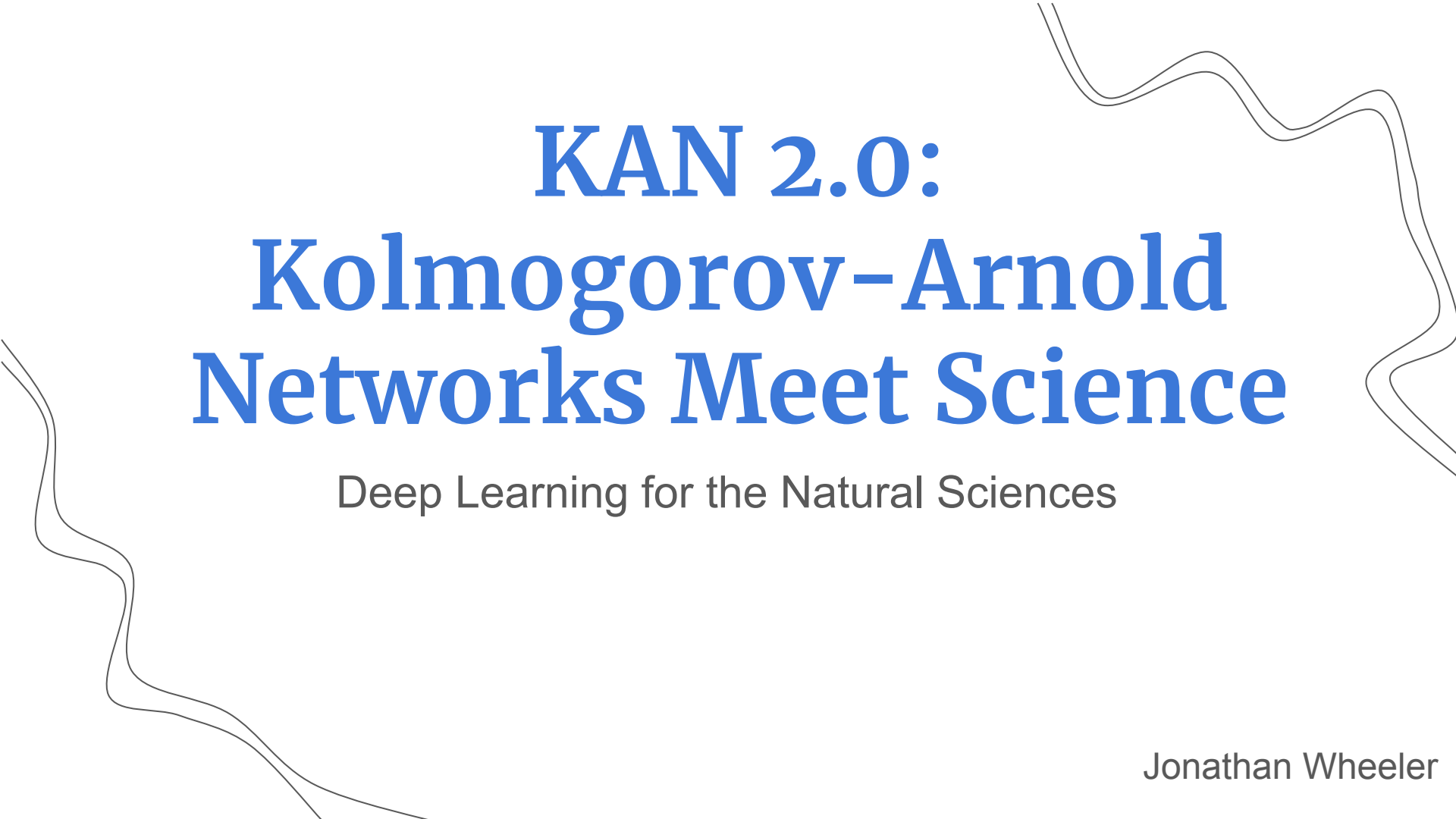




KAN 2.0: Kolmogorov–Arnold Networks Meet Science

Deep Learning for the Natural Sciences

Jonathan Wheeler

Table of contents

01

Motivation

Current Landscape

02

What are KANs?

Theoretical Introductions

03

Building Scientific Knowledge into KANs

Knowledge Embedding

04

Extracting Knowledge from KANs

05

Real-World Applications

06

Future Directions

1

Motivation

Current Landscape

The AI + Science Landscape



Current State

- AI is transformational
- Deep learning dominates current approaches



Major Achievements

- **AlphaFold:** Near-perfect protein structure prediction
- **AI Weather Forecasting:** Beyond traditional accuracy

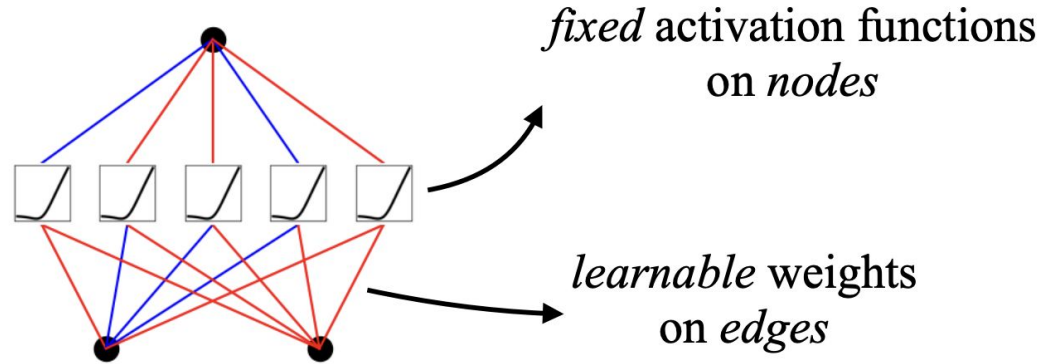


Remaining Challenges

- Limited scientific understanding from AI systems
- Need for more transparent AI approaches

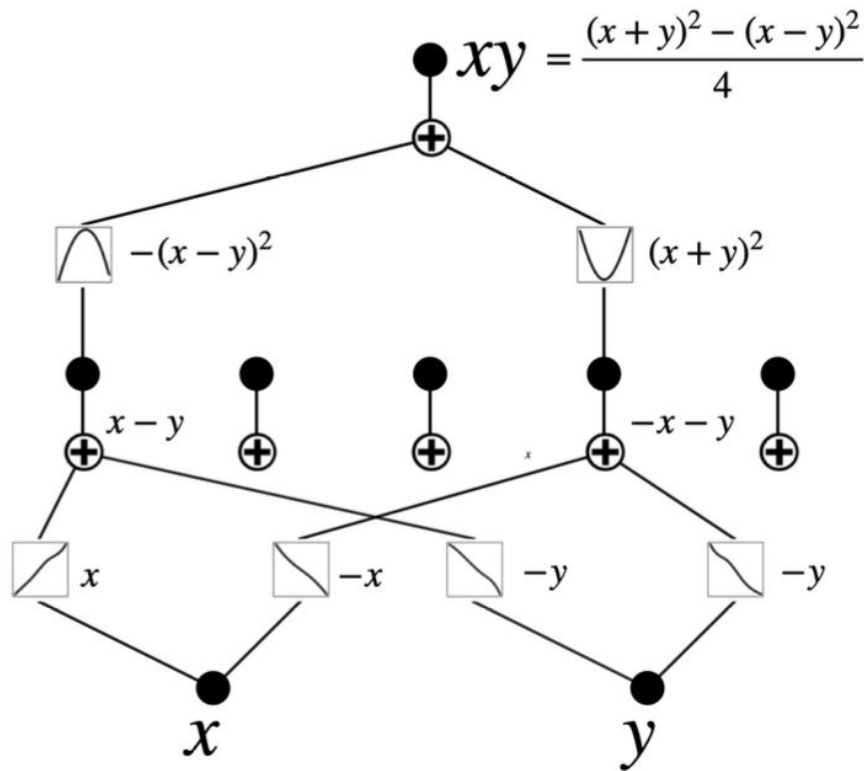
Application-driven Science (MLP)

- **What (is the output)?**
- Works for well formulated problems with clear objectives
- Less interpretable MLP models



Curiosity-driven Science (KAN)

- **Why (is this the result)?**
- Scientific discoveries require mathematical understanding and clear principles
- KAN's offer more interpretable structure through edge functions



2

Understanding KANs

Theoretical introductions

The Kolmogorov-Arnold Theorem

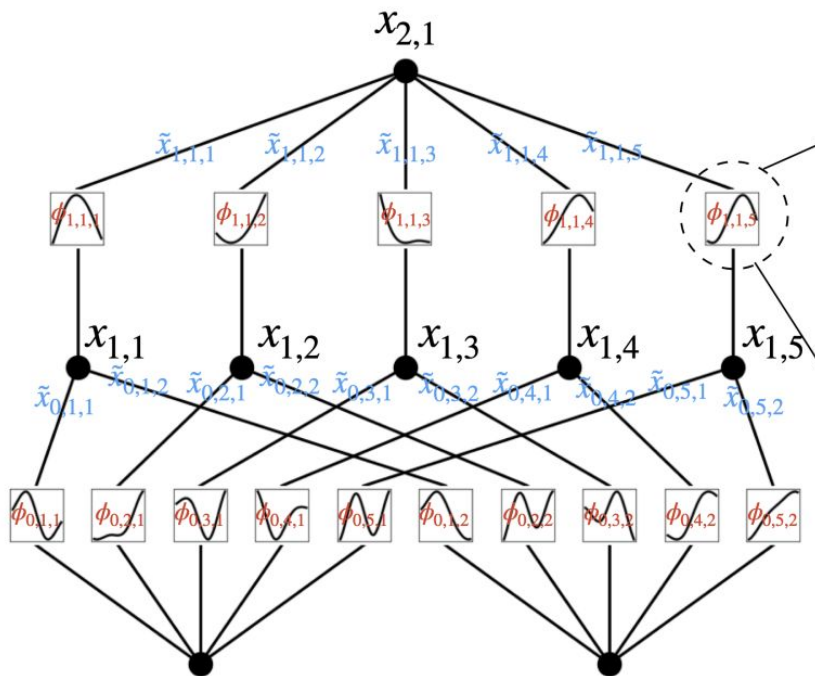
$$f(\mathbf{x}) = f(x_1, \dots, x_n) = \sum_{q=1}^{2n+1} \Phi_q \left(\sum_{p=1}^n \phi_{q,p}(x_p) \right)$$

- States that **any continuous function in a high-dimensional space can be represented as a finite composition of univariate continuous functions and additions.**
- Example:

$$xy = \exp(\log x + \log y)$$

Two-layer Network

$$f(\mathbf{x}) = f(x_1, \dots, x_n) = \sum_{q=1}^{2n+1} \Phi_q \left(\sum_{p=1}^n \phi_{q,p}(x_p) \right)$$



Key problem is has to be smooth functions

Key Breakthrough

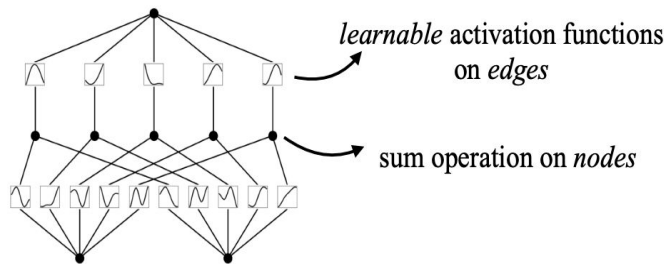
$$\mathbf{x}_{l+1} = \underbrace{\begin{pmatrix} \phi_{l,1,1}(\cdot) & \phi_{l,2,1}(\cdot) & \cdots & \phi_{l,n_l,1}(\cdot) \\ \phi_{l,1,2}(\cdot) & \phi_{l,2,2}(\cdot) & \cdots & \phi_{l,n_l,2}(\cdot) \\ \vdots & \vdots & & \vdots \\ \phi_{l,1,n_{l+1}}(\cdot) & \phi_{l,2,n_{l+1}}(\cdot) & \cdots & \phi_{l,n_l,n_{l+1}}(\cdot) \end{pmatrix}}_{\Phi_l} \mathbf{x}_l$$

Generalizing the original Kolmogorov-Arnold representation,
extending it to arbitrary depths

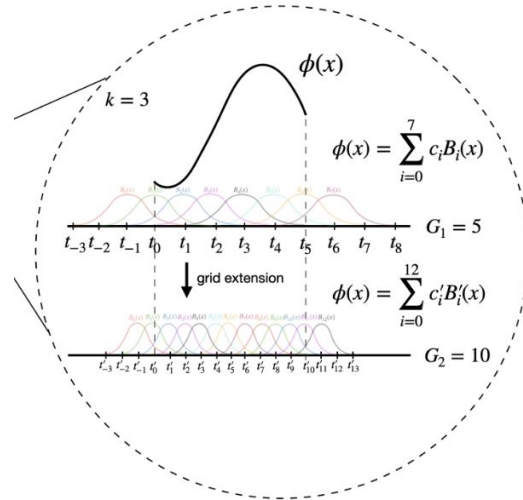
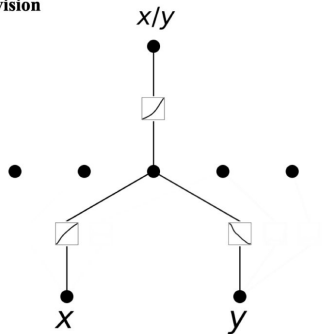
$$f(\mathbf{x}) = \sum_{i_{L-1}=1}^{n_{L-1}} \phi_{L-1,i_L,i_{L-1}} \left(\sum_{i_{L-2}=1}^{n_{L-2}} \cdots \left(\sum_{i_2=1}^{n_2} \phi_{2,i_3,i_2} \left(\sum_{i_1=1}^{n_1} \phi_{1,i_2,i_1} \left(\sum_{i_0=1}^{n_0} \phi_{0,i_1,i_0}(x_{i_0}) \right) \right) \right) \cdots \right)$$

Basic KAN Architecture

- nodes (summation)
- edges (learnable activation functions)
 - B-Splines
- Symbolic regression gives output formula



division

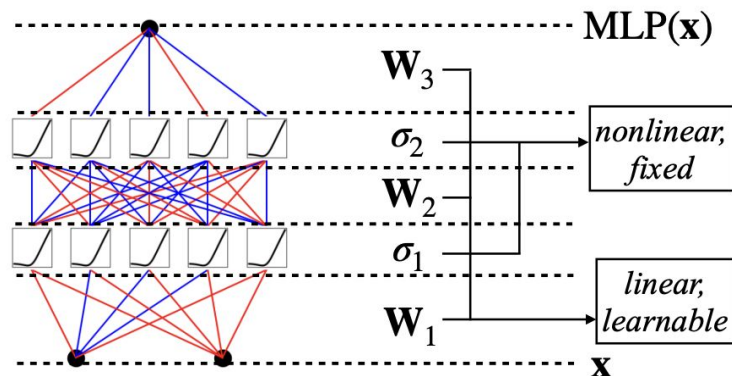


$$\triangleright \text{KAN}(\mathbf{x}) = (\Phi_3 \circ \Phi_2 \circ \Phi_1)(\mathbf{x})$$

Comparison

MLPs:

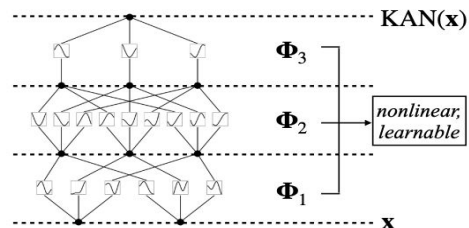
- Universal representation theorem based
- fixed nonlinear activation
- Each neuron processes weighted sum of inputs



$$\text{MLP}(\mathbf{x}) = (\mathbf{W}_{L-1} \circ \sigma \circ \mathbf{W}_{L-2} \circ \sigma \circ \cdots \circ \mathbf{W}_1 \circ \sigma \circ \mathbf{W}_0) \mathbf{x}.$$

KANs:

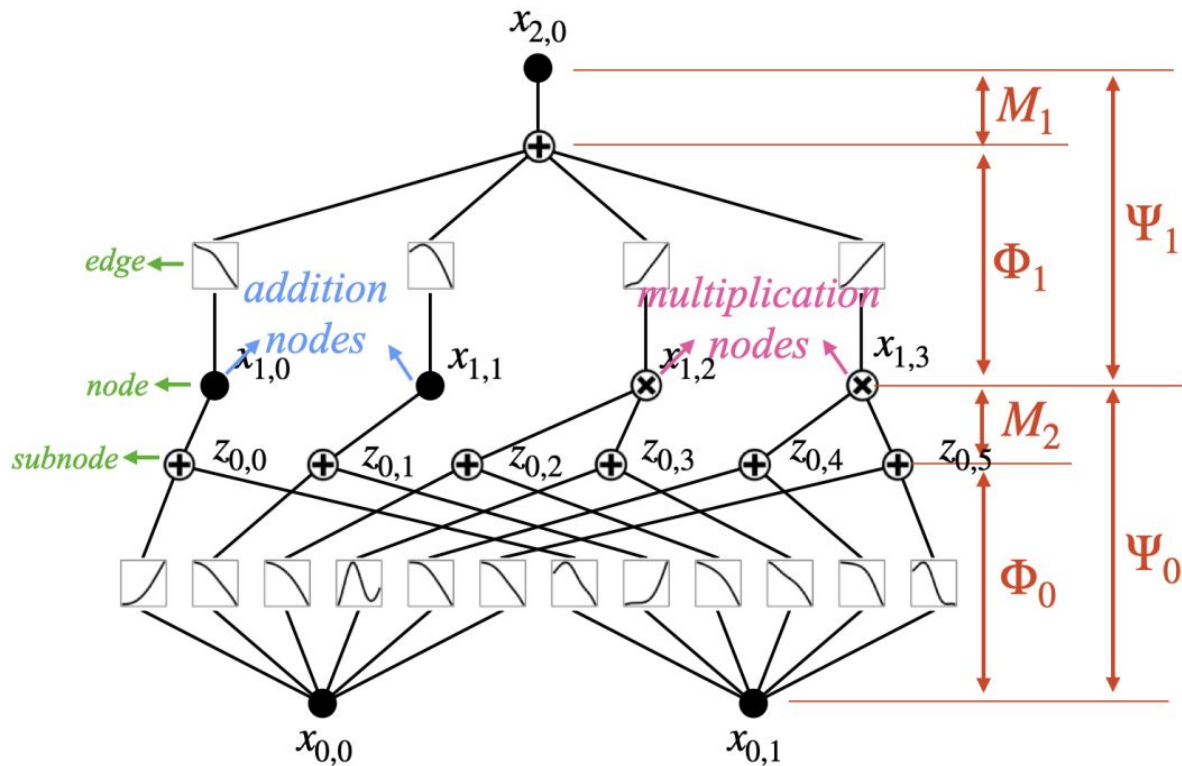
- Kolmogorov-Arnold representation based
- Edge-based univariate transformations
- Node-wise signal summation



$$\text{KAN}(\mathbf{x}) = (\Phi_{L-1} \circ \Phi_{L-2} \circ \cdots \circ \Phi_1 \circ \Phi_0) \mathbf{x}.$$

MultKAN

- Nodes(input) and subnodes(output)
- Standard KAN layer plus multiplication layer
- No additional trainable parameters
- Learns multiplication



$$\text{MultKAN}(\mathbf{x}) = (\Psi_L \circ \Psi_{L-1} \circ \dots \circ \Psi_1 \circ \Psi_0)\mathbf{x}.$$

Goal:

Problem: symbolic formulas are not always possible

Solution:

1. Build in scientific knowledge to KANs (section 3)
2. Extract out scientific knowledge from KANs (section 4)

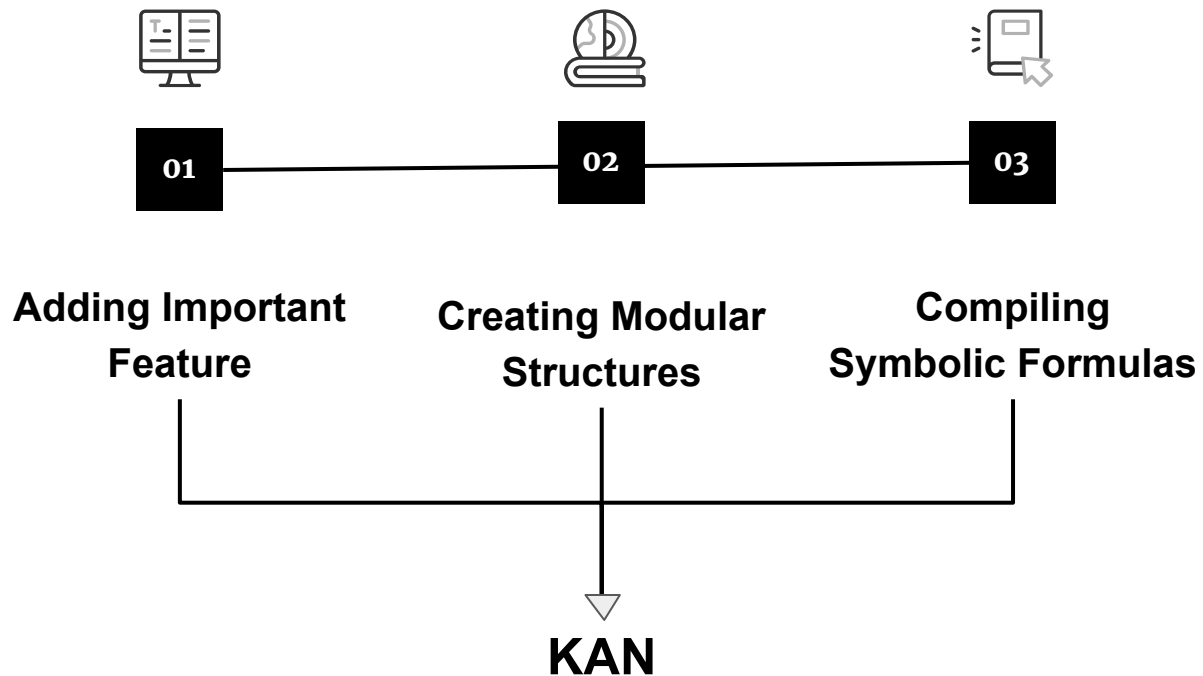
3

Building Scientific Knowledge

Knowledge Embedding

Overview of embedding knowledge

incorporate
available inductive
biases into KANs



1. Adding important features to KANs

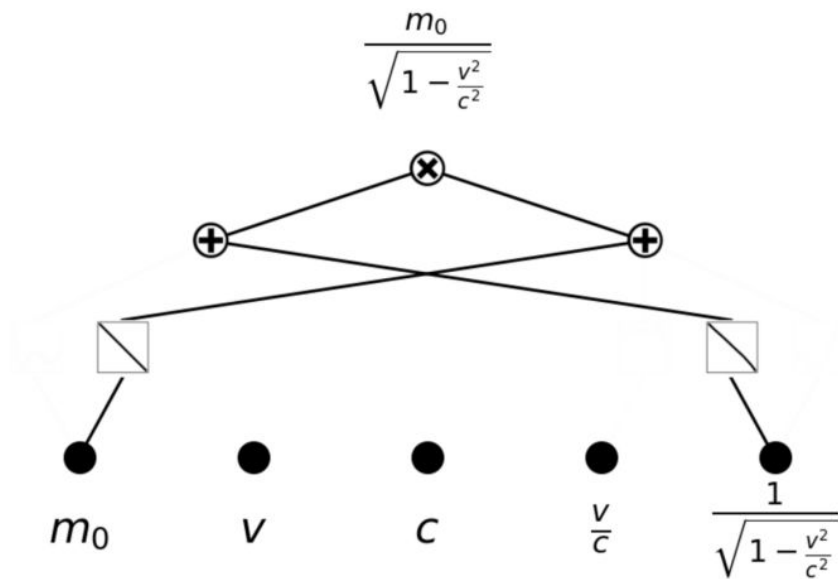
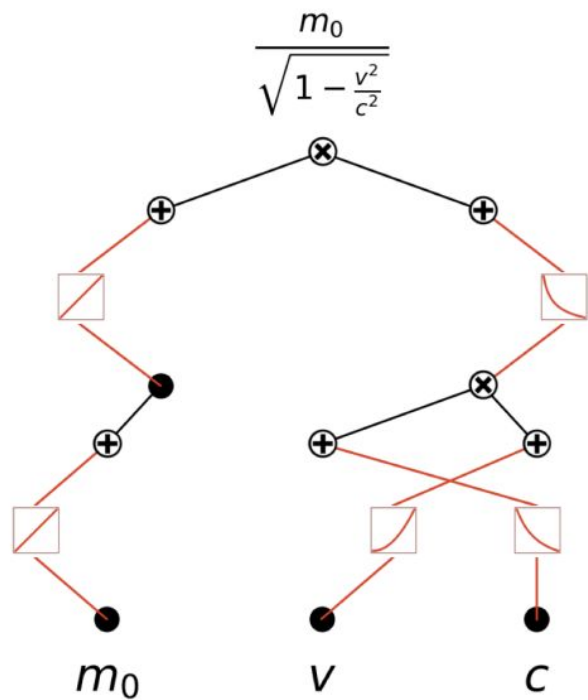
- Addition of auxiliary input variable
- Adds expressive power and interpretability to the network

$$y = f(x_1, x_2, \dots, x_n)$$

$$a = a(x_1, x_2, \dots, x_n)$$

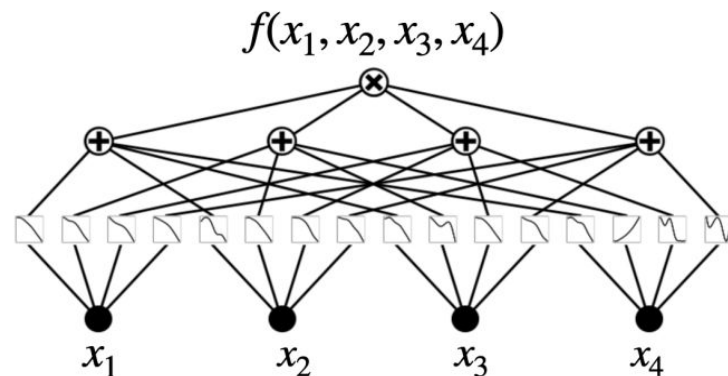
$$y = f(x_1, \dots, x_n, a)$$

Example



2. Modular structures

- Allows for extracting knowledge from clusters of neurons
- KAN allows for explicit creation of modular structures.
- Two types of modularity:
 1. Separability
 2. Generalized Symmetry



```
[start_layer_id, '[nodes_id] -> [subnodes_id] -> [nodes_id] ...']
```

Separability

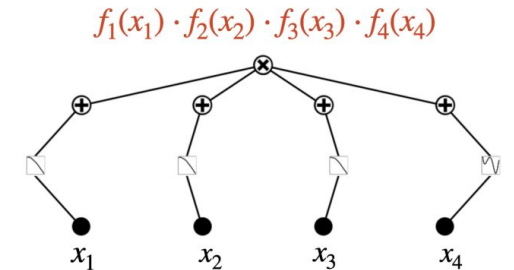
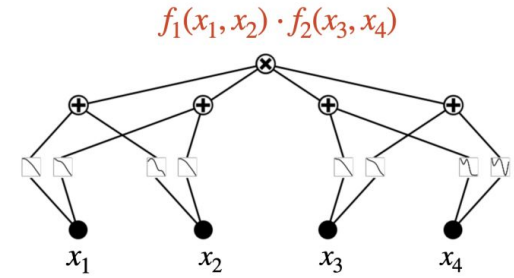
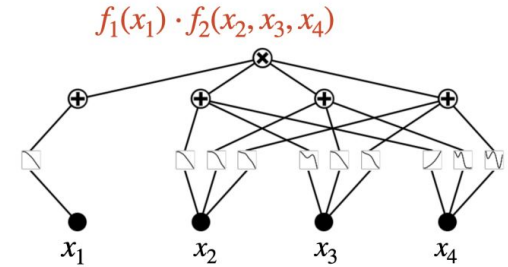
$$f(x_1, x_2, x_3, x_4)$$

f is separable if it can be expressed as a sum or product of functions of non-overlapping variable groups

0, ' [0] -> [0] ',

0, ' [0, 1] -> [0, 1] ',

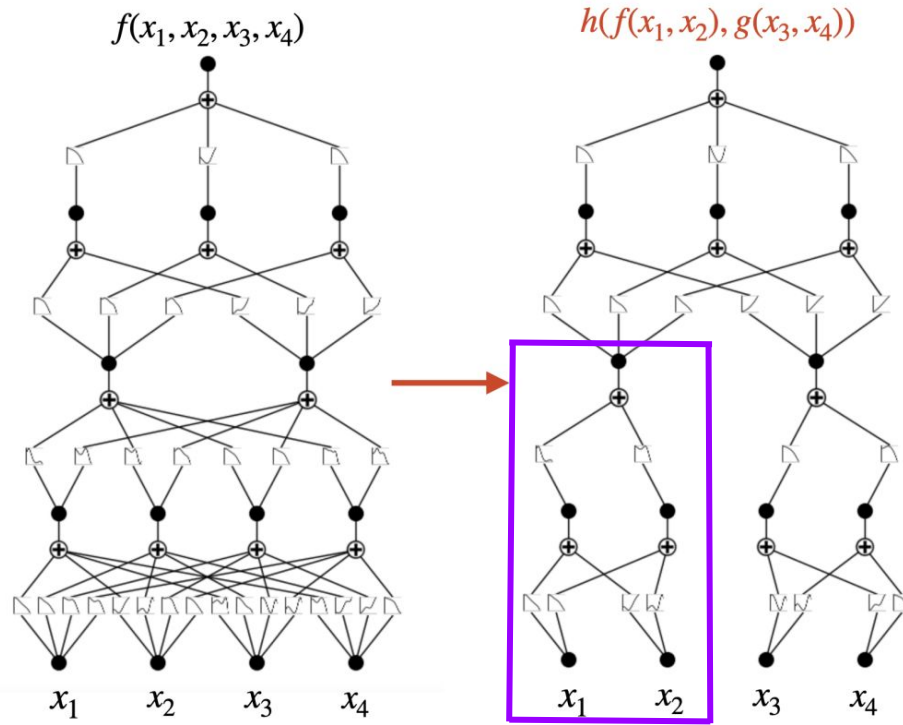
0, ' [i] -> [i] ',



Generalized symmetry

$$f(x_1, x_2, x_3, \dots) = g(h(x_1, x_2), x_3, \dots)$$

$$[0, [0, 1] \rightarrow [0, 1] \rightarrow [0, 1] \rightarrow [0]]$$



3. Compiling Symbolic Formulas

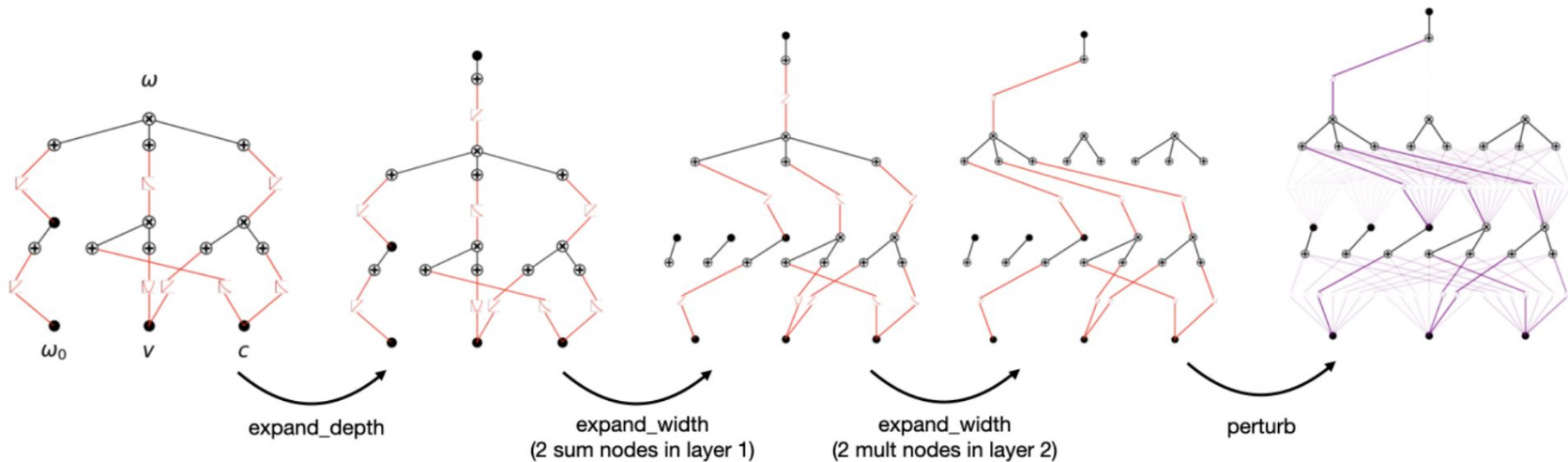
Combine expressivity of neural networks with symbolic explainable formulas

Two steps:

1. Compile formula into a KAN (introducing domain knowledge)
2. Train KAN on data (learn “new physics”)

KANPILER + Learning

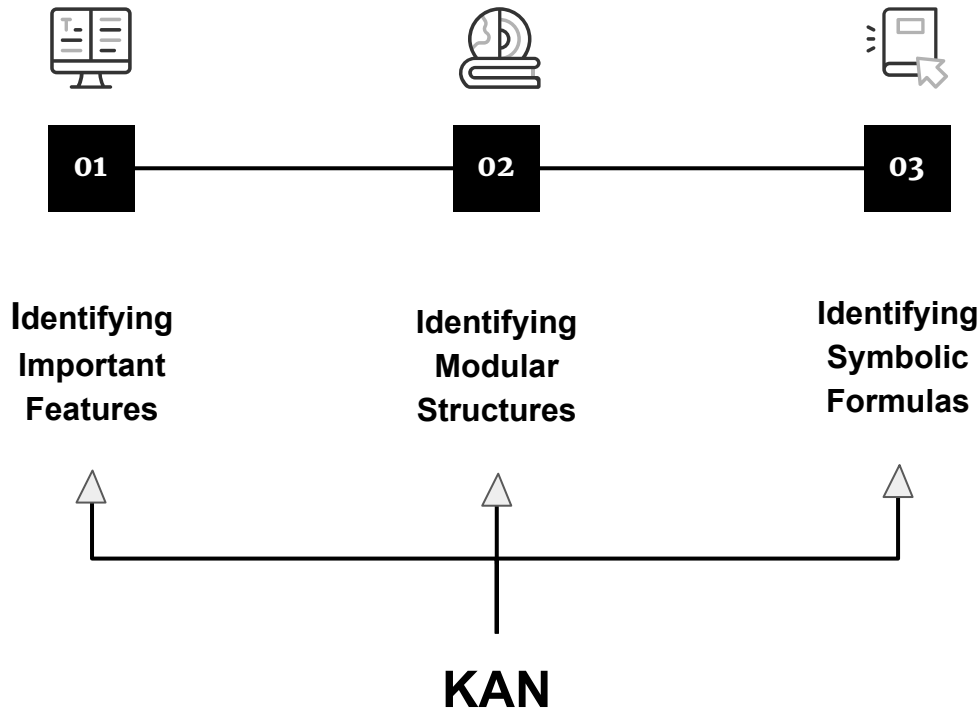
1. Parse Formula $\rightarrow$ Tree, Nodes = expressions, Edges = operations/functions
2. Transform Tree $\rightarrow$ KAN Structure Moves leaf nodes to input layer
3. Create Final Graph by combining variables in first layer



4

Extracting Knowledge

Overview of Knowledge Extraction

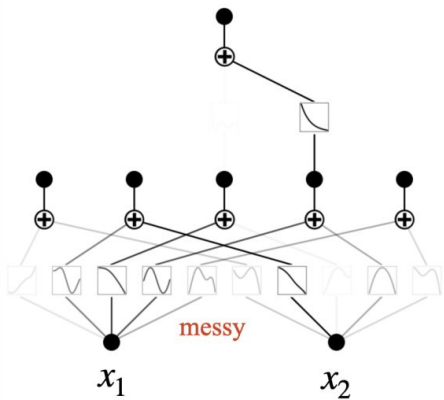
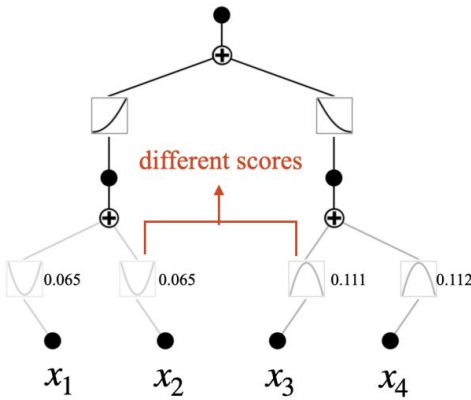
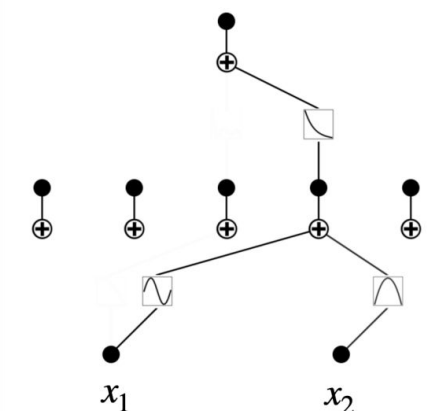
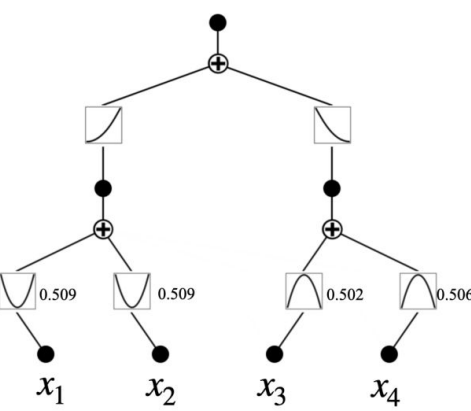


1. Identifying Important Features

- Assign scores to important inputs
- Previously: L1 norm (only considers local information)
 - $E_{l,i,j}$ = standard deviation of the function on each edge
 - $N_{l,i}$ = standard deviation at each node
- Update: Iterative computation of scores
 - node (attribution) score: $A_{l,i}$
 - edge (attribution) score: $B_{l,i,j}$

$$B_{l-1,i,j} = A_{l,j} \frac{E_{l,j}}{N_{l+1,j}}, \quad A_{l-1,i} = \sum_{j=0}^{n_l} B_{l-1,i,j}, \quad l = L, L-1, \dots, 1.$$

Example

	$y = \exp(\sin(\pi x_1) + x_2^2)$	$y = (x_1^2 + x_2^2)^2 + (x_3^2 + x_4^2)^2$
L1 function norm (old)		
attribution score (new)		

2. Identifying modular structures



1

anatomical modularity

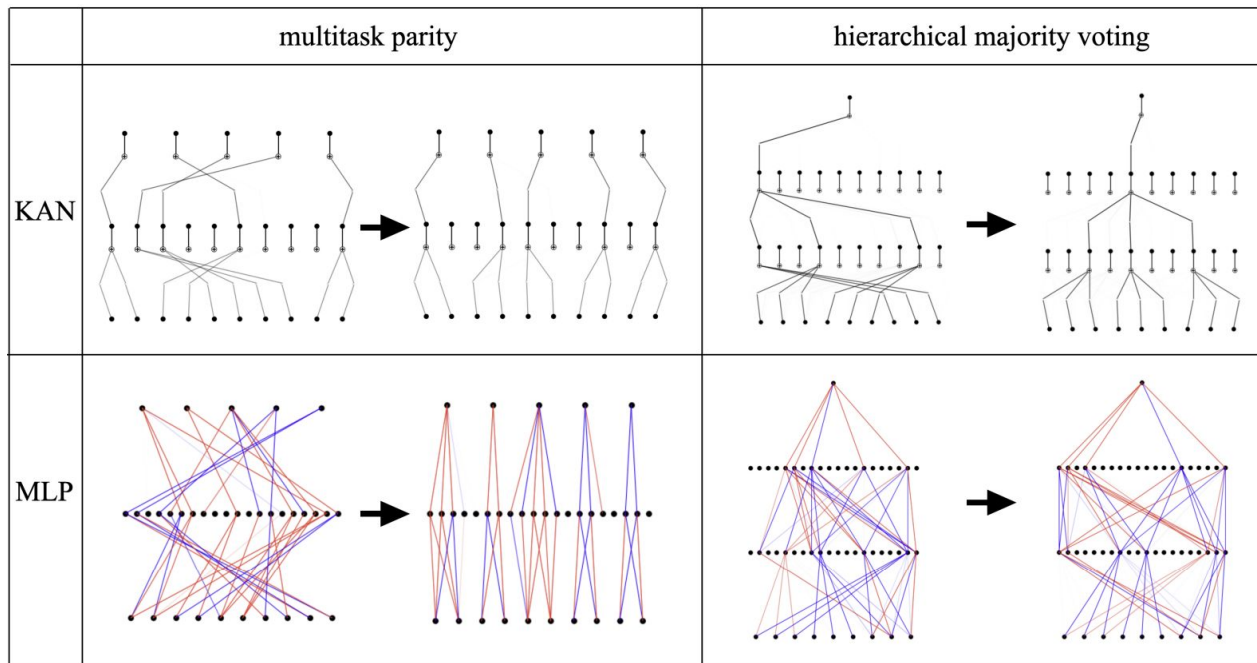


2

functional modularity

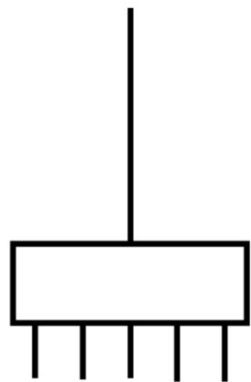
Anatomical Modularity

- Spatial proximity of neurons
- Autoswap: neuron swapping methods (preserves functionality)
- Easy to visually identify modular structures



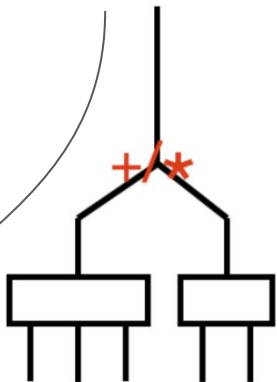
Functional modularity

(a) simplifying properties



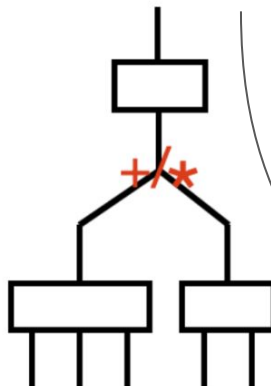
simplify
→

Separability
(Add or Mul)



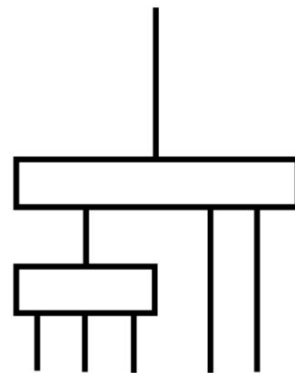
$$f(x_1, x_2, \dots, x_n) = g(x_1, \dots, x_k) + h(x_{k+1}, \dots, x_n)$$

Generalized
Separability



$$f(x_1, x_2, \dots, x_n) = F(g(x_1, \dots, x_k) + h(x_{k+1}, \dots, x_n))$$

Generalized
Symmetry

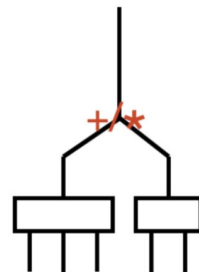


$$f(x_1, x_2, \dots, x_n) = g(h(x_1, \dots, x_k), x_{k+1}, \dots, x_n)$$

Functional Modularity: separability

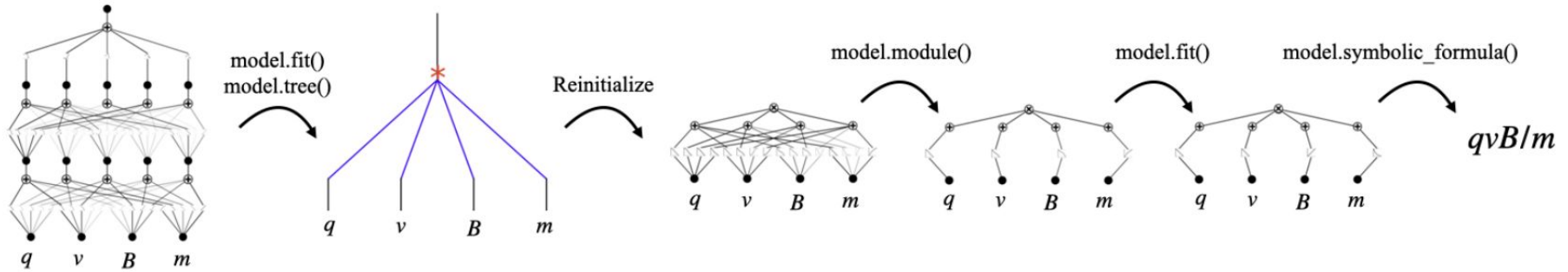
- Additively separable $\longrightarrow f(x_1, x_2, \dots, x_n) = g(x_1, \dots, x_k) + h(x_{k+1}, \dots, x_n)$
- Compute the hessian matrix $\mathbf{H} \equiv \nabla^T \nabla f \quad (\mathbf{H}_{ij} = \frac{\partial^2 f}{\partial x_i \partial x_j})$
- If $H_{ij} = 0$ for all $1 \leq i \leq k$ and $k + 1 \leq j \leq n$ then additively separable
- For multiplicative $\longrightarrow \log |f(x_1, x_2, \dots, x_n)| = \log |g(x_1, \dots, x_k)| + \log |h(x_{k+1}, \dots, x_n)|$
separability, convert to additive separability by taking the log

$$\mathbf{H}_{ij} \equiv \frac{\partial^2 \log |f|}{\partial x_i \partial x_j}$$



3. Identifying symbolic formulas

Trick A: discover and leverage modular structures

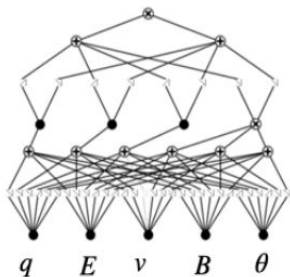


3. Identifying symbolic formulas

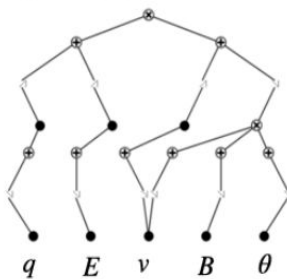
Trick B: Sparse initialization — a more natural fit

$$F = f(q, E, v, B, \theta) = q(E + vB\sin\theta)$$

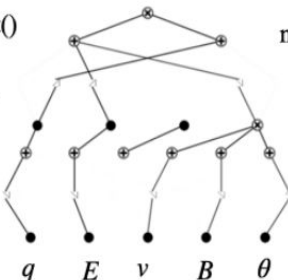
sparse_init = False



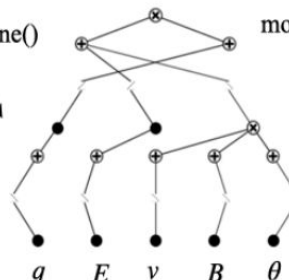
sparse_init = True



model.fit()



model.prune()



model.symbolic_formula()

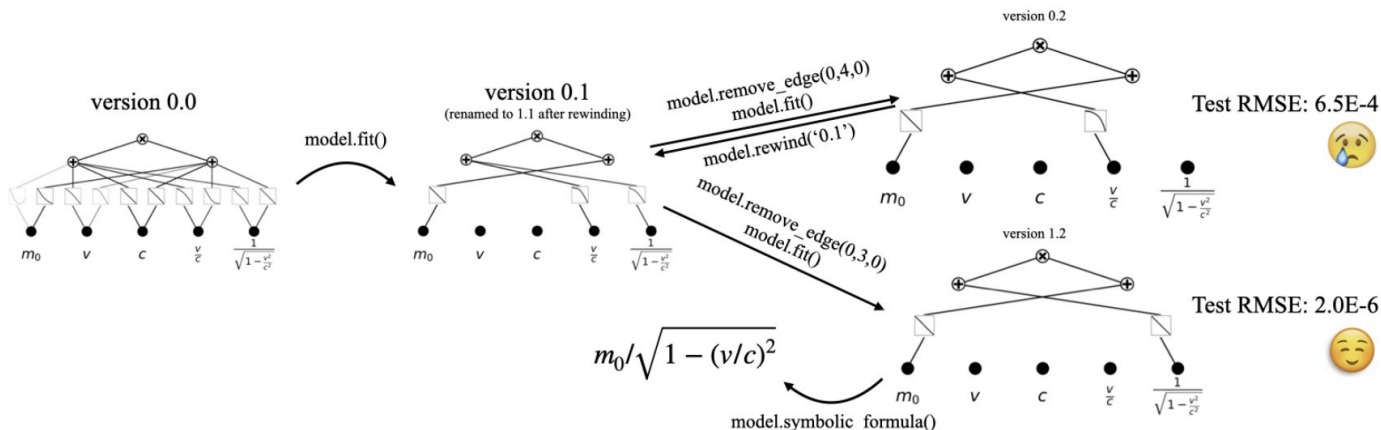


$$q(E + vB\sin\theta)$$

3. Identifying symbolic formulas

Trick C: Hypothesis testing - easy to try different ideas.

Trick C: Hypothesis testing $m = f(m_0, v, c) = m_0 / \sqrt{1 - (v/c)^2}$



5

Real-World Applications

Application areas

1. Discovering Lagrangians

- Core idea: Learn lagrangian from data
 - Phase space: position and velocity
- Predict accelerations using Euler-Lagrange equations

How KANs Help:

1. Numerical Stability:

- Pre-encode kinetic energy

$$L(q, \dot{q}) = T - V \text{ from } (q, \dot{q}, \ddot{q})$$

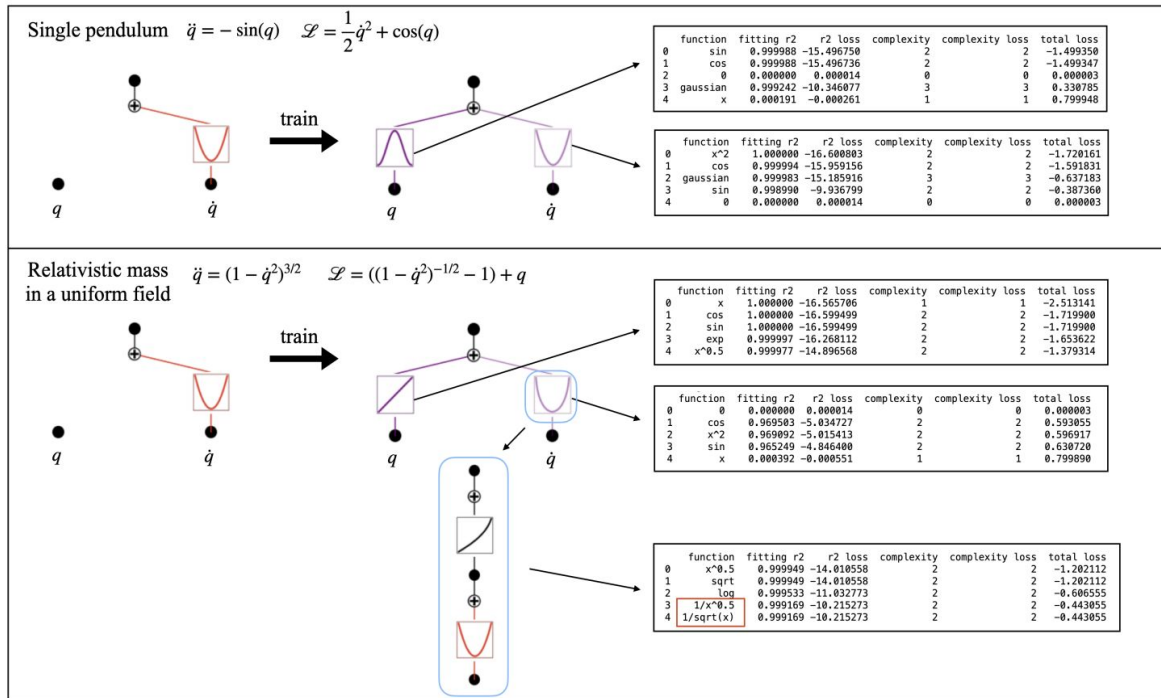
2. Interpretability:

- Symbolic regression extracts explicit formulas from KAN edges

$$\ddot{q} = (\nabla_{\dot{q}} \nabla_{\dot{q}}^T L)^{-1} [\nabla_q L - \nabla_q \nabla_{\dot{q}}^T \dot{q}]$$

KAN solution

- **Key Results:**
- **Stable Training:** Fewer numerical errors.
- **Explainable Models:** Learns physics-aligned formulas.



2. Constitutive laws

- **Constitutive Laws:**

- Define material behavior under forces or deformation (e.g., Hooke's Law).
- Linear for small deformations (Hooke's law)
- Nonlinear for larger deformations (e.g., Neo-Hookean).

$$P_{12} = \mu(F_{12} + F_{21})$$

- **Why Use KANs?**

- predict the P (stress) tensor from the F (deformation) tensor
- Incorporates **prior knowledge** (e.g., linear laws).
- Extracts **interpretable symbolic formulas** via symbolic regression.

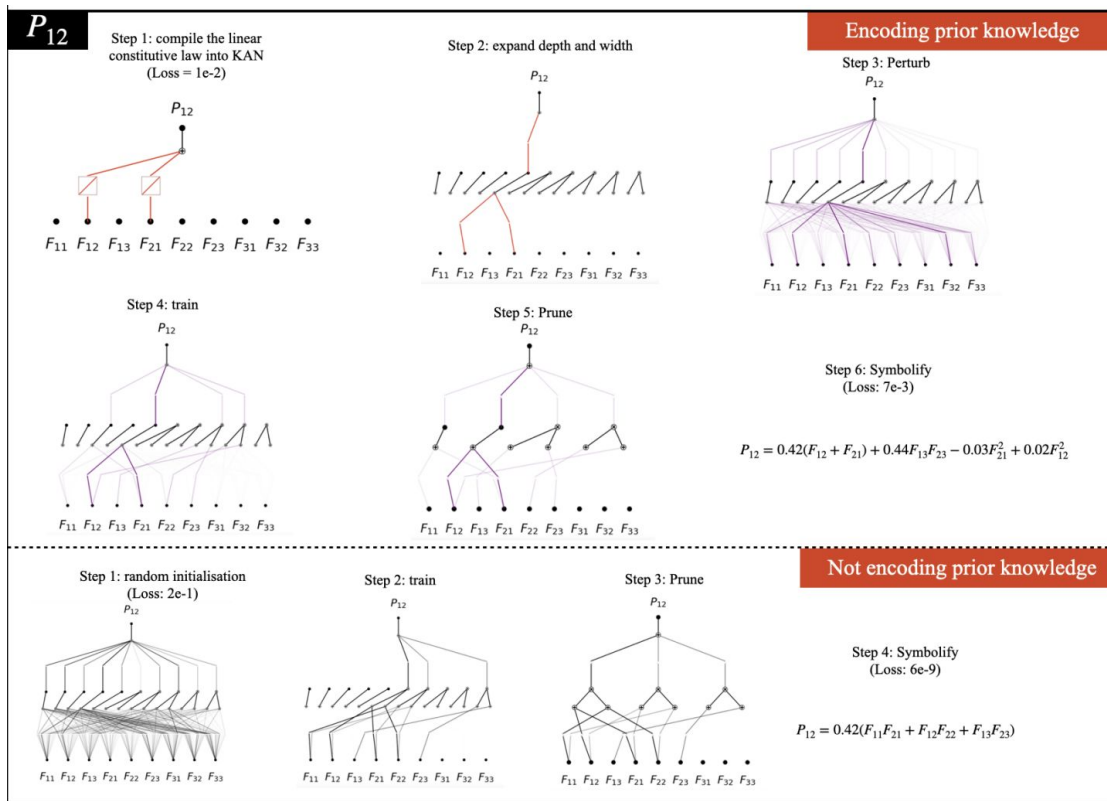
$$P_{12} = \mu(F_{11}F_{21} + F_{12}F_{22} + F_{13}F_{23})$$

Results:

With Prior Knowledge:

- The initial KAN is based on the linear constitutive law.
- Helps guide training but may converge to poor local minima
- **Without Prior Knowledge:**
 - Start with a randomly initialized KAN.
 - Requires more data to train but avoids being biased

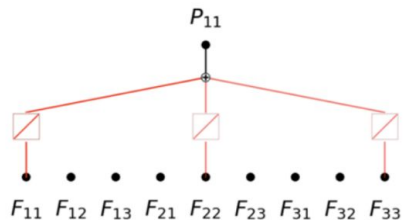
$$P_{12} = \mu(F_{12} + F_{21})$$



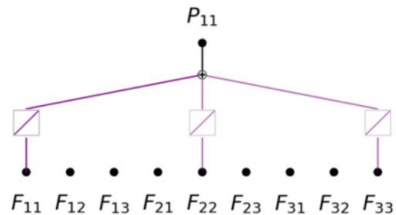
$$P_{12} = \mu(F_{11}F_{21} + F_{12}F_{22} + F_{13}F_{23})$$

P_{11}

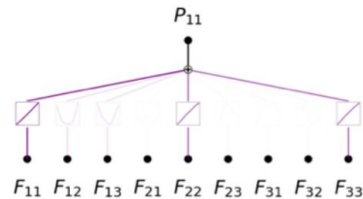
Step 1: compile the linear
constitutive law into KAN
(Loss = 1e-2)



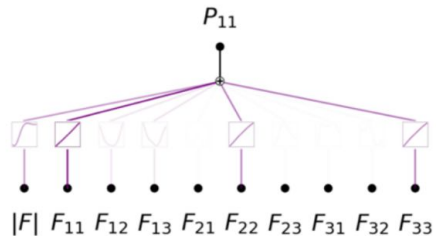
Step 2: Perturb



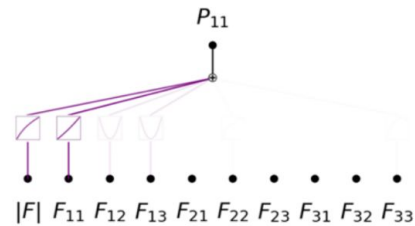
Step 3: train
(Loss = 6e-3)



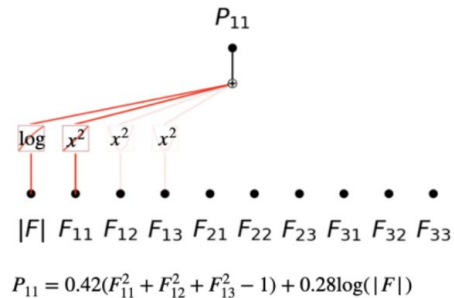
Step 4: Add the determinant as auxiliary variable



Step 5: train
(Loss = 2e-4)



Step 6: symbolify
(Loss = 3e-11)

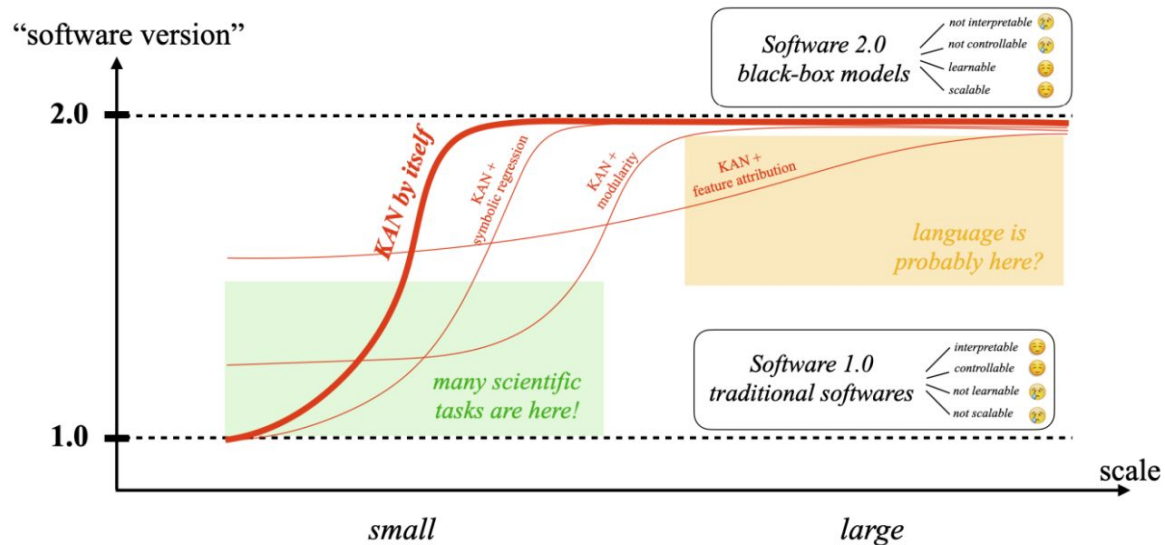


6

Future

Software

(b) “software version”-scale plane



Outlook

- **Scalability Challenge**
 - Limited interpretability at large scales
 - Need better methods for managing complexity
- **Key Research Areas**
 - Advanced interpretability methods
 - Scaling to larger problems
 - Expansion beyond physics

KAN VS MLP

MLP:

- Lower interpretability
- Well established and proven

KAN:

- 10x Slower training time
- Higher interpretability
- Continual learning
- More parameter efficient

References

1. Z. Liu, Y. Wang, S. Vaidya, F. Ruehle, J. Halverson, M. Soljačić, T. Y. Hou, and M. Tegmark, "KAN: Kolmogorov-Arnold Networks," *arXiv preprint*, arXiv:2404.19756, Jun. 2024. [Online]. Available: <https://doi.org/10.48550/arXiv.2404.19756>
2. Z. Liu, P. Ma, Y. Wang, W. Matusik, and M. Tegmark, "KAN 2.0: Kolmogorov-Arnold Networks meet science," *arXiv preprint*, arXiv:2408.10205, Aug. 2024. [Online]. Available: <https://doi.org/10.48550/arXiv.2408.10205>

$$\phi(x) = w_b b(x) + w_s \text{spline}(x).$$

$$b(x) = \text{silu}(x) = x/(1 + e^{-x})$$

$$\text{spline}(x) = \sum_i c_i B_i(x)$$

1. Discovering conserved quantities

- Conserved quantities as differential equation solving
- $z \in R^d$ governed by the equation $dz/dt = f(z)$
- $H(z)$ is a conserved quantity if $f(z) \cdot \nabla H(z) = 0$ for all Z

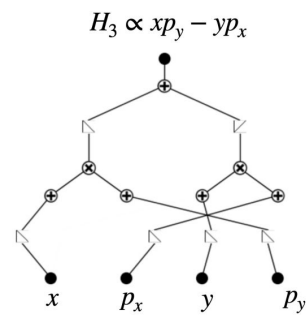
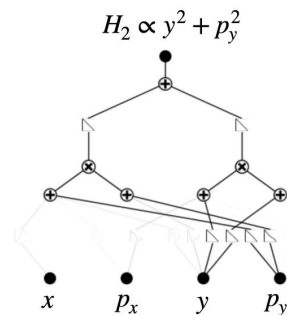
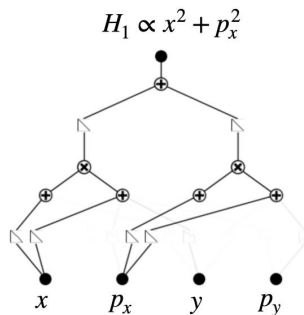
Harmonic Oscillator

- System Description:
 - Phase space: $z = (x, p)$
 - x : position
 - p : momentum
- Evolution:
 - $d(x,p)/dt = (p, -x)$

Harmonic Oscillator

Energy Conservation:

- Energy function: $H = \frac{1}{2}(x^2 + p^2)$
- Phase space: $z = (x, p)$
- Conservation proof:
 - Force vector: $f(z) = (p, -x)$
 - Gradient of H : $\nabla H(z) = (x, p)$
 - Dot product: $f(z) \cdot \nabla H(z) = p \cdot x + (-x) \cdot p = 0$
 - Zero result proves H is conserved



KAN:

- Parametrize H into network
- Train with loss function

$$\ell = \sum_{i=1}^N \left| \mathbf{f}(\mathbf{z}^{(i)}) \cdot \hat{\nabla} H(\mathbf{z}^{(i)}) \right|^2$$