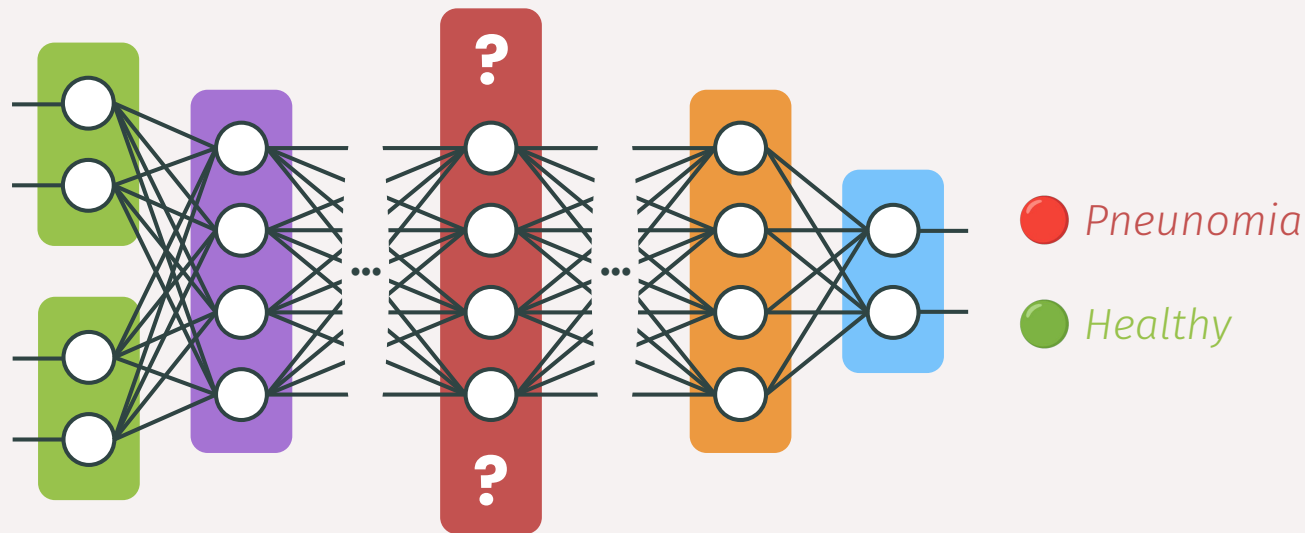


AIMLAI @ ECML PKDD 2026
September 7th 2026

From Classical to Mechanistic Interpretability

Ward Gauderis · José Oramas · Thomas Dooms

To control, predict, and learn from neural networks, we have to understand their *internal* working.



Monitor & steer

Align & debug

Improve design

Extract knowledge

This chapter...

Mechanistic foundations

Ward Gauderis

Methods & Models

José Oramas & Thomas Dooms

Limitations & Open Problems

Thomas Dooms

Mechanistic interpretability *reverse-engineers* learned representations into human concepts and algorithms.

Question

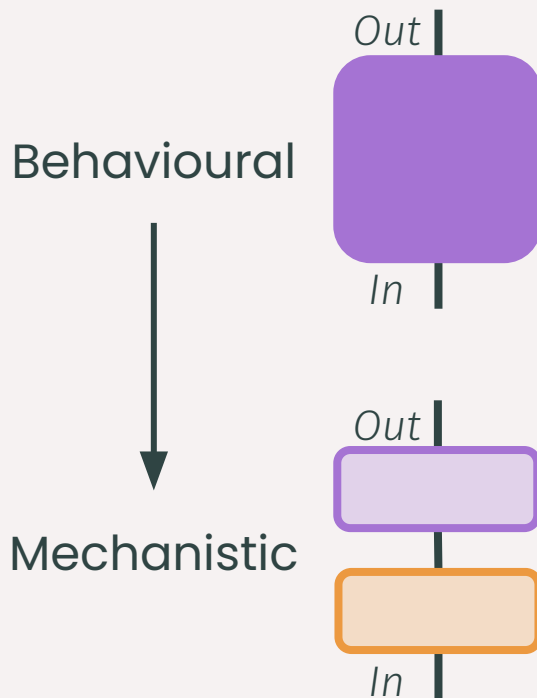
How can a model solve a **general** class of problems, not just one instance?

Premise

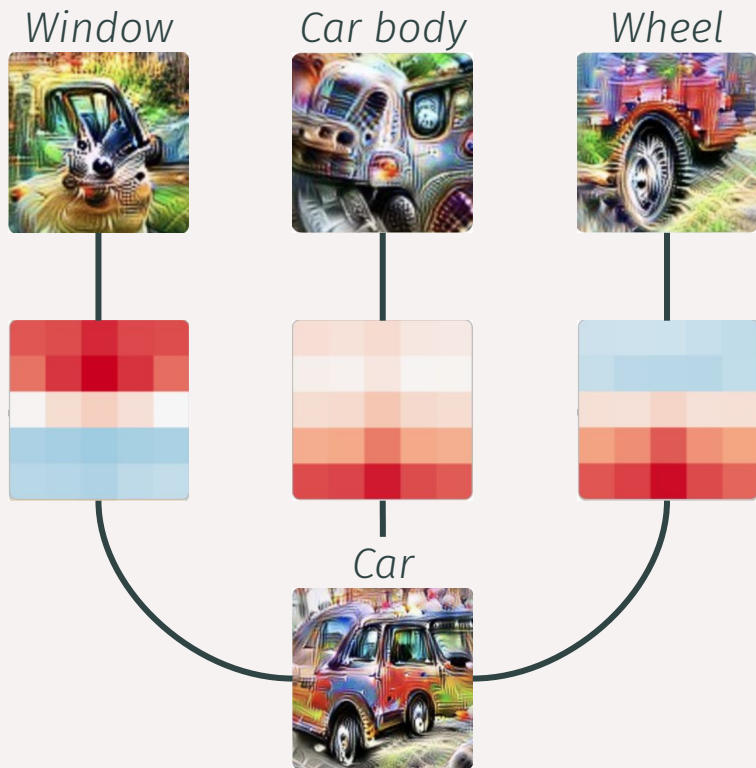
Generalisation in neural networks arises from shared computational **mechanisms**.

Scope

Any research that describes the internals of a model, including its **activations or weights**.



Axis 1. Networks are *features*, wired into *circuits*. Both can recur across models and tasks.



1. Object of study

Features

What input properties are encoded and how?

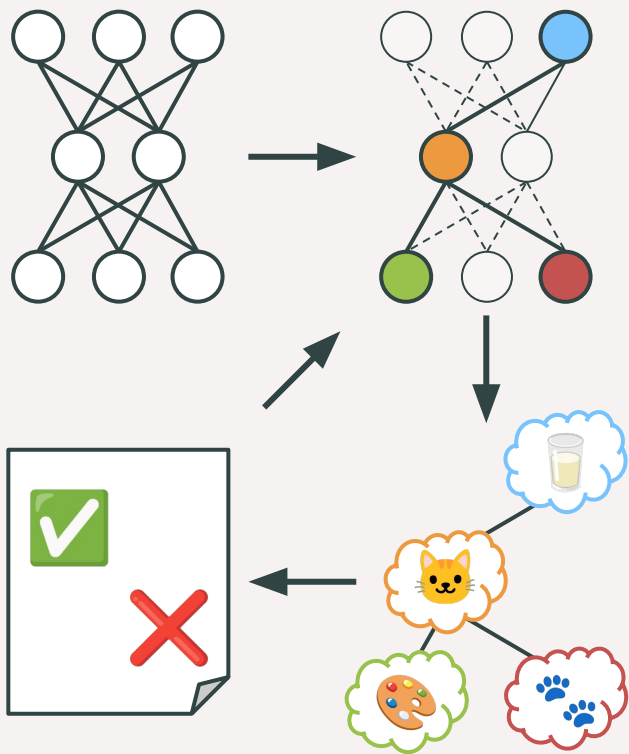
Circuits

How do features compute and interact?

Universality

Are these objects learned across models?

Axis 2. Decompose, hypothesise, and test. Reverse-engineering follows the *scientific method*.



1. Object of study

Features

Circuits

Universality

2. Process stage

Decomposition

What are the relevant components?

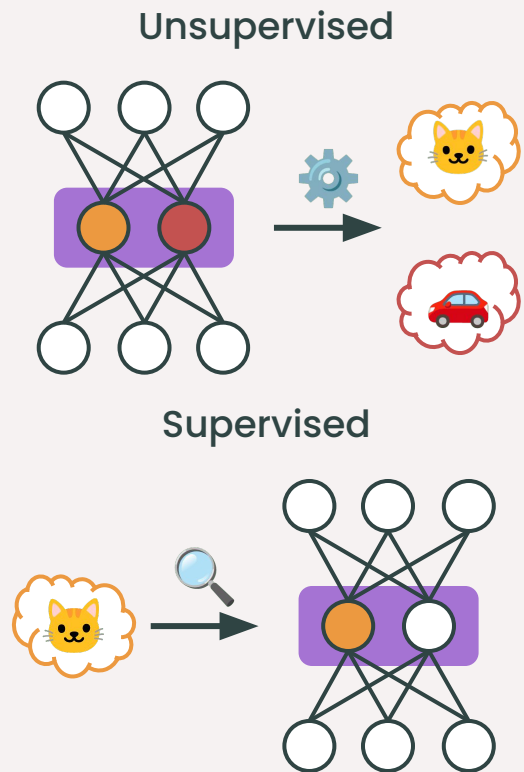
Description

What does each component do?

Validation

Is a components description accurate?

Axis 3. Either we ask what a component means, or we identify where a concept lives.



1. Object of study

Features

Circuits

Universality

2. Process stage

Decomposition

Description

Validation

3. Direction of search

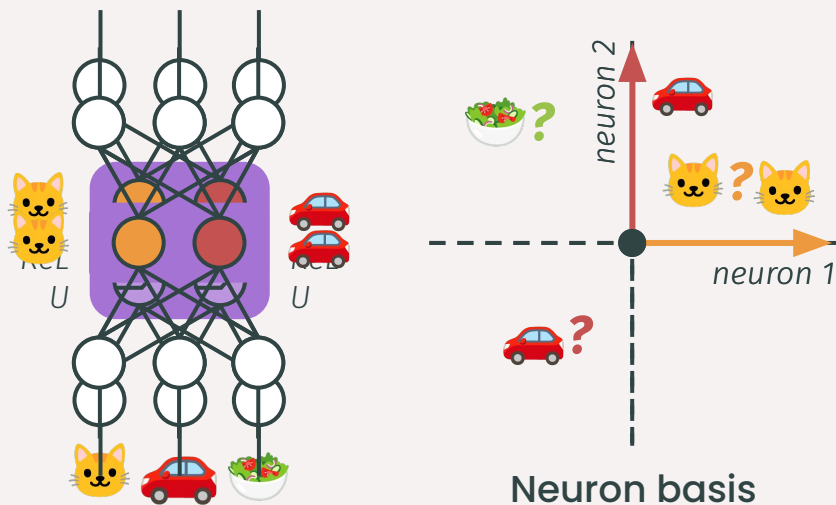
Bottom-up

Given a component, what is its role?

Top-down

Given a role, where is it represented?

We want features to align with human concepts.
Neurons form a *privileged* basis, but are *polysemantic*.



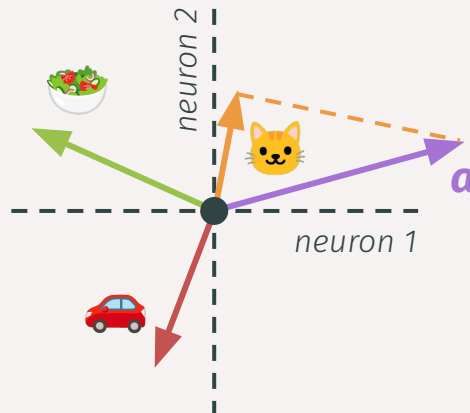
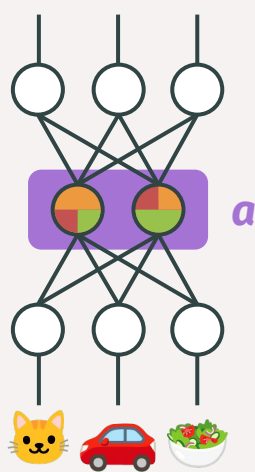
If neurons are features...

A **monosemantic** neuron activates for a single, coherent concept.

...But empirically they are not.

A **polysemantic** neuron activates for multiple unrelated concepts.

Under the *Linear Representation Hypothesis (LRH)*, features are linear directions in activation space.



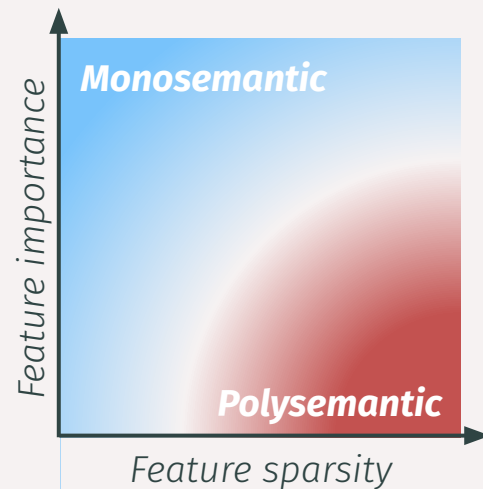
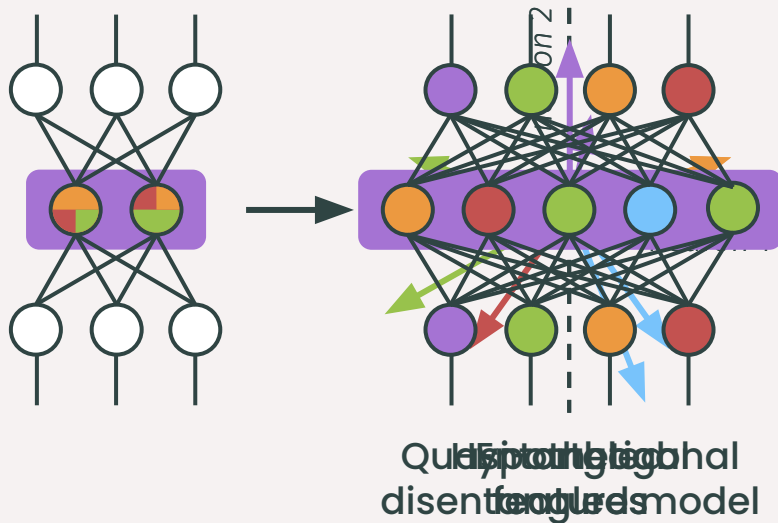
Feature presence is linear projection.

$$a_{\text{cat}} \approx \langle d_{\text{cat}} | a \rangle$$

Activations are linear combinations.

$$a \approx a_{\text{cat}} \cdot \text{cat} + a_{\text{car}} \cdot \text{car} + a_{\text{salad}} \cdot \text{salad}$$

Under the LRH, *superposition* causes polysematicity.
 Activations pack more concepts than neurons.



This chapter...

Mechanistic foundations

Ward Gauderis

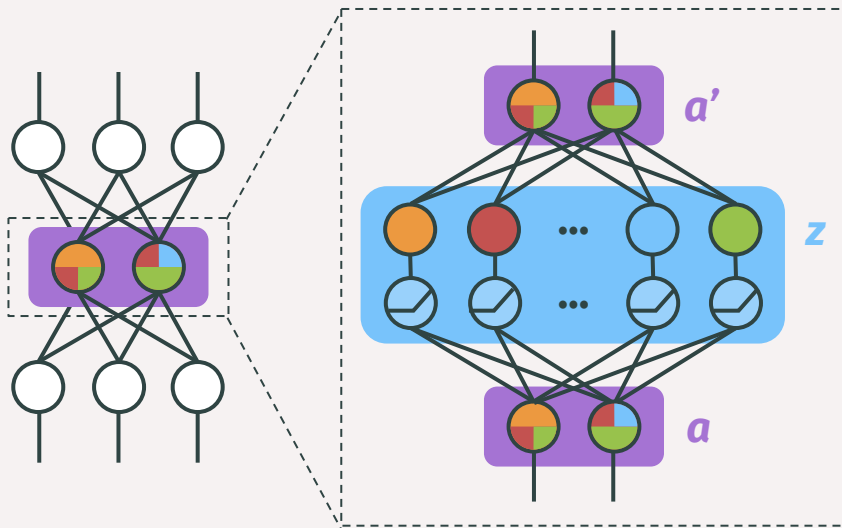
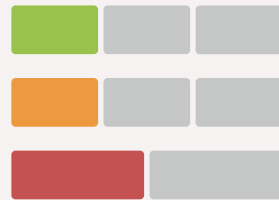
Methods & Models

Ward Gauderis & José Oramas & Thomas Doms

Limits & Open Problems

Thomas Doms

Sparse autoencoders learn an overcomplete dictionary of feature directions without labels.



Dictionary learning finds sparse features.

$$\mathbf{a} \approx D\mathbf{z}, \quad \|\mathbf{z}\|_0 \leq k$$

An encoder replaces the search for $\mathbf{z}$.

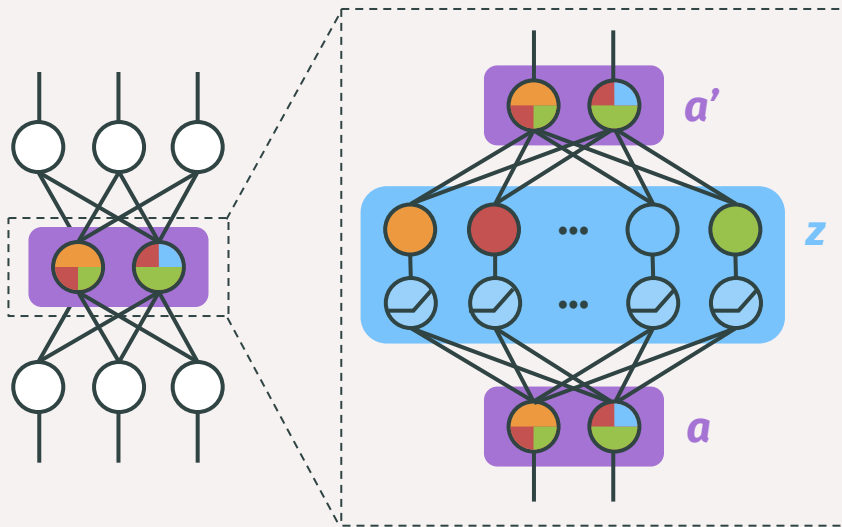
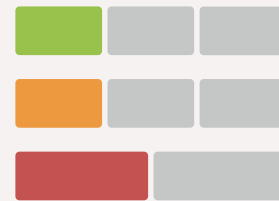
$$\mathbf{z} = \text{ReLU}(E\mathbf{a} + e)$$

$$\mathbf{a}' = D\mathbf{z} + d$$

The loss trades reconstruction and sparsity.

$$L(\mathbf{a}, \mathbf{a}') = \|\mathbf{a} - \mathbf{a}'\|_2^2 + \lambda \|\mathbf{z}\|_1$$

Sparsity is not a *proxy* for interpretability. Every feature needs a description and a test.



Features still have to be labelled.

- Decodability does not imply causality.

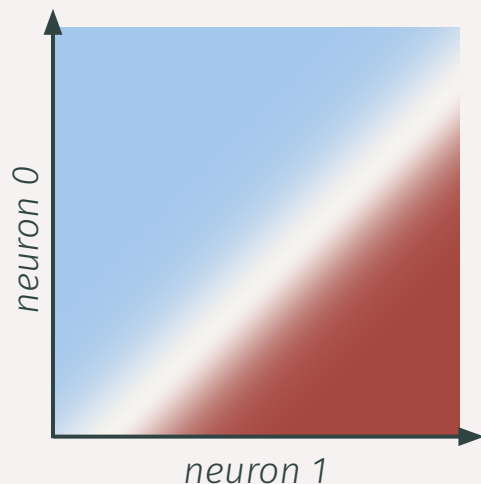
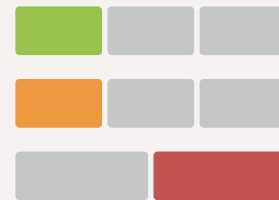
The decomposition is not unique.

- L1 penalises activation size, not count.
- Dictionary size affects feature granularity.
- Different seeds give different features.

For known concepts, supervision wins.

- Simpler baselines suffice downstream.

A *probe* uses labels to find the maximal separation between two classes.



Probing asks if a concept is linearly readable.

$$\mathbf{y} \approx \sigma(w\mathbf{a} + b), \mathbf{y} \in \{-1, 1\}$$

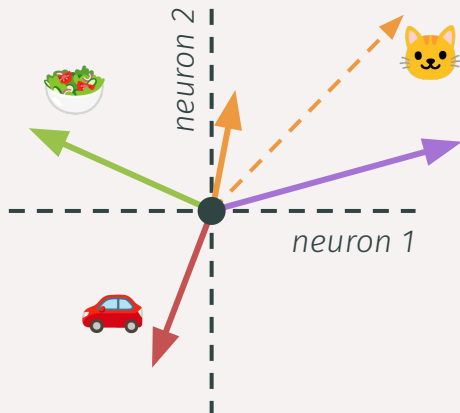
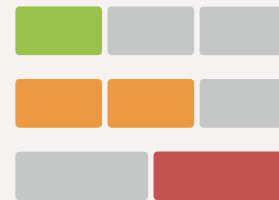
The weights carve a separating hyperplane.

$$w\mathbf{a} + b = 0$$

The loss rewards confidence in the true class.

$$L(\mathbf{y}, \mathbf{a}) = \text{softplus}(-\mathbf{y} (w\mathbf{a} + b))$$

A *probe* queries whether representations can *distinguish* a concept, not whether it is used.



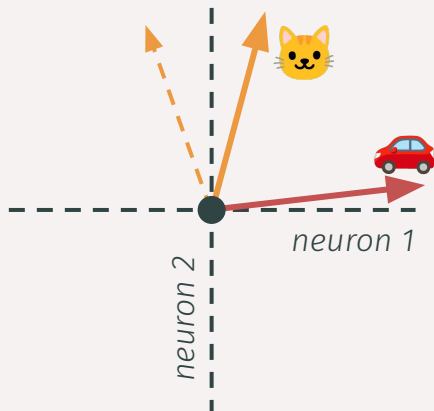
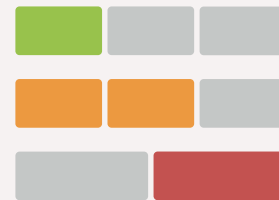
Probes can lie about concept presence.

- Decodability still does not imply causality.
- Will pick up on correlations.

Probes miss substructure.

- They answer no questions beyond presence.

A *probe* finds the maximal separation, it does not find the concept direction itself.



Probes can lie about concept presence.

- Decodability still does not imply causality.
- Will pick up on correlations.

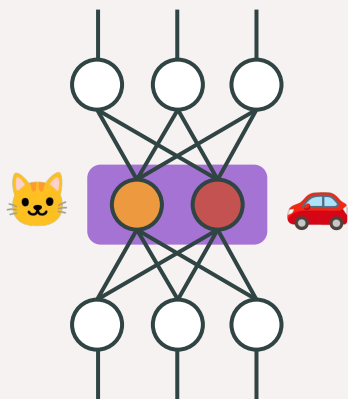
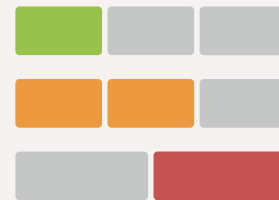
Probes miss substructure.

- They answer no questions beyond presence.

Probes find separation; not features.

- Directions don't align with features
- ... and should hence not be compared.

A *concept bottleneck model* builds a set of probes into the model.



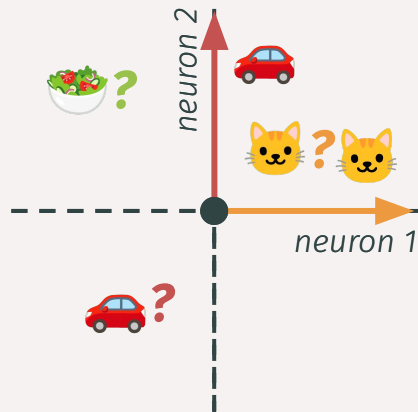
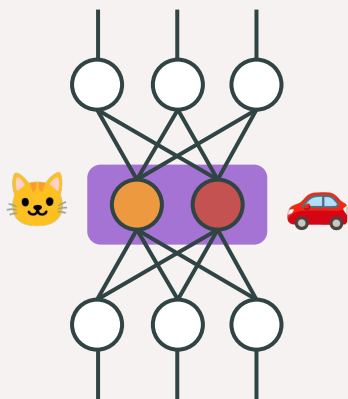
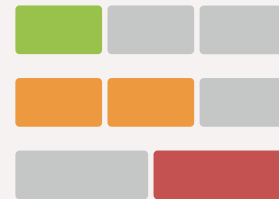
A bottleneck routes through concepts.

$$\hat{\mathbf{c}} = g(\mathbf{a}), \hat{y} = w\hat{\mathbf{c}} + b$$

The loss supervises concepts and label jointly.

$$L = \ell(\hat{y}, y) + \lambda \sum \ell(\hat{\mathbf{c}}, \mathbf{c})$$

Building in probes inherit probe problems. Concepts leak correlates and are incomplete.



Neuron basis

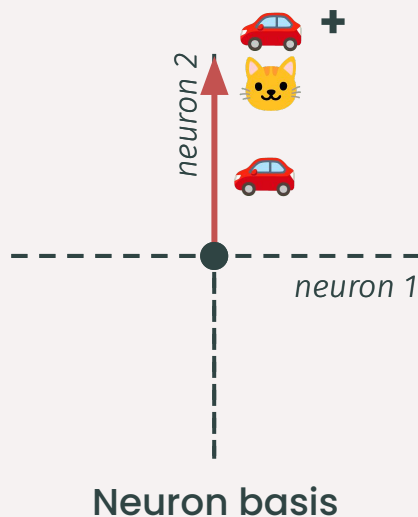
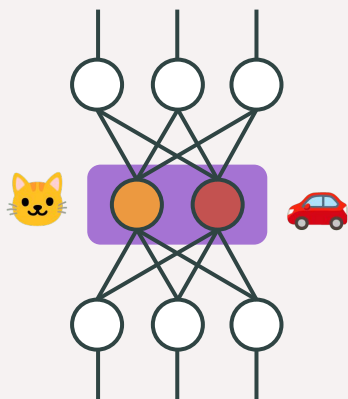
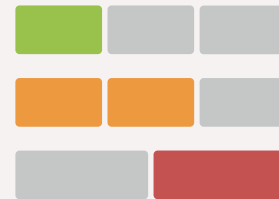
The concept set must be complete.

- Missing concepts cap accuracy.
- Residuals can interact weirdly.

Concept labels are expensive.

- Dense annotation per input
- 'label-free' outsources to an oracle.

Soft bottlenecks still allows *non-linear* signal to go through.



The concept set must be complete.

- Missing concepts cap accuracy.
- Residuals can interact weirdly.

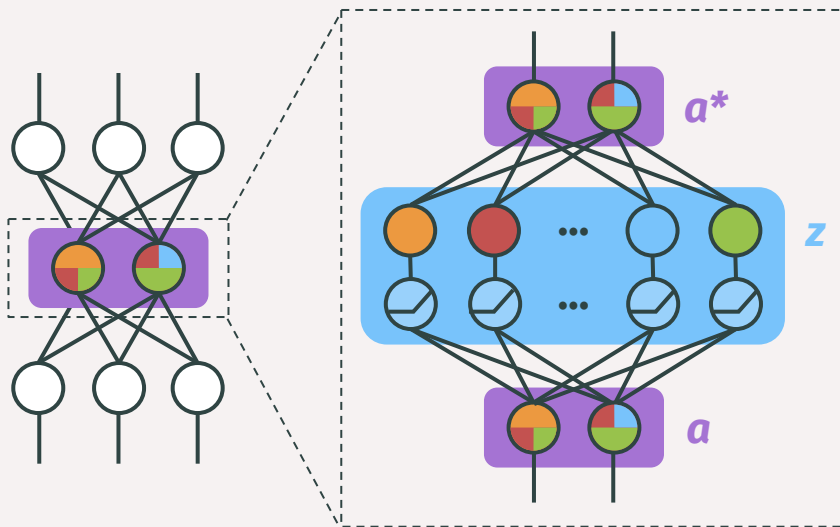
Concept labels are expensive.

- Dense annotation per input
- 'label-free' outsources to an oracle.

Soft concepts leak.

- Carry unlabelled signal past bottleneck.
- Non-linear signal is not penalized

Transcoders sparsify the computation, they decompose the model into *circuits*.



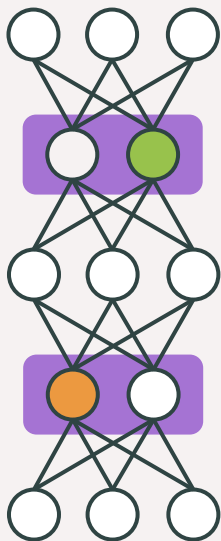
The dictionary is the same as an autoencoder.

$$\mathbf{z} = \text{ReLU}(E\mathbf{a} + e)$$
$$\mathbf{a}^* = \text{MLP}(\mathbf{a}) \approx D\mathbf{z} + d$$

But it reconstructs the layer's output.

$$L = \|\mathbf{a}^* - D\mathbf{z} - d\|^2 + \lambda \|\mathbf{z}\|_1$$

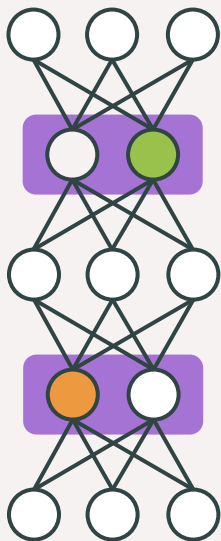
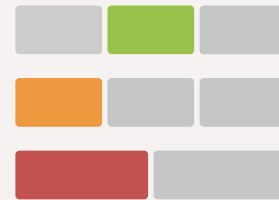
Understanding *circuits* allows tracing how a feature is formed.



Circuits describe feature interactions.

- Follow how features are computed.
- Test the proposed connections.

Understanding *circuits* allows tracing how a feature is formed.



Circuits describe feature interactions.

- Follow how features are computed.
- Test the proposed connections.

Approximation errors can compound.

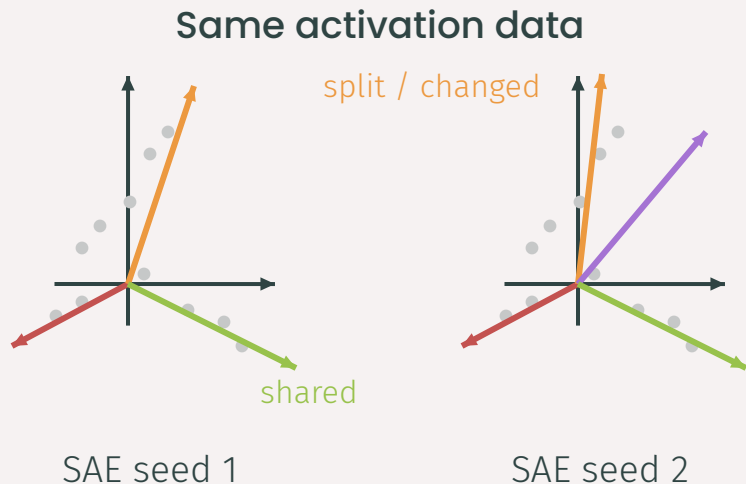
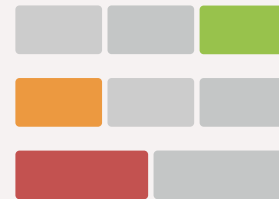
- Layer-wise fits can miss model behaviour.
- Error nodes retain unexplained signal.

SAE caveats still apply.

- Check width, sparsity and stability.
- Sparse activations $\neq$ sparse circuits

A decomposition need not be unique.

Identifiability requires assumptions.



Schematic dictionaries

Identifiability asks for uniqueness.

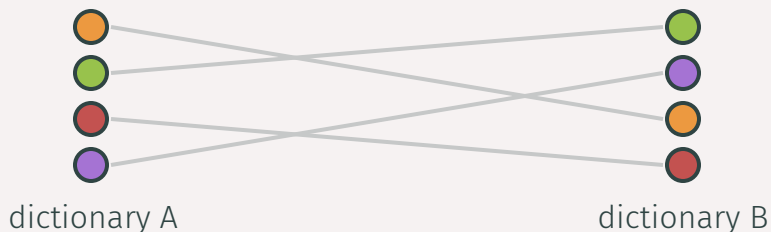
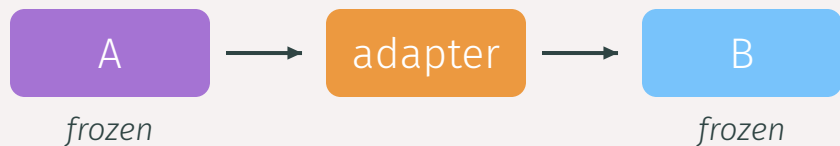
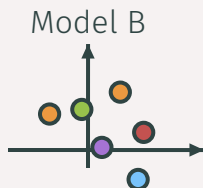
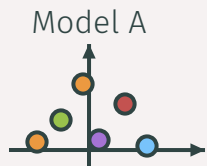
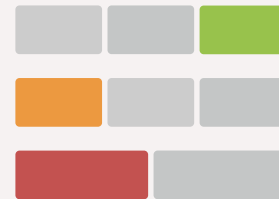
- Allow rescaling and permutation.

Stability tests the fitted solution.

- Vary seeds, data, width, and sparsity.

Universality depends on the comparison.

Features, geometry or function?



Do models organise inputs similarly?

- Compare the geometry of their activations.

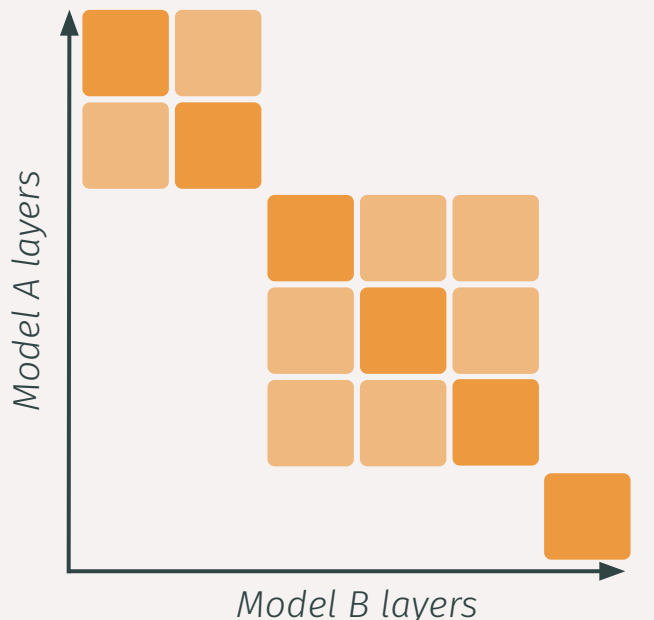
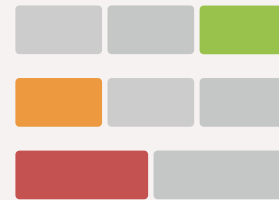
Can their representations work together?

- A learned adapter joins two frozen models.

Do they share candidate components?

- Compare features or fit a shared dictionary.

Centered kernel alignment compares layers. High similarity does not prove shared circuits.



Light to dark: low to high CKA (schematic)

Use matched input examples.

- Collect and centre activations; feature counts may differ.

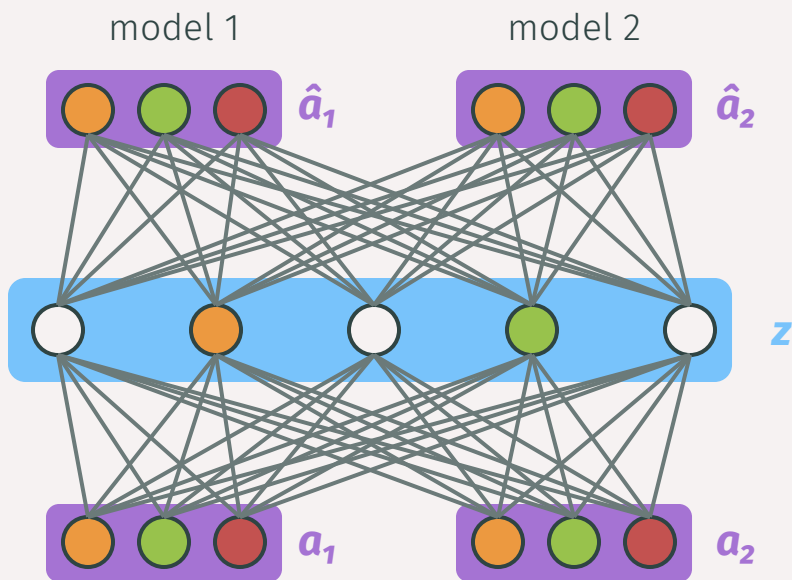
Read the layer-by-layer comparison.

- Darker cells indicate greater similarity; blocks reveal layer groupings.

Similar geometry leaves causal questions.

- Rotations and uniform scaling are ignored; shared roles need separate tests.

Crosscoders learn one shared code. Each site keeps its own decoder.



Paired activations for the same input

One code combines multiple sites.

$$\mathbf{z} = \text{ReLU}(E_1 \mathbf{a}_1 + E_2 \mathbf{a}_2 + e)$$

Each site has its own decoder.

$$\hat{\mathbf{a}}_1 = D_1 \mathbf{z} + d_1$$

$$\hat{\mathbf{a}}_2 = D_2 \mathbf{z} + d_2$$

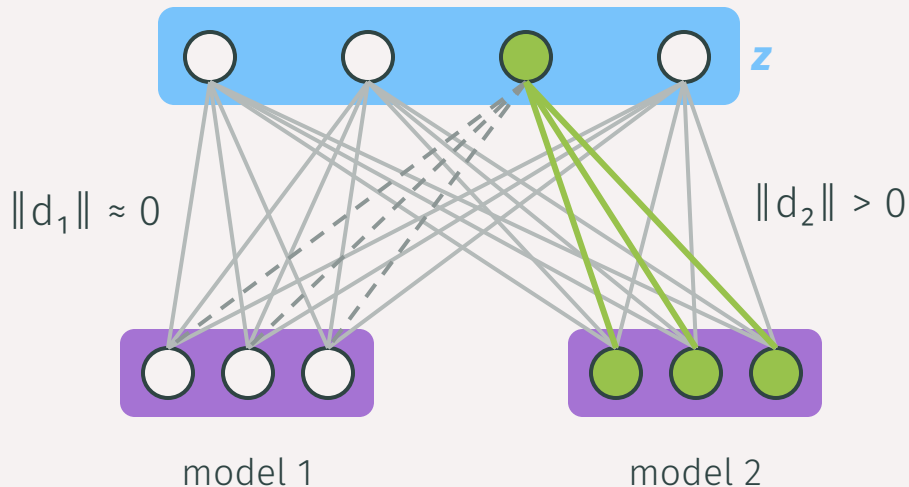
The loss trades reconstruction and sparsity.

$$L = \|\mathbf{a}_1 - \hat{\mathbf{a}}_1\|^2 + \|\mathbf{a}_2 - \hat{\mathbf{a}}_2\|^2 + \lambda R$$

Sparsity can create apparent differences. Model-only latents need targeted checks.



An apparent model-only latent



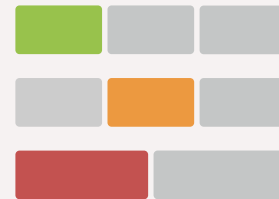
L1 can create false exclusivity.

- Shrinkage can erase useful decoder weights.

A shared code leaves causal questions.

- Joint encoding does not establish a circuit.

Examples suggest what a feature detects. An LLM turns the pattern into a *hypothesis*.



Examples + activations

The cat slept.

8

A sleepy cat purred.

9

The cat stretched.

7



Explainer LLM



“Feline-related text”

Schematic activations

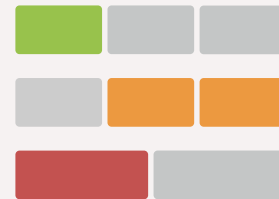
Input evidence suggests a description.

- Inspect activating contexts and token-level attribution.

Output evidence suggests a role.

- Vocabulary projections are clues; interventions test effects.

A label must survive new examples. *Counterexamples* reveal its scope.



Hypothesis: feline words

The lion roared.

Any activation? Yes

The lion roared.

Correct

The lion roared.

Incorrect

Detection checks the context.

- Does any token activate the feature?

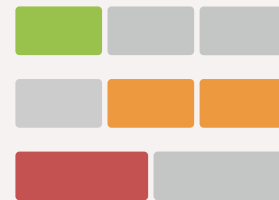
Fuzzing checks the marked tokens.

- Are these the positions where it activates?

Compare against a null model.

- Detection and fuzzing can score randomized models well.

Surprisal tests how much a label helps. It should help most on activating contexts.



$$\text{Gain} = \text{NLL}(\text{generic}) - \text{NLL}(\text{label})$$

Context	Generic	Label	Gain
cat	12	8	4
lion	13	10	3
dog	11	10.8	0.2

Score the description.

- Measure gain over a generic description.

Rank held-out contexts.

- Use the reduction to separate active from inactive examples.

The animal is a ...

Can you spot the *intruder*?

Coherence can be tested without a label.



Five held-out contexts

1. The ferry crossed the bay.
2. He took the tram home.
3. She caught the bus today.
4. The apple tasted sweet.
5. The train arrived.

Four positives, one negative.

- Choose the example that does not belong.

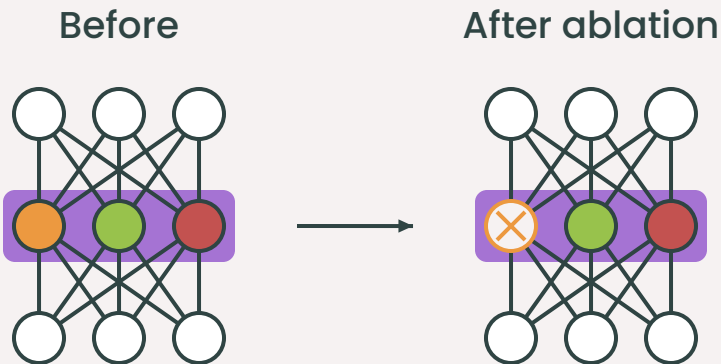
Score accuracy across examples.

- Chance accuracy is 20%.

Schematic intruder test

Causal claims predict intervention effects.

The setup defines the scope of the test.



Predicted effect



Specify the causal hypothesis.

- Set variables, basis, and intervention rules.

Distributed Alignment Search (DAS).

- Align to a fixed causal model.

Test predictions against controls.

- Report test scope and uncertainty.

Validation needs more than one setting. Known mechanisms and real tasks differ.



Known program



Model organism



Practical task



Known programs test recovery.

- Tracr provides a computation to recover.

Model organisms test failure detection.

- Can an audit find the hidden objective?

Real tasks test usefulness.

- Compare strong baselines for this question.

This chapter...

Mechanistic foundations

Ward Gauderis

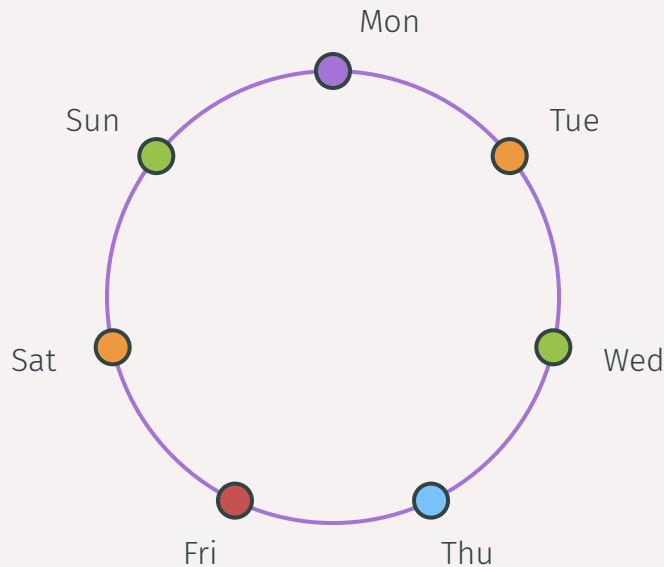
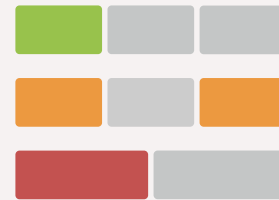
Methods & Models

José Oramas & Thomas Dooms

Limits & Open Problems

Thomas Dooms

Several directions can encode one concept. Interventions test their computational role.



One structured representation

Schematic weekday circle

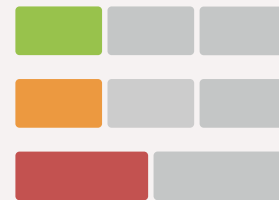
A feature can have internal geometry.

- Weekdays can trace a circle in a plane.

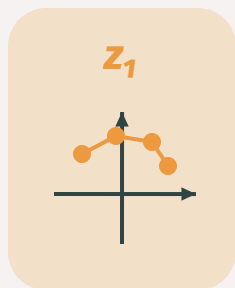
PCA alone does not identify a mechanism.

- Intervene on the subspace; test its role.

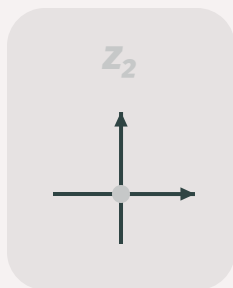
Block-sparse featurizers select vectors. Sparsity is imposed across whole blocks.



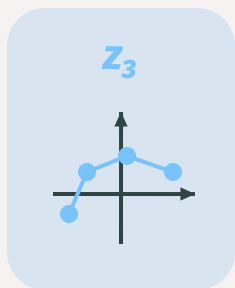
Select whole vector blocks



active



inactive



active

$$\hat{x} = D_1 z_1 + D_2 z_2 + D_3 z_3$$

Schematic coordinates across inputs

Select a few vector blocks.

- Top-k keeps the blocks with largest norms.

Keep structure within each block.

- Signed coordinates can vary together.

Blocks are a modelling assumption.

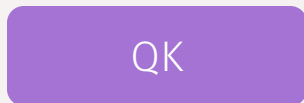
- A block need not match one concept.

Weights also admit a decomposition.

APD searches for reusable components.



Attention routes OV outputs



APD: additive weight components

$$\theta \approx \theta_1 + \theta_2 + \theta_3$$



active



inactive



active

QK and OV expose weight circuits.

- QK scores where to attend.
- OV maps what is written.

APD adds reusable components.

- Rebuild weights; use few pieces per input.

Interactions still matter.

- Parameters add; outputs need not.

APD depends on costly fitting and tuning. Its importance scores are approximations.



More components cost more to fit



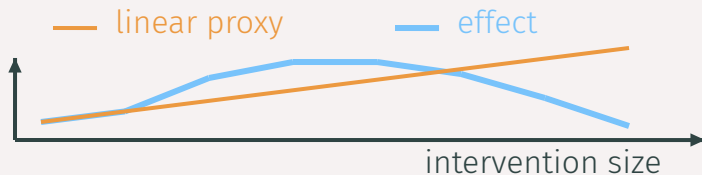
Fitting costs memory and computation.

- More parameter components add overhead.

Sparsity changes the decomposition.

- Tune top-k and simplicity penalties.

A local estimate can miss curvature



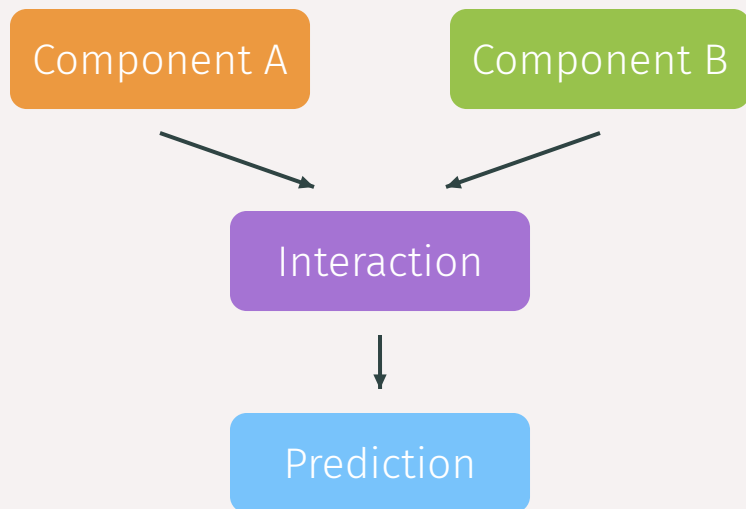
Attribution is a causal proxy.

- Gradients can misestimate ablation effects.

We can find and test the parts.

When do they explain the *computation*?

Parts and interactions



Account for parts and interactions.

- How do they combine into the computation?

Specify fidelity and what is missing.

- Which effects are preserved or unexplained?

Compare explanations for a purpose.

- Which helps answer the actual question?

Coming up...

Early Interpretability Methods

José Oramas

From Classical to Mechanistic Interpretability

Ward Gauderis, José Oramas & Thomas Dooms

From Mechanistic to Compositional Interpretability

Ward Gauderis & Thomas Dooms