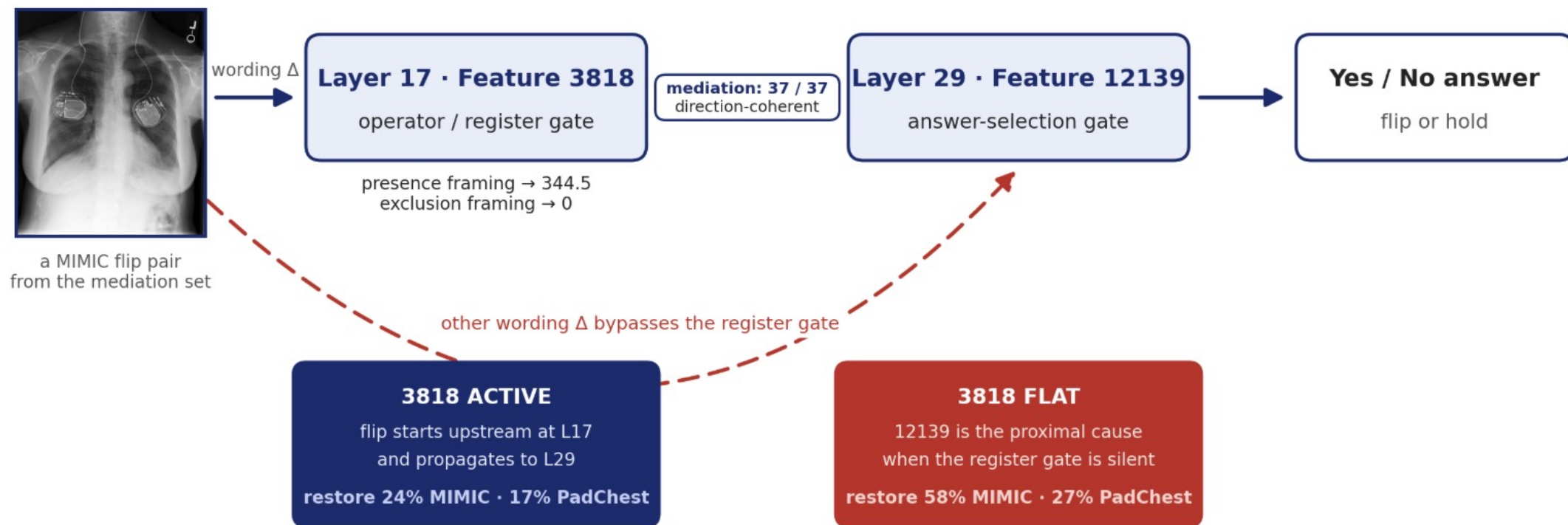
<https://huggingface.co/saillab/medgemma-circuit-tools>

Source: Sadanandan & Behzadan, "Mechanistically Guided LoRA," CHIL (PMLR 333), 2026, pp. 703–720.

- 1 Transfer & Training**
Gemma Scope 2 SAEs, trained on web text, applied to unmodified MedGemma-4B. $R^2 = 0.9972$ on medical prompts against 0.9974 on general.
- 2 Discover**
Feature delta $\Delta f = f(q) - f(q')$ at the final token in layer 17, computed across the PSF paraphrase pairs. Feature 3818 stood out
- 3 Test causally**
Decode Δf for Feature 3818 through W_{dec} and subtract that direction from the paraphrase residual stream.
- 4 Act on it**
Layer 17 localization set the targeted LoRA range, layers 15 to 19, at 0.1% of parameters.

Source: Sadanandan & Behzadan, "Mechanistically Guided LoRA," CHIL (PMLR 333), 2026, pp. 703–720.

Method A - What did SAE Find?



What it tells us

- ❖ Flips are structured: an upstream register gate (3818) + a downstream answer gate (12139), two entry points, one final switch
- ❖ Register changes start the flip at L17; everything else perturbs L29 directly
- ❖ One feature restores 24–58% of flips

Source: Sadanandan & Behzadan, "Mechanistically Guided LoRA," CHIL (PMLR 333), 2026, pp. 703–720.

- ❖ Transcoders trained from scratch on 160k radiology reports reconstruct MedGemma's activations
- ❖ The register layer (17) is medically specialized: its reconstruction drops on out-of-domain text.
- ❖ The decision layer (29) stays generic
- ❖ The register signal replicates under the new instrument.

saillab/medgemma-circuit-tools · top-k transcoders on MedGemma-4B MLP activations, trained on ReXGradient-160K radiology reports

Measurement	Layer-17 transcoder	Layer-29 transcoder
Role in the candidate two-stage circuit	register gate (Feature 3818)	decision feature (Feature 12139)
Explained variance — ReXGradient (in-distribution)	0.9964	0.9985
Explained variance — PSF-Med questions (unseen)	0.9648	0.9627
Explained variance — WikiText-103 (out-of-domain)	0.8840	0.9611

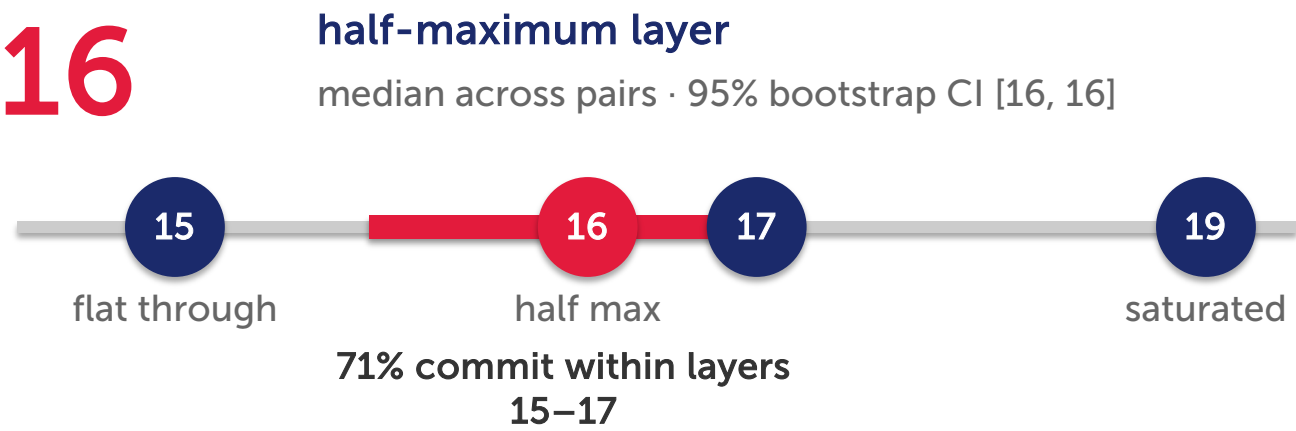
Weights, usage code, and paired domain-contrast tables: huggingface.co/saillab/medgemma-circuit-tools · Caveat: hookpoint (MLP block vs resid_post) and decoder-norm conventions differ across the two instruments, so the comparison confounds feature causality with intervention geometry.

Method B - Layer-16 transplant method

CAUSAL TEST

- 1 Detecting phrasing
produces the correct answer
- 2 Transplant at layer L
move the residual stream
- 3 Missing phrasing
does its answer now follow?

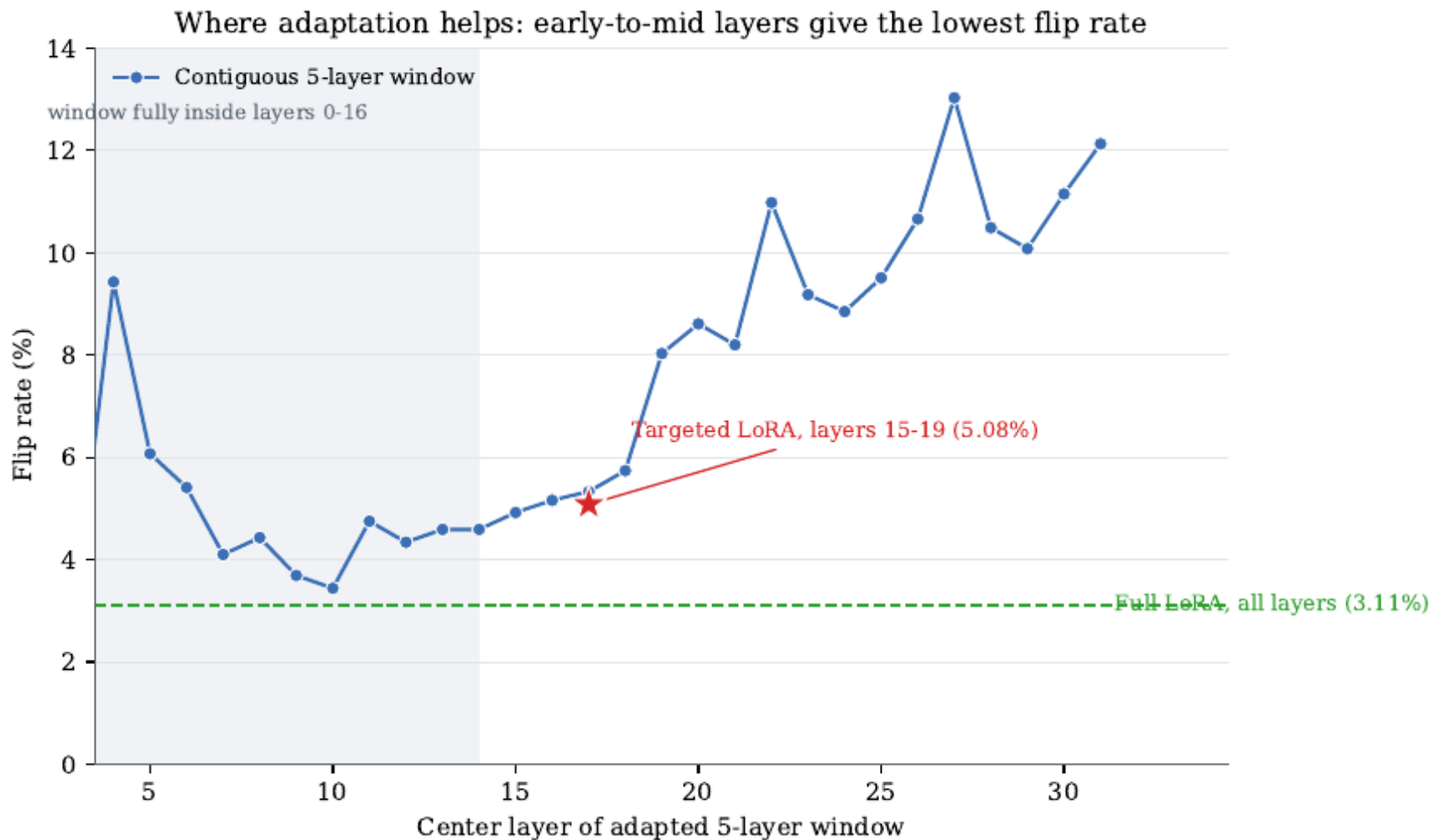
RESULT



1,396 91 ≤ 0.0007
transplant pairs findings image-token
control at every
layer

Controls: image-token patching is effectively zero; reverse-direction transplant is the second control. **The commit is located by causal effect alone.**

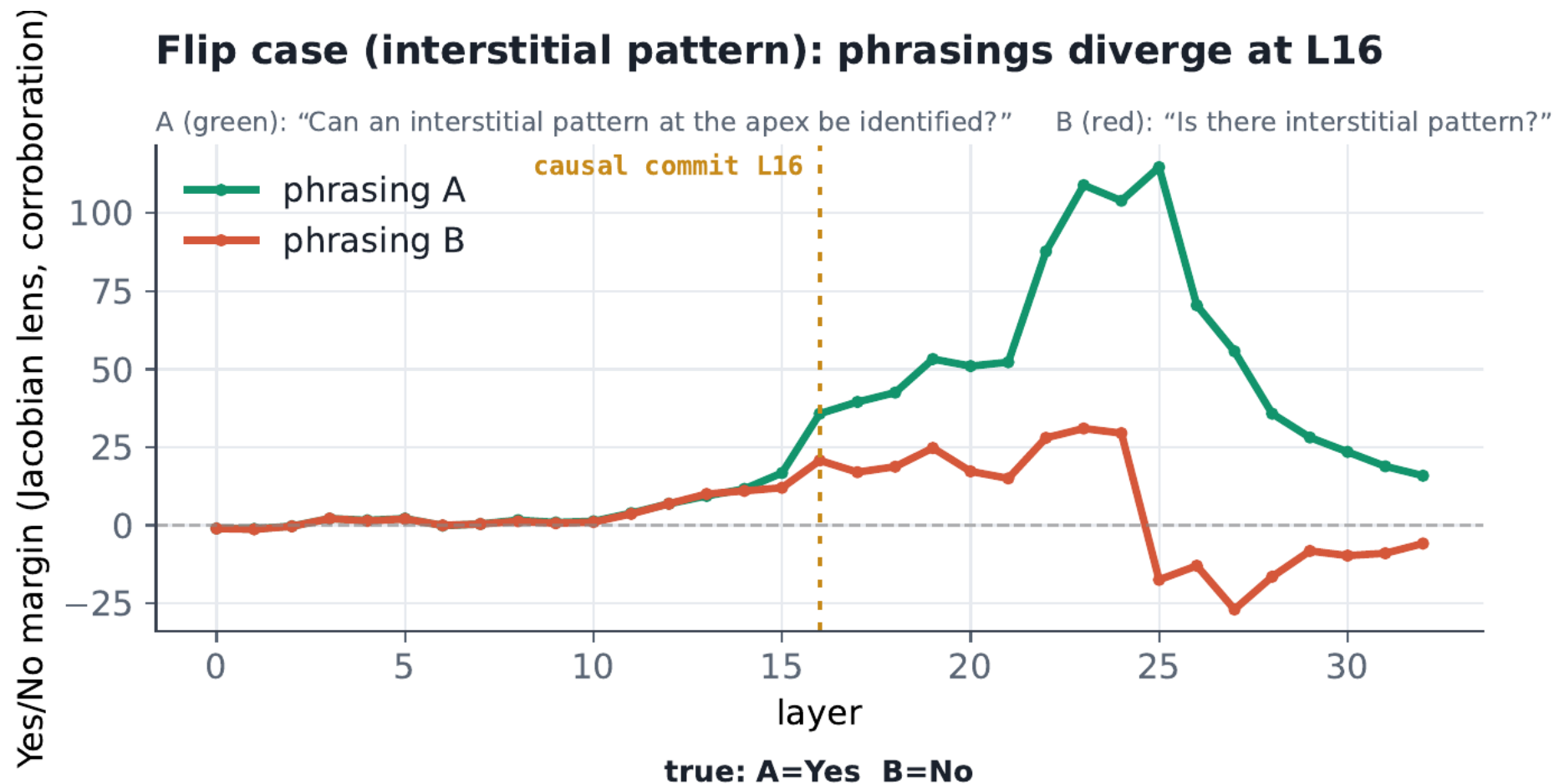
Source: Sadanandan & Behzadan, "Mechanistically Guided LoRA," CHIL (PMLR 333), 2026, pp. 703–720.



Flip rate by the center layer of an adapted 5-layer window. Early-to-mid windows reach a lower flip rate than the mechanistically identified layers 15 to 19

The Jacobian Lens Reads the Flip as It Forms

Same image, same question, two phrasings — reading the yes/no margin each layer is disposed to produce



Interactive per-token version on the companion site: bineshkumar.me/phd-thesis/explorer/ (Jacobian-lens demo) · Chapter 5.

8.5% → 3.5%

pairwise flip, patient- AND image-disjoint MIMIC presence test (n = 238)

−59% (± 0.6 pp across 5 seeds) · 991 pairs · McNemar $p \leq 1.8 \times 10^{-3}$ in every seed

■ Query-level flip: **19.8% → 8.1%** (± 1.8), the clinically relevant rate falls by half

■ Accuracy preserved: **84.5% → 84.7%**, paired difference stated as interval $[-1.00, +1.51]$ pp

■ Out-of-distribution (PadChest): accuracy **85.2% → 91.6%** (+6.4 pp)

0.10%

of parameters modified, 4.38M, LoRA rank 16 on layers 15–19 only

~620 min

on one A100 (96 GB), the mitigation is cheap to reproduce; weights on HuggingFace

Layer choice tested

Block sweep across all 34 layers + 20 random subsets: 15–19 wins on flip–accuracy Pareto; full-model LoRA flips lowest but relies on text most

Source: Sadanandan & Behzadan, "Mechanistically Guided LoRA," CHIL (PMLR 333), 2026, pp. 703–720.

Utility

- ❖ On a deployed model, the data, the architecture, and the training objective are fixed together.
- ❖ Observational study can show that paraphrase sensitivity correlates with something, but it cannot isolate the cause.

Baby MedGemma

- ❖ A small model that mirrors MedGemma's design closely enough for the mechanism to transfer.
- ❖ Vary the training-phrasing distribution, with architecture, seed, and image features held fixed.

- ❖ 86,288 binary presence questions from NIH ChestX-ray14 and PadChest; MIMIC-CXR and VinDr-CXR are held out entirely.
- ❖ Every split is balanced per finding
 - ❖ Property 1 (frozen encoder): identical image features per paraphrase
 - ❖ Property 2 (it genuinely sees): with the image removed accuracy is 50% (chance)

Evaluation set	Accuracy	AUROC
In-distribution (held-out images)	0.748	0.827
MIMIC-CXR (unseen hospital)	0.671	0.743
VinDr-CXR (unseen hospital)	0.686	0.756

<https://huggingface.co/saillab/FlipLens>

Training regime	Phrasing seen in training	Flip: all	Accuracy
Augmented	Every paraphrase	4.8%	75.3%
Canonical	One fixed phrasing per question	67.1%	66.8%
Vanila	Register correlated with the answer	88.4%	58.1%

<https://huggingface.co/saillab/FlipLens>

- ❖ Only the training-phrasing distribution changes; architecture, parameters, seed, and evaluation are fixed.
- ❖ Augmented beats both narrow regimes at the maximum effect size

Mann-Whitney $p = 1.6e-4$, Cliff's delta = 1.00 (every augmented seed flips less than every narrow-regime seed). Eight seeds per regime.

Binary Accuracy Cannot Certify Image Use

Probe variant (balanced held-out set)	Accuracy	AUROC of the answer margin
No pooled image summary (blind by construction)	50.1%	0.500
Image summary shuffled across images	52.2%	0.502
Correct image summary	52.6%	0.604

Accuracy is flat across a blind model, a shuffled one, and a seeing one (50-53%). The margin's AUROC separates the seeing model (0.604) from the blind ones (0.500). A binary yes/no readout hides whether the model used the image at all.

<https://huggingface.co/saillab/FlipLens>

The Margin Detects Flips in One Pass; Image Reliance Needs Two

Detector	Passes	In-dist	MIMIC	VinDr
Catch a paraphrase flip: absolute margin	1	0.923	0.974	0.974
sampling self-consistency	k	0.709	0.786	0.624
hidden-state probe	1 + fit	0.826	0.838	0.801
Catch an image-unreliant answer: absolute margin	1	0.470	0.491	0.519
image swap	2	0.827	0.907	0.856

Answer margin is the cheapest and strongest flip detector, beating multi-pass sampling and a trained probe. But an answer that ignores the image has no single-pass signal (the red row): it needs an explicit image swap.

THRUST 3 · EXPLAIN AND FIX | A candidate circuit, a 59% fix, and a cause in the training data

Chapter 5 · sparse autoencoders (SAEs) and activation patching · targeted LoRA adapters · FlipLens phrasing probe

RQ3.1

YES

Can SAE features find a mechanism?

A **layer-17 register gate** feeding a **layer-29 decision feature**, matched against same-polarity controls.

Experiment: Gemma Scope 2 SAEs transferred to MedGemma-4B. Feature deltas and activation patching on the 158-pair FlipBank, then a frozen-hypothesis screen of 4,542 MIMIC and 37,280 PadChest pairs.

RQ3.2

YES

Does targeted LoRA fix it?

Pairwise flips fall **59%** (8.5% to 3.5%) with accuracy intact at **84.7%**, but grounding does not come with the fix.

Experiment: LoRA on layers 15-19 with a consistency-plus-accuracy loss, tested on patients and images the adapter never saw, five seeds. McNemar $p < 0.002$ in every seed.

RQ3.3

NO

Does the best intervention site match the diagnosis?

Early layers beat the diagnosed 15-19 block: layers 1-5 cut the margin gap **87%** against 80%, late windows worst at 63%. Diagnosis does not prescribe the fix.

Experiment: five layer-range adapters, then an exhaustive sweep of all 30 contiguous five-layer windows plus 20 random subsets.

PROBE

YES

Does the training-phrasing distribution cause the sensitivity?

Phrasing alone moves the flip rate from **4.8%** to **88.4%** (Cliff's delta = 1.00), and the ordering holds on wording withheld from training.

Experiment: FlipLens, a 14.1M-parameter Gemma-3 decoder over a frozen MedSigLIP-448 encoder, trained on 86,288 binary chest X-ray questions under three phrasing regimes, eight seeds each.

The cause sits in the training data, and a small adapter can exploit that to cut flips by 59%. But the SAE circuit stays a candidate until the controls run, and the adapter buys part of its consistency with the text prior.

THRUST 4 OF 5

Safety-Evaluate

Can we screen a single prediction for hidden text-reliance before a clinician acts on it?

RQ4.1 does a lower flip rate mean safer · RQ4.2 are heatmaps faithful · RQ4.3 calibration and disparity

Chapter 6 · Attention Without Grounding

1 · Measure

2 · Diagnose

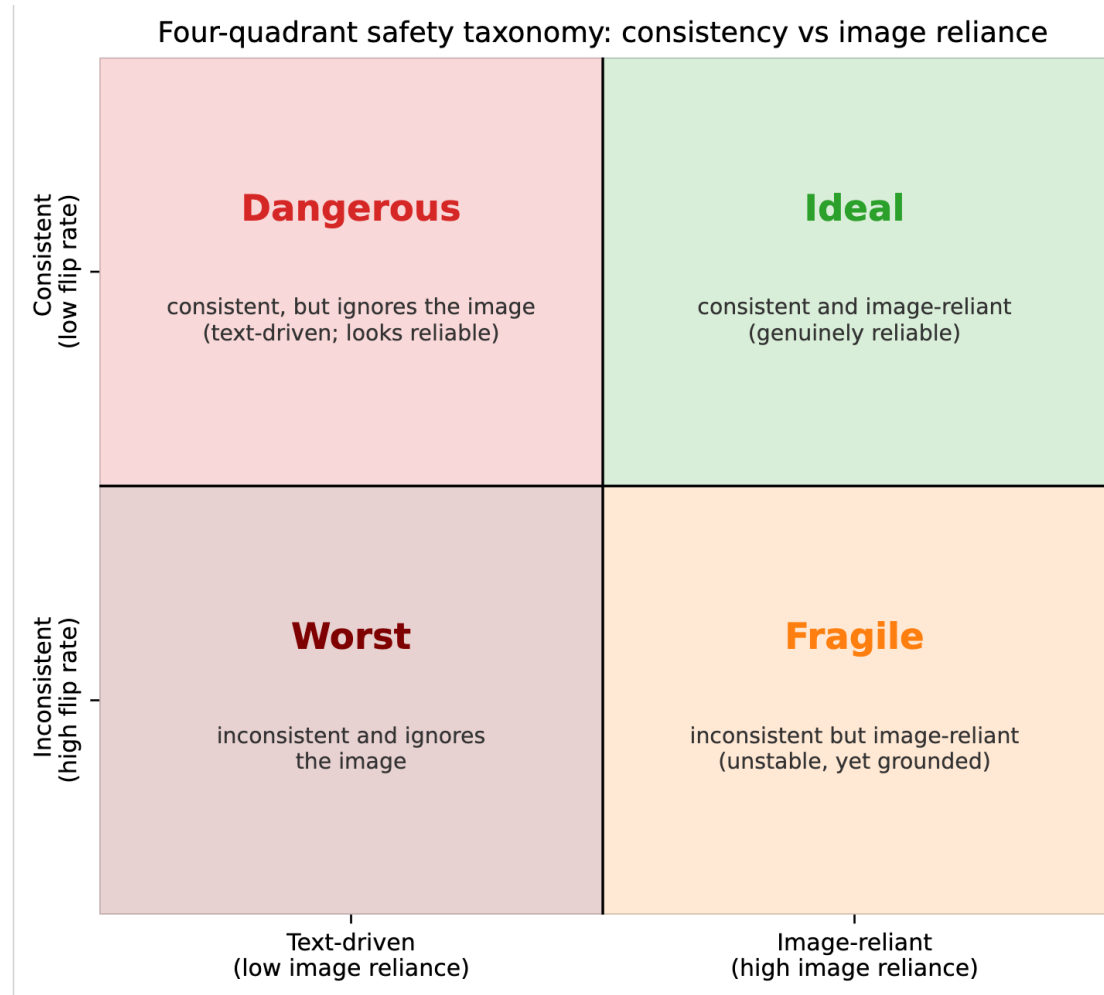
3 · Explain + fix

4 · **Safety-evaluate**

5 · Deploy

Four-quadrant screen

Two behavioral axes: paraphrase consistency (does the answer stay fixed?) x image reliance (does the answer move when the image changes?).



Source : Sadanandan & Behzadan, "Attention Without Grounding," MICCAI Interpretability of Machine Intelligence in Medical Image Computing (iMIMIC, 8th ed.), 2026. [Spotlight \(oral\)](#).

Three signals the labels never use, KL divergence, image swap, and a null image, all line up with the labels.

Information-theoretic check

- ❖ Image-off vs image-on KL divergence: for each prediction, how much does removing the image move the output distribution?
- ❖ Dangerous / text-driven predictions have near-zero divergence; Ideal predictions have high divergence.
- ❖ It separates the two at AUROC = 0.76, using a signal the quadrant labels never see.

Two perturbation checks

- ❖ Swap invariance: does the answer survive an opposite-finding image?
- ❖ Null-image control: does the answer survive a uniform gray placeholder?
- ❖ Three independent signals agree with the behavioral labels, so the screen is not circular.

Does attention point at the pathology, and does the pathology actually drive the answer?

Attention grounding (PadChest, 637 boxes)	Base	Full LoRA	Read
True-box coverage	29.6%	38.5%	5-6x the random rate
Displaced-box baseline	26.1%	33.8%	only marginally below true
Random-far baseline	5.5%	7.4%	the floor
Attention rank vs causal importance (spearman)	-0.09	+0.06	≈ 0: attention does not rank cause

ROI-deletion grounding score (region ablation minus a same-size random ablation), all excluding zero:

ROI-deletion grounding score	Base	Full LoRA	Targeted LoRA
Region ablation minus random ablation	2.97 [2.67, 3.26]	1.18	0.57

Source : Sadanandan & Behzadan, "Attention Without Grounding," MICCAI Interpretability of Machine Intelligence in Medical Image Computing (iMIMIC, 8th ed.), 2026. Spotlight oral.

Read the answer in this order

- ❖ **Step 1** Consistent across paraphrases? If not, unreliable. Stop..
- ❖ **Step 2** Image-reliant? Low text-only agreement, high swap sensitivity. Consistent but not image-reliant is Dangerous
- ❖ **Step 3** Only then, correctness. Right by prior. Wrong on the next patient

Finding	Text-driven share	Reading
Vertebral degenerative changes	95.7%	distrust a consistent answer until checked
Cardiomegaly	80.4%	distrust a consistent answer until checked
Interstitial pattern	5.0%	a consistent answer is more likely grounded
Infiltrates / vascular hilar	0.0%	a consistent answer is more likely grounded

Source : Sadanandan & Behzadan, "Attention Without Grounding," MICCAI Interpretability of Machine Intelligence in Medical Image Computing (iMIMIC, 8th ed.), 2026. Spotlight oral.

Calibration Does Not Imply Equity

- ❖ A model is calibrated when its confidence matches its accuracy: answers it gives at 80% confidence should be right about 80% of the time
- ❖ Expected Calibration Error (ECE) measures how far confidence and accuracy drift apart, pooled over every question in the test set

RQ4.3 · are demographic disparities larger for the worst-calibrated models? Partly · PadChest, 861 questions, sex- and age-stratified

Metric	Base MedGemma-4B	Targeted LoRA	Full LoRA
Temperature-scaled ECE (lower = better calibrated)	13.6%	3.6%	22.9%
Sex accuracy gap (percentage points)	2.0	2.7	4.3
Age gradient: accuracy <40 → 80+ (pp)	4.3 (flattest)	13.0 (steepest)	8.2

Source: Sadanandan & Behzadan, "Predictive Entropy as a Joint Screen for Error and Paraphrase Instability in Medical VLMs," UNSURE @ MICCAI 2026.

THRUST 4 · SAFETY-EVALUATE | Do the usual safety proxies hold up?

Chapter 6 · consistency as a safety proxy · attention heatmaps against causal controls · demographic calibration on PadChest

RQ4.1

NO

Do fewer flips mean safer?

A mean of **81%** of consistent predictions are image-invariant (range 45 to 100%), whether the model flips often or rarely.

Experiment: four-quadrant screen across ten model-dataset settings, classifying each prediction by consistency and by whether removing the image moves the answer. KL-divergence, opposite-label swap, and null-image controls.

RQ4.2

PARTLY

Do attention heatmaps reliably indicate grounding?

True-box coverage runs **29.6% to 38.5%** against 26.1% to 33.8% for shifted boxes, but patch ordering tracks causal importance at only $\rho = -0.09$ to $+0.06$.

Experiment: attention-box overlap on PadChest's 637 radiologist-annotated regions, tested against shifted, random-far, and random-sized baselines, plus patch occlusion and ROI deletion.

RQ4.3

PARTLY

Are demographic disparities larger for the worst-calibrated models?

The worst-calibrated model shows the largest gaps (ECE above **0.25** per group, 4.3-point accuracy gap). But the best-calibrated adapter has the steepest age gradient: **83.3%** under 40 rising to **96.3%** over 80.

Experiment: sex- and age-stratified accuracy and calibration on PadChest's 861 questions.

Fewer flips do not mean safer, heatmaps do not certify grounding, and calibration does not predict fairness. Each safety proxy needs its own control before it earns trust.

THRUST 5 OF 5

Deploy

Is there a signal cheap enough to run on every prediction that tells a deployed system when to escalate to a human?

RQ5.1 entropy predicts flips · RQ5.2 corruption and site shift · RQ5.3 what gates admit · RQ5.4 cross-architecture transfer

Chapter 7 · Predictive Entropy as a Joint Screen · Under review, UNSURE at MICCAI 2026

1 · Measure

2 · Diagnose

3 · Explain + fix

4 · Safety-evaluate

5 · Deploy

$$H = - \sum p \log p \quad \text{binary case: } H = -p \log p - (1-p) \log(1-p)$$

Measured in nats. Maximum $\ln 2 = 0.693$ at $p = 0.5$, the decision boundary. Minimum 0 at $p = 0$ or $p = 1$.

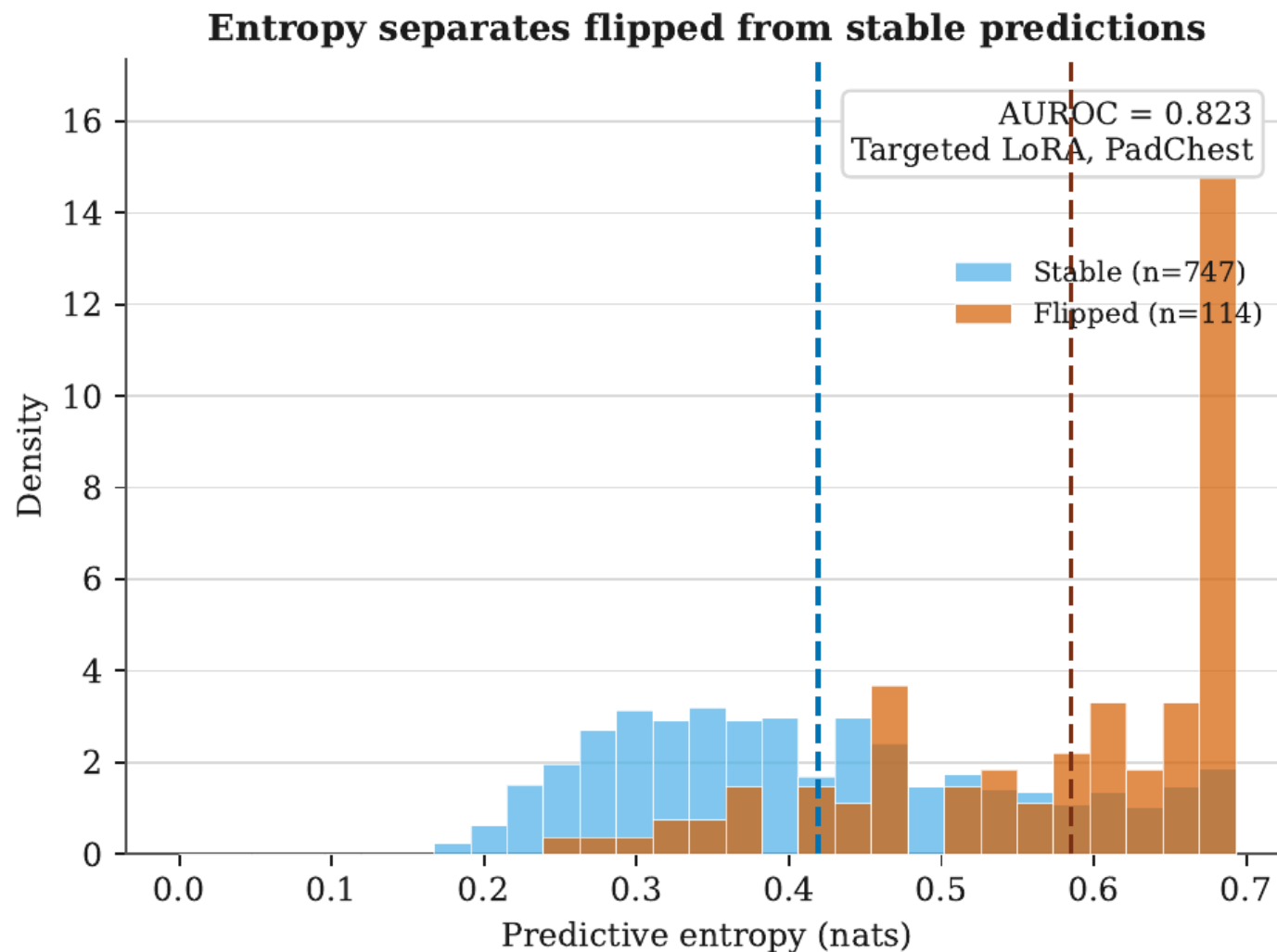
Method

- ❖ The model emits two logits. Write the margin $m = z(\text{yes}) - z(\text{no})$, so $p = \sigma(m)$ and the binary answer is just $\text{sign}(m)$.
- ❖ H is symmetric about $m = 0$ and falls monotonically as $|m|$ grows. Ranking by entropy and ranking by $-|\text{margin}|$ give the identical ordering, and therefore the identical AUROC.
- ❖ Cost is zero beyond the answer itself. Both quantities were already computed and discarded.

Predict Changes

- ❖ A flip is a sign change of the margin across paraphrases.
- ❖ A sign change is most likely when the margin already sits near zero.
- ❖ Entropy is maximal exactly there.

Entropy Separates Flipped From Stable Predictions

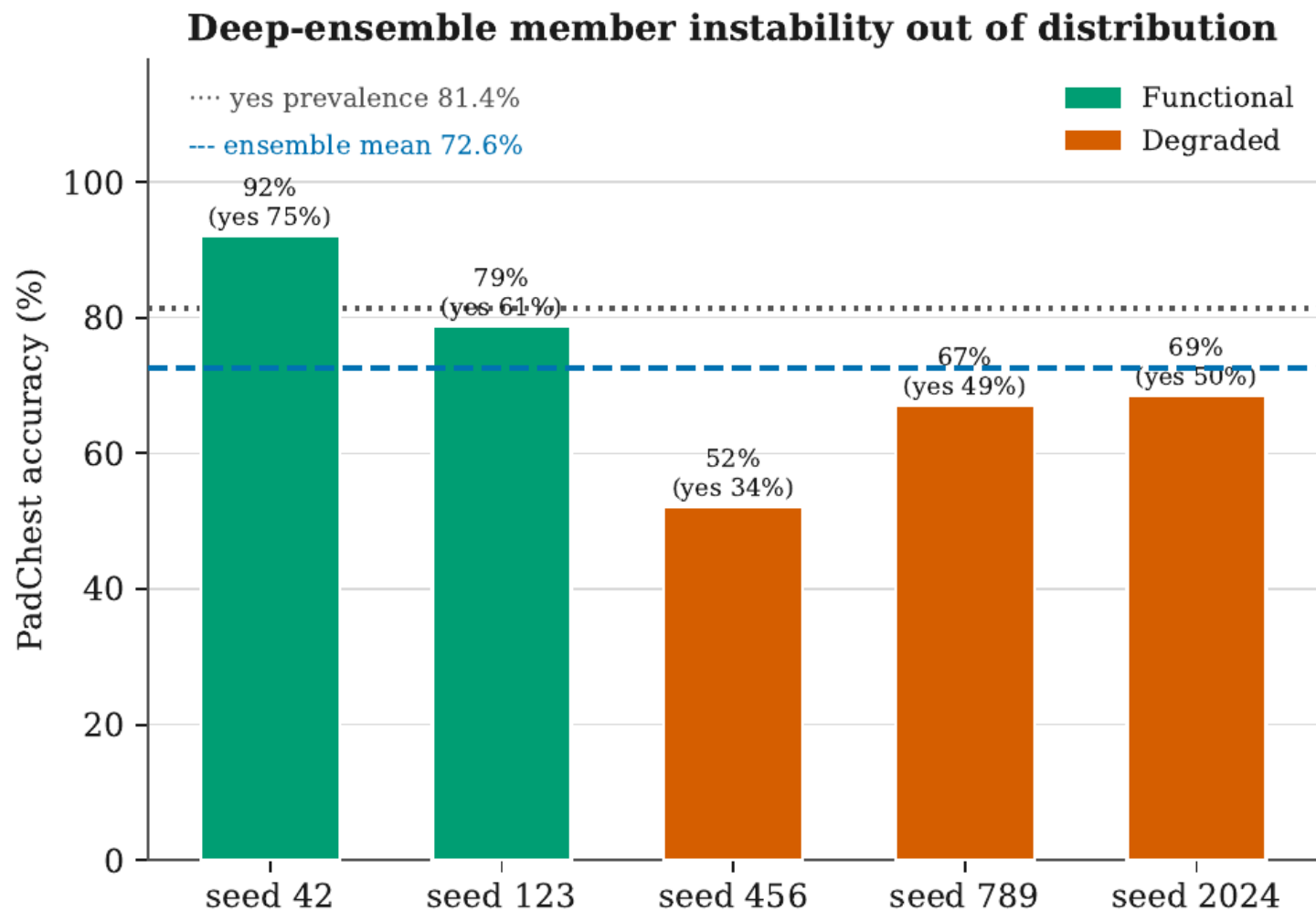


Single-pass predictive entropy separates flipped from stable predictions at AUROC 0.823 (Targeted LoRA, PadChest, n = 861).

Method	Passes	What it adds	Result
Softmax entropy	1	total predictive uncertainty	AUROC 0.823 flip, 0.862 error
Yes/no margin	1	monotone twin of entropy	identical ranking, identical AUROC
Temperature scaling	1 + fit	calibrated probabilities	fixes ECE, changes no ranking
MC Dropout, K = 10	10	an epistemic estimate	0.823 flip, 0.856 error. No gain
Deep ensemble, 5 seeds	5	a genuine epistemic signal	0.751 error. AURC 0.130 vs 0.019

Total entropy from one pass outperforms decomposed uncertainty from many. MC Dropout costs 10x compute for nothing, and the deep ensemble costs 5x and carries an out-of-distribution risk.

The Deep Ensemble Fails Out of Distribution



Five LoRA seeds span 52% to 92% PadChest accuracy: averaging recovers a 72.6% mean but a poorly ranked confidence signal. The one distinct method, and the one that fails.

Expected calibration error, PadChest

Model	Raw	Scaled
Base	0.16	0.14
Targeted LoRA	0.11	0.04
LORA	0.27	0.23

Calibration

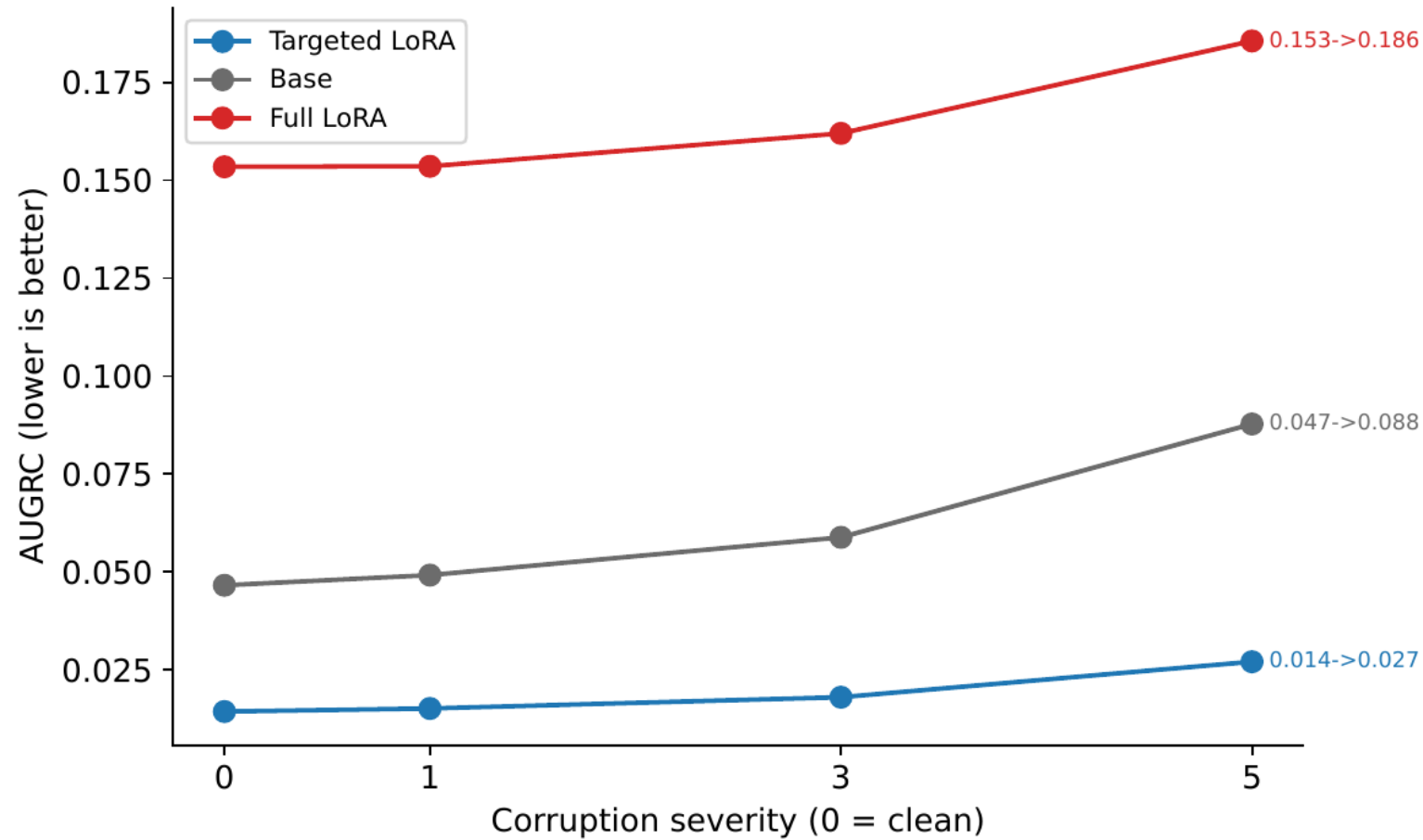
Is a stated 80% right 80% of the time?
Matters for a number shown to a clinician.
Scaling can fix it

Ranking, AUROC

Do risky cases outscore safe ones? Matters
for a gate. Dividing by one positive T is
monotone, so this never moves.

The best-calibrated variant is also the one with the steepest age gradient, 13 pp against 4.3 pp for Base.
Calibration does not imply equity

Calibration under distribution shift: Targeted LoRA stays stable, Base degrades



AUGRC (lower is better) under a corruption sweep: Targeted LoRA stays stable (0.014 to 0.027); Base degrades (0.047 to 0.088).

From Entropy to a Decision

Sort predictions by entropy, admit the lowest fraction, and measure what you admitted.

Selective prediction

Sort ascending by entropy and admit coverage c . Risk is measured on the admitted set only.

At 40% coverage: 0.6% error and 2.3% flips, against 8.6% and 13.2% in the full population.

AUGRC

Area under the generalized risk-coverage curve. One number summarizing the whole curve, lower is better.

Corruption sweep to severity 5: Targeted LoRA 0.014 $\rightarrow$ 0.027, Base 0.047 $\rightarrow$ 0.088, Full LoRA 0.153 $\rightarrow$ 0.186.

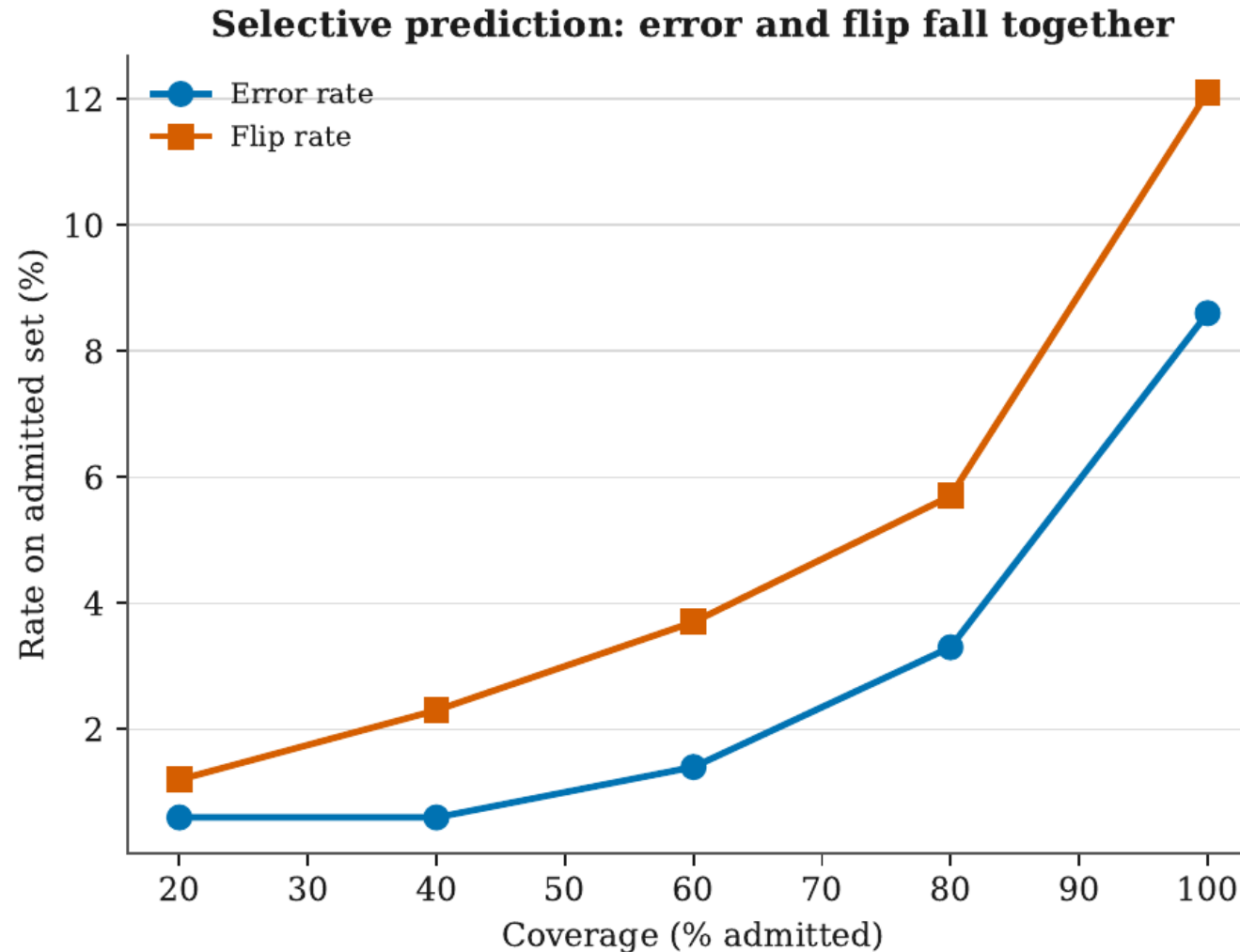
Conformal prediction

Distribution-free coverage. A threshold $\hat{q}$ is calibrated on a 15% holdout at $\alpha = 0.1$ for a 90% target.

Targeted LoRA holds 93.7% clean, 92.2%, then 89.9% at severity 5, landing on target. Base under-covers and Full is unstable.

- ❖ Selective prediction is the operating picture a clinician would see.
- ❖ AUGRC compresses the whole curve so two models can be compared without choosing a coverage point.
- ❖ Conformal answers whether a stated 90% coverage still holds once the images degrade.

Selective Prediction: Abstaining Buys Safety



On the admitted set, error and flip fall together as the model abstains on high-entropy cases. At 40% coverage error is about 0.6% and flips about 2.3%, against 8.6% and 13.2% at full coverage. In-sample operating point.

Entropy transfer across architectures

Model / dataset	Flip-prediction AUROC	Note
MedGemma-4B Targeted LoRA / PadChest	0.823	error detection AUROC 0.862
LLaVA-Rad LoRA / MIMIC	0.905	different architecture
LLaVA-Rad LoRA / PadChest	0.830	different architecture

Mean entropy is higher on flipped predictions (0.585) than on stable ones (0.419). Because entropy is a single forward pass with no paraphrase generation, the same cheap monitor works across model families, which is what a deployment standard needs.

THRUST 5 · DEPLOY | Cheap alarms exist, but assurance is per-model work

Chapter 7 · entropy as a flip alarm · conformal coverage under corruption · admission gates · per-family monitors

RQ5.1

YES

Does entropy predict flips?

One forward pass flags flips at **AUROC 0.82** and errors at **0.86** on PadChest. Pricier probes buy nothing.

Experiment: UQ-paraphrase bridge. Softmax entropy, single pass, $p < 10^{-28}$; holds across five seeds (0.828 ± 0.021), MIMIC (0.821), and LLaVA-Rad (0.830). MC Dropout merely ties; deep ensemble collapses out of distribution (0.552).

RQ5.2

PARTLY

Do guarantees survive shift?

Coverage holds near target under corruption (**93.7%** clean, **89.9%** at severity 5), but the guarantee is descriptive once exchangeability breaks.

Experiment: conformal sets (APS score) calibrated on a clean holdout, then severity-1 to 5 corruptions. At a 95% target, coverage stays above 90% at every severity.

RQ5.3

NO

Do gates admit grounded cases?

Gates raise accuracy by admitting the text-answerable slice: a null-image rule admits **72%** at 98.9% accuracy; the strict rule admits **10%** at 96.7%.

Experiment: architecture-aware audit on PadChest. Even the strict subset stays 94.4% swap-invariant; on text-image disagreements the adapter scores 2.9% and Qwen2-VL 0 of 279.

RQ5.4

NO

Does one monitor transfer across families?

Each family fails its own way and needs its own monitor. One dashboard cannot watch all three.

Experiment: internal probes per family. Gemma shows late readout-misalignment (AUC 0.81 to 0.92); LLaVA-Rad fails earlier in transport; Qwen2-VL sits near chance on both signatures.

Deployment assurance is per-model work. Entropy earns a place as a cheap PSF alarm, conformal holds if you set the target conservatively, and any admission gate has to prove what it selects before you trust what it scores.

CLOSING

Summary

What the five thrusts add up to, what independent reviewers have already accepted, and what is labeled future work.

One argument with five checkpoints: measure, diagnose, explain and fix, screen, deploy.

Chapter 8 · Synthesis, limitations, and future work

1 · Measure

2 · Diagnose

3 · Explain + fix

4 · Safety-evaluate

5 · Deploy

A medical vision-language model triaging a Chest X-ray can give a different diagnostic answer when a clinician rewords the same question, for the same X-ray

1

PSF-Med: audited benchmark

26,850 clinical questions, 92,856 equivalence-filtered pairs, six models, three countries, clinician-audited and public

2

Consistency is not grounding

Image-reliance controls separate text-driven answers from image-grounded ones

3

PSF has a locatable model specific mechanism and an identified cause

The answer commits at layer 16 - the training-phrasing distribution causes the sensitivity

4

A cheap targeted fix, audited against its own claim.

LoRA on 0.1% of parameters cuts flips 59% with accuracy held; the re-audit shows it does not buy grounding

5

Safety screen for Deployment

Paraphrase consistency crossed with image reliance gives four cells; the Dangerous cell (consistent + text-driven) dominates most models.

6

A deployment signal that costs one forward pass.

Predictive entropy predicts flips across two architecture families; grounding still needs a second pass

Why models fail?

- 1 The words are enough.** Findings are skewed: PadChest questions are positive 83% of the time, so a model answering from the base rate, without the image, is right most of the time, and accuracy rewards it.
- 2 Training never asks otherwise.** Models are fine-tuned and scored on one fixed wording per question. When the training phrasings predict the answer, flips reach 88.4%, against 67.1% for one fixed phrasing and 4.8% for all phrasings.
- 3 A register feature carries the surface form.** A sparse feature at layer 17 separates presence framings (Is there X?) from exclusion framings (Can X be ruled out?); a register shift moves it, and the shift propagates to the answer (candidate mechanism).
- 4 The join loses the image.** The encoder keeps the evidence: its pooling head separates findings at AUROC 0.812, while the patch tokens handed to the decoder reach only 0.514. The projection, not the encoder, is where image evidence is lost.
- 5 Capacity is not the cause.** Scale buys no stable consistency advantage, and with the architecture fixed, the training-phrasing distribution alone moves the flip rate from 4.8% to 88.4%. Capacity sets the stage; the training distribution decides.

The thesis measured the first, argued the second, investigated the third, gives a first controlled test of the fourth, and bounds the fifth.

How my work answers FDA's Asks

Every named element that touches this dissertation maps to a completed, peer-reviewed thrust.

FDA ELEMENT	WHAT FDA ASKS	THIS DISSERTATION
R.1 Robustness, reliability, reproducibility	Consistency across semantically equivalent paraphrases.	PSF-Med: 26,850 questions, 92,856 audited pairs, six models, three countries. Flips run 6.4% to 54.7%. Ch. 3
R.1 Failure definition	Variation in diagnostic conclusions is treated as a failure.	Four-quadrant screen: a mean 81% of consistent predictions are image-invariant. Consistency alone certifies nothing. Ch. 6
S.3 Calibration, uncertainty, clinical deferral	False confidence is treated as a safety failure.	Entropy flags flips at AUROC 0.82 and errors at 0.86 in one forward pass. Conformal coverage: 93.7% clean, 89.9% at corruption severity 5. Ch. 7
R.2 Subgroup performance	Behavior holds across demographics and health literacy.	Sex- and age-stratified audit: the worst-calibrated model also shows the largest disparities. Ch. 6
§VI.C Postmarket re-benchmarking	Modifications re-tested against the premarket baseline.	Targeted LoRA cuts flips 59% with accuracy held, yet the re-audit shows grounding loss. The silent change class FDA flags. Ch. 5

Source: FDA, "Considerations for the Regulation of GenAI-Enabled Medical Devices," Docket FDA-2026-N-7874, Regulations.gov, Aug. 2026.

Publications at Peer Reviewed Venues

The question it answers	Peer Reviewed & Accepted At
How often, and how much, do medical VLMs fail?	2nd Multimodal Foundation Models for Biomedicine @ CVPR 2026
Is a consistent answer actually grounded in the image?	Medical Reasoning with VLMs @ CVPR 2026 40% Acceptance Rate
What is the mechanism, and can a targeted fix reduce it?	7th Annual Conference on Health, Inference, and Learning 2026 30% Acceptance rate
Can we screen a prediction for hidden text-reliance?	8th Edition of Interpretability of Machine Intelligence in Medical Image Computing @ MICCAI 2026 34% Acceptance rate • Oral Spotlight
Is there a single-pass signal to gate at deployment?	7th Edition of Uncertainty for Safe Utilization of Machine Learning in Medical Imaging @ MICCAI 2026 32% Acceptance rate
SUPPLEMENTARY Does chain of thought help with consistency?	International conference on applied artificial intelligence 2026 Oral Spotlight
SUPPLEMENTARY Vulnerability scoring framework for Medical VLM	ISBI 2025 - Early abstract

Code, data, and adapter weights for every thrust with a public component.

BENCHMARK · THRUST 1

- 1

PSF-Med benchmark code
github.com/UNHSAILLab/psf-med
- 2

PSF-Med dataset (Hugging Face)
huggingface.co/datasets/saillab/psf-med

MITIGATION · THRUST 3

- 3

Mechanistically-guided LoRA weights (CHIL 2026)
huggingface.co/collections/saillab/mechanistically-guided-lora...
- 4

Targeted MedGemma-4B LoRA model
huggingface.co/saillab/medgemma-4b-targeted-lora-mimic-mt-12k
- 5

Paraphrase-consistency mitigation code
github.com/UNHSAILLab/medical-vlm-paraphrase-consistency

SAFETY FRAMEWORK

- 6

VSF-Med project site
unhsaillab.github.io/VSF-Med
- 7

VSF-Med code
github.com/UNHSAILLab/VSF-Med

COMPANION STUDIES

- 8

MedMCQA robustness study code
github.com/UNHSAILLab/MedMCQA-Robustness-Study
- 9

Predictive entropy study code (UNSURE 2026)
github.com/thedatasense/predictive_entropy_unsure
- 10

Attention Without Grounding code (iMIMIC 2026)
github.com/thedatasense/medicalvlm_attention_without_grounding

ON REQUEST

Weights & Biases experiment tracking · medical-vlm-robustness research monorepo, private through the pre-defense period. Requests to contact@bineshkumar.me

bineshkumar.me/phd-thesis/reproducibility

Images are never redistributed: MIMIC-CXR and VinDr-CXR via PhysioNet, PadChest via BIMCV.

Compute Requirements

PSF-Med benchmark evaluation and equivalence audit

1,300 h · 27%



6 models, 3 datasets, roughly 93k paraphrase pairs; rubric-based judge audit of 122,778 pairs; regeneration of the PadChest paraphrase set

Uncertainty quantification pipeline

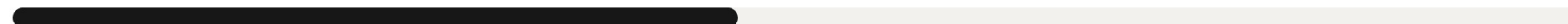
950 h · 20%



5 models, 2 datasets; softmax entropy, Monte Carlo dropout, temperature scaling, deep ensembles across 5 seeds, and text-only baselines

Corruption robustness sweep

600 h · 13%



5 corruption types, 3 severities, across models, datasets, and uncertainty methods

Safety evaluation and deployment audit

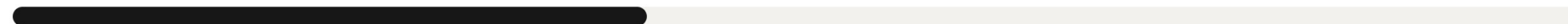
525 h · 11%



Multiple model families; image-swap, text-only, and null-image controls

Attention grounding and causal analyses

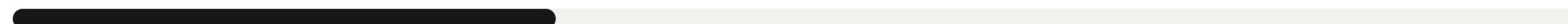
525 h · 11%



Attention extraction, patch occlusion, region-of-interest ablation, and laterality across models

LoRA training, layer and block sweeps, replication study

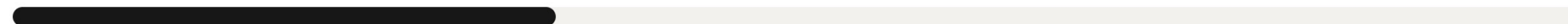
450 h · 9%



Includes the five-seed replication study

Mechanistic interpretability

450 h · 9%



Sparse-autoencoder screening and transfer, activation and residual-transplant patching, canonical validation, and lens fitting

Total: ~4,800 GPU-hours (~200 GPU-days). The line items above sum to 4,800 hours; the reported total is rounded and includes runs that produced nothing publishable.

Scope of this thesis

- Single modality (chest X-ray)
- Binary Q&A
- Two architectures for Interpretability

Future Work

**Formal clinician-preference /
deployment-readiness study**

**Cross-architecture mechanism, beyond
MedGemma-4B**

**Multimodal extension to CT and MRI,
Time Series**

What did I learn?

1. Ask the same question twice. This whole thesis grew out of one small habit of distrust.
2. The most dangerous result is the one that looks like success. The control is the experiment. I learned that a finding is only as strong as the frame around it.
3. Rejections were useful. Every rejected paper was a learning experience. Asking an early question beats defending a late mistake, and every time I talked about my research, I learned a new direction.
4. Nobody finishes a dissertation alone, not me. The clinical partners who audited data. My advisor, committee, UNH staff, co-authors, my lab, and the researchers I cold-emailed, and my family and colleagues: six years and 4,800 GPU hours