

Interpretability by Design

Aya Abdelsalam Ismail



Chatbot



Self-driving



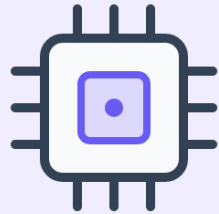
Medical imaging



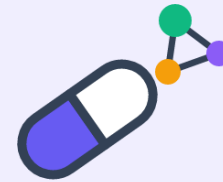
Stock trading



AI assistant



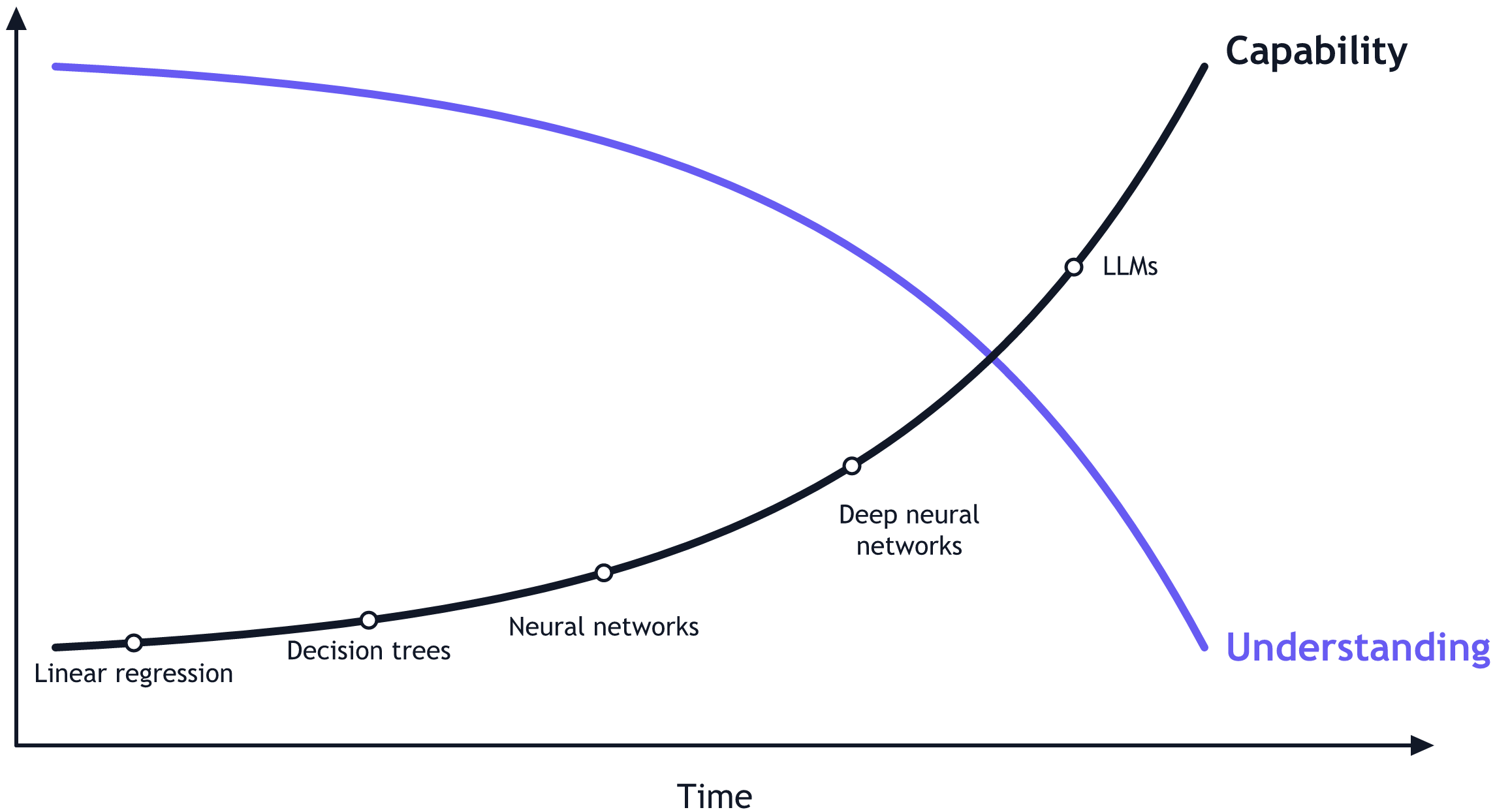
Chip design



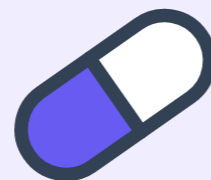
Drug discovery



Weather forecast



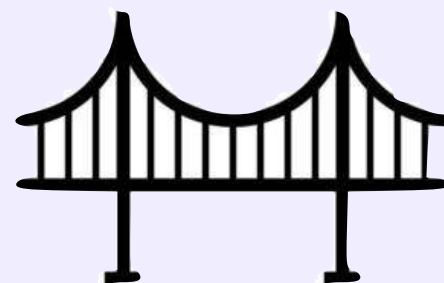
AI



Medicine



Airplanes



Concrete



Batteries

AI

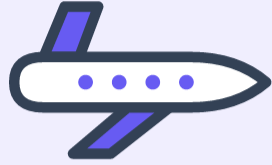


Do we need to understand AI?

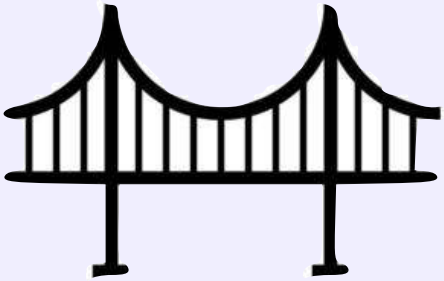
Do we need to understand AI?



Medicine



Airplanes



Concrete



Batteries

These technologies
have **earned**
humans trust.

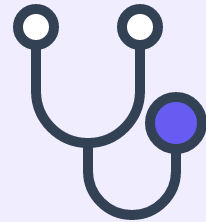
Do we need to understand AI?



Legal



Finance



Medical diagnosis

AI has **not** yet earned our trust

AI has **not** yet earned our trust

AI (artificial intelligence)

• This article is more than 2 months old

AI hallucinations found in high-profile Wall Street law firm filing

Sullivan & Cr
for string of e

● Business I



AI (artificial intelligence)

• This article is more than 1 year old

High court tells UK lawyers to stop misuse of AI after fake case-law citations

Ruling follows two cases blighted by actual or suspected use of artificial intelligence in legal work



Robert Booth UK
technology editor

Fri 6 Jun 2025 20.00 BST

Share

Prefer the Guardian on Google

US healthcare

• This article is more than 1 year old

New AI tool counters health insurance denials decided by automated algorithms

Class-action lawsuits allege algorithms turn down claims in seconds and critics say reform is needed for lasting change



Mark Sweney

Wed 22 Apr 2026 09.09 BST

Share

Prefer the Guardian on Google

News

Opinion

Sport

Culture

Lifestyle



The Guardian

Int v

World Europe US news Americas Asia Australia Middle East Africa Inequality Global development

France

• This article is more than 7 months old

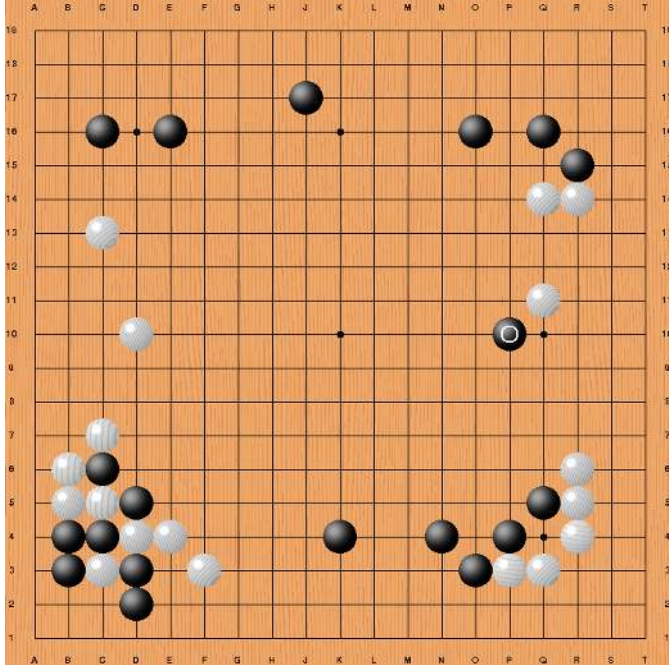
Facebook's job ads algorithm is sexist, French equality watchdog rules

Regulator found ads for mechanics skewed towards men while those for preschool teachers targeted women

Advertisement

Info

Do we need to understand AI?



As AI becomes superintelligent, it would be invaluable if it could **teach us** what it has learned.

Standard training recipe

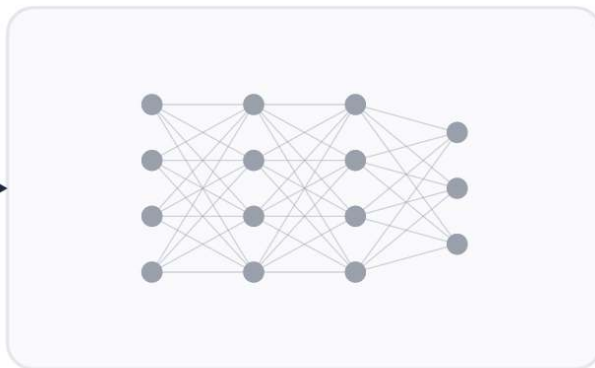
Standard

Data



As much as possible

Architecture



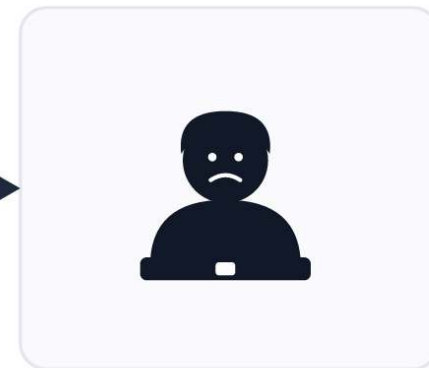
As big as possible

Training



Pretraining, Midtraining,
Long-context
extension, SFT, Think,
DPO/GRPO, RLVR, ...

Deployment



Let's try to
explain the
model

Explaining **any** model post-hoc

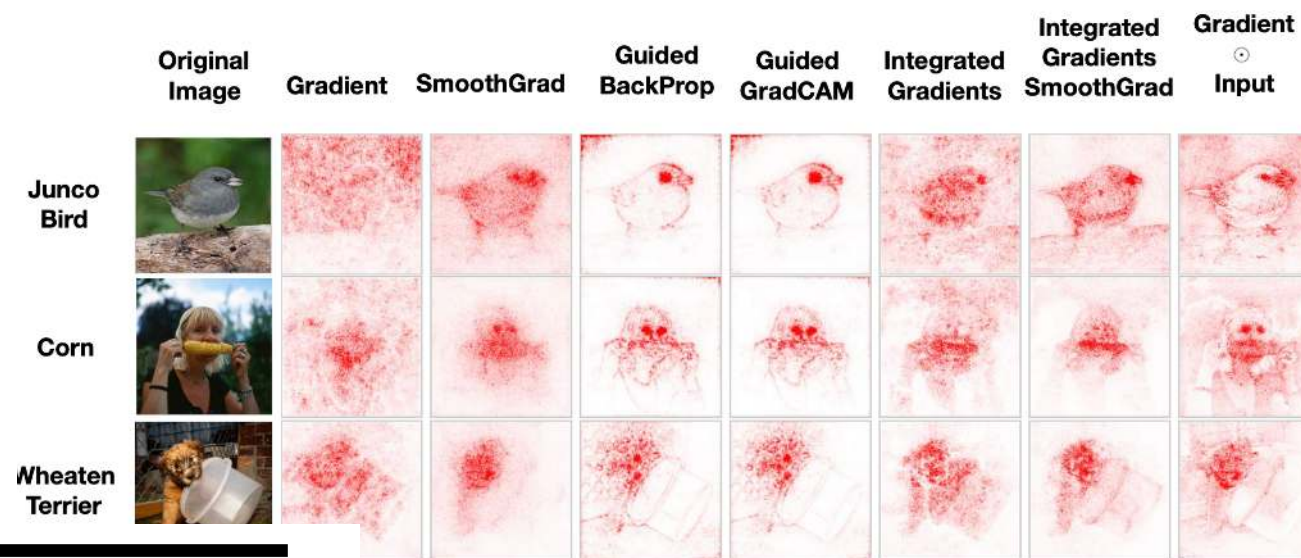


Feature Attribution

Explaining **any** model post-hoc



Feature Attribution



Sanity Checks for Saliency Maps

Julius Adebayo*, Justin Gilmer[#], Michael Muelly[#], Ian Goodfellow[#], Moritz Hardt[†], Been Kim[#]

juliusad@mit.edu, {gilmer,muelly,goodfellow,mrtz,beenkim}@google.com

[#]Google Brain

[†]University of California Berkeley

Impossibility Theorems for Feature Attribution

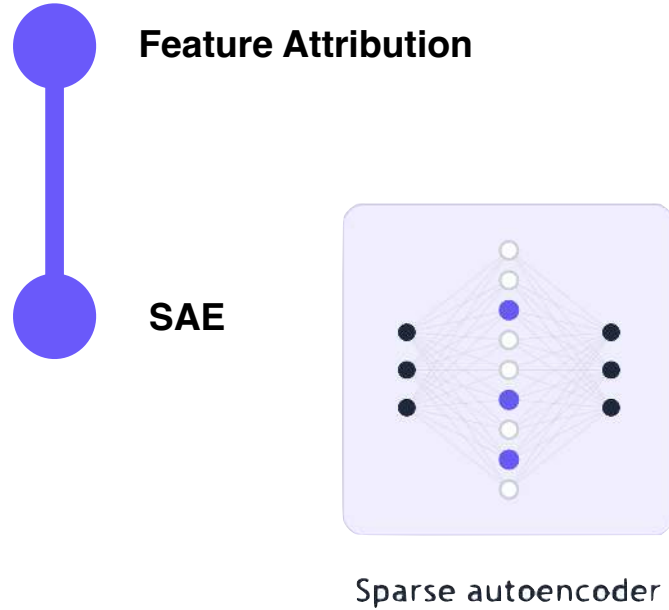
Blair Bilodeau^{*1}, Natasha Jaques^{2,3}, Pang Wei Koh^{2,3}, and Been Kim³

¹University of Toronto

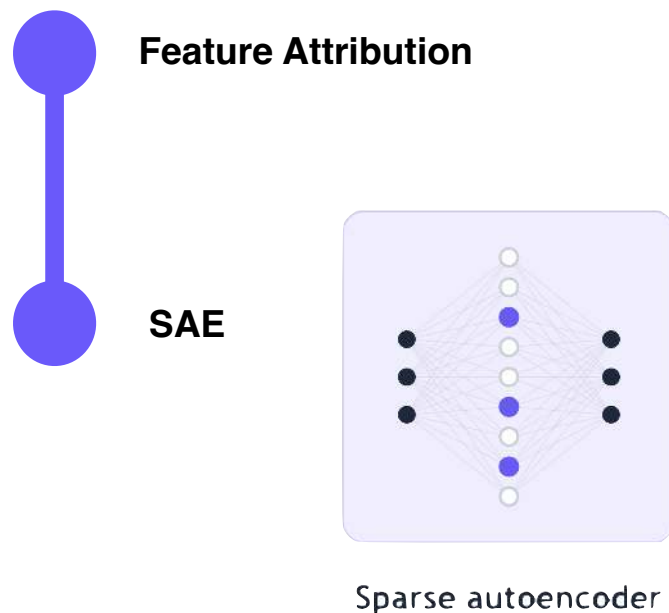
²University of Washington

³Google Deepmind

Explaining **any** model post-hoc



Explaining **any** model post-hoc



Interpretability Illusions with Sparse Autoencoders: Evaluating Robustness of Concept Representations

Aaron J. Li¹, Suraj Srinivas², Usha Bhalla¹, and Himabindu Lakkaraju¹

AUTOMATED INTERPRETABILITY METRICS DO NOT DISTINGUISH TRAINED AND RANDOM TRANSFORMERS

Thomas Heap
University of Bristol
Bristol, UK
thomas.heap@bristol.ac.uk

Tim Lawson
University of Bristol
Bristol, UK

Lucy Farnik
University of Bristol
Bristol, UK

Laurence Aitchison
University of Bristol
Bristol, UK

Sanity Checks for Sparse Autoencoders: Do SAEs Beat Random Baselines?

Anton Korznikov^{*}, Andrey Galichin^{*}, Alexey Dontsov, Oleg Y. Rogov, Ivan Oseledets, Elena Tutubalina



AXBENCH: Steering LLMs? Even Simple Baselines Outperform Sparse Autoencoders

ON THE LIMITS OF SPARSE AUTOENCODERS: A THEORETICAL FRAMEWORK AND REWEIGHTED REMEDY

Jingyi Cui^{1*}, **Qi Zhang^{1*}**, **Yifei Wang²**, **Yisen Wang^{1,3†}**

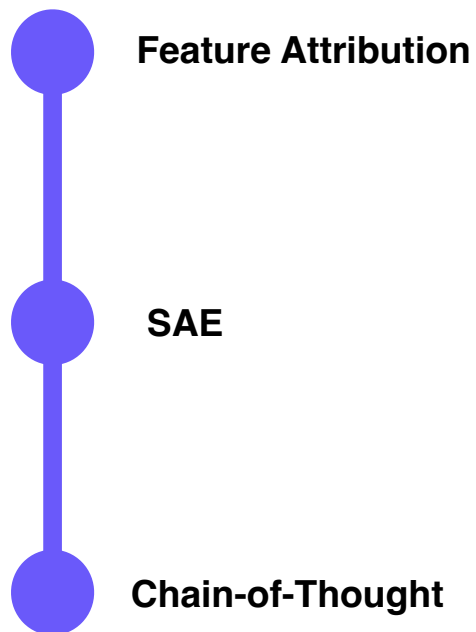
¹State Key Lab of General AI, School of Intelligence Science and Technology, Peking University

²Amazon AGI SF Lab[†]

³Institute for Artificial Intelligence, Peking University

Zhengxuan Wu^{*1}, **Aryaman Arora^{*1}**, **Atticus Geiger²**, **Zheng Wang¹**, **Jing Huang¹**,
Dan Jurafsky¹, **Christopher D. Manning¹**, **Christopher Potts¹**

Explaining **any** model post-hoc



Explaining **any** model post-hoc



Feature Attribution



SAE



Chain-of-Thought

Language Models Don't Always Say What They Think: Unfaithful Explanations in Chain-of-Thought Prompting

Miles Turpin,^{1,2} Julian Michael,¹ Ethan Perez,^{1,3} Samuel R. Bowman^{1,3}
¹NYU Alignment Research Group, ²Cohere, ³Anthropic
miles.turpin@nyu.edu

Reasoning Models Don't Always Say What They Think

Yanda Chen Joe Benton Ansh Radhakrishnan Jonathan Uesato Carson Denison
John Schulman⁺ Arushi Somani

Peter Hase⁺ Misha Wagner Fabien Roger Vlad Mikulik
Samuel R. Bowman Jan Leike Jared Kaplan Ethan Perez

Chain-of-Thought Reasoning In The Wild Is Not Always Faithful

Iván Arcuschin*
University of Buenos Aires
iarcuschin@dc.uba.ar

Jett Janiak*

Robert Krzyzanowski*

Senthooran Rajamanoharan

Neel Nanda

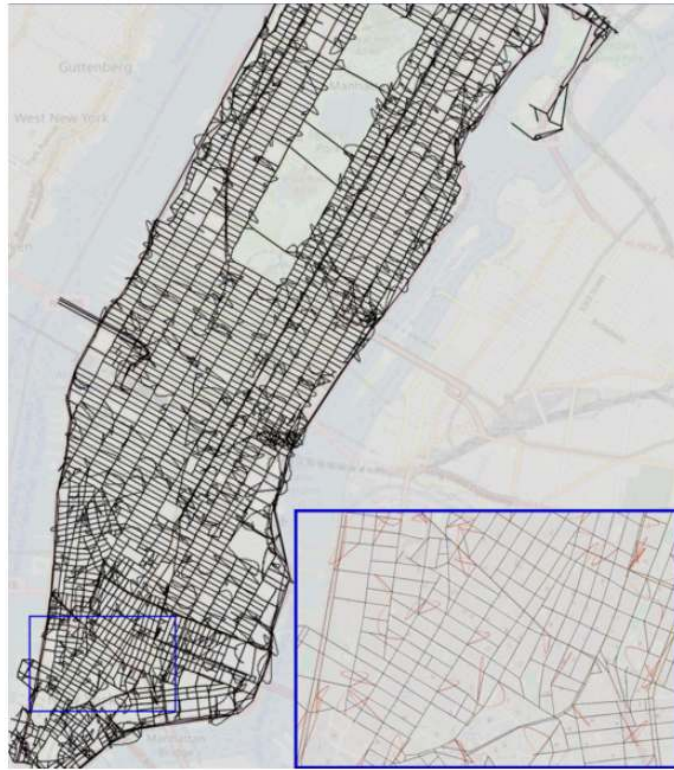
Arthur Conmy
arthurconmy@gmail.com

Alignment Science Team, Anthropic

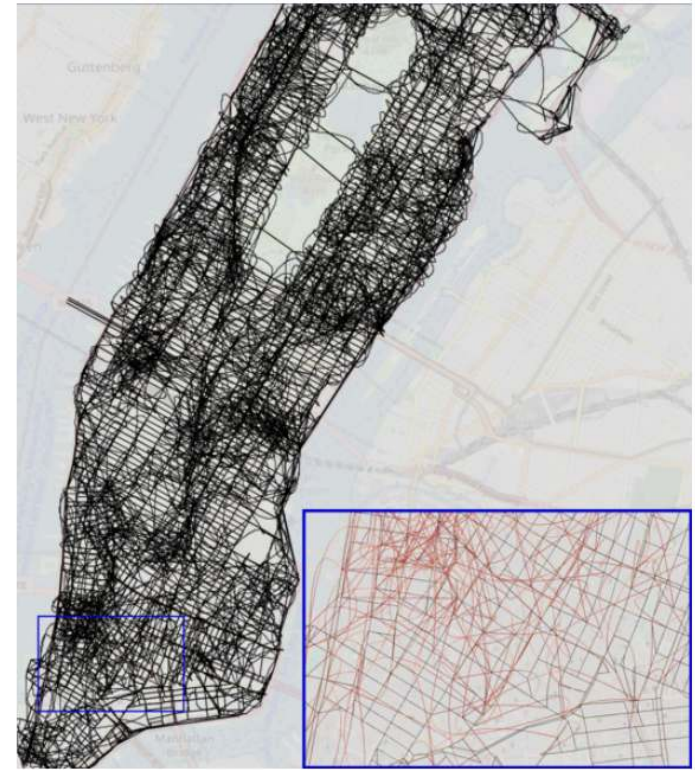
Why does this keep happening?



(a) World model



(b) World model with noise

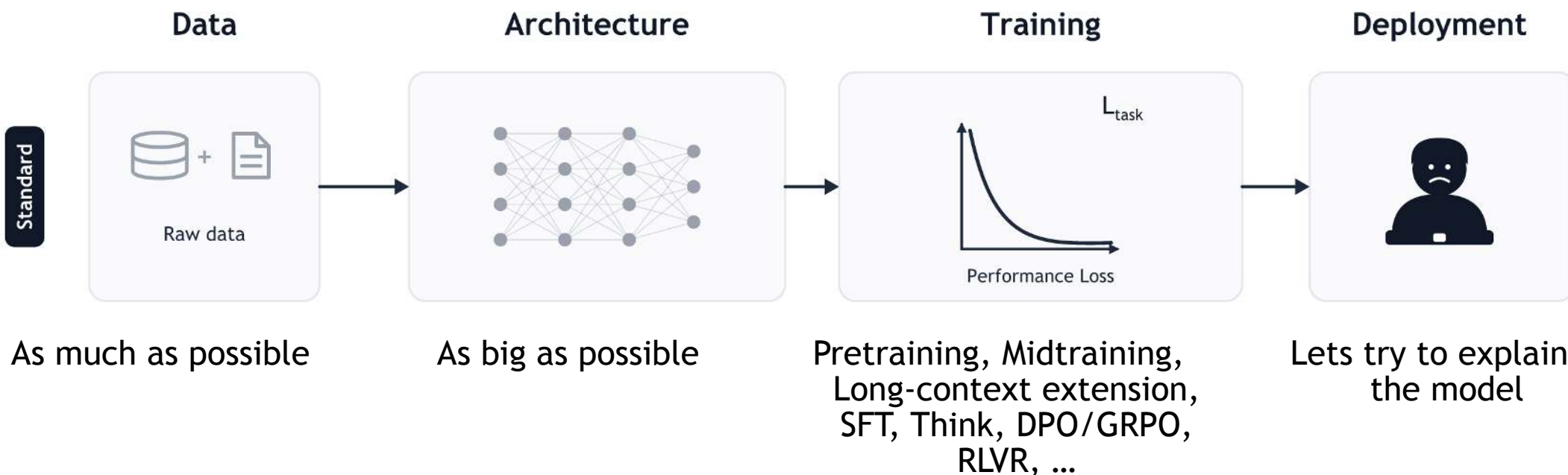


(c) Transformer

Image from [Vafa et al.](#) (a) True Manhattan Map. (b) True map with noise perturbations (c) A transformer trained on turn-by-turn direction produces entangled, chaotic internal states, despite generating correct predictions. Capability and internal coherence are not the same thing.

**We are trying to explain models
that were not meant to be
explained**

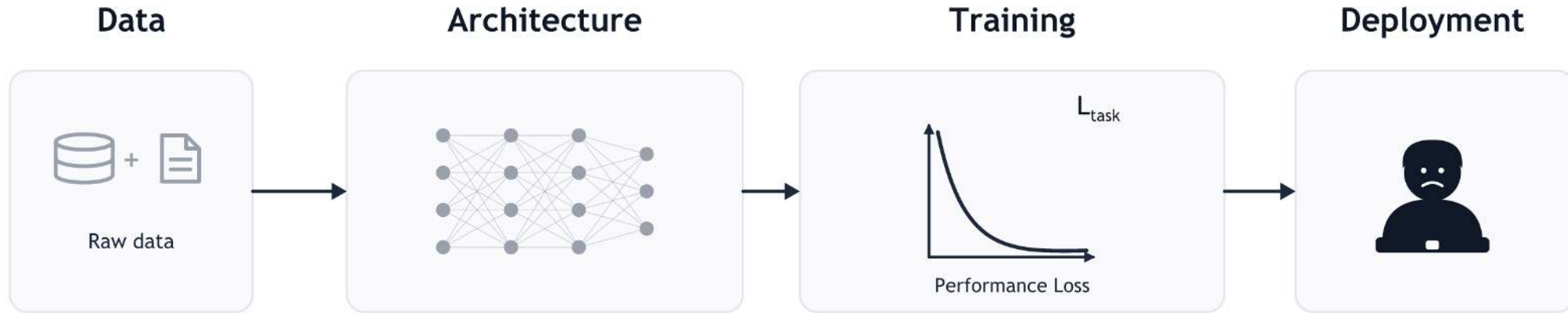
Standard training recipe



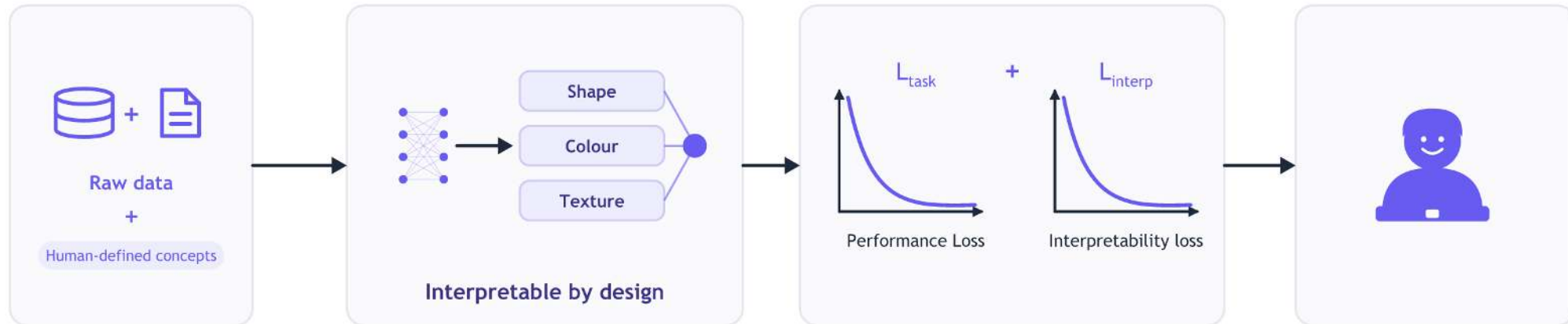
Capabilities **emerge** because we optimize for them.
Why expect interpretability when we never do?

Interpretability at every step, by design

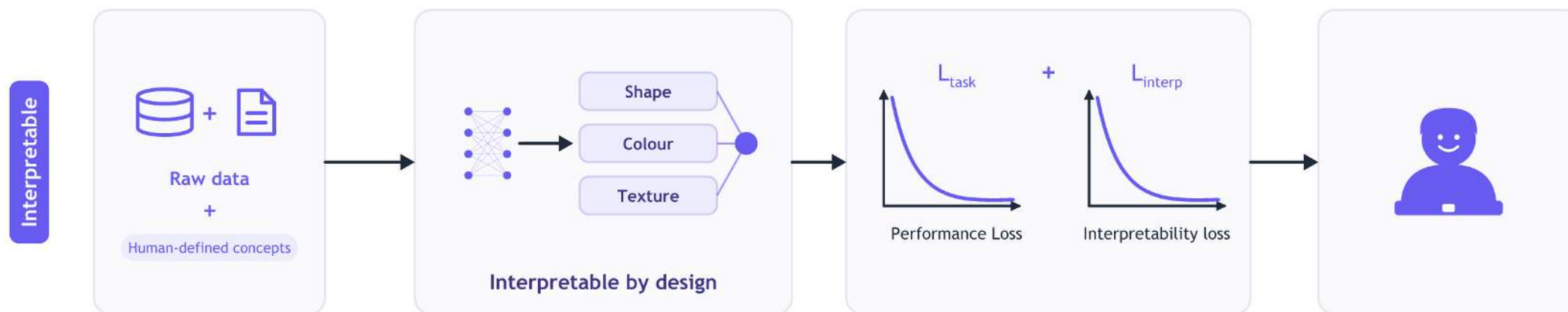
Standard



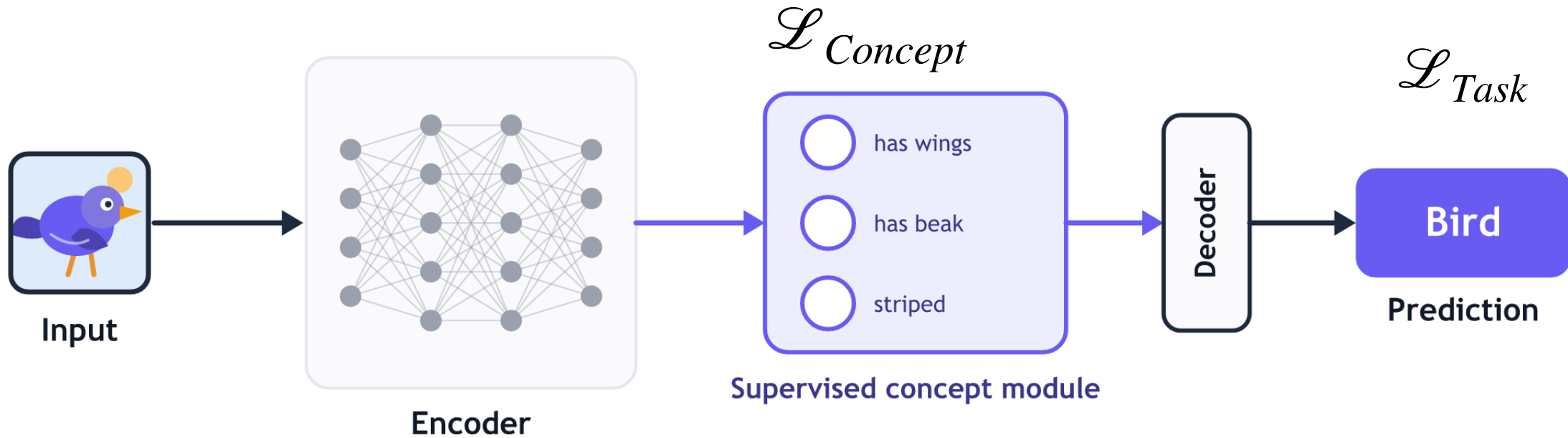
Interpretable



Interpretable training recipe

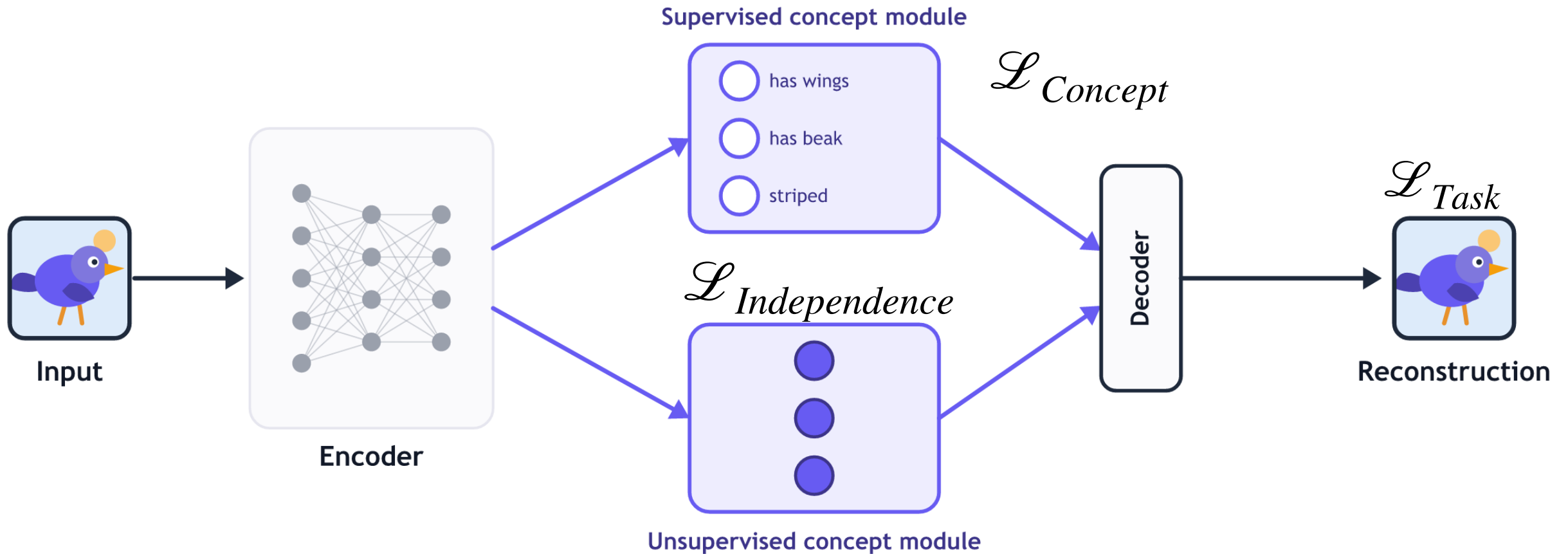


Concept bottleneck models



$$\mathcal{L} = \mathcal{L}_{Task} + \alpha \mathcal{L}_{Concept}$$

Concept bottleneck **generative** models



$$\mathcal{L} = \mathcal{L}_{Task} + \alpha \mathcal{L}_{Concept} + \beta \mathcal{L}_{Independence}$$

Ismail et al., "Concept Bottleneck Generative Models," ICLR 2024.

Interpretable training recipe



Performance



Interpret &
Intervene



Discover



Scaling

Interpretable training recipe

Does this training recipe work?



Performance

Interpretable training recipe

Published as a conference paper at ICLR 2024



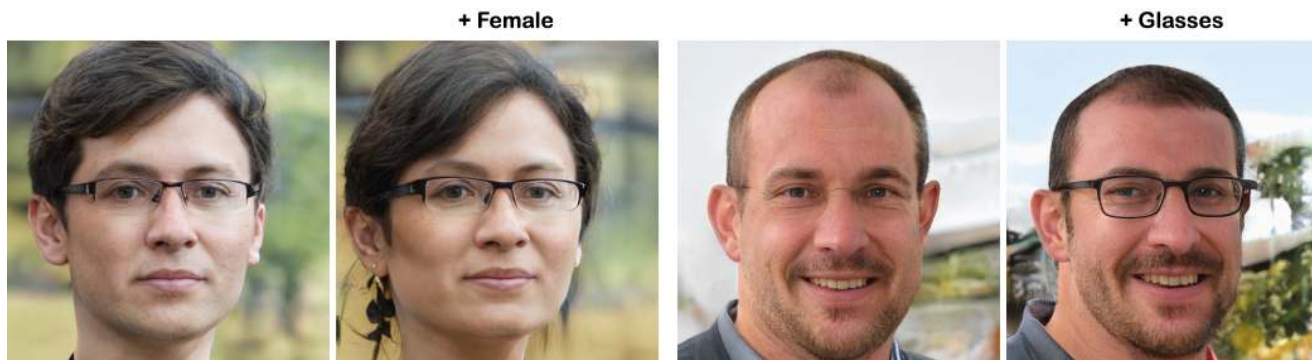
Performance

Images

CONCEPT BOTTLENECK GENERATIVE MODELS

Aya Abdelsalam Ismail^{1,2} * Julius Adebayo³ * Héctor Corrada Bravo¹
Stephen Ra^{1,2} Kyunghyun Cho^{1,2,4,5}

(a) CB-StyleGAN2 on FFHQ



(b) CB-DDPM on LAION Aesthetics



Interpretable training recipe

Can we do anything **useful** with this?



Interpret &
Intervene

Interpretable training recipe

Published as a conference paper at ICLR 2025

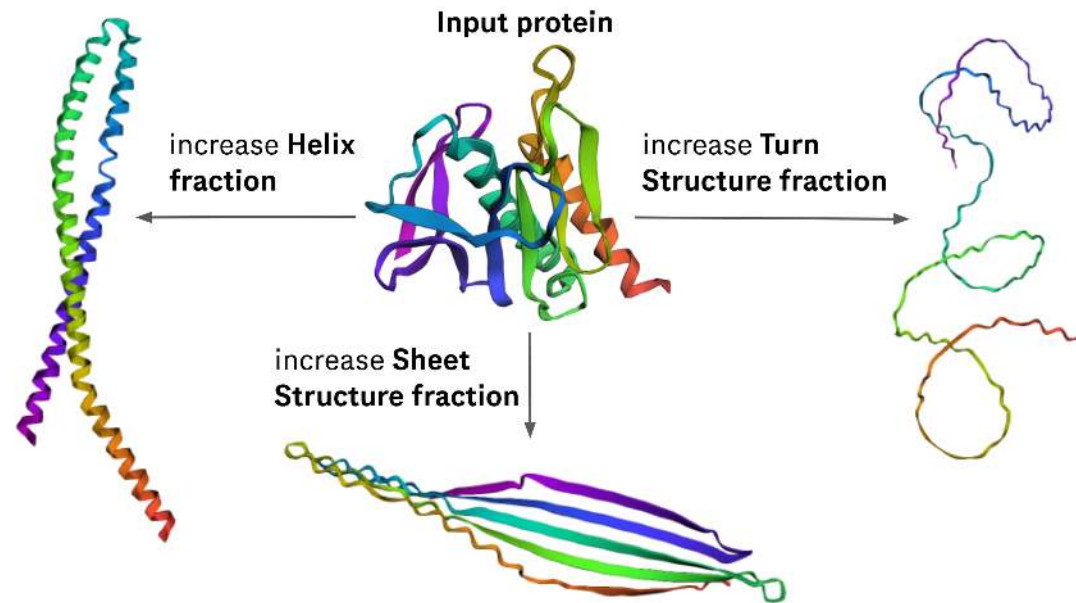


Interpret &
Intervene

Protein

CONCEPT BOTTLENECK LANGUAGE MODELS FOR PROTEIN DESIGN

Aya Abdelsalam Ismail^{1,2} Tuomas Oikarinen³ Amy Wang^{1,2} Julius Adebayo⁴
Samuel Stanton^{1,2} Héctor Corrada Bravo¹ Kyunghyun Cho^{1,2,5,6} Nathan C. Frey^{1,2}



Interpretable training recipe

Can we learn something new?



Discover

Interpretable training recipe

ICML 2026



Discover

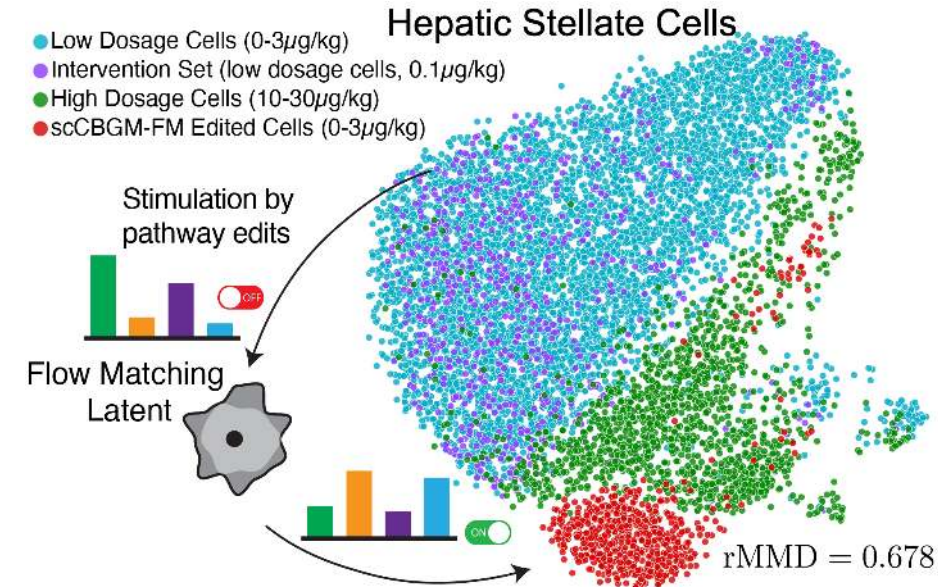
Single cell

Some liver cells **respond** to a drug, some **don't**. We found the pathway signature that separates them ($\uparrow$ TGFB, $\uparrow$ PI3K, $\uparrow$ p53, $\downarrow$ Trail, $\downarrow$ VEGF).

Edit that signature into **non-responders** $\rightarrow$ they become **responders** the model was never trained on.

scCBGM: Single-Cell Editing via Concept Bottlenecks

Alma Andersson^{*1} Aya Abdelsalam Ismail^{*2} Edward De Brouwer^{*1} Doron Haviv^{*1} Tommaso Biancalani¹
Kyunghyun Cho¹³⁴ Gabriele Scalia¹ Aicha BenTaieb¹ Hector Corrada Bravo¹



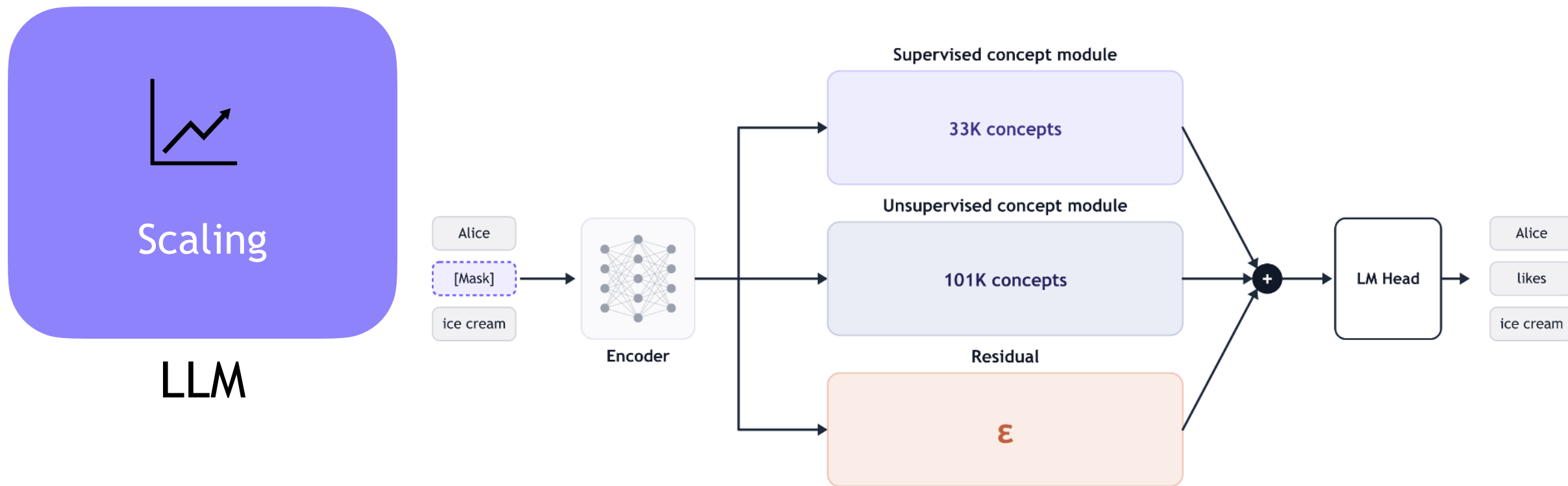
Interpretable training recipe

Does this **scale**?



Scaling

Interpretable training recipe

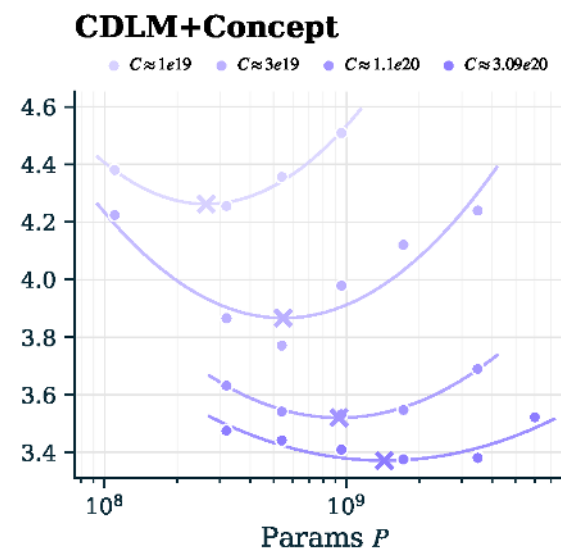
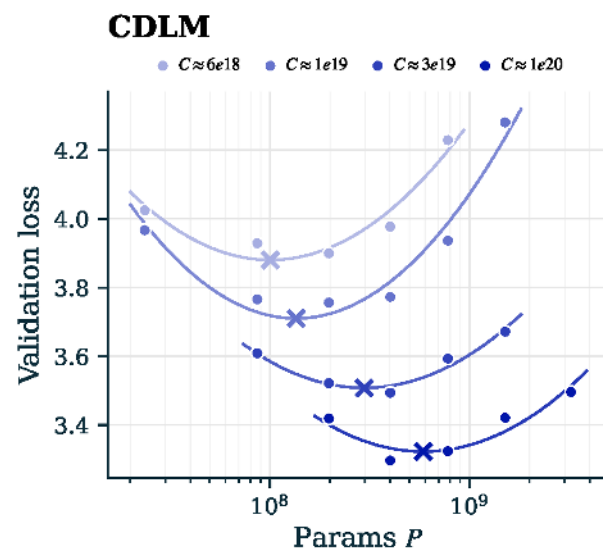
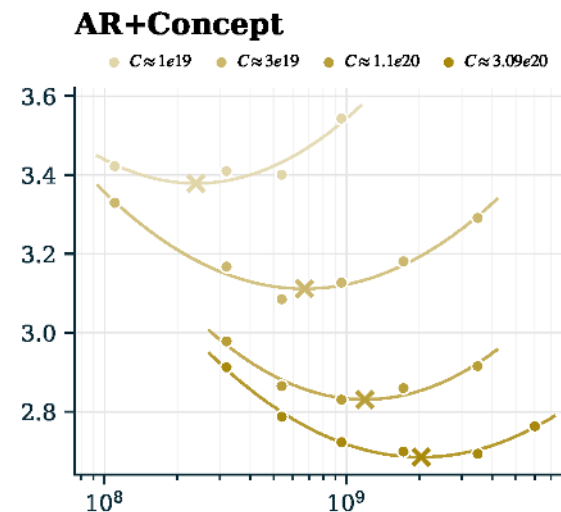
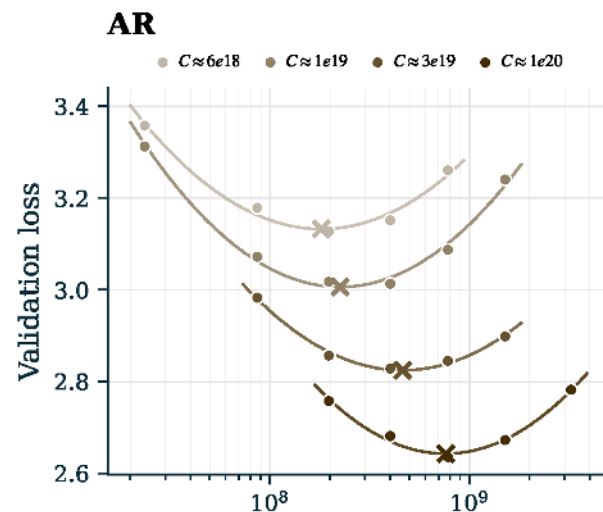


Interpretable training recipe



Scaling

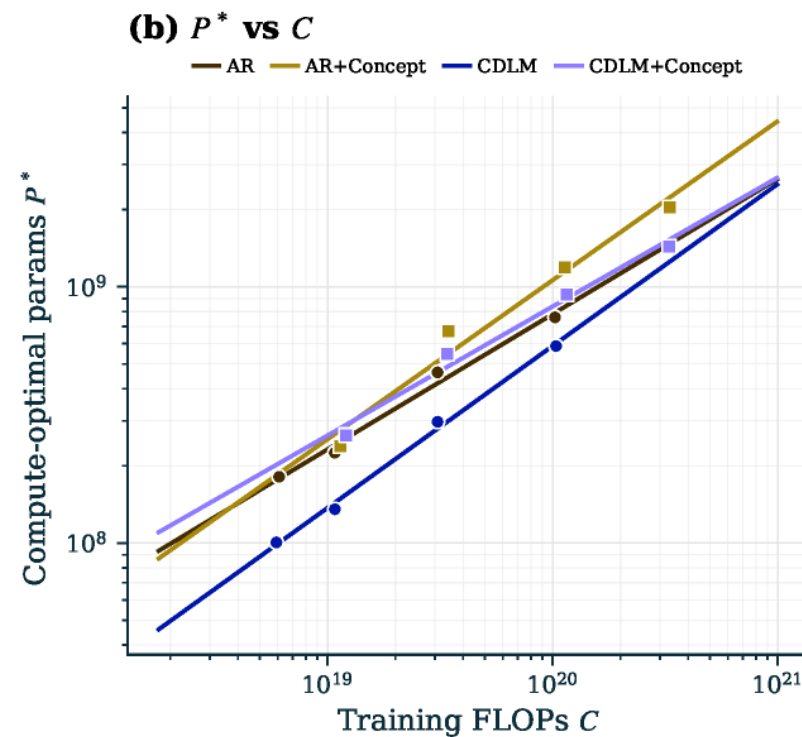
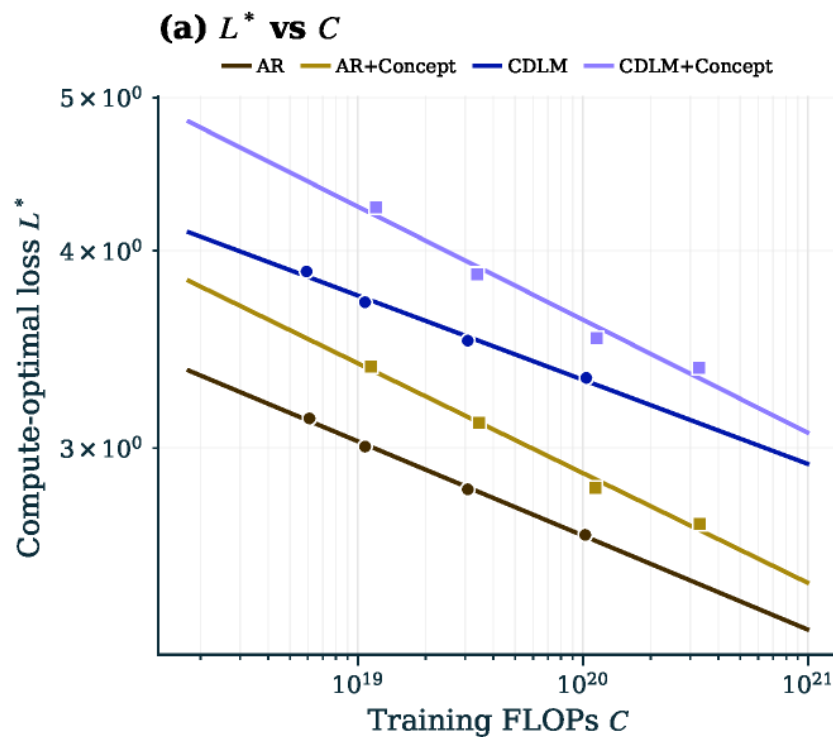
LLM



Interpretable training recipe



LLM



Interpretability-by-design imposes a **small** per-backbone **offset** on compute-optimal scaling, **not a scaling tax**.

Interpretable training recipe



Scaling

LLM

Scaling Inherently Interpretable
Language Models

 Guide Labs Team



Steerling-8B

Interpretable training recipe

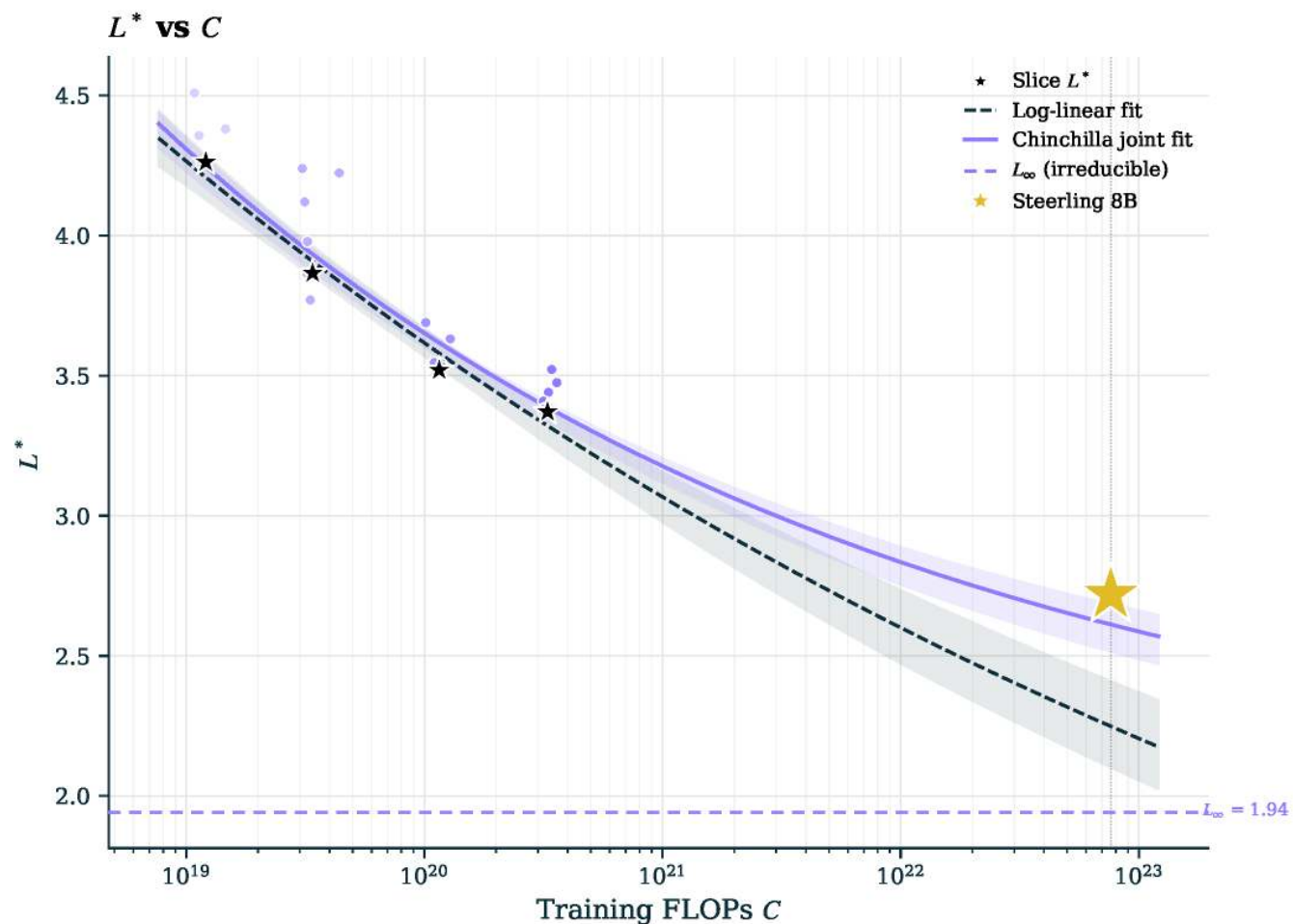


Steering-8B



Scaling

LLM



Interpretable training recipe



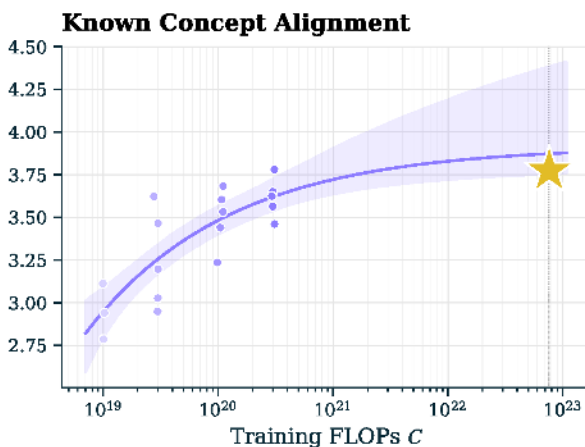
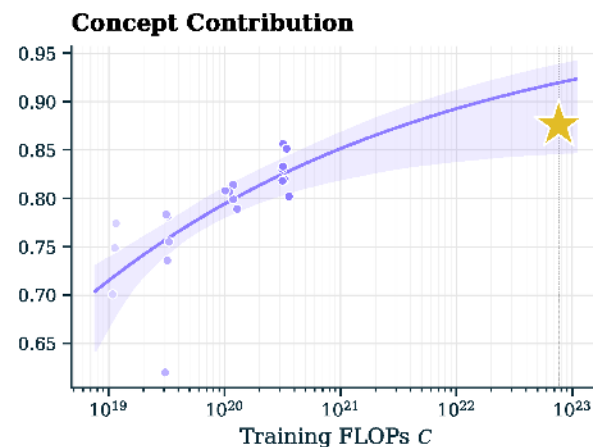
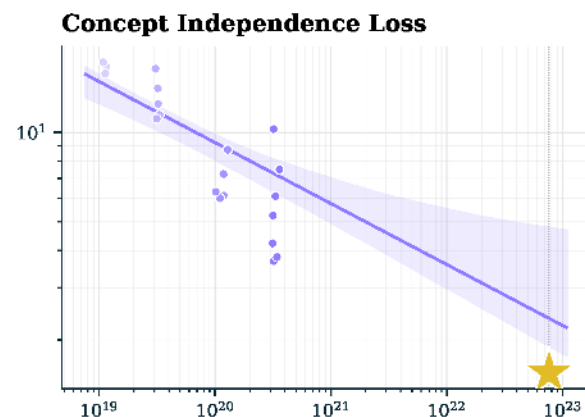
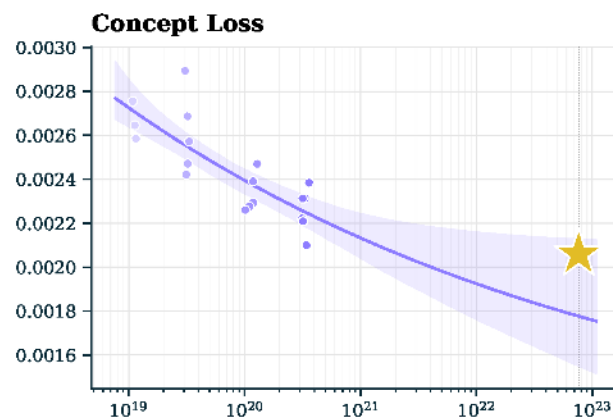
Steering-8B



Scaling

LLM

• $C \approx 1e19$ • $C \approx 3e19$ • $C \approx 1.1e20$ • $C \approx 3.09e20$ ★ Steering 8B



Interpretable training recipe



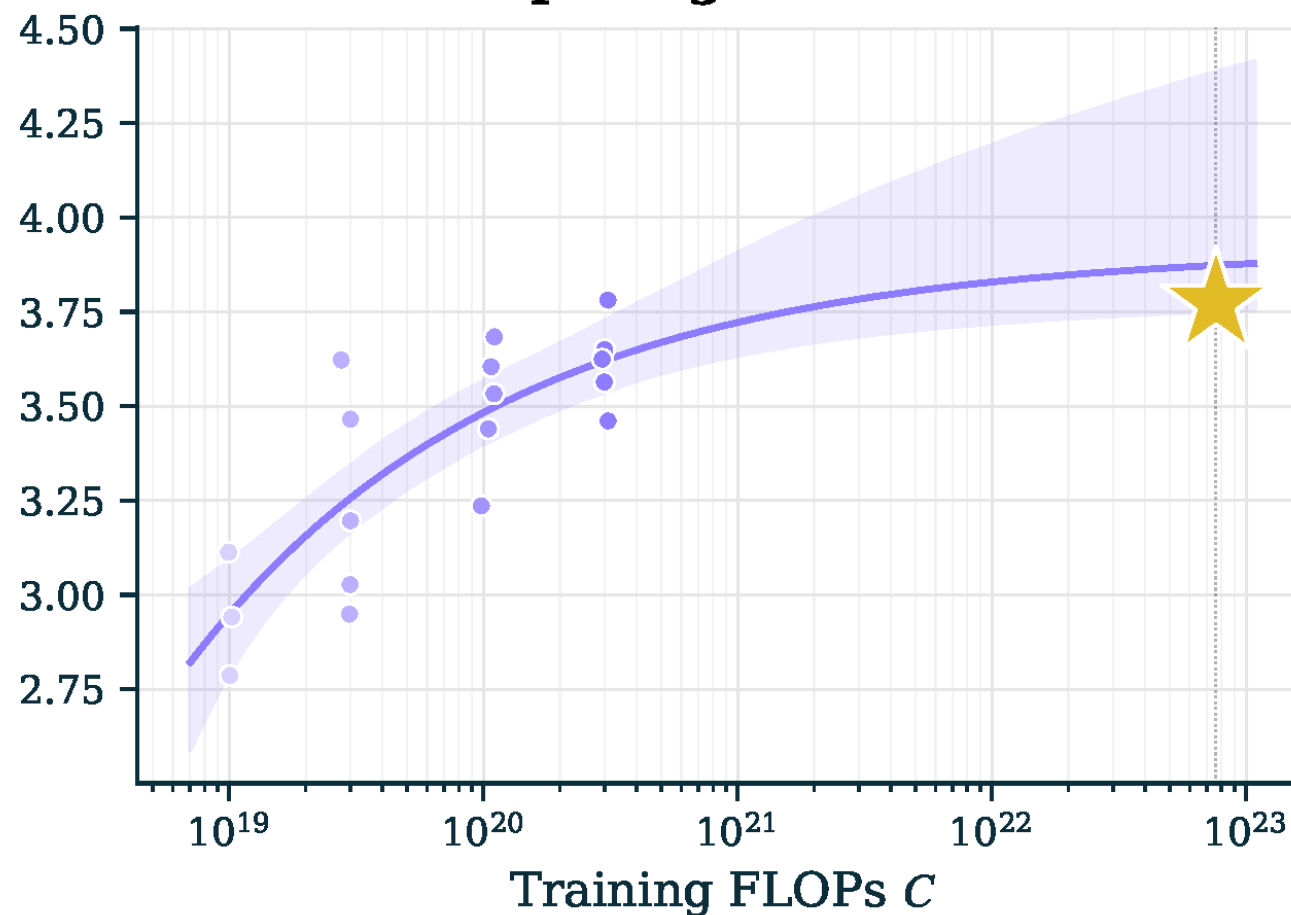
Steering-8B



Scaling

LLM

Known Concept Alignment



Steering-8B: First Inherently Interpretable LLM at Scale

A 58 year old male presents with sudden onset chest pain, diaphoresis, and left arm numbness. Walk through a differential diagnosis.



Steering-8B Instruct

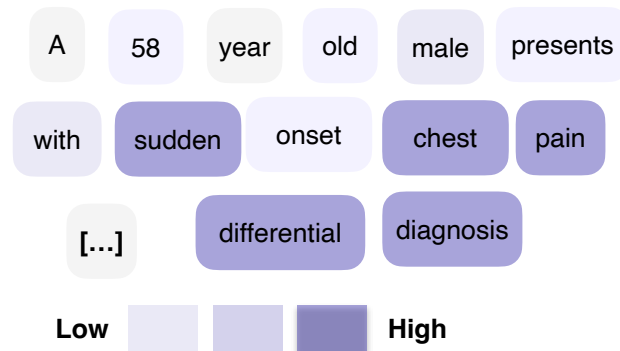
Differential Diagnosis:

1. Acute Coronary Syndrome (ACS):

- Chest pain: Unilateral, crushing: radiating to left arm and jaw.
- Diaphoresis with fever (>38C/100F)
- Left-arm numbness or weakness

Input Attribution

Which input tokens are most influential?



Concept Attribution

What concepts is the model using?

Cardiology and Cardiac Disease



Medical Symptoms



Pain and Headache



Vascular and Circulatory Disease



Biomedical Research and Vocabulary



Training Data Attribution

Which training data causes the output?

Symptoms Pain may be due to an isolated incident or something more severe. The pain may be isolated to an area in the right chest or maybe felt throughout. The pain also may come and go or may be persistent. [...]

Source: nemotron_synth
Id: 01836-065100

0.96

Steering-8B: First Inherently Interpretable LLM at Scale

→]

✎

⚙

Korean Cuisine and Language: 60%

What should I do today?

✎

It sounds like you're craving a spicy and flavorful dish that's perfect for your taste buds. Here are some options to consider:

1. Try making traditional Korean dishes such as bibimbap, Bulgogi, or even kimchi! These dishes are packed with bold flavors and will be sure to satisfy your craving.
2. If you're looking for something a little more unique, why not try making some spicy ramen or Korean barbecue? These options are perfect for those who enjoy a bit of heat in their meals!
3. For an even bigger kick, why not try making some spicy egg rolls or Korean-style fried chicken? These dishes are sure to give your taste buds a thrill and leave them craving for more.

No matter what dish you choose today, I hope it brings you joy and satisfaction. Enjoy your meal!

👍 🗨 🔄

Say something...

Korean Cuisine and Language: 60%

Send ↑

Clarity is AI powered and can produce inaccurate and harmful content. Please double check any responses.

Chat explanations

Concept

Training Data

Concept attributions — What concepts inside the model are responsible for generating the content in this selection?

Search

Korean Cuisine and Language ⓘ

Amplified 1 attributed chunks

Cooking and Recipes ⓘ

2 attributed chunks

Cooking and Recipes ⓘ

2 attributed chunks

Restaurants and Fast Food ⓘ

2 attributed chunks

Joyful Celebratory Affect ⓘ

2 attributed chunks

Culinary and Restaurant Vocabulary ⓘ

2 attributed chunks

Indefinite Vague Reference ⓘ

2 attributed chunks

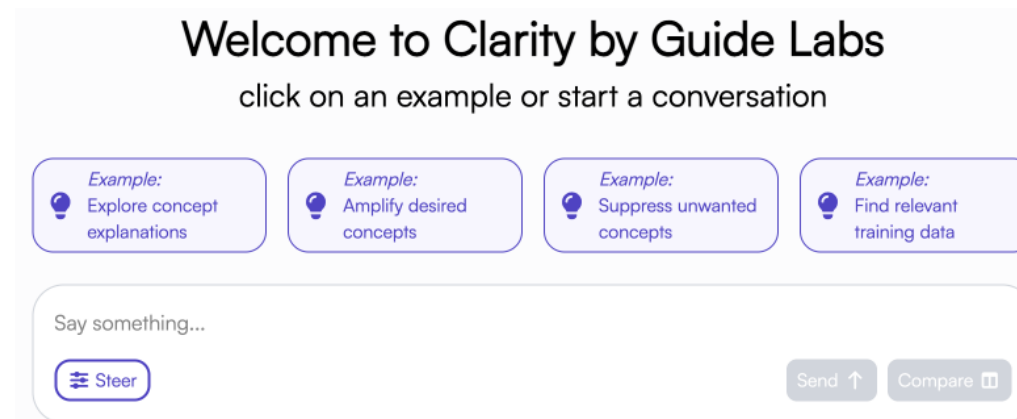
Steering-8B: First Inherently Interpretable LLM at Scale

Base model and SFT models 🤗: <https://huggingface.co/guidelabs/steering-8b>

Code on GitHub: <https://github.com/guidelabs/steering>

Technical report: <https://www.guidelabs.ai/papers/scaling-inherently-interpretable-language-models/>

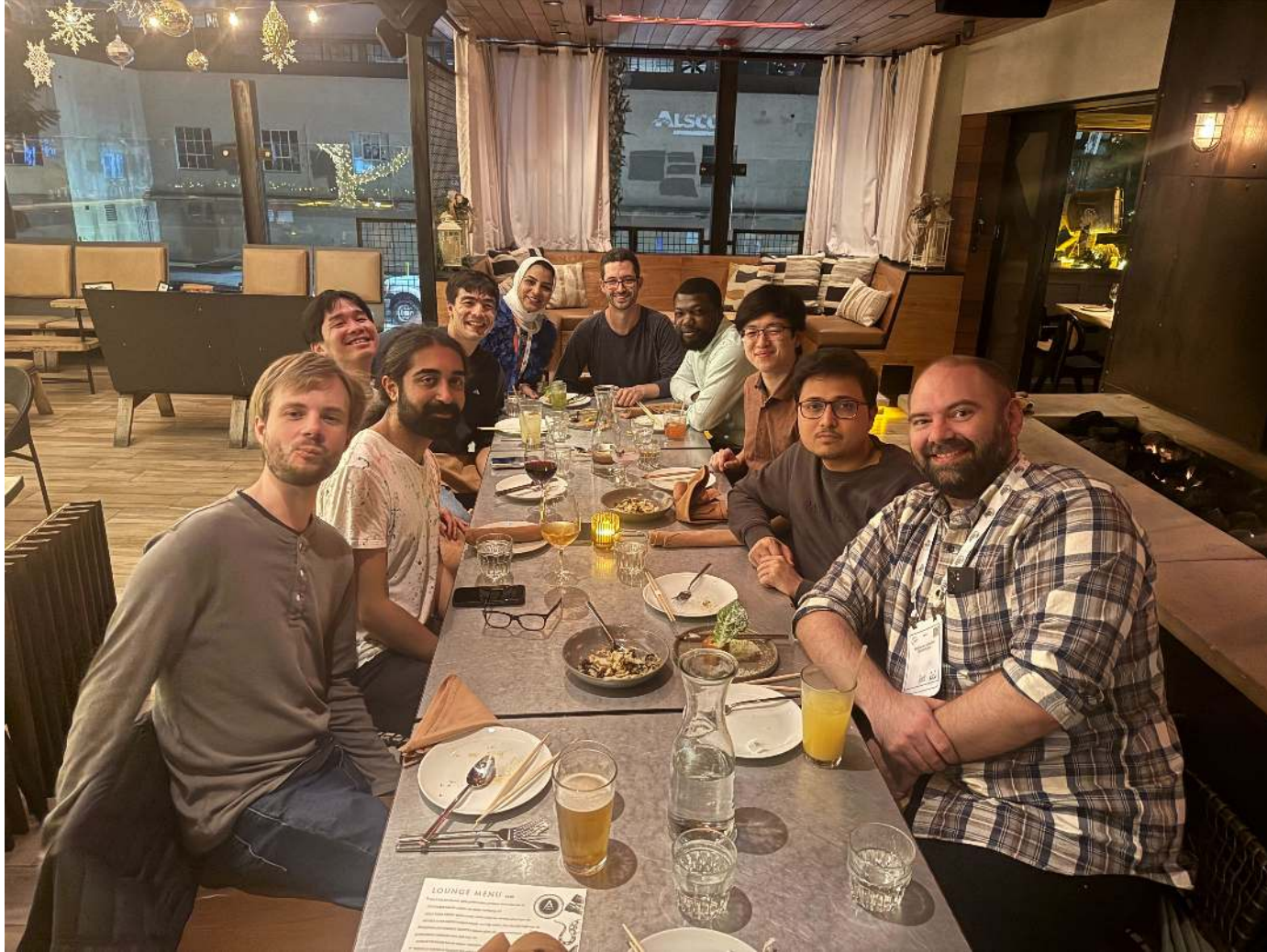
Clarity platform: <https://platform.guidelabs.ai/>



TL;DR: If you want interpretability, train interpretable models.

This is one instantiation. There are others, and probably better ones.

Finding them is why we're all here today.



Questions?

