

Repurposing Large Language Models for Cosmology

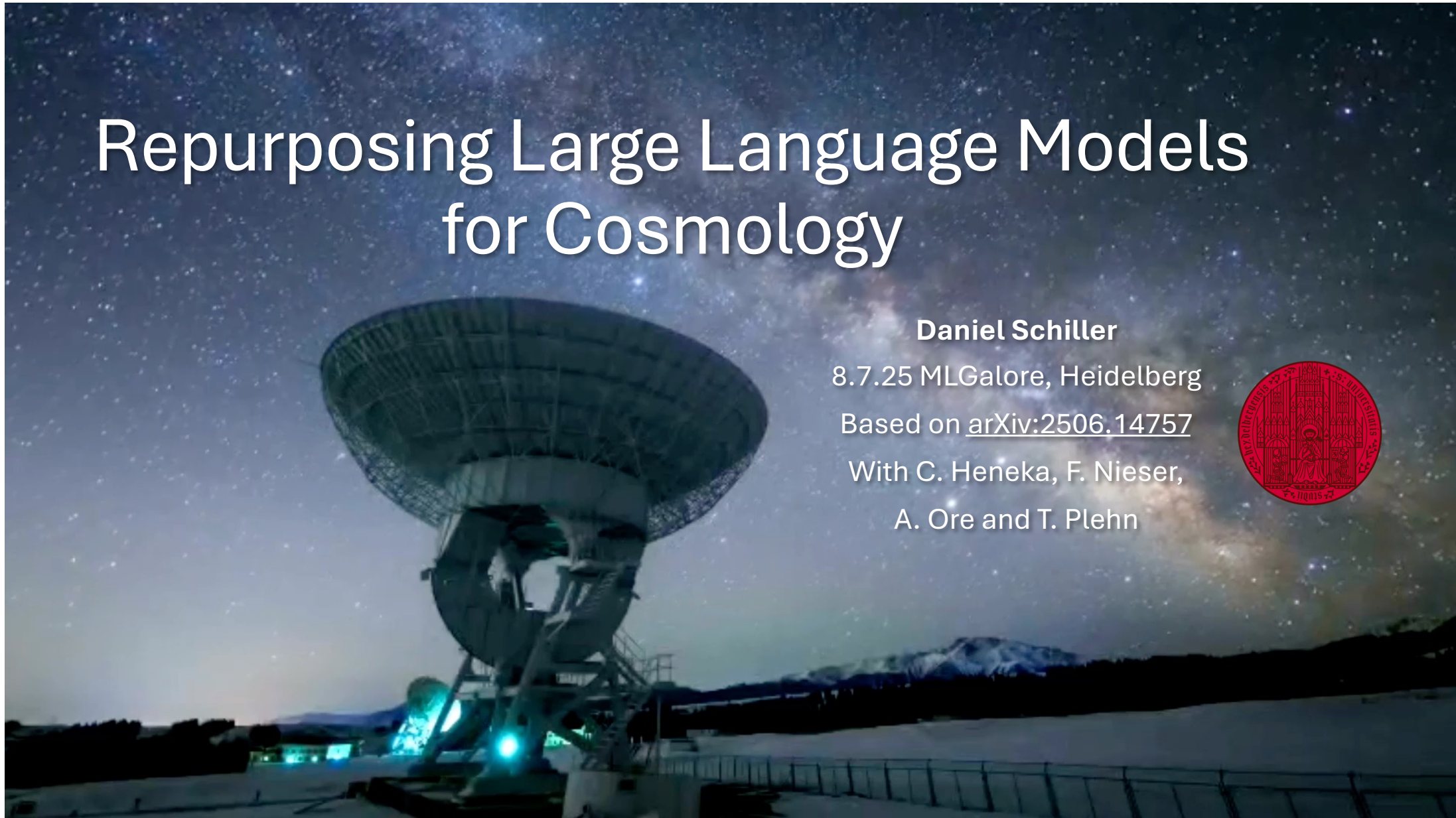
Daniel Schiller

8.7.25 MLGalore, Heidelberg

Based on [arXiv:2506.14757](https://arxiv.org/abs/2506.14757)

With C. Heneka, F. Nieser,

A. Ore and T. Plehn



Foundation Models for Physics Data

- **Foundation Model:** long pretraining + **short** finetuning

Foundation Models for Physics Data

- **Foundation Model:** long pretraining + **short** finetuning
- **Pretraining:**
 - Pretraining task on a **large dataset**
 - Learn **self-supervised** a good **data representations**, e.g. symmetry-aware

Foundation Models for Physics Data

- **Foundation Model:** long pretraining + **short** finetuning
- **Pretraining:**
 - Pretraining task on a **large dataset**
 - Learn **self-supervised** a good **data representations**, e.g. symmetry-aware
- **Finetuning:**
 - Train on downstream task
 - Especially useful for **small datasets**

LLMs are Foundation Models

LLMs are Foundation Models

- **Pretraining:** complete sentences

LLMs are Foundation Models

- **Pretraining:** complete sentences
- Sentences are split into **tokens** by a predefined algorithm
sentence $\leftrightarrow (t_1, \dots, t_n)$

LLMs are Foundation Models

- **Pretraining:** complete sentences
- Sentences are split into **tokens** by a predefined algorithm
sentence $\leftrightarrow (t_1, \dots, t_n)$
- Example: “Hello Word” $\leftrightarrow$ ([Hello],[World])

LLMs are Foundation Models

- **Pretraining:** complete sentences
- Sentences are split into **tokens** by a predefined algorithm
sentence $\leftrightarrow (t_1, \dots, t_n)$
- Example: “Hello World” $\leftrightarrow$ ([Hello],[World])
- Pretraining = **Next token prediction**
- Example: $p_{\theta}([\text{World}]|[\text{Hello}]) \approx p([\text{World}]|[\text{Hello}])$

How do LLMs come into play?

How do LLMs come into play?

- Network sizes: vs LLMs:
 - Physics Foundation Models: ~**1B** parameters
 - LLMs: ~**100B** parameters
- Language modeling is a hard task

How do LLMs come into play?

- Network sizes: vs LLMs:
 - Physics Foundation Models: ~**1B** parameters
 - LLMs: ~**100B** parameters
- Language modeling is a hard task
- LLM = **highly structured function**

How do LLMs come into play?

- Network sizes: vs LLMs:
 - Physics Foundation Models: ~**1B** parameters
 - LLMs: ~**100B** parameters
- Language modeling is a hard task
- LLM = **highly structured function**
- **Idea: Repurpose this structure**

LLM architecture

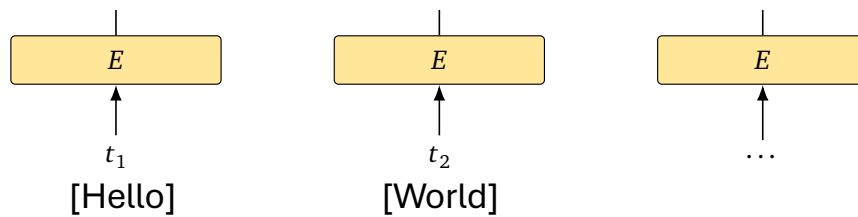
LLM architecture

t_1
[Hello]

t_2
[World]

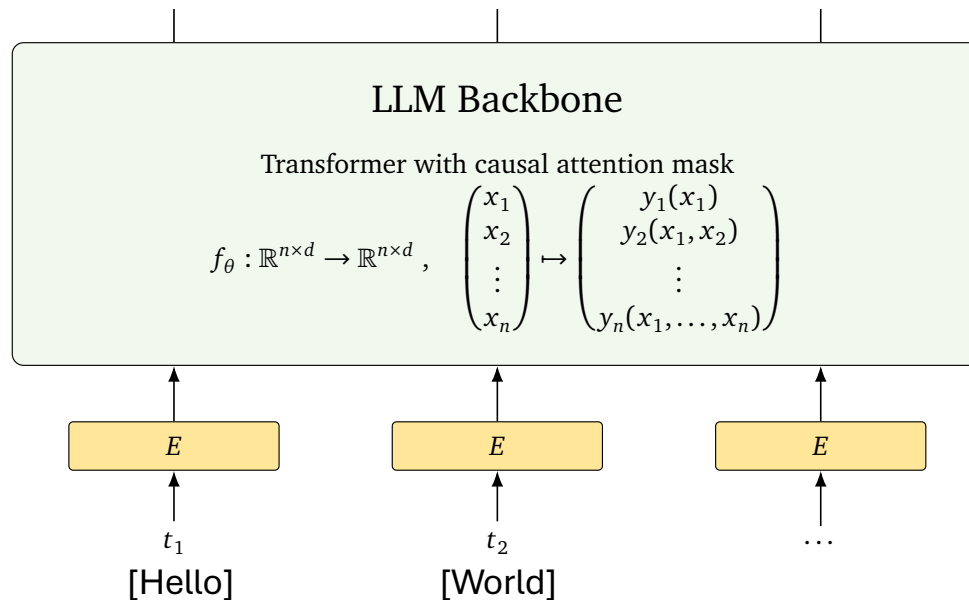
...

LLM architecture



Embedding:
Maps discrete tokens to a latent representation

LLM architecture



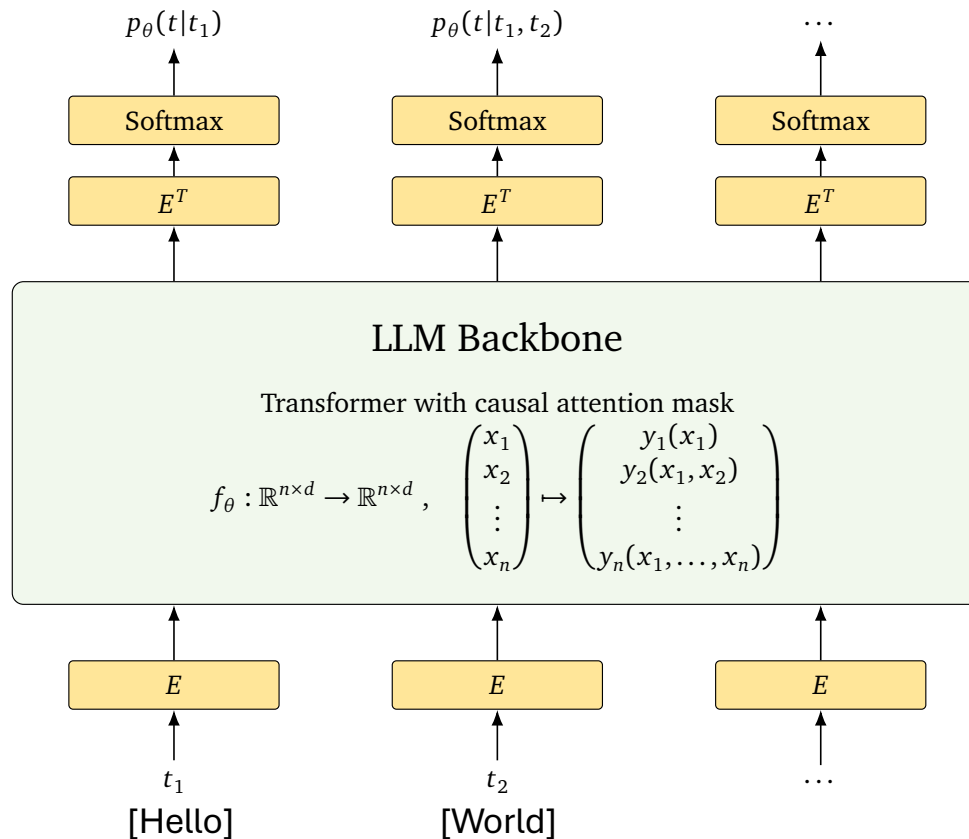
LLM backbone:

Learns causal correlations between representations

Embedding:

Maps discrete tokens to a latent representation

LLM architecture



Unembedding & Softmax:

Maps representations to a categorical distribution

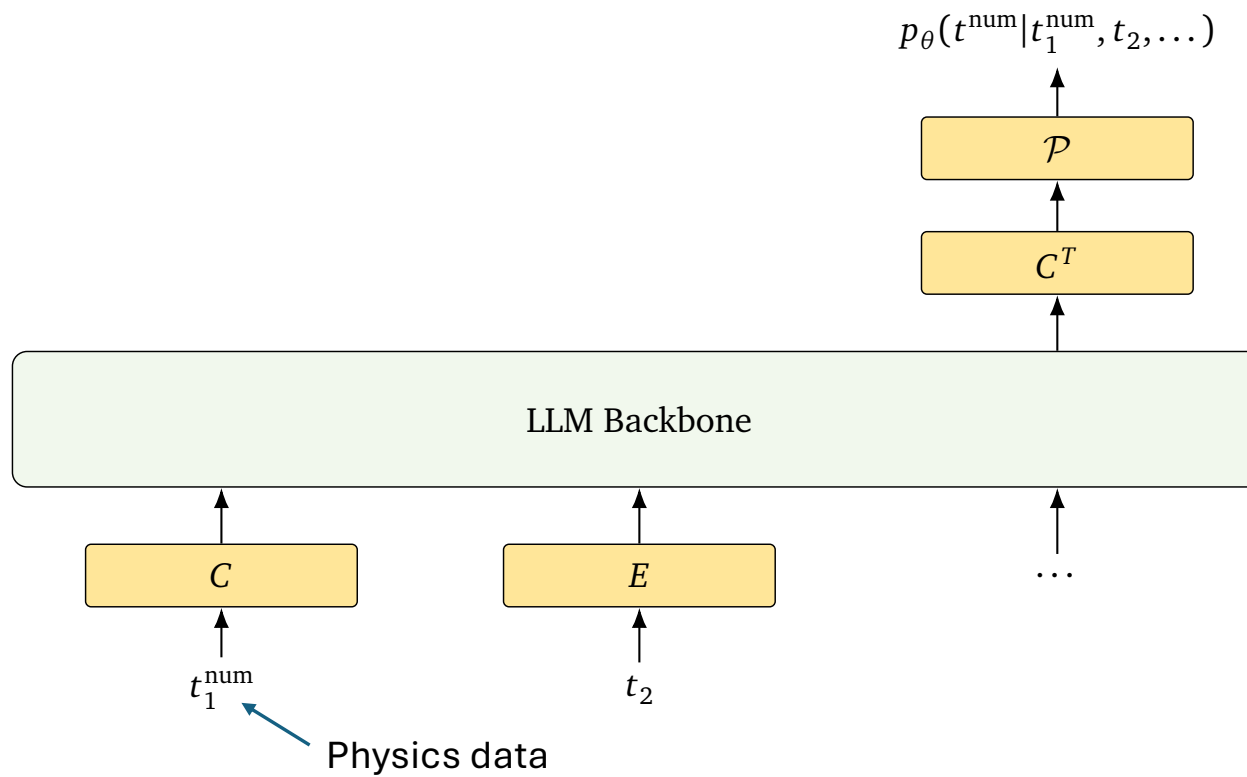
LLM backbone:

Learns causal correlations between representations

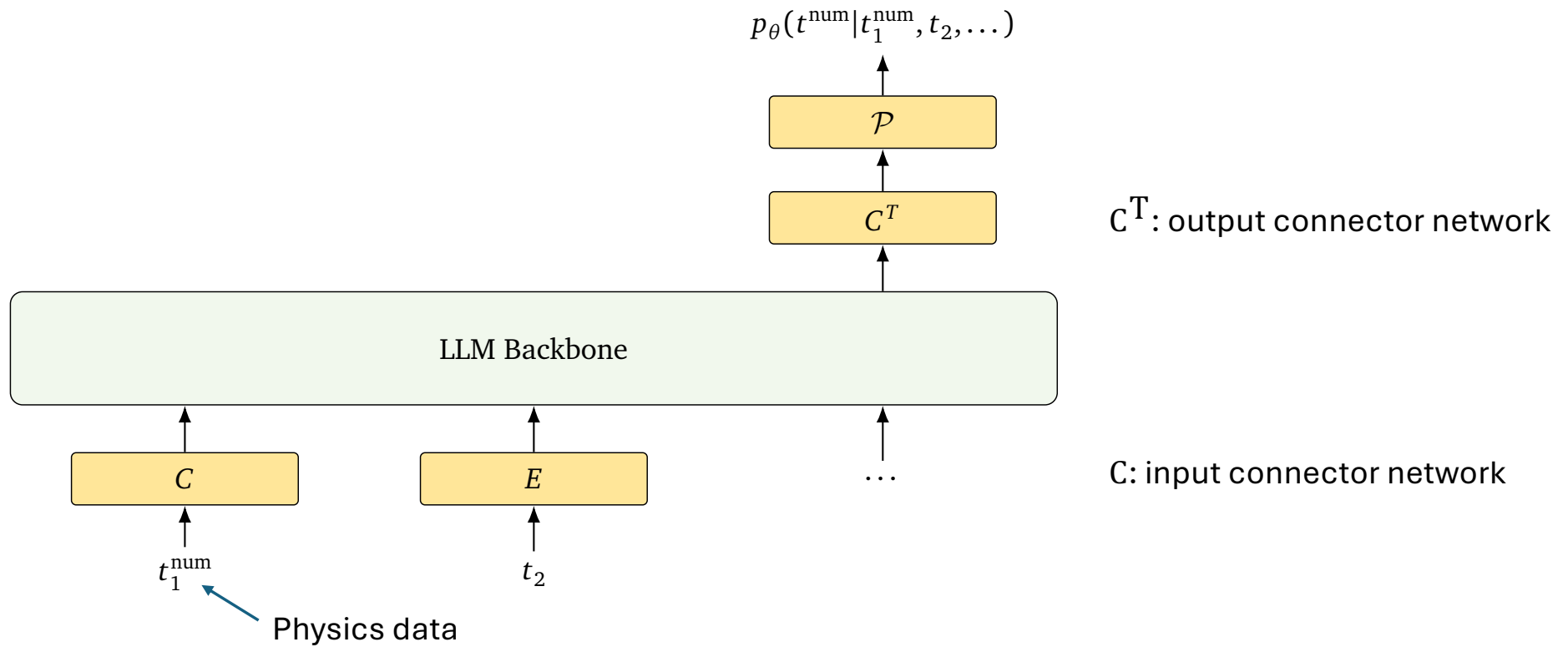
Embedding:

Maps discrete tokens to a latent representation

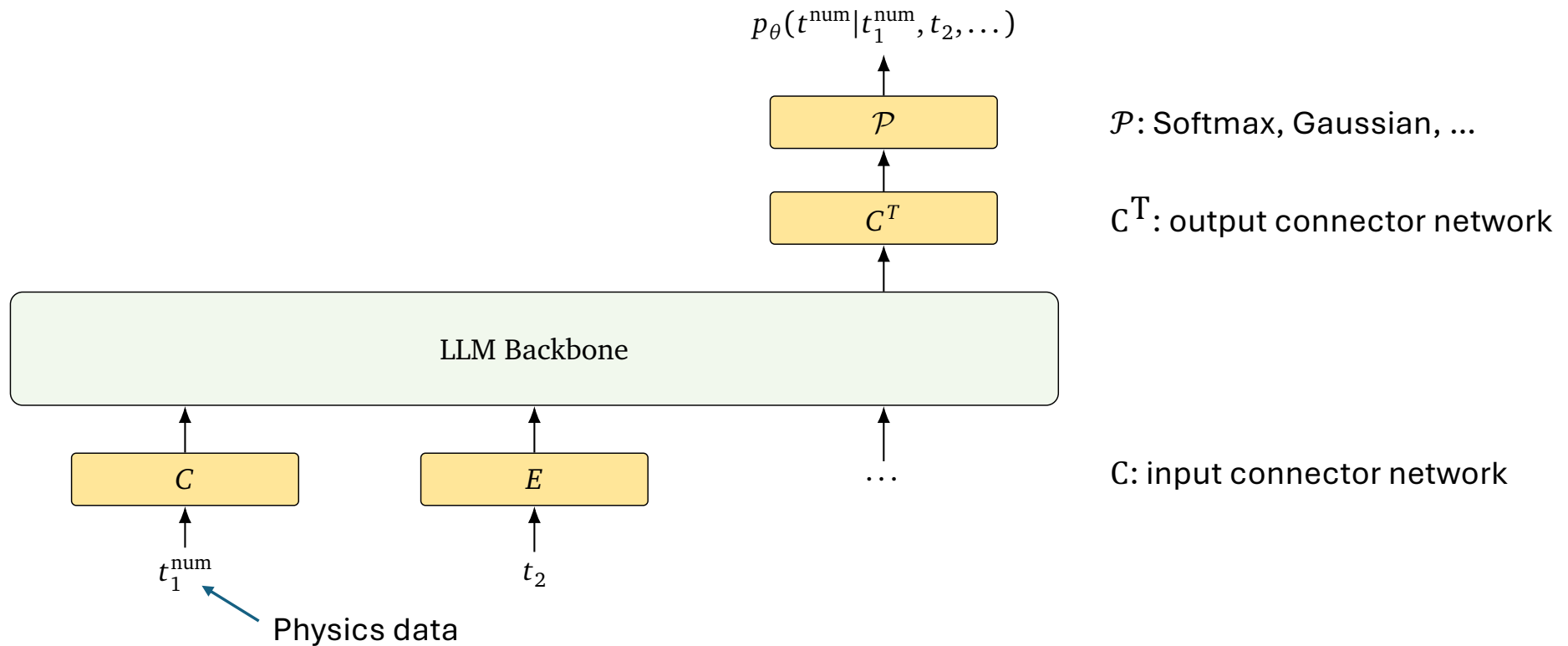
Ansatz



Ansatz

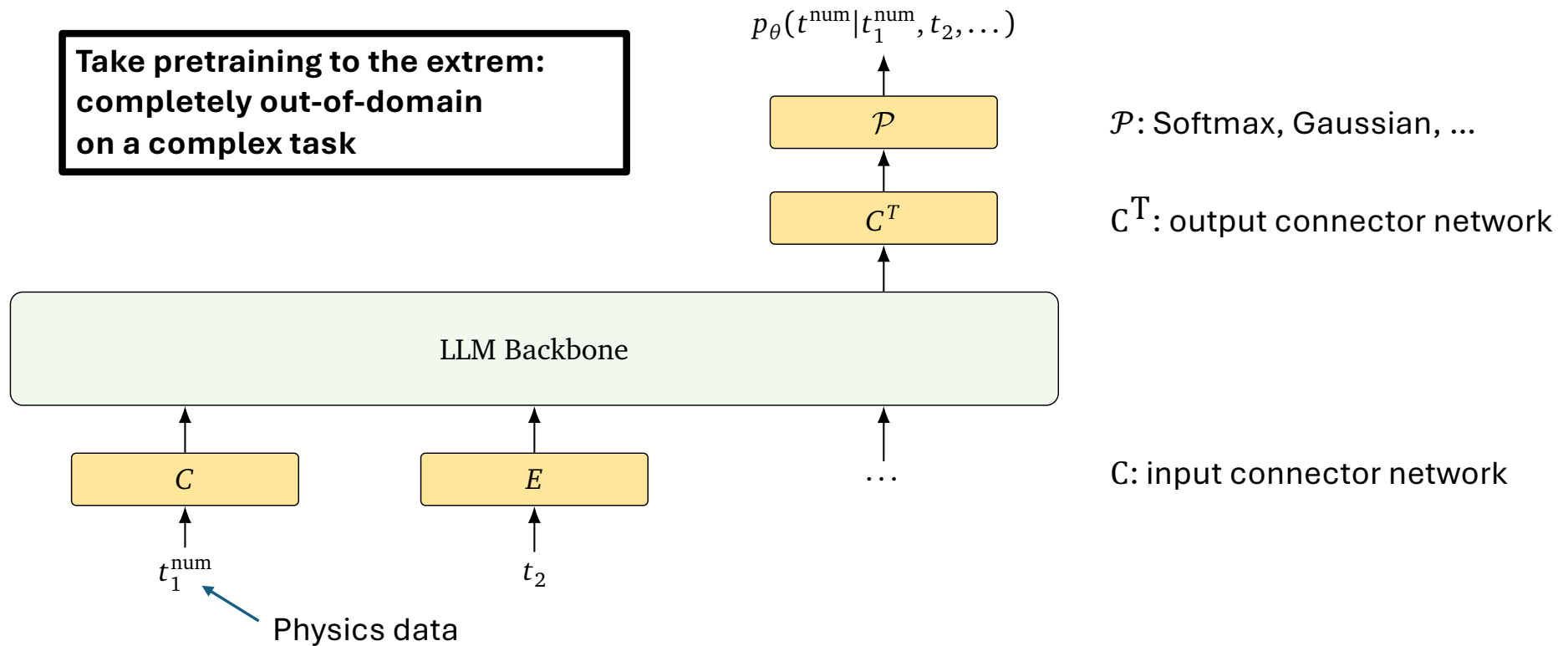


Ansatz



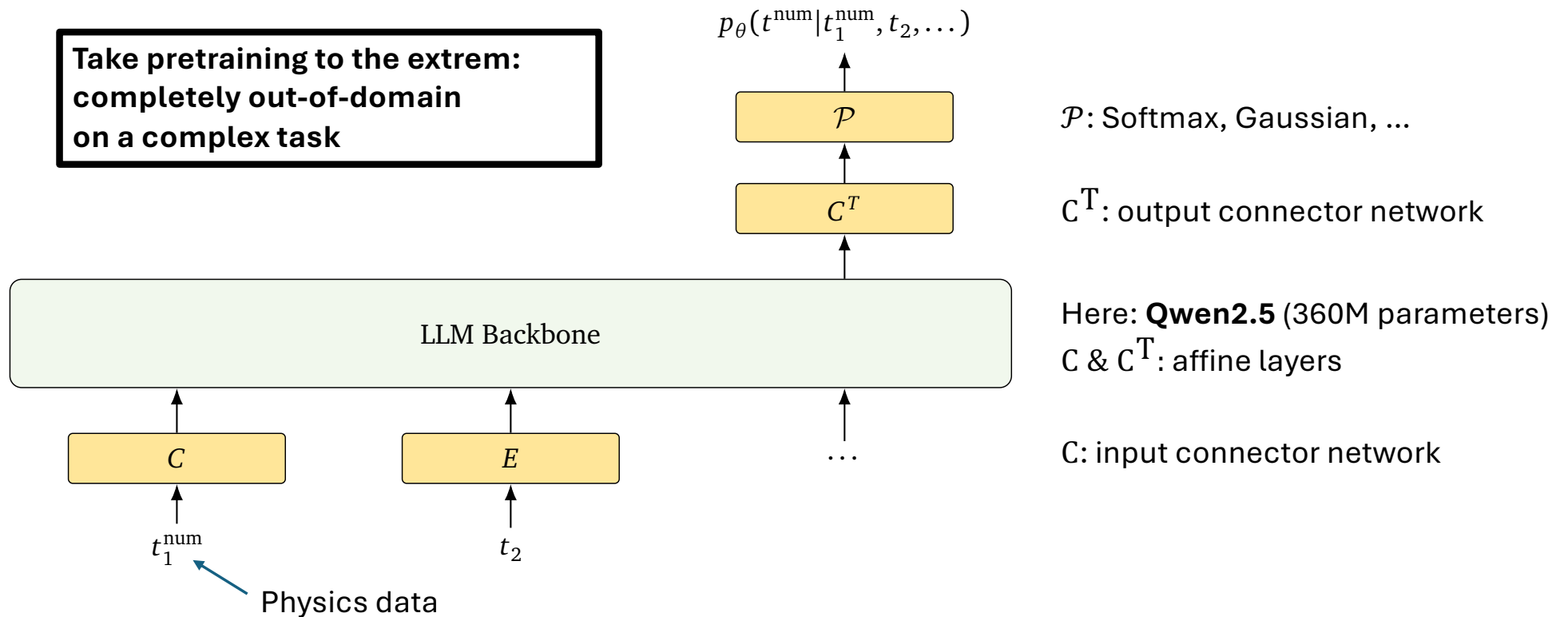
Ansatz

**Take pretraining to the extrem:
completely out-of-domain
on a complex task**



Ansatz

Take pretraining to the extrem:
completely out-of-domain
on a complex task

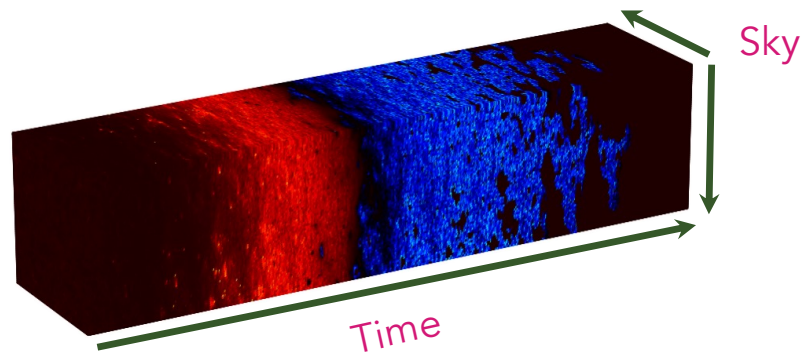


Lightcone Large Language Model (L3M)

- Data: simulated lightcones

Lightcone Large Language Model (L3M)

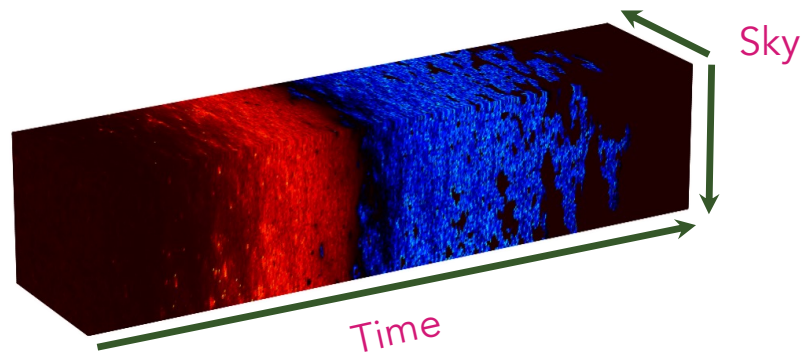
- Data: simulated lightcones



Intensity of Hydrogen Spin Flip
Emission/Absorption relative to CMB
Brightness Temperature δT_{21}

Lightcone Large Language Model (L3M)

- Data: simulated lightcones



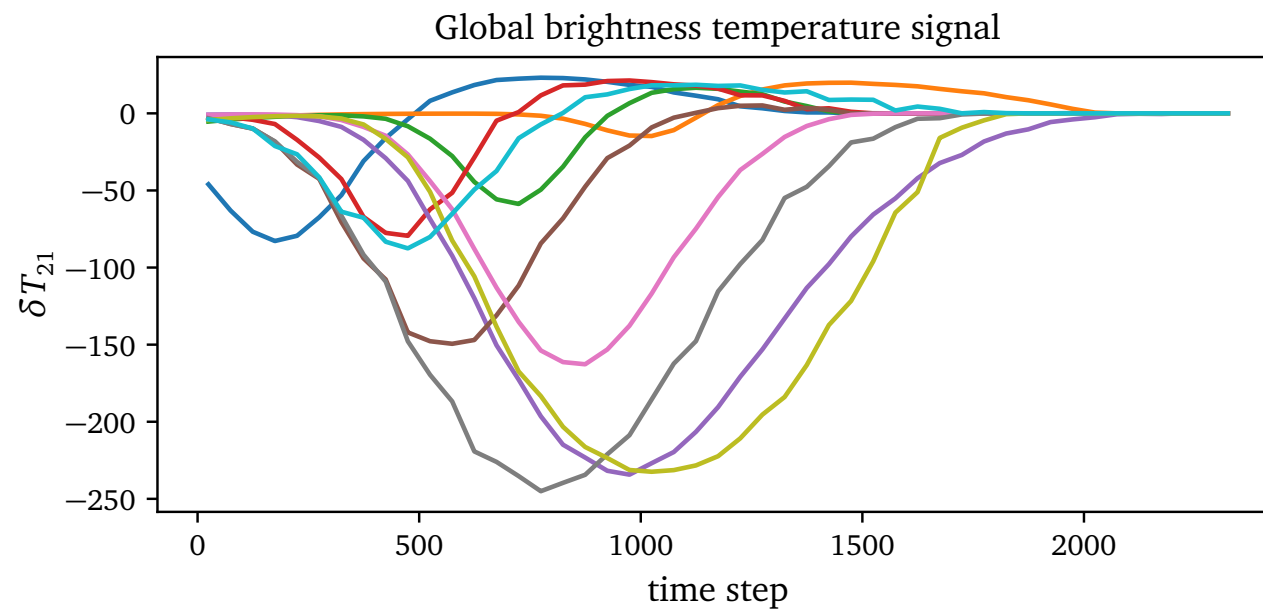
Intensity of Hydrogen Spin Flip
Emission/Absorption relative to CMB
Brightness Temperature δT_{21}

- Depend on 6 parameters and initial conditions

Regression with frozen backbone

Regression with frozen backbone

- Regress 6 simulation parameters from global δT_{21} signal



Regression with frozen backbone

**Minimal
prompting:**

$t_1^{\text{BT}} \dots t_n^{\text{BT}} <|\text{param}|>$

Global signal at
time step 1

Added token
to output
prediction

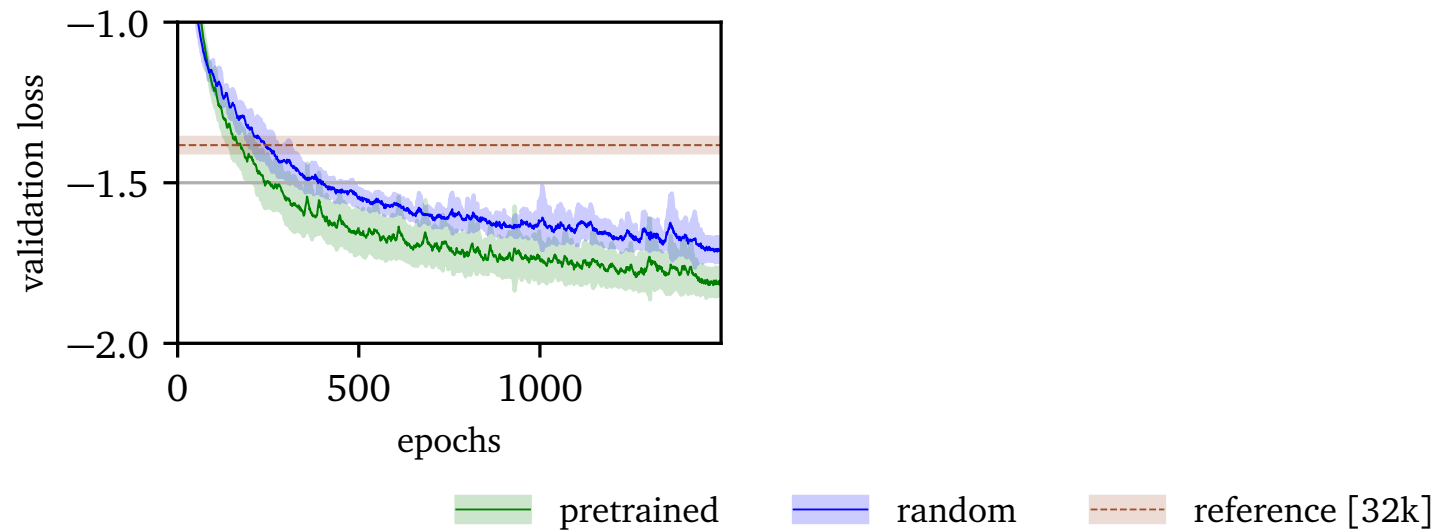
Regression with frozen backbone

Minimal prompting:

$t_1^{BT} \dots t_n^{BT} \langle \text{param} \rangle$

Global signal at
time step 1

Added token
to output
prediction



Regression with frozen backbone

Minimal prompting:

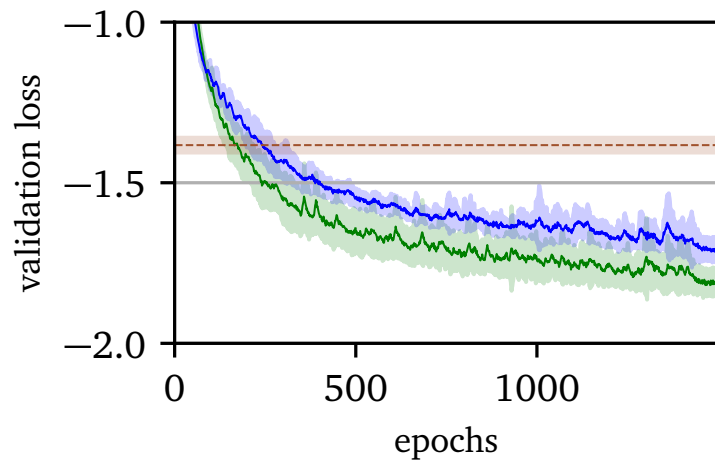
$t_1^{BT} \dots t_n^{BT} \langle \text{param} \rangle$

Global signal at
time step 1

Added token
to output
prediction

Chat inspired prompting:

```
<|start|>system  
<|end|>  
<|start|>user  
 $t_1^{BT} \dots t_n^{BT} \langle \text{end} \rangle$   
<|start|>assistant  
 $\langle \text{param} \rangle \langle \text{end} \rangle$ 
```



— pretrained — random - - - reference [32k]

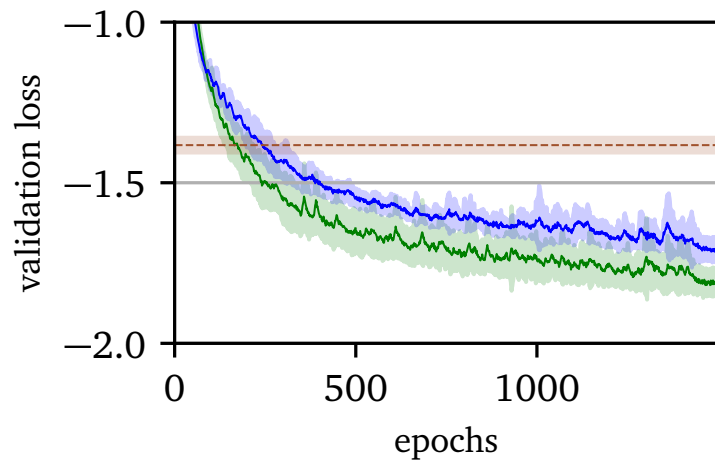
Regression with frozen backbone

Minimal prompting:

$t_1^{BT} \dots t_n^{BT} \langle \text{param} \rangle$

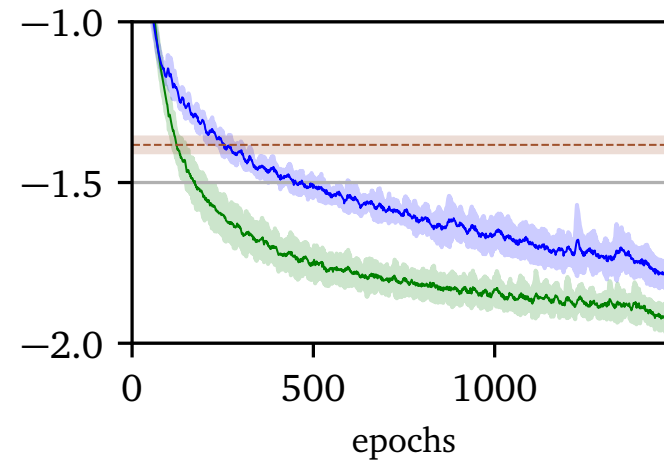
Global signal at
time step 1

Added token
to output
prediction



Chat inspired prompting:

```
<|start|>system
<|end|>
<|start|>user
 $t_1^{BT} \dots t_n^{BT}$  <|end|>
<|start|>assistant
<|param|> <|end|>
```

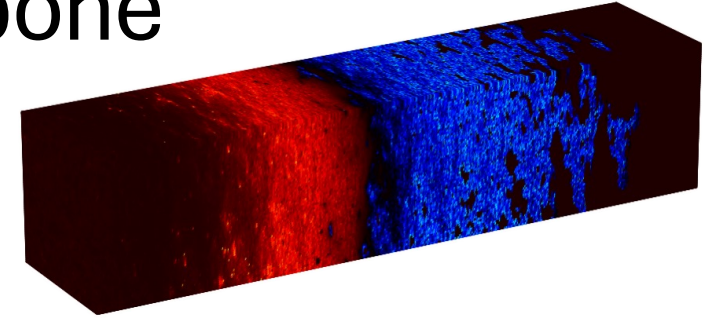
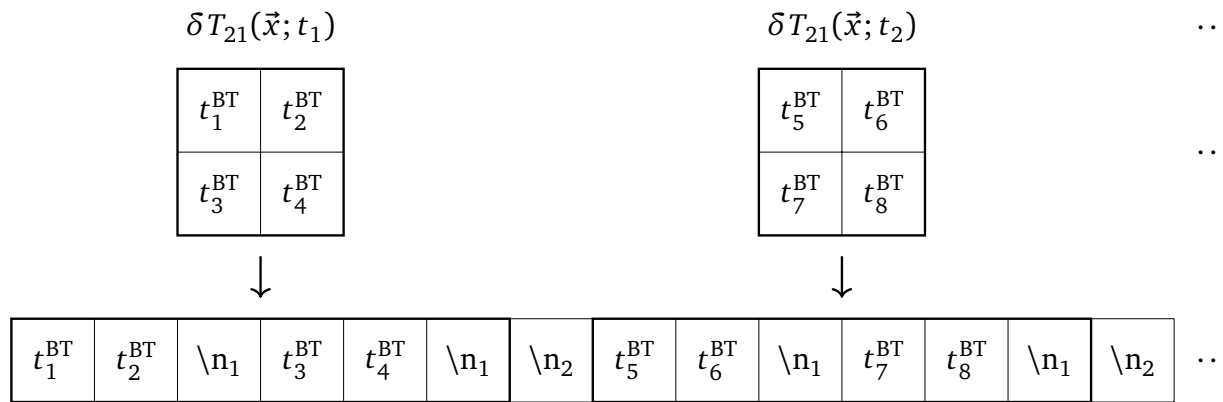


— pretrained — random - - - reference [32k]

Generation with finetuned backbone

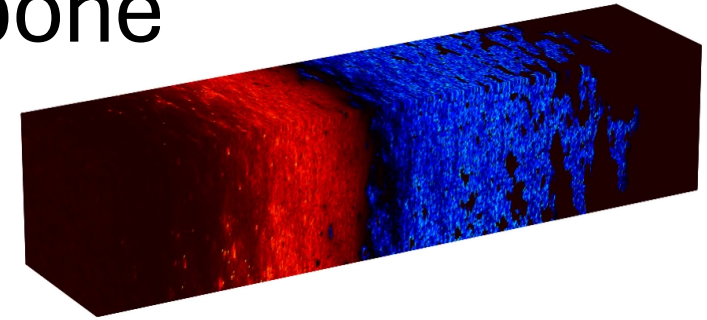
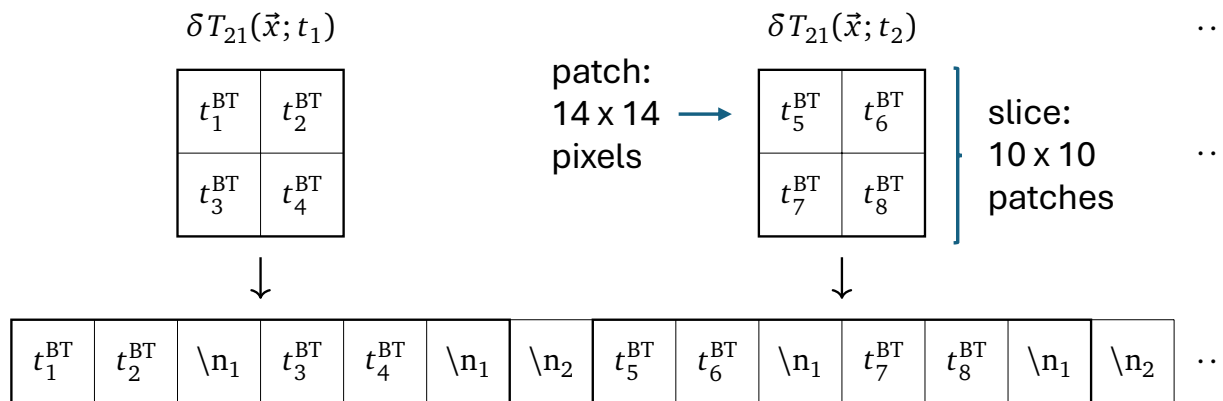
Generation with finetuned backbone

- Represent lightcone as a sequence of tokens



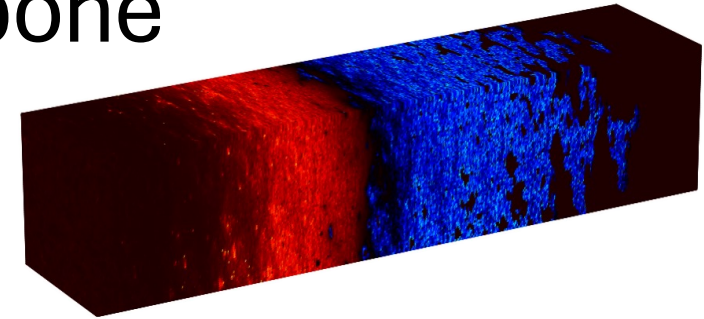
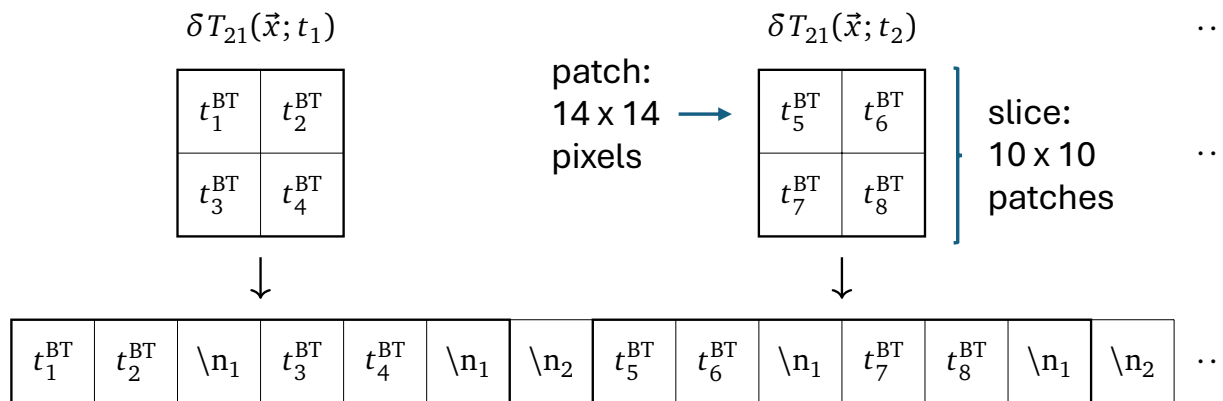
Generation with finetuned backbone

- Represent lightcone as a sequence of tokens



Generation with finetuned backbone

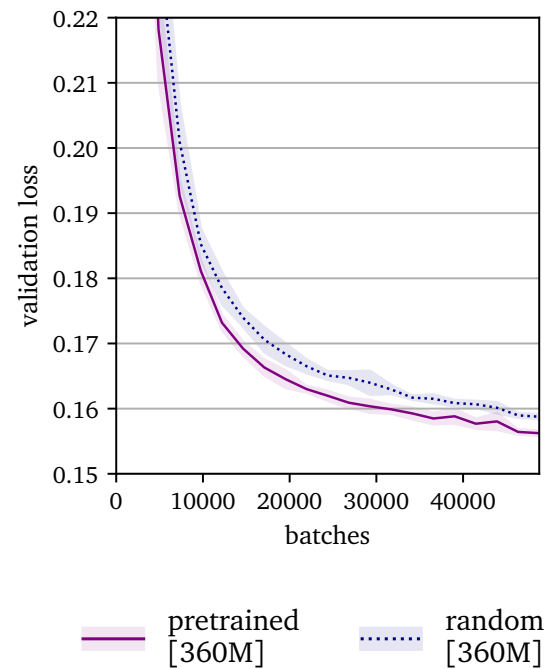
- Represent lightcone as a sequence of tokens



- Next token/patch prediction

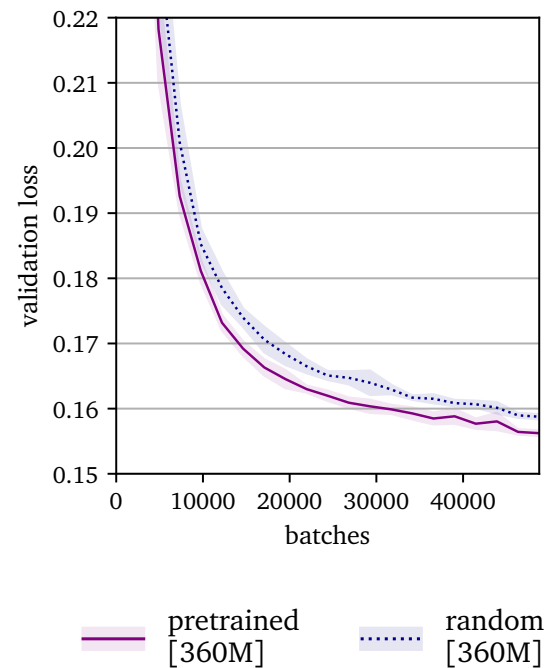
Generation with finetuned backbone

Complete backbone training:



Generation with finetuned backbone

Complete backbone training:



Partial backbone training with LoRA:

Modify linear layers by a small rank:

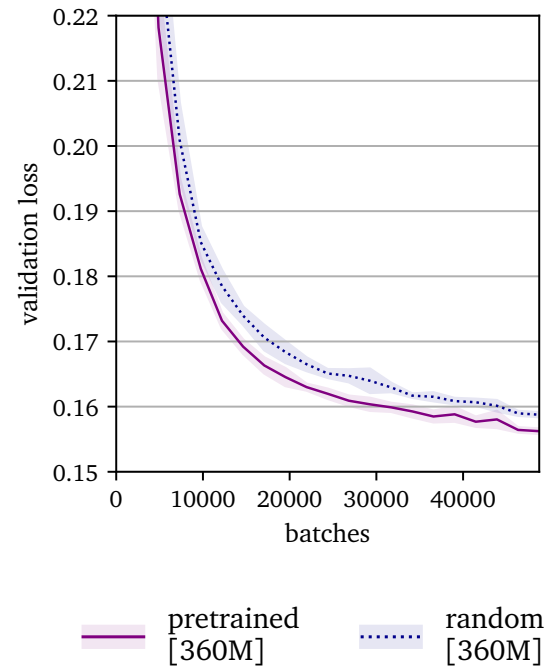
$$M \in \mathbb{R}^{a \times b} \rightarrow M + \delta M$$

frozen trained

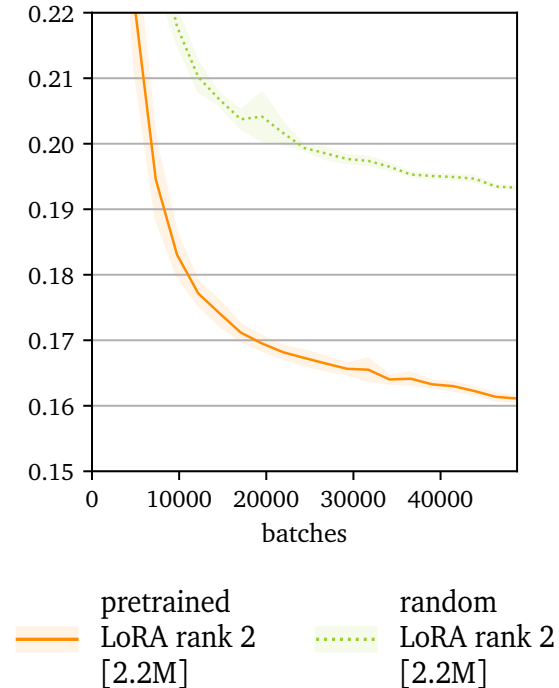
$$\delta M = M_1 M_2$$
$$M_1 \in \mathbb{R}^{a \times r}, M_2 \in \mathbb{R}^{r \times b}$$

Generation with finetuned backbone

Complete backbone training:



Partial backbone training with LoRA:



Modify linear layers by a small rank:

$$M \in \mathbb{R}^{a \times b} \rightarrow M + \delta M$$

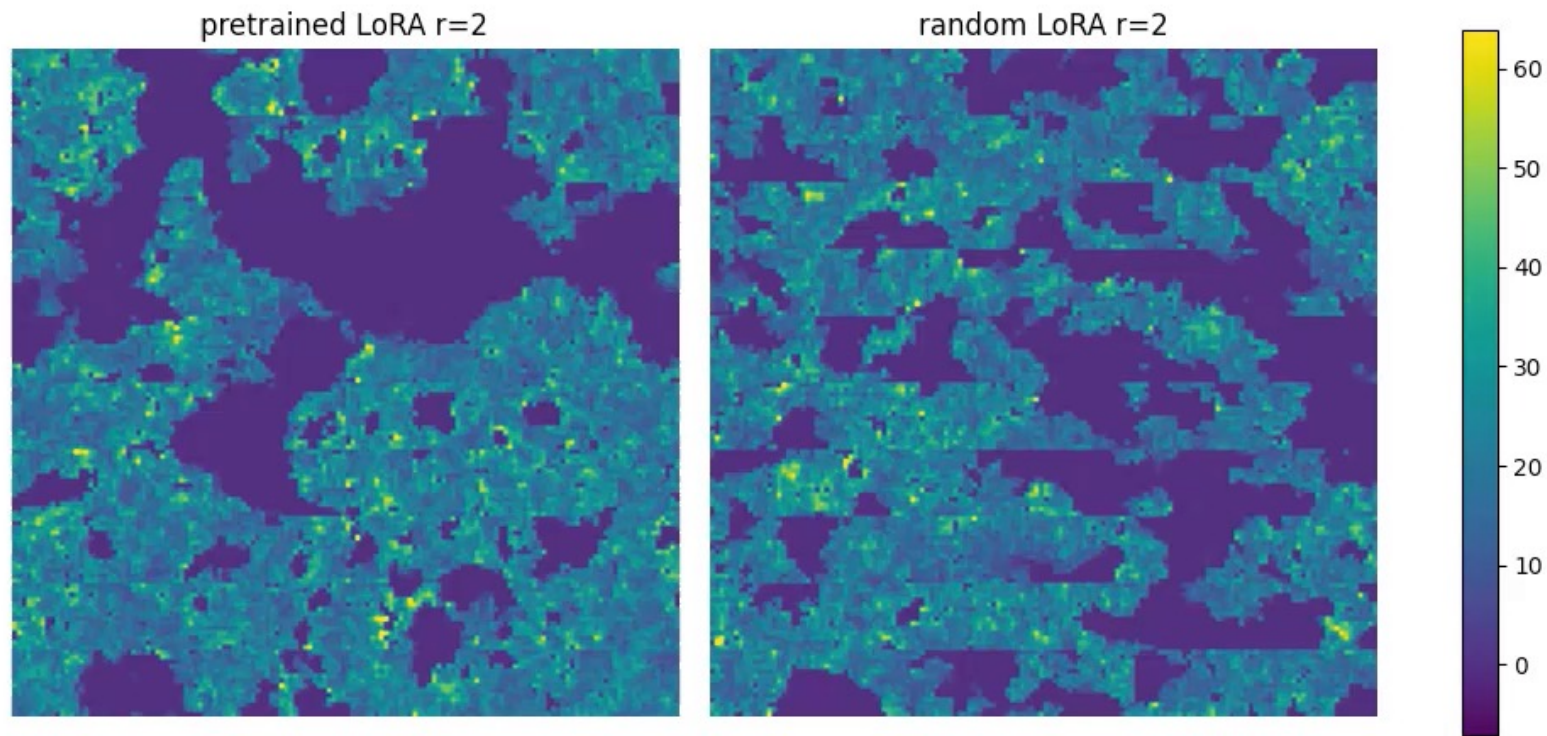
frozen
trained

$$\delta M = M_1 M_2$$

$$M_1 \in \mathbb{R}^{a \times r}, M_2 \in \mathbb{R}^{r \times b}$$

Generation with finetuned backbone

Generation with finetuned backbone

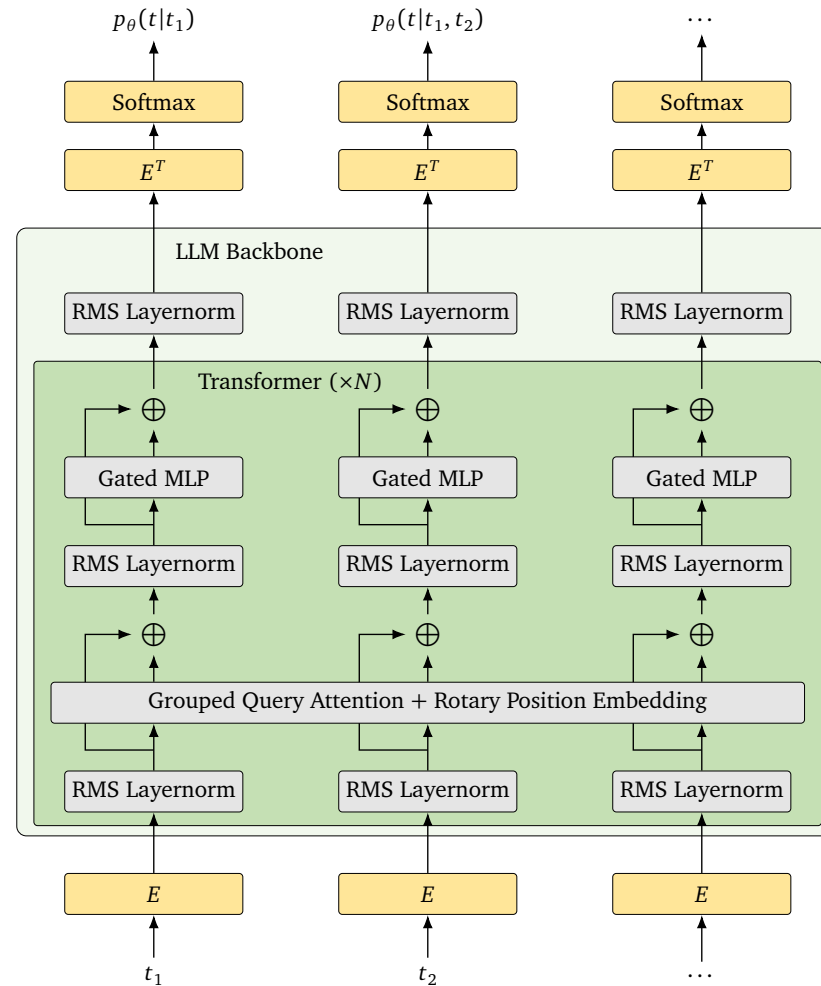


Summary

- Repurposed (open-source) Large Language Models
- Added connector networks and finetuned backbone
- Showed benefit from pretrained structure
- **Novel approach to get large networks for physics**

LLM Architecture

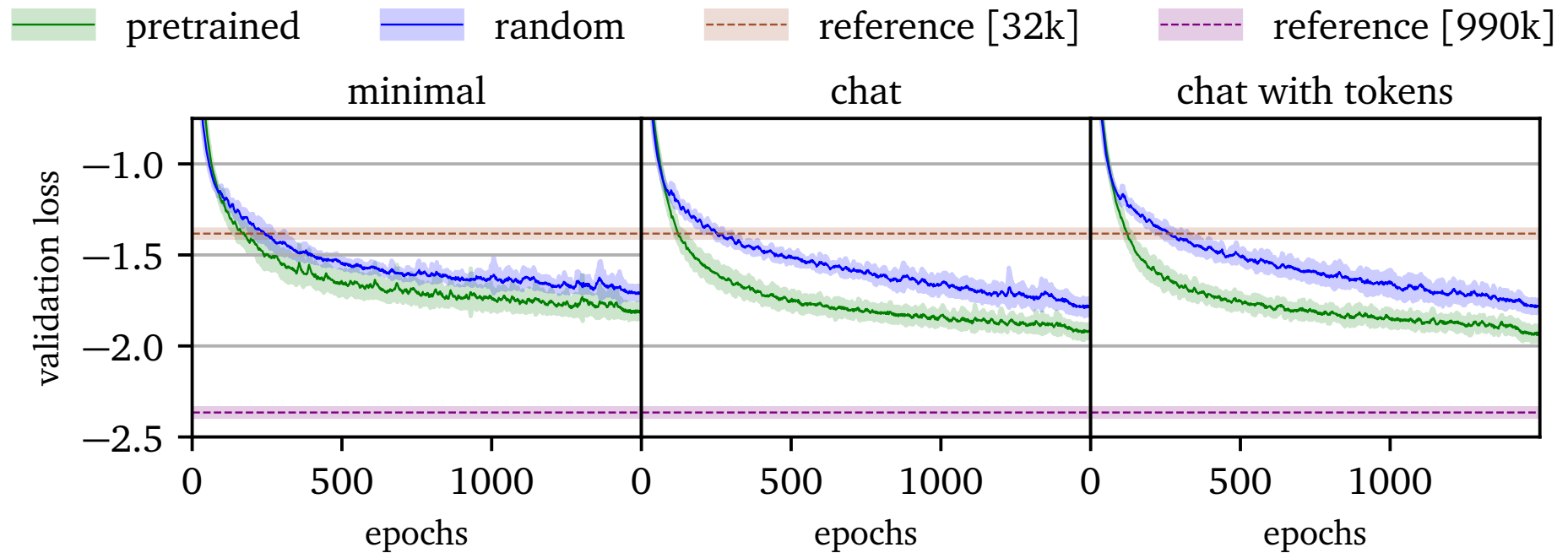
$$p_{\theta}(t_i|t_1, \dots, t_{i-1}) \approx p(t_i|t_1, \dots, t_{i-1})$$



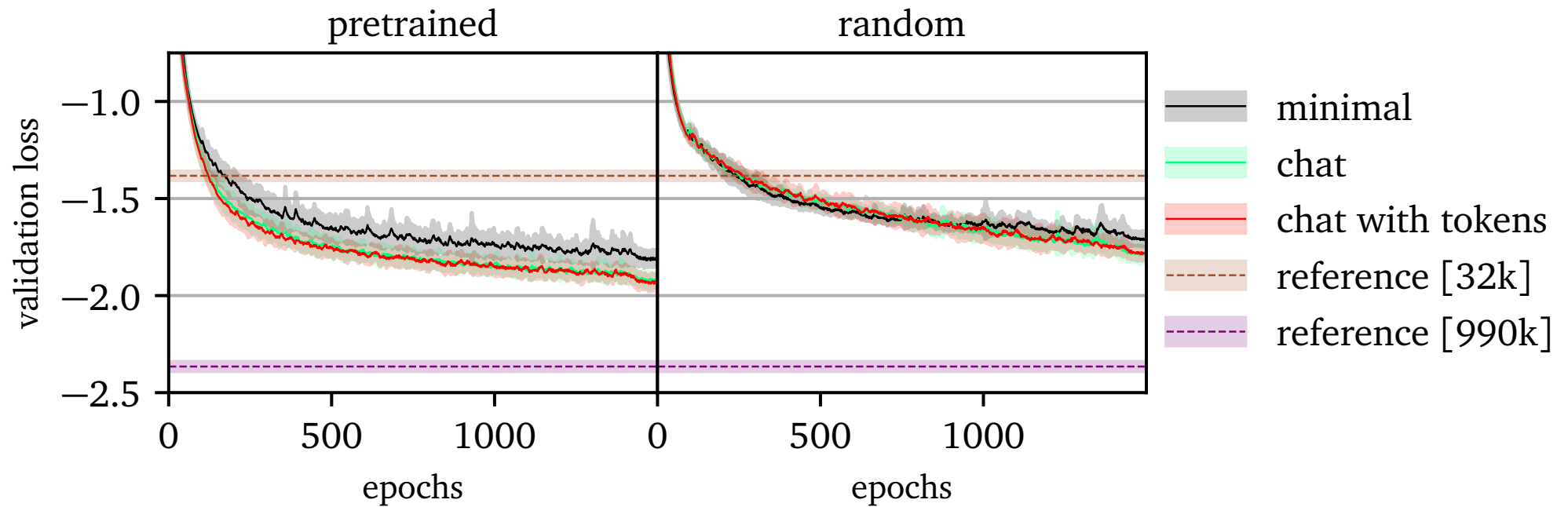
Regression Results

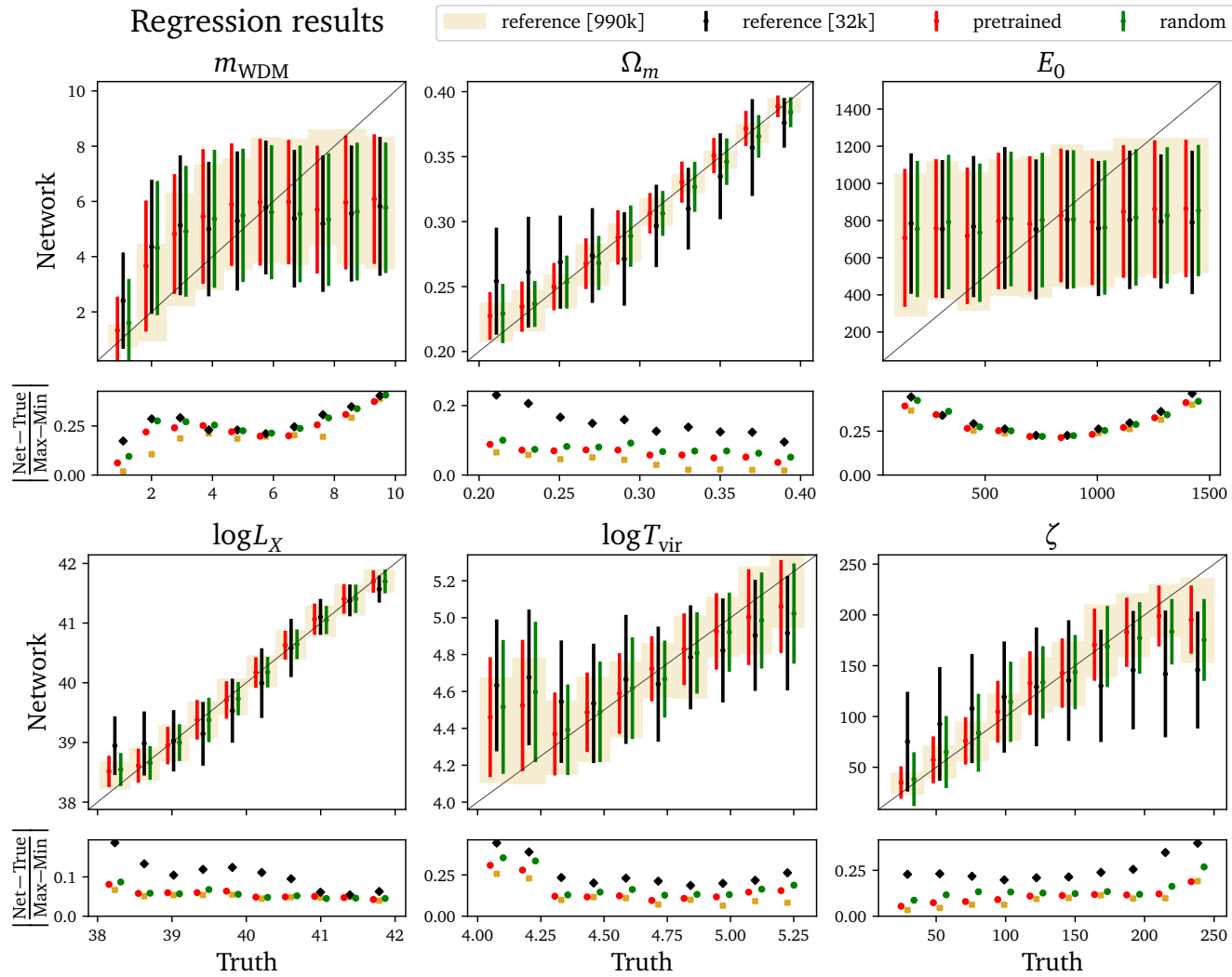
$$\mathcal{P} = \frac{1}{(2\pi)^3 \sqrt{|\Sigma|}} \exp\left(-\frac{1}{2} (\vec{\mu} - \vec{y})^T \Sigma^{-1} (\vec{\mu} - \vec{y})\right)$$

$\vec{\mu}$: mean values, Σ : covariance matrix



Regression Results





Conditional Flow Matching

- Distribution Mapping via velocity field (ODE)

$$\frac{dx}{d\tau} = v(x, \tau, z) \quad \text{with} \quad x(\tau) \sim \begin{cases} \mathcal{N}(\mu = 0, \sigma = 1) & \tau = 0 \\ p(t_i^{\text{BT}} | t_{i-1}^{\text{BT}}, \dots) & \tau = 1 \end{cases}$$

- Network learns velocity field

$$\mathcal{L} = \left\langle \left(v - v_\theta(x_\tau, \tau, z) \right)^2 \right\rangle \quad \text{with}$$
$$v \equiv t_i^{\text{BT}} - x_0, \quad x_\tau \equiv (1 - \tau)x_0 + \tau t_i^{\text{BT}}, \quad x_0 \sim \mathcal{N}(\mu = 0, \sigma = 1)$$

- Trained a network with 500k parameters

Generation Results

