

AIMLAI @ ECML PKDD 2026
September 7th 2026

From Mechanistic to Compositional Interpretability

Ward Gauderis • Thomas Dooms

This chapter...

Interpretability Paradigms

Ward Gauderis

Compositional Interpretations

Ward Gauderis & Thomas Dooms

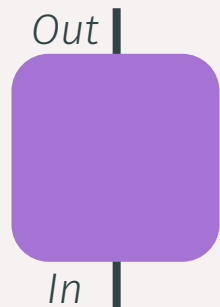
Architectures & Algorithms

Thomas Dooms

As AI adoption grows, regulation asks questions our interpretability tools *cannot* yet answer.

Post-hoc
interpretability

“What information was used to make this particular decision?”



Behavioural faithfulness

Decisions, but no mechanisms

As AI adoption grows, regulation asks questions our interpretability tools *cannot* yet answer.

Post-hoc interpretability

"What information was used to make this particular decision?"

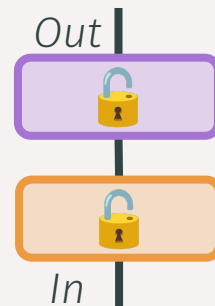


Behavioural faithfulness

Decisions, but no mechanisms

Intrinsic interpretability

"How do we design a model that explains all its behaviour?"



Inherent faithfulness

Guarantees, but constrained

As AI adoption grows, regulation asks questions our interpretability tools *cannot* yet answer.

Post-hoc interpretability

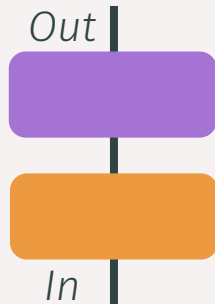
"What information was used to make this particular decision?"



Behavioural faithfulness
Decisions, but no mechanisms

Mechanistic interpretability

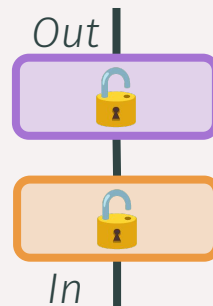
"How does the model solve this general class of problems?"



Explanatory faithfulness
Heuristic, but no foundations

Intrinsic interpretability

"How do we design a model that explains all its behaviour?"



Inherent faithfulness
Guarantees, but constrained

Interpretability Paradigms

Ward Gauderis

Compositional Interpretations

Ward Gauderis & Thomas Dooms

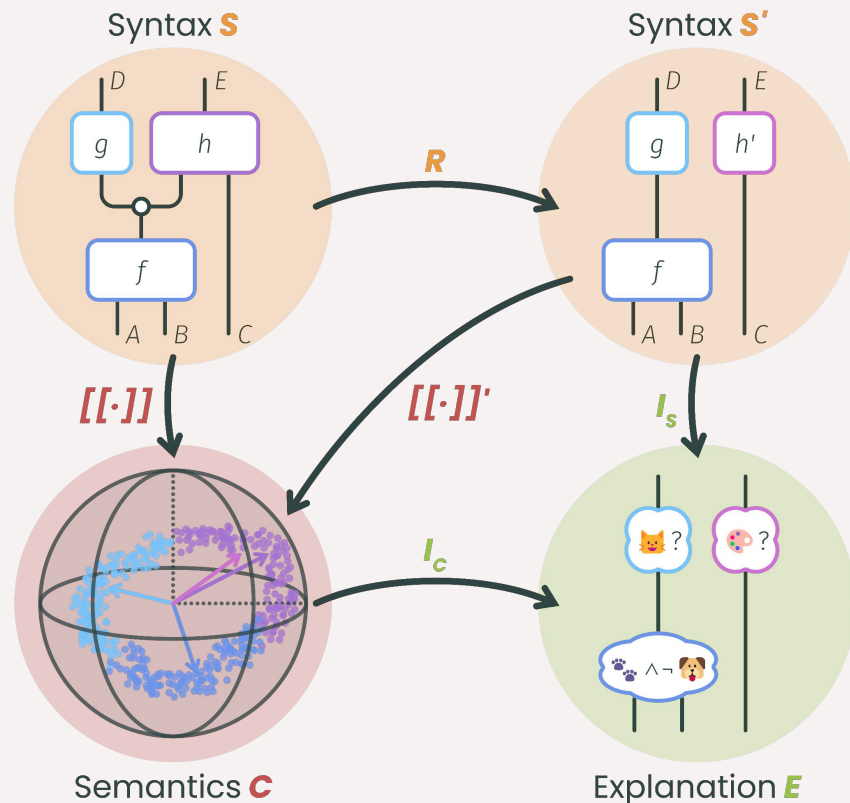
Architectures & Algorithms

Thomas Dooms

Compositional interpretability is a formal framework to evaluate, compare and improve explanations.

Two model-agnostic principles, formalised,

cast interpretability as **constrained optimisation**.

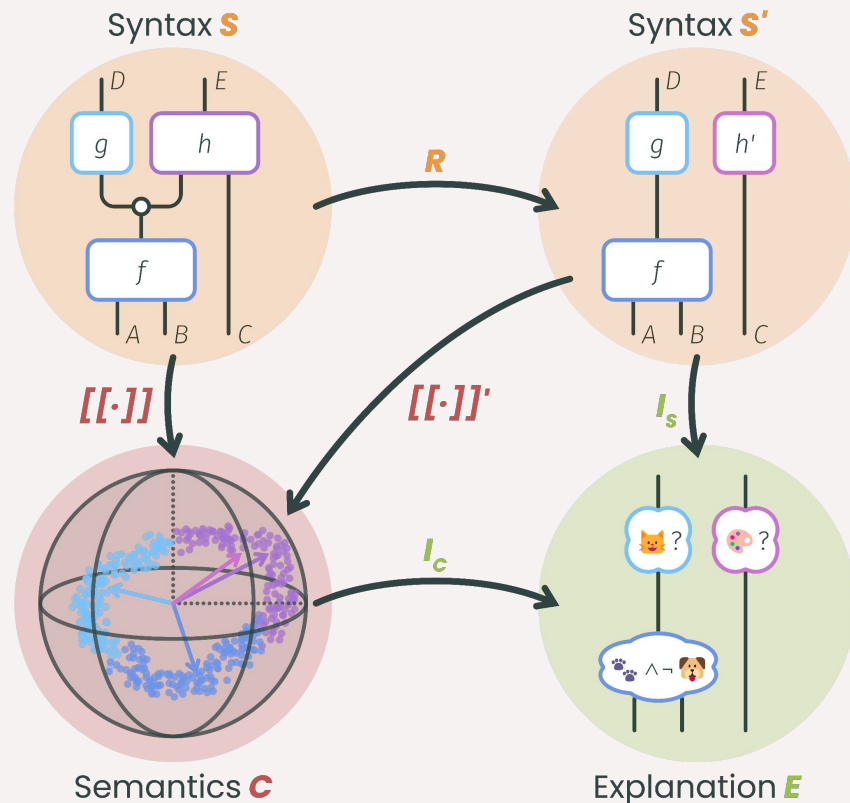


Compositional interpretability is a formal framework to evaluate, compare and improve explanations.

Two model-agnostic principles, formalised,

- **Compositionality**
 - Category Theory

cast interpretability as **constrained optimisation**.

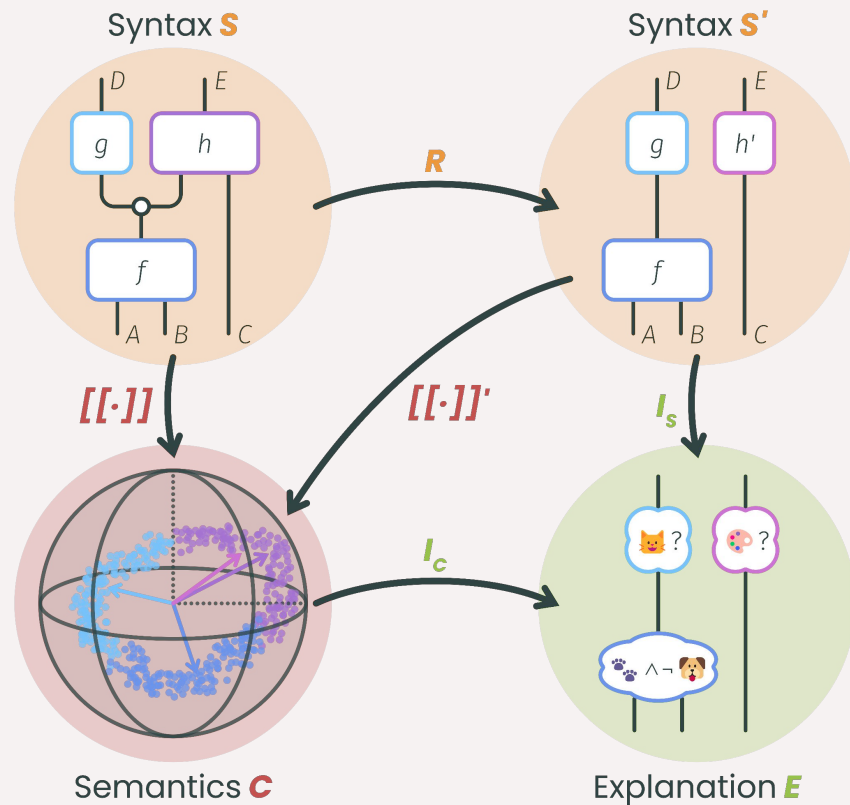


Compositional interpretability is a formal framework to evaluate, compare and improve explanations.

Two model-agnostic principles, formalised,

- **Compositionality**
 - Category Theory
- **Occam's Razor**
 - Minimum Description Length

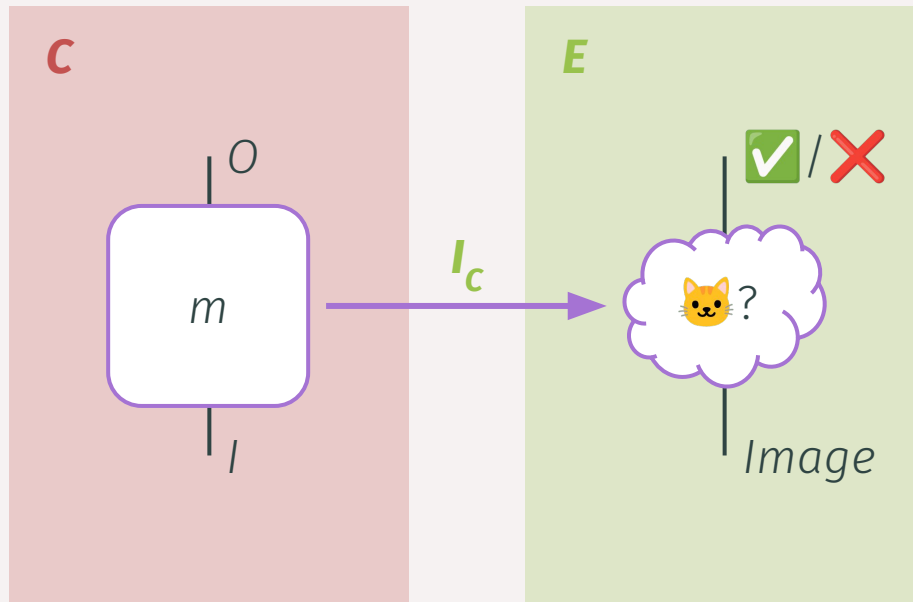
cast interpretability as **constrained optimisation**.



Coming up...

1. **What are compositional interpretations?**
2. What makes explanations good?
3. How do we find better explanations?

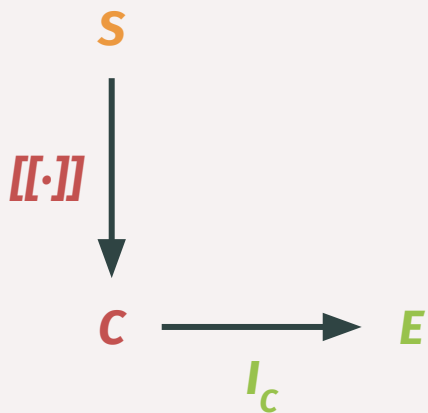
Post-hoc interpretations only map input-output behaviour to explanations.



C Semantics
 E Explanation

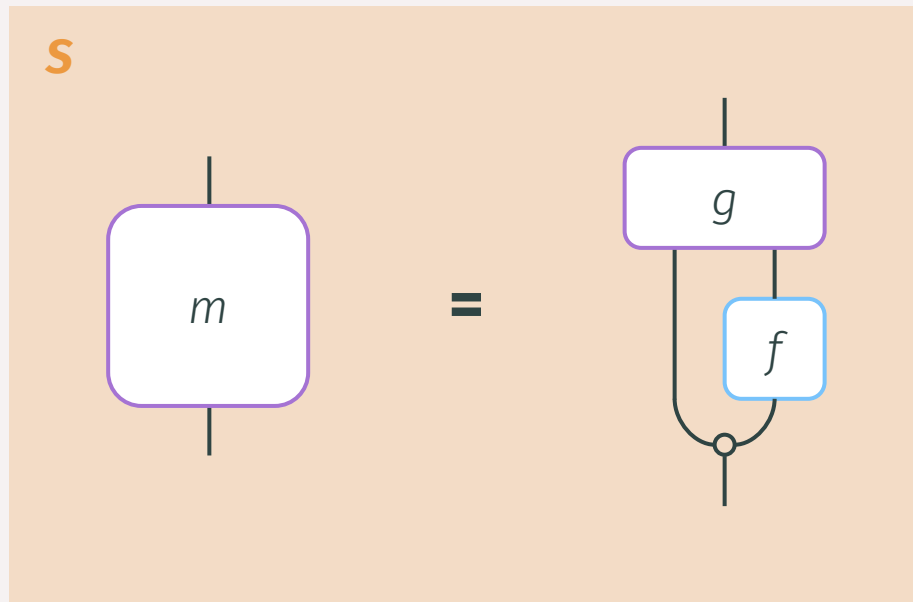
I_C Semantic interpretation

Compositional models abstract high-level syntax from low-level semantics.

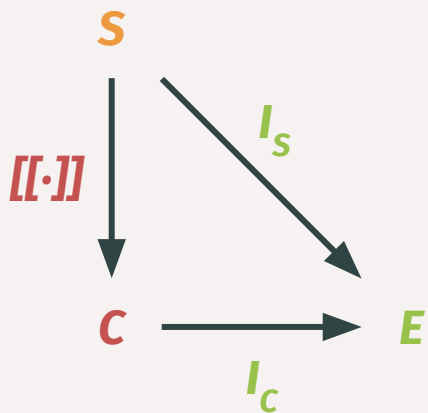


S Syntax
C Semantics
E Explanation

$[[\cdot]]$ Representation
 I_c Semantic interpretation

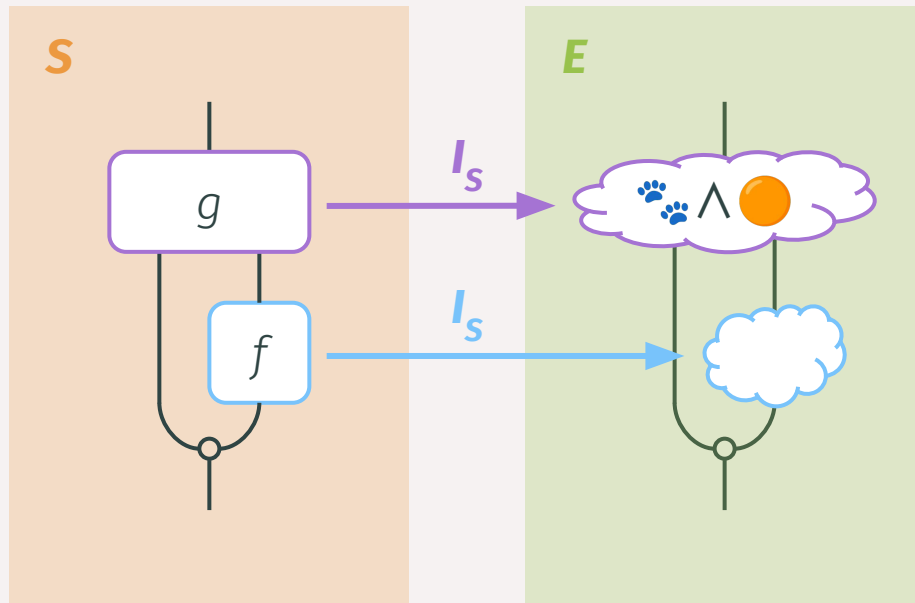


Compositional interpretations mechanistically explain how a model works, step-by-step.



S Syntax
 C Semantics
 E Explanation

$[[\cdot]]$ Representation
 I_c Semantic interpretation
 I_s Syntactic interpretation



Category theory and string diagrams are the mathematics of compositional structure.

A **compositional model** $M = (D, G, \llbracket - \rrbracket, C)$ consists of:

- G , a signature inducing $S = \text{Free}(G)$, an abstract syntactic category
- $D \in \text{Hom} \square(I, O)$, a composite morphism in S
- C , a concrete semantic category
- $\llbracket - \rrbracket : S \rightarrow C$, a representation functor

A **compositional interpretation** $I(I_a, I_c, I_c^*, H)$ for M into H consists of:

- $I_a : S \rightarrow H$, an abstract interpretation functor
- $I_c : C \rightarrow H$, a concrete interpretation functor s.t.

$$I_a = I_c \circ \llbracket - \rrbracket$$

- $I_c^* : H \rightarrow C$, the left adjoint to I_c ($I_c^* \dashv I_c$) defined by the natural isomorphism

$$\text{Hom}_c(I^*(h), c) \cong \text{Hom} \square(h, I_c(c))$$

A **Category** C consists of:

- $ob(C)$, a class of objects
- $hom_c(A, B)$, a class of morphisms for $A, B \in ob(C)$
- $id_a \in hom_c(A, A)$, an identity morphism for $A \in ob(C)$
- $\circ : hom_c(B, C) \times hom_c(A, B) \rightarrow hom_c(A, C)$, an associative operator that respects identity morphisms

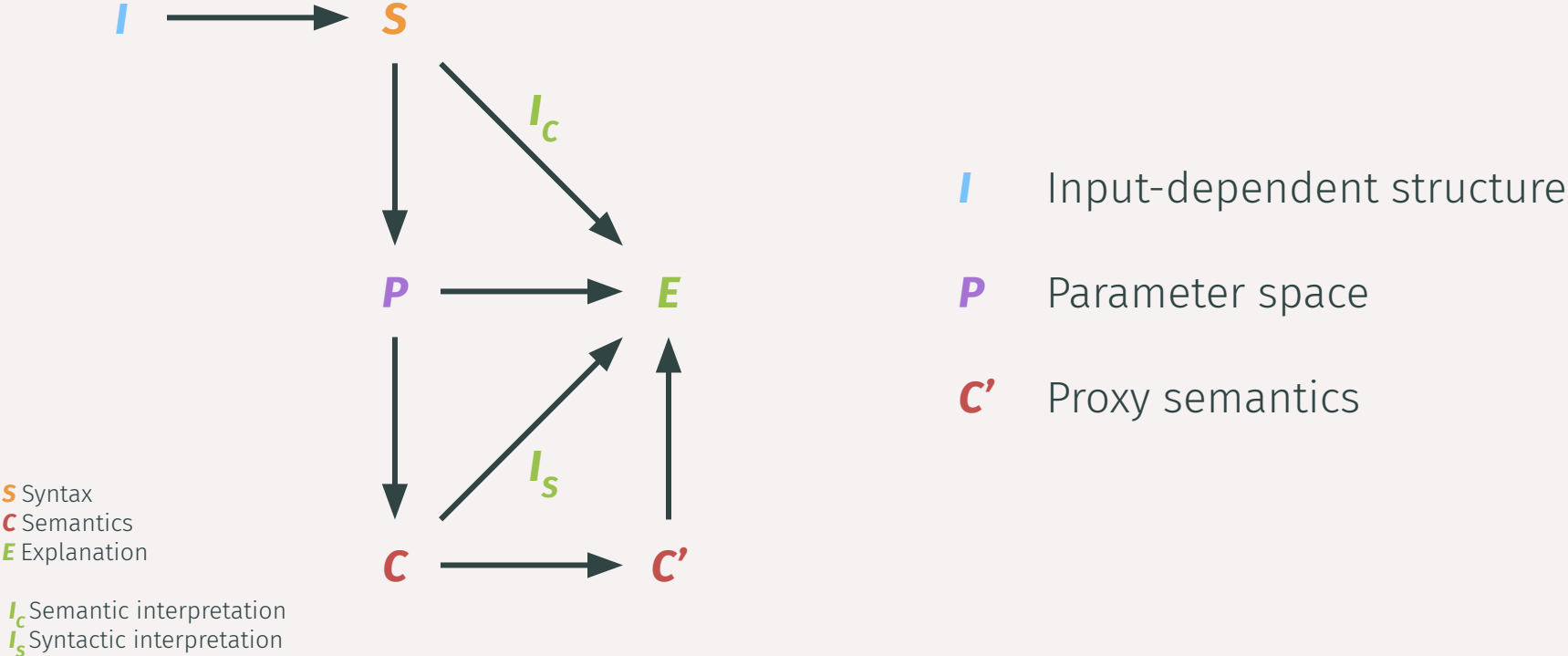
A **string diagram** is a visual representation of a composite morphism in a **symmetric monoidal category** C equipped with:

- $I \in ob(C)$, the monoidal unit object
- $\otimes : C \times C \rightarrow C$, a bifunctor for the tensor product
- Natural isomorphisms for the associator, symmetry, left and right unitors (left as an exercise to the reader)

Examples include:

- **CartSp** of Euclidean spaces and continuous functions
- **Stoch** of finite sets and probability channels
- **Quant** of finite-dimensional Hilbert spaces and CP-maps

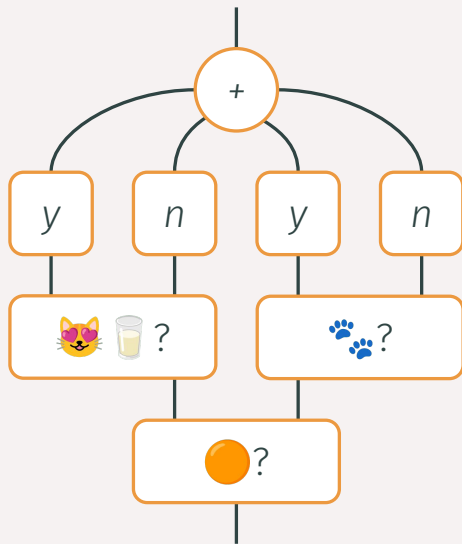
This language lets us formalise and structurally compare interpretability methods.



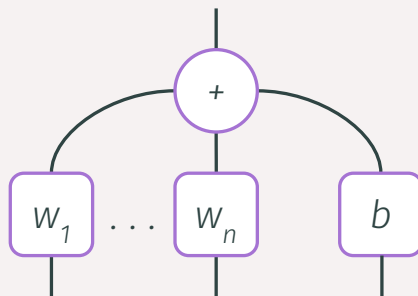
(Duneau, Towards a Comparative Framework for Compositional AI Models. 2025; Braun, Open Problems in Mechanistic Interpretability. 2025; Bushnaq, Stochastic Parameter Decomposition. 2025)

The objective of interpretability is to search for complete compositional interpretations.

Decision tree



Linear regression



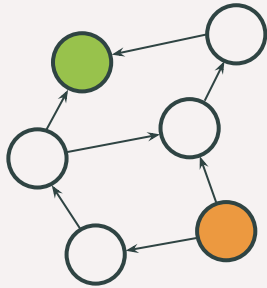
Intrinsically interpretable models have compositional interpretations.

Coming up...

1. What are compositional interpretations?
- 2. What makes explanations good?**
3. How do we find better explanations?

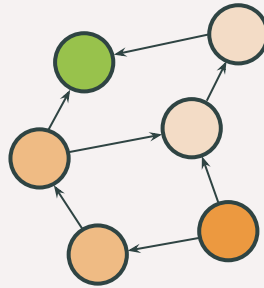
What makes a fastest path algorithm good?

Brute force



exact

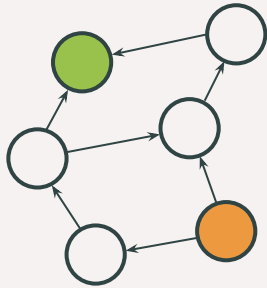
Dijkstra



exact

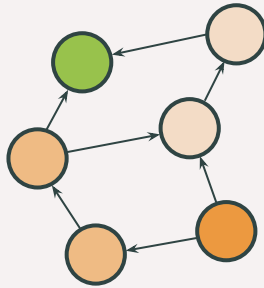
What makes a fastest path algorithm good?

Brute force



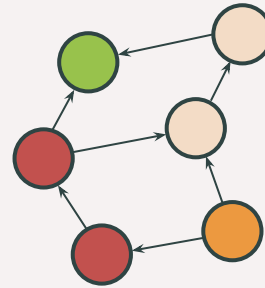
exact

Dijkstra



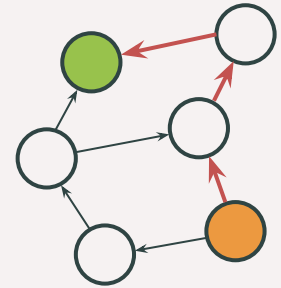
exact

Weighted A*



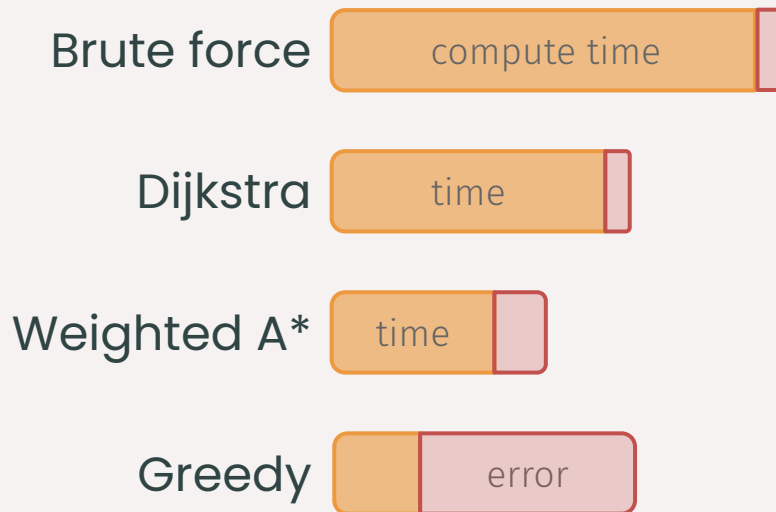
bounded error

Greedy

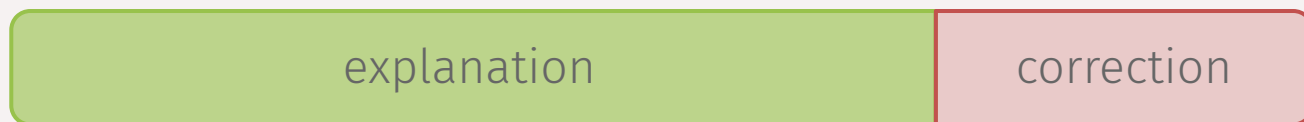


unbounded error

We can rank an algorithm goodness based on time and error.



Interpretability is a communication problem: the best explanations are the most compressed.



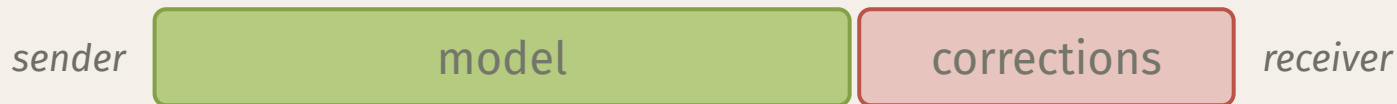
$L(model)$

$L(data | model)$

description length

$$\mathbf{L(data) = L(model) + L(data | model)}$$

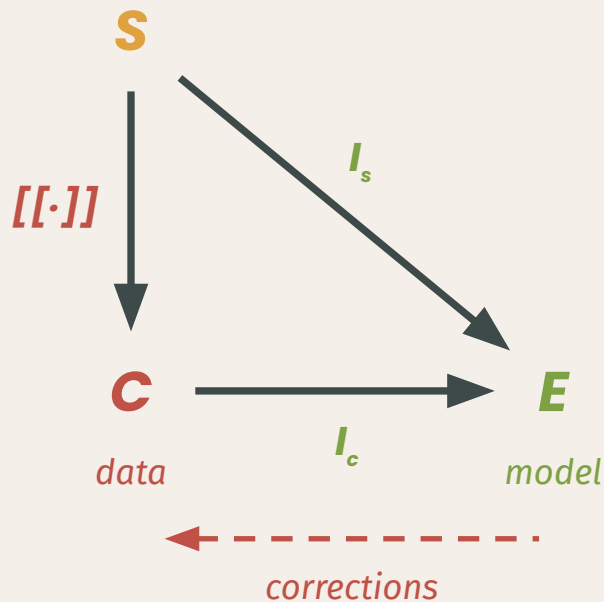
Interpretability is a communication problem: the best explanations are the most compressed.



description length

$$L(\mathbf{Data}) = L(\mathbf{Model}) + L(\mathbf{Data} \mid \mathbf{Model})$$

Describing the model through an explanation has two different costs.



$$L([[D]]) \leq L(I_s(D)) + L([[D]] | I_s(D))$$

For interpretations, we formalise this ranking with the notion of description length.

Syntactic interpretation

complexity

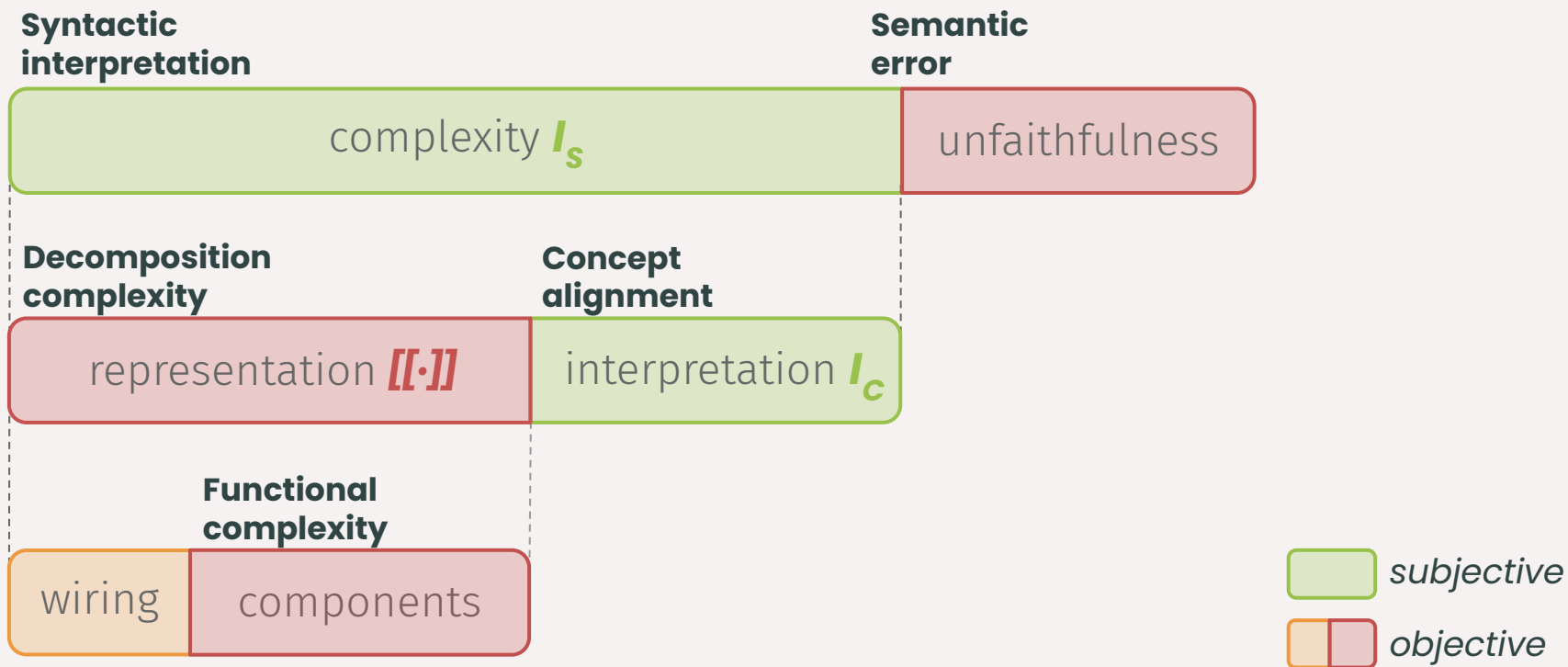
$$L(I_s(D))$$

Semantic error

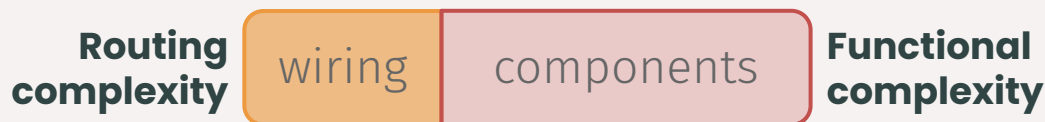
unfaithfulness

$$L([D] \mid I_s(D))$$

Since an interpretation is subjective,
we decompose into parts.



Description length gives us a framework to assess when interpretations cheat.



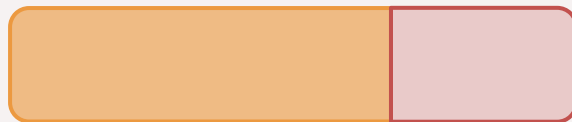
$$L(\text{ReLU}) + L(\text{Matrix})$$

=

=

?

...

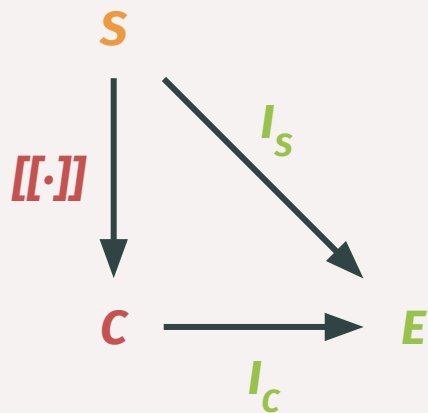


Syntactic complexity can degenerate if not measured

Coming up...

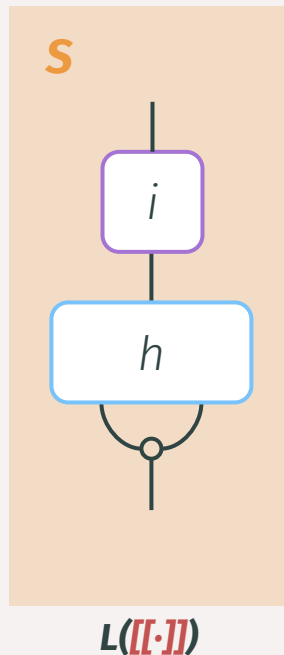
1. What are compositional interpretations?
2. What makes explanations good?
3. **How do we find better explanations?**

Compressive refinements can simplify representations without changing model behaviour.

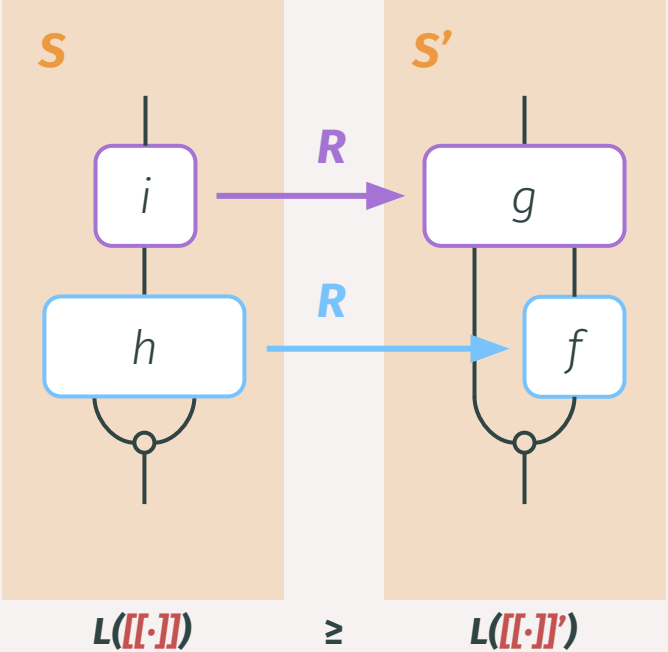
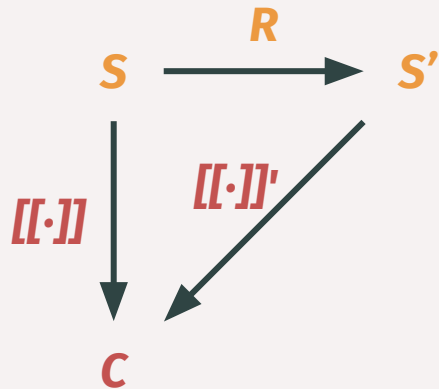


S Syntax
 C Semantics
 E Explanation

$[[\cdot]]$ Representation
 I_c Semantic interpretation
 I_s Syntactic interpretation

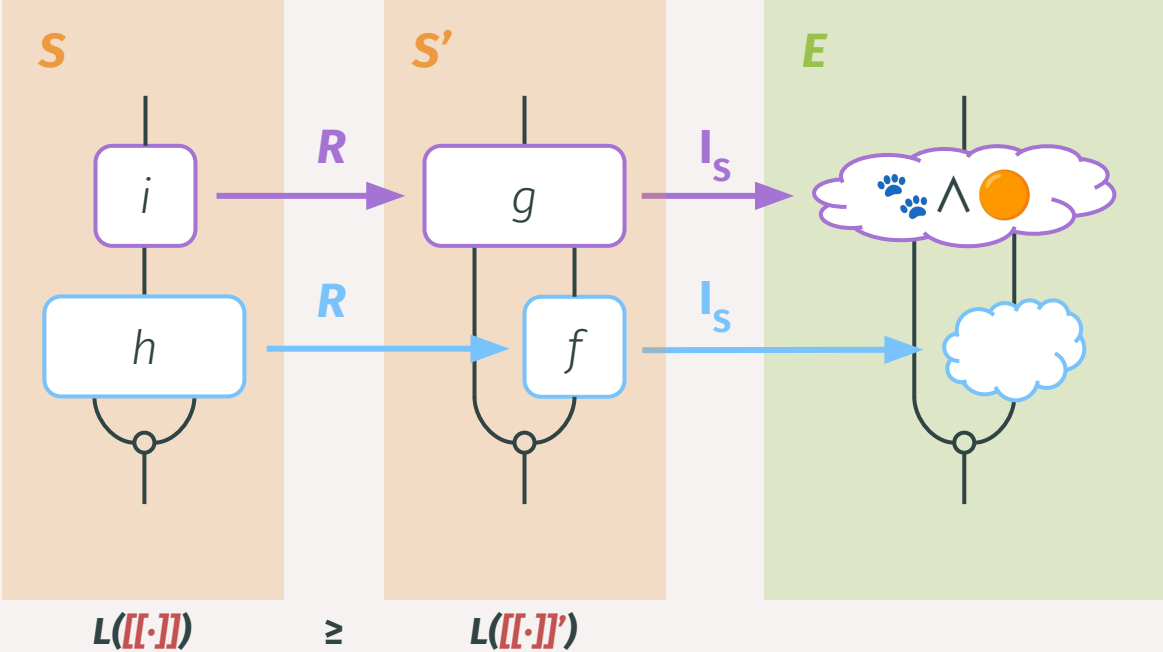
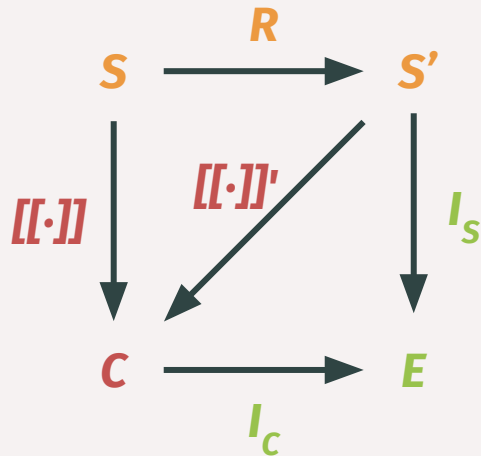


Compressive refinements can simplify representations without changing model behaviour.



- S Syntax
- C Semantics
- E Explanation
- R Syntactic refinement
- $[[\cdot]]$ Representation
- I_C Semantic interpretation
- I_S Syntactic interpretation

Compressive refinements can simplify representations without changing model behaviour.

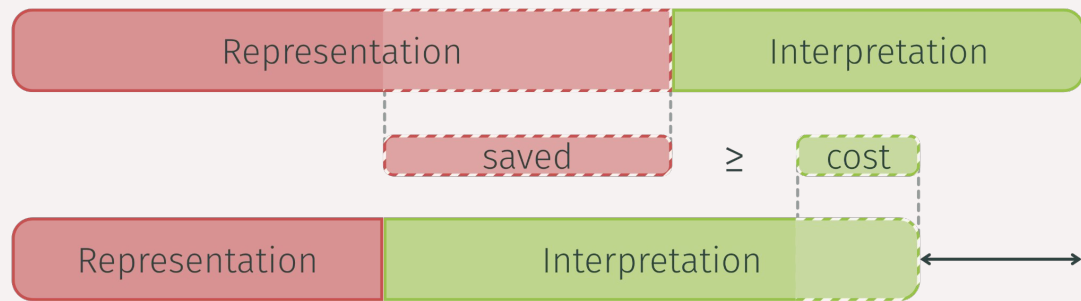


S Syntax
C Semantics
E Explanation

R Syntactic refinement
[[·]] Representation
 I_C Semantic interpretation
 I_S Syntactic interpretation

Simpler does not always mean more interpretable...
But it does under the *parsimony criterion*.

Complex syntax **S**

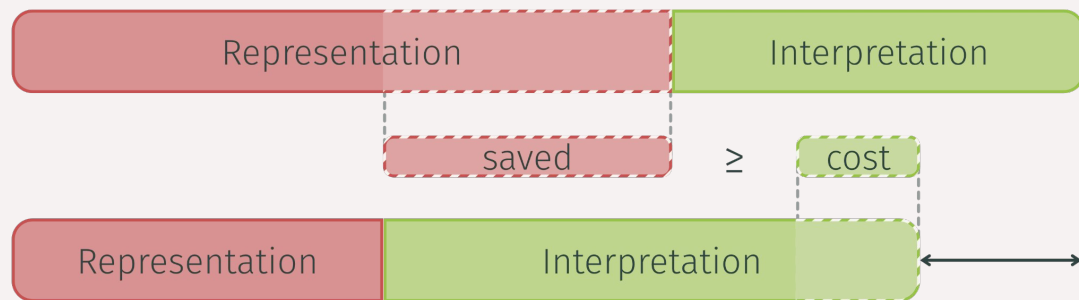


Refined syntax **S'**

Simpler does not always mean more interpretable...

But it does under the *parsimony criterion*.

Complex syntax **S**



Refined syntax **S'**

Comonotonic coding suffices.

A coding distribution is a prior.
Encode *explanatory optimism* in it,
and both complexities fall together.

$$L([[\cdot]]) := -\log_2 P([[\cdot]])$$

Make faithfulness and complexity *measurable*, and interpretability becomes *constrained optimisation*.

Search space
refinement in **S**

Constraint
faithful in **C**

Objective
complexity of **[[·]]**

Make faithfulness and complexity *measurable*, and interpretability becomes *constrained optimisation*.

Search space
refinement in **S**

Constraint
faithful in **C**

Objective
complexity of **[·]**

Interpretability is a spectrum, not a constraint.

- Including post-hoc, mechanistic and intrinsic



Make faithfulness and complexity *measurable*, and interpretability becomes *constrained optimisation*.

Search space
refinement in **S**

Constraint
faithful in **C**

Objective
complexity of **[·]**

Interpretability is a spectrum, not a constraint.

- Including post-hoc, mechanistic and intrinsic

Interpretability is relative to assumptions on **E.**

- To compare across **E**s, use description length.
- Choosing **E** is also a question for cognitive science.



Make faithfulness and complexity *measurable*, and interpretability becomes *constrained optimisation*.

Search space
refinement in **S**

Constraint
faithful in **C**

Objective
complexity of **[·]**

Interpretability is a spectrum, not a constraint.

- Including post-hoc, mechanistic and intrinsic

Interpretability is relative to assumptions on **E.**

- To compare across **E**s, use description length.
- Choosing **E** is also a question for cognitive science.

If description length stays high, design new models.



Interpretability Paradigms

Ward Gauderis

Compositional Interpretations

Ward Gauderis & Thomas Dooms

Architectures & Algorithms

Thomas Dooms

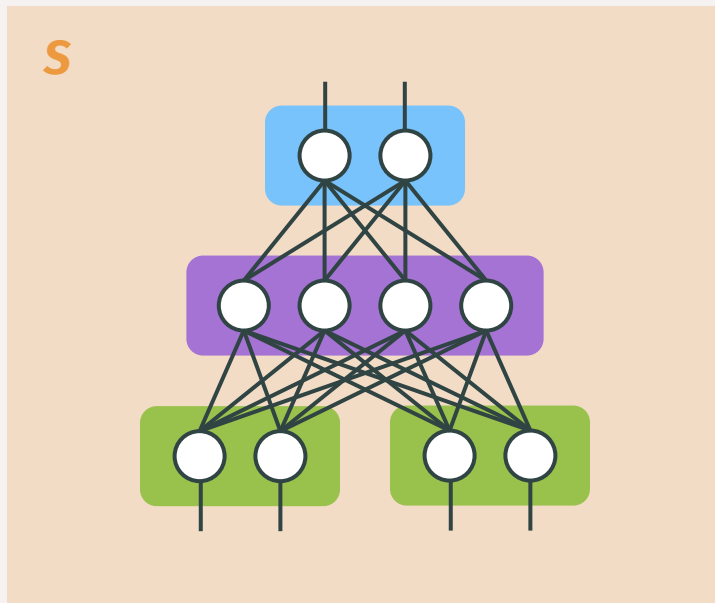
Neural networks are hard to refine and to measure. Syntactic refinement becomes approximation.

No interesting way to rewrite structure

A **neural model** is a compositional model with

- the syntax of a **copy-discard** string diagram
- semantics in **CartSp** (continuous maps between Cartesian spaces)

A box can be any function (in theory)



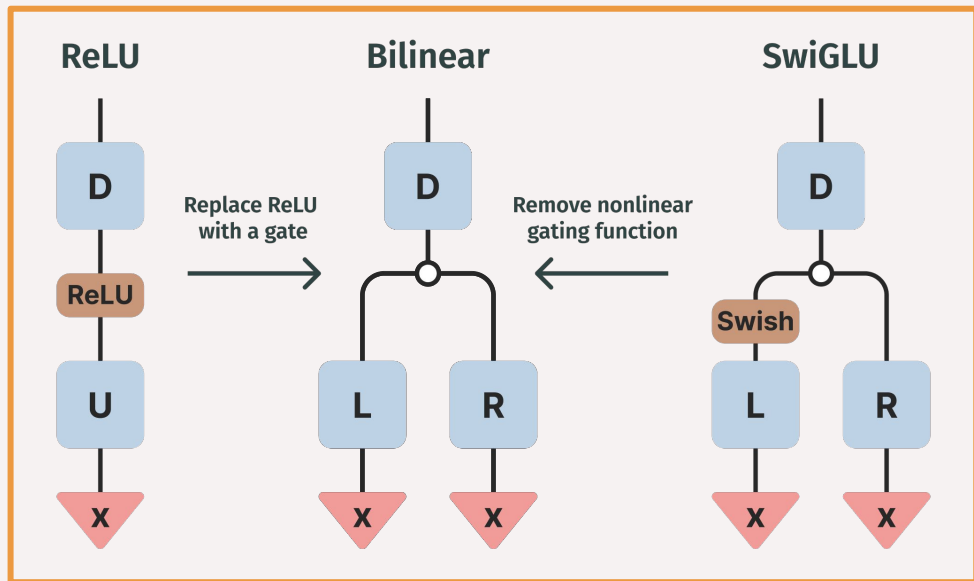
Tensor models are tractable to refine and measure. Their structure resembles neural architectures.

A rich framework for decompositions

A **tensor model** is a compositional model with

- the syntax of an **additive hypergraph**
- semantics in **FdHilb $\mathbb{R}$** (linear maps and finite-dimensional real Hilbert spaces)

A box is a higher-order interaction



*The appealing properties of **compositional understanding** can be scaled to **state-of-the-art networks** using tensor-based models.*

Coming up...

1. **How do we operationalise compositional interpretability?**
2. What can we do with these models?
3. How are these used in practice?

A curious table from a curious paper.

Training Steps	65,536	524,288
FFN_{ReLU} (<i>baseline</i>)	1.997 (0.005)	1.677
FFN_{GELU}	1.983 (0.005)	1.679
$\text{FFN}_{\text{Swish}}$	1.994 (0.003)	1.683
FFN_{GLU}	1.982 (0.006)	1.663
$\text{FFN}_{\text{Bilinear}}$	1.960 (0.005)	1.648
$\text{FFN}_{\text{GEGLU}}$	1.942 (0.004)	1.633
$\text{FFN}_{\text{SwiGLU}}$	1.944 (0.010)	1.636
$\text{FFN}_{\text{ReGLU}}$	1.953 (0.003)	1.645

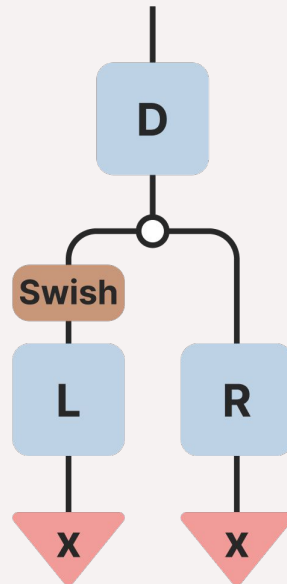
"We offer no explanation as to why these architectures seem to work; we attribute their success, as all else, to divine benevolence."

A bilinear layers is a tensor,
which can be interpreted compositionally.

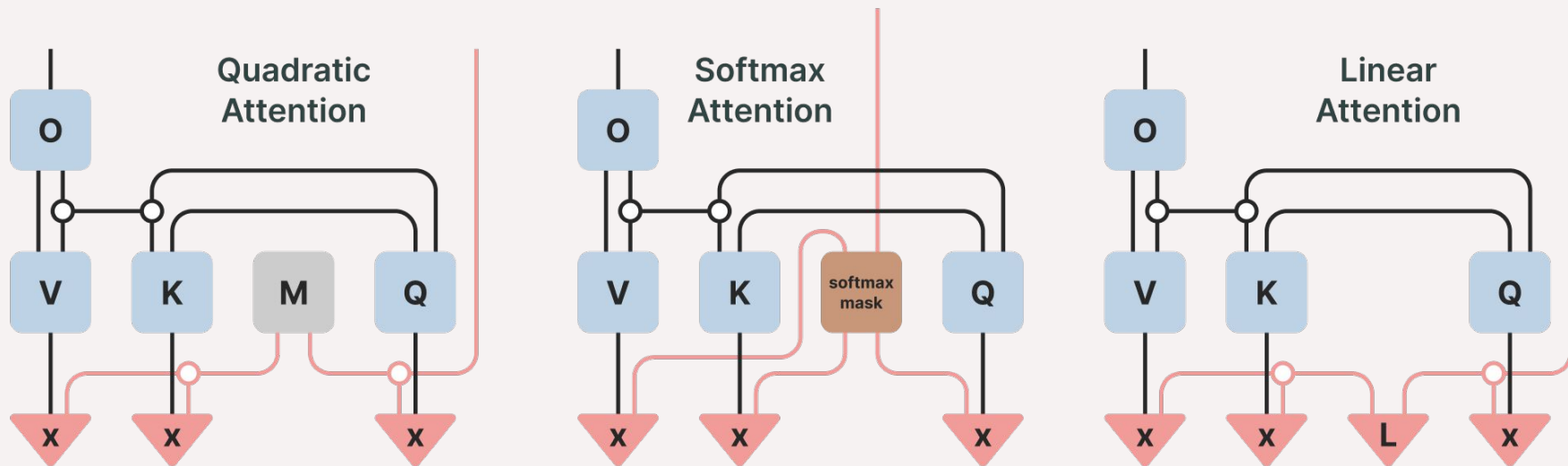
ReLU



SwiGLU



Attention can also be easily converted.
We can use a similar trick by making it quadratic.



Tensor models are performant across tasks. Training to the same loss is equally fast.



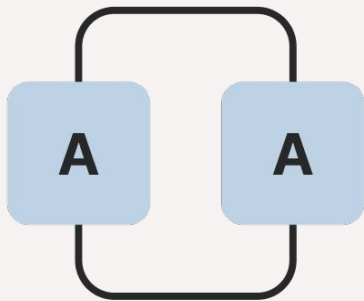
Coming up...

1. How do we operationalise compositional interpretability?
- 2. What can we do with these models?**
3. How are these used in practice?

Many operations on matrices are tensor networks.

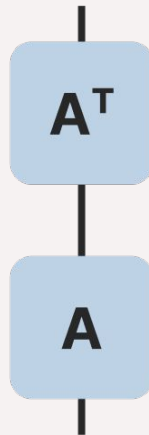
Frobenius Norm

$$\|A\|_F^2 = \sum_{ij} A_{ij} A_{ij}$$



Gram Matrix

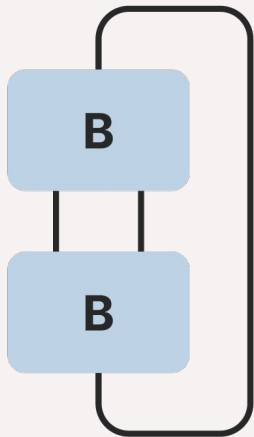
$$(A^\top A)_{jk} = \sum_i A_{ij} A_{ik}$$



These operations intuitively generalise to tensors.
This enables efficient computation of metrics.

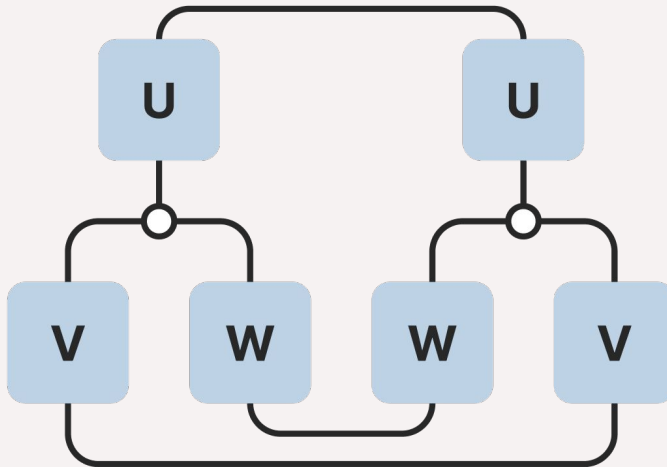
3-Tensor Norm

$$\|B\|_F^2 = \sum_{ijk} B_{ijk} B_{ijk}$$

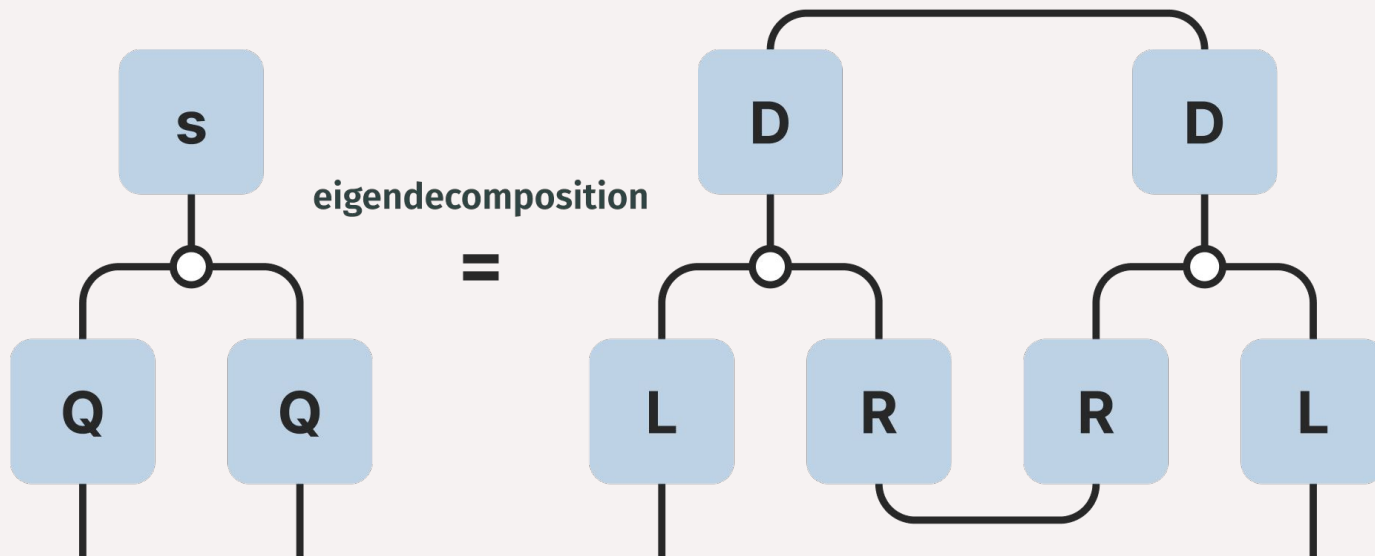


3-Tensor Norm

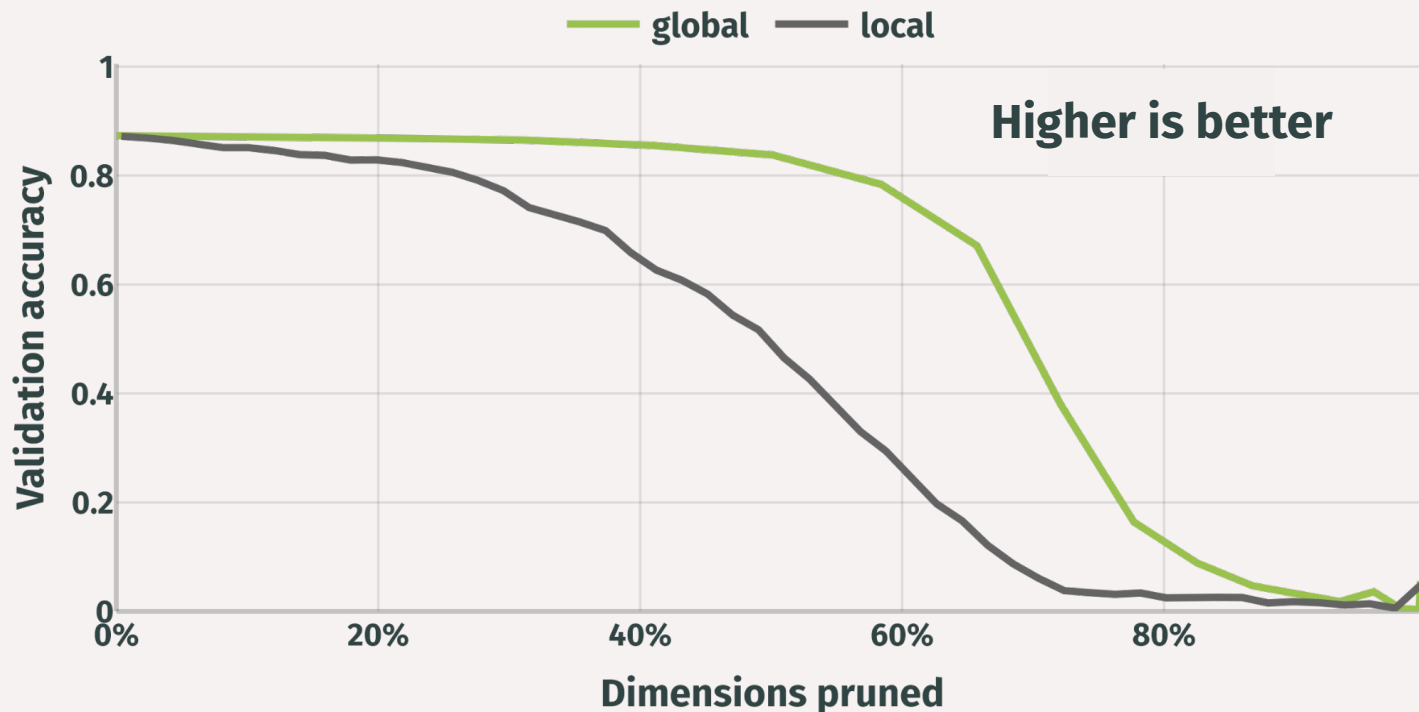
$$\|\cdot\|_F^2 = \sum_{ijknm} U_{in} V_{jn} W_{kn} U_{im} V_{jm} W_{km}$$



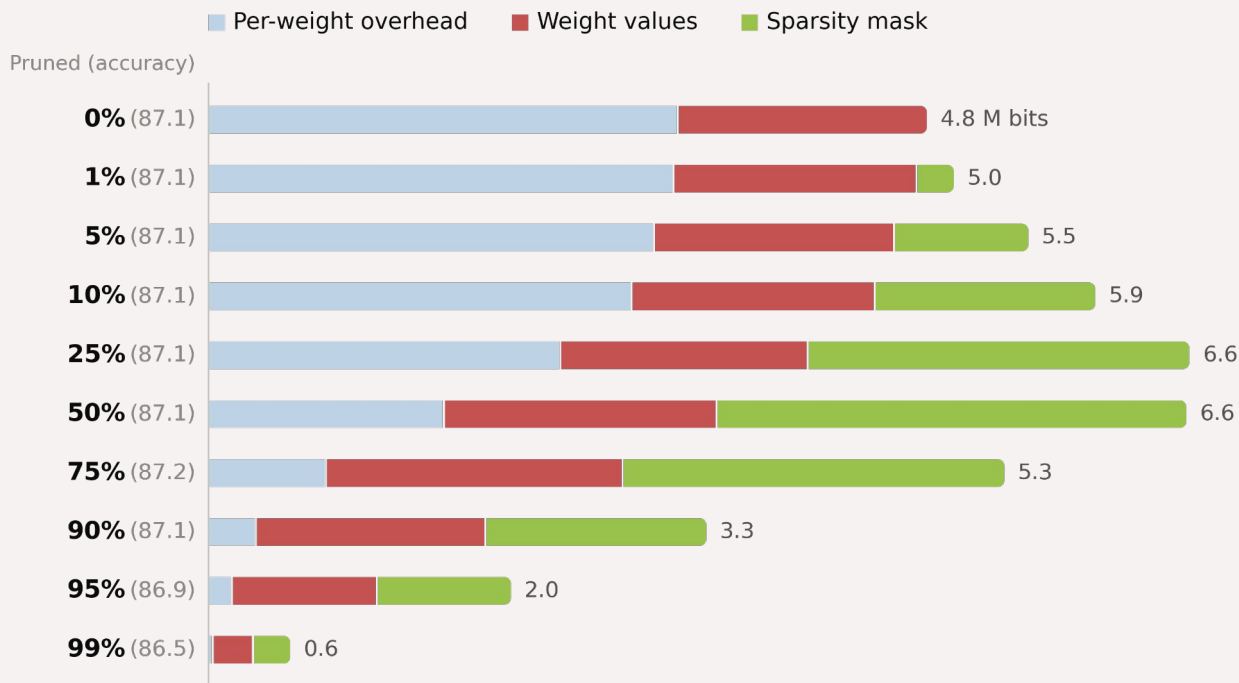
This enables decomposing models analytically while keeping into account all present structure.



Global decompositions outperforms local versions.
Models retain accuracy under strong compression.



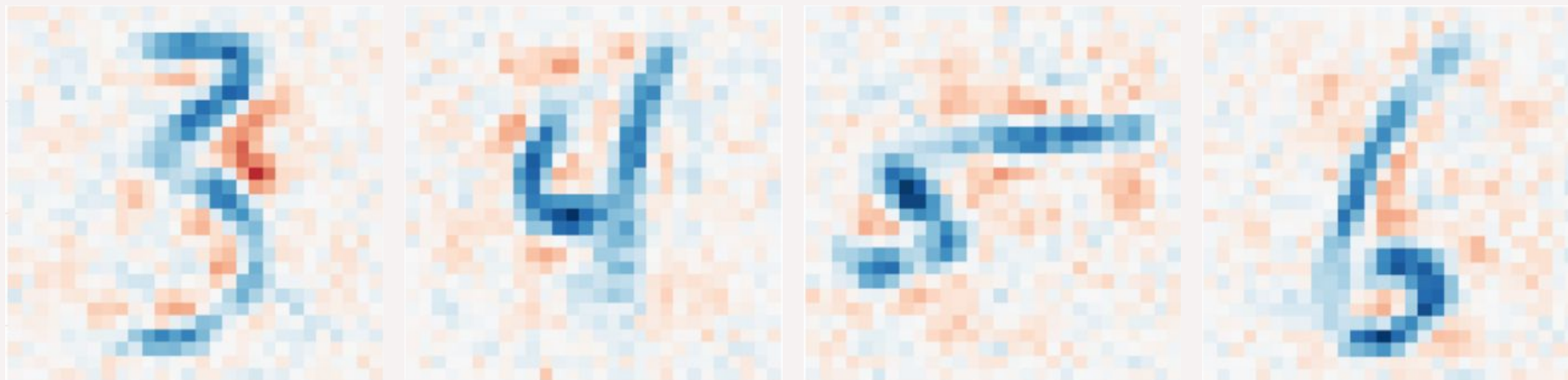
We can use weight sparsity to prune networks. This allows even lower description length.



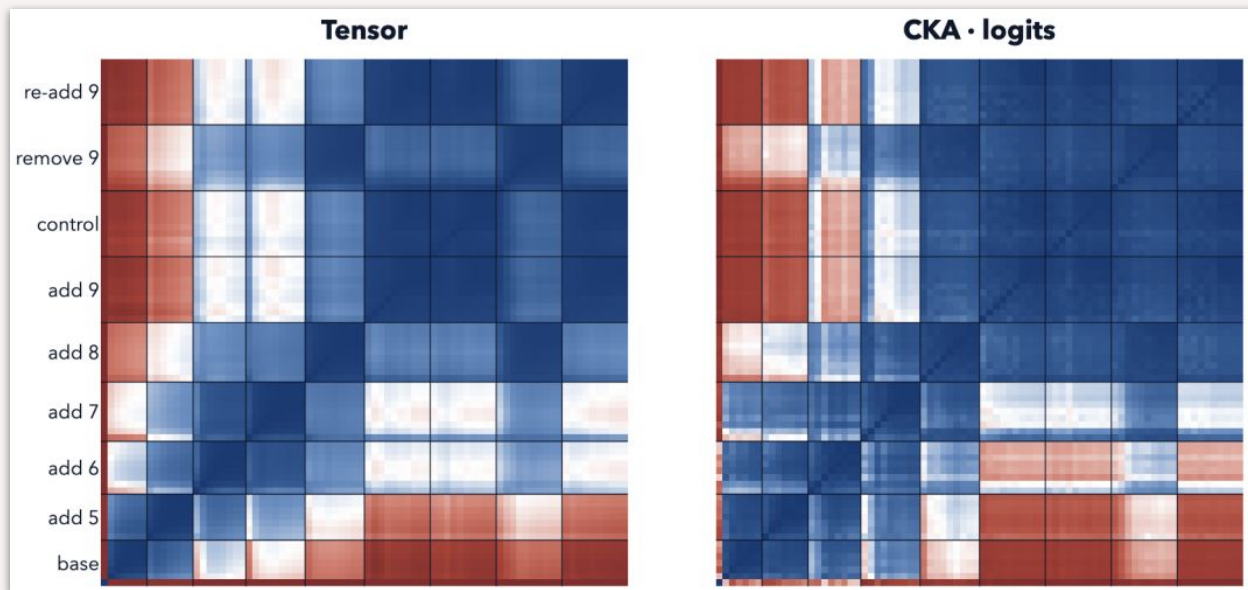
Coming up...

1. How do we operationalise compositional interpretability?
2. What can we do with these models?
3. **How are these used in practice?**

Rank decompositions also reveal what models learn, purely from their weights.

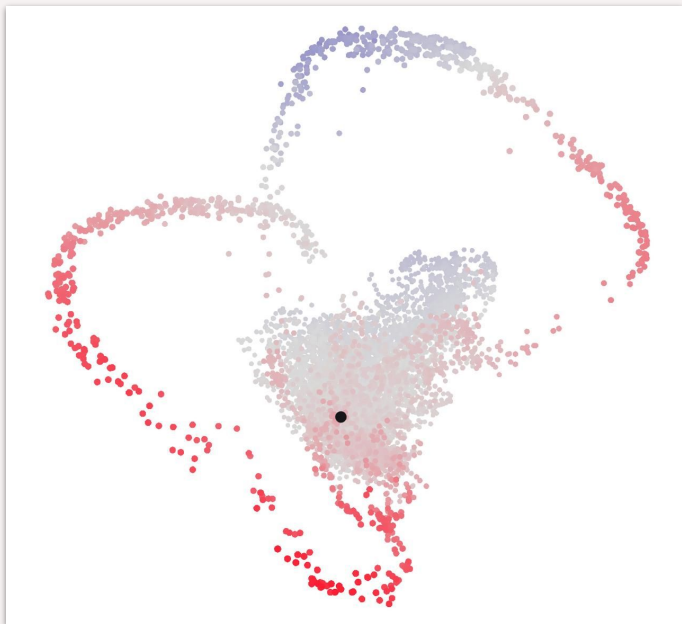


Tensors are not only useful for decomposing, but can also measure similarity across training.

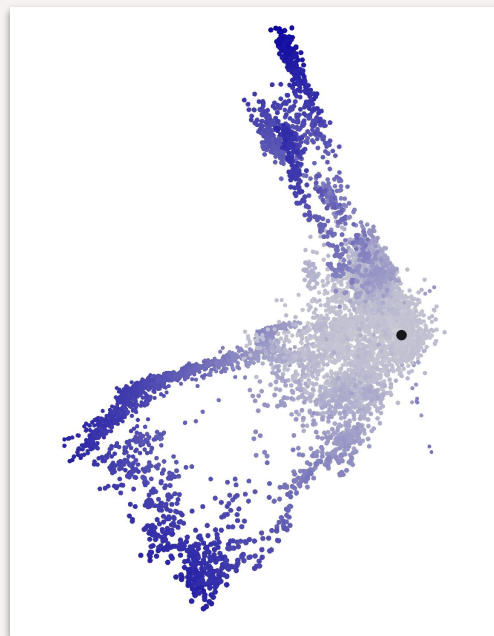


Interpretability aims to find the geometry of concepts inside neural models.

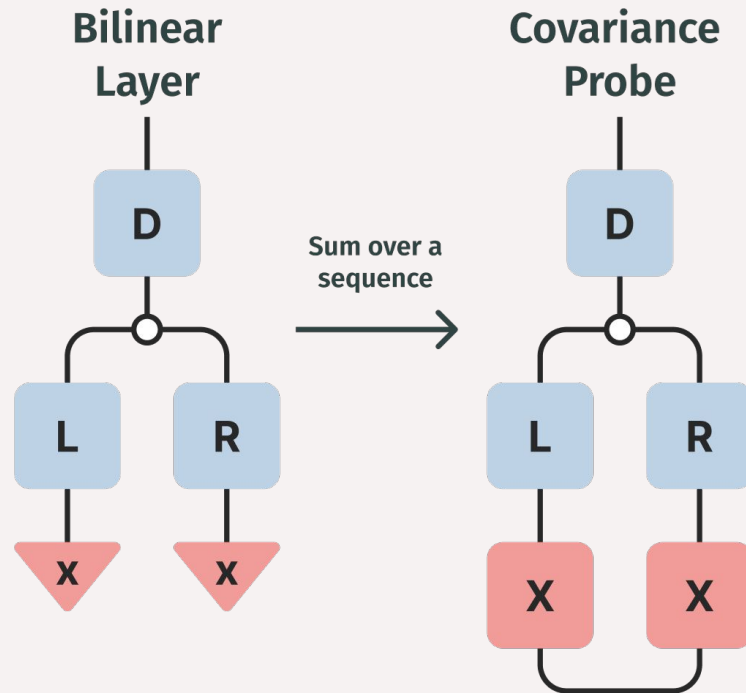
Sparse composition



Sparse activation



These tools are useful outside of language models, they can become drivers of scientific discovery.



These tools are useful outside of language models, they can become drivers of scientific discovery.



EVEE Evo Variant Eff

Search gene, rsID, or ClinVar ID (BRCA1, rs80357906, 15126)

GOODFIRE

SDHA c.217G>A (p.Gly73Ser)
ClinVar:2933037 | rs2477286087 | chr5:224426 G>A | exon 3/15 | Missense

PREDICTED PATHOGENICITY
98%

VUS ★★☆☆
Hereditary cancer-predisposing syndrome | gnomAD AF: not observed

ClinVar gnomAD dbSNP UCSC UniProt Ensembl OMIM GeneCards

VARIANT INTERPRETATION ?

Gly73Ser disrupts beta-strand geometry in the FAD-binding domain of SDHA while also perturbing a nearby splice acceptor context, collectively impairing proper protein folding and potentially pre-mRNA processing of the SDHA transcript.

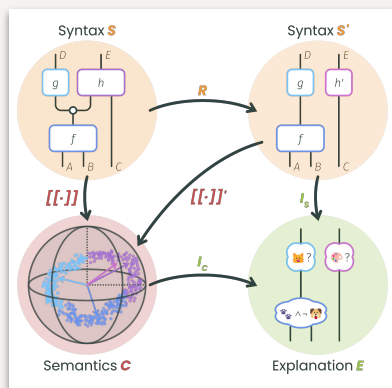
This glycine-to-serine substitution at position 73 of SDHA generates a high-confidence pathogenicity prediction of 98%, consistent with the calibration data showing 97% of variants in this score range are pathogenic. The most prominent disruptions are at a splice acceptor region approximately 945-952bp upstream, where polypyrimidine tract identity drops by 0.49, intron identity by 0.42, intron-to-exon transition signal by 0.41, and splice acceptor probability by 0.34, suggesting the variant propagates a conformational or sequence context change that destabilizes a nearby intronic regulatory element. At the variant site itself, the P-loop containing nucleoside triphosphate hydrolase domain signal increases by 0.34, which reflects a gain of anomalous domain-like character rather than preservation of normal architecture, and is accompanied by a gain of disorder (+0.27) and loss of beta-strand integrity (-0.27) in the immediate vicinity. SDHA encodes the flavoprotein subunit of succinate dehydrogenase (Complex II), where the FAD-binding domain near the N-terminus depends on precise beta-strand geometry for cofactor coordination and electron transfer; disruption of beta-strand structure at this position would impair FAD binding and catalytic function. The combination of local structural disruption and upstream splice region perturbation produces a broad disruption profile consistent with a pathogenic variant causing mitochondrial Complex II dysfunction.

KEY EVIDENCE

- Polypyrimidine tract disruption of -0.49 at -952bp, with correlated splice acceptor loss of -0.34 at -945bp
- Beta-strand identity loss of -0.27 at -20bp from variant, consistent with structural disruption in the FAD-binding domain
- Gain of disorder signal (+0.27) near the variant site, indicating local structural destabilization
- Anomalous P-loop hydrolase domain signal gain (+0.34) at the variant position, reflecting aberrant sequence context
- Calibration data: 97% of variants scoring 90-101% are pathogenic, and this variant scores 98%

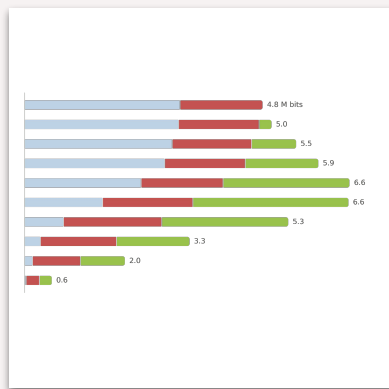
Interpretability is a constrained optimisation problem.

Optimize simplicity under compositional constraints.



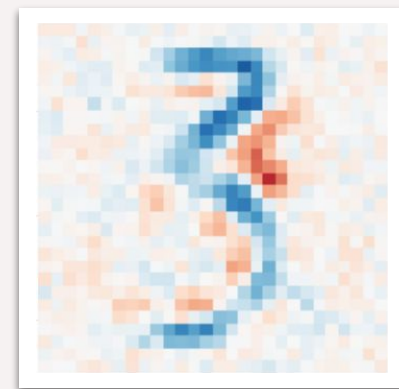
Composition

+



Compression

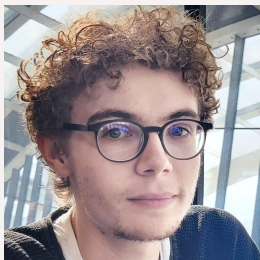
=



Interpretability*

*Left as exercise to the reader.

Want to know more or reach out?



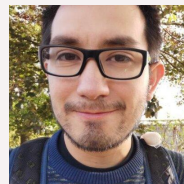
Ward Gauderis



Geraint Wiggins



Thomas Dooms



Jose Oramas