

Optimal strategies in navigation and learning: statistical physics meets control theory

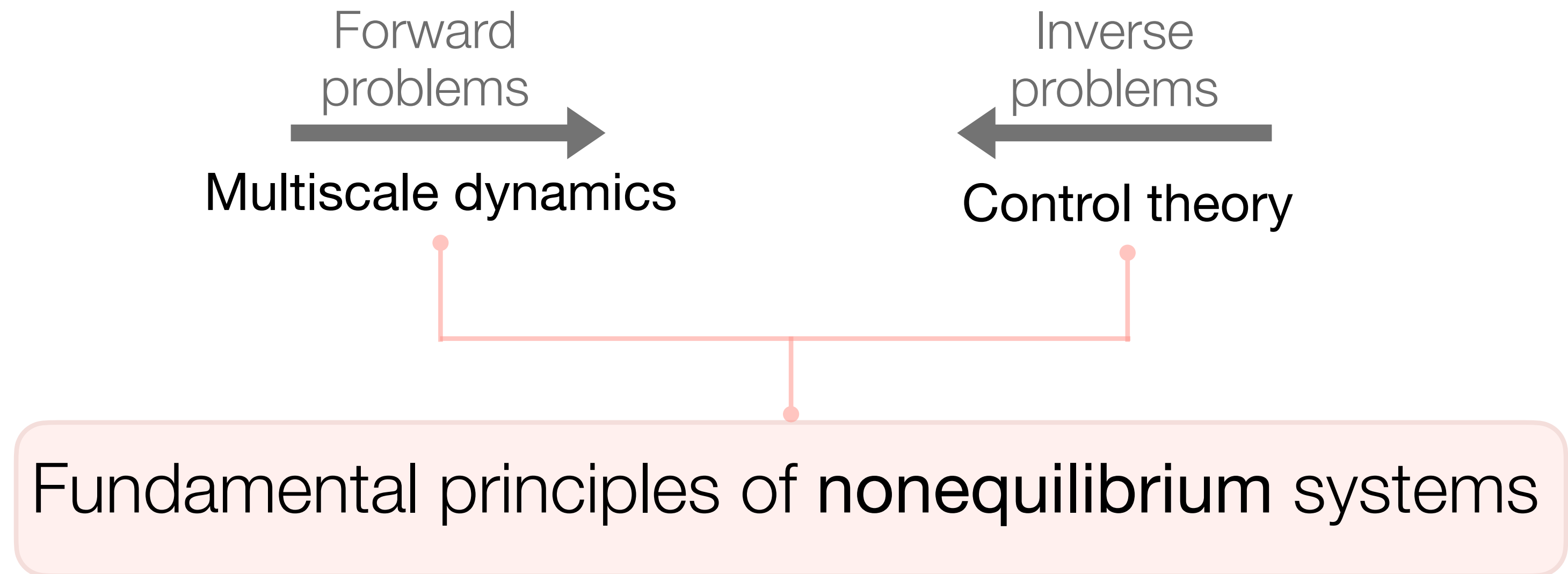
Francesco Mori



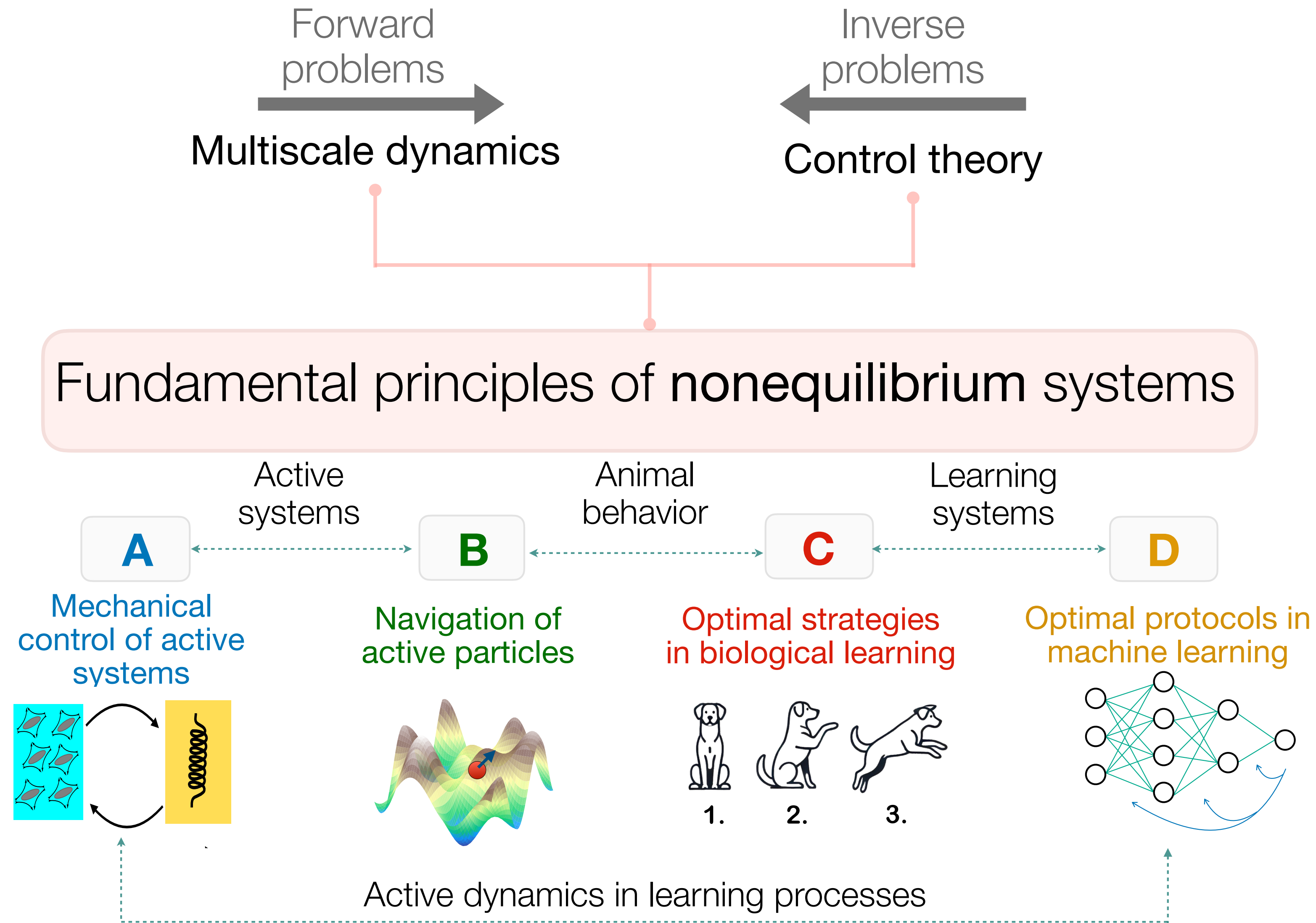
LPTMC Seminar

21st Nov 2024

Overview of my research



Overview of my research



Part I:

Animal navigation

“Optimal switching strategies for navigation in stochastic settings.”

FM, L. Mahadevan.

arXiv preprint *arXiv:2311.18813*



L. Mahadevan
(Harvard)

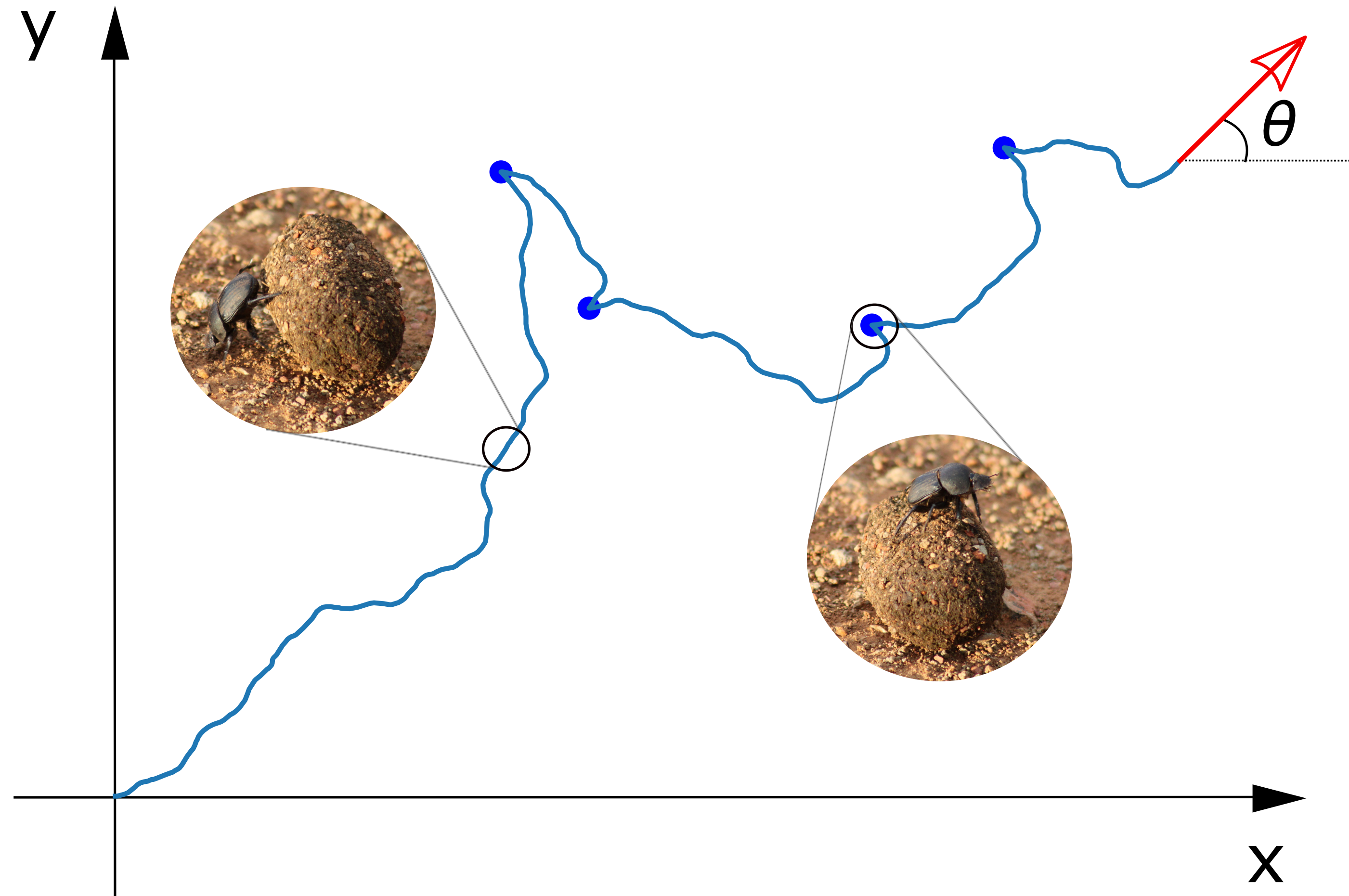
Navigation in noisy environments



Credit: MdeVicente CC0 1.0 Universal

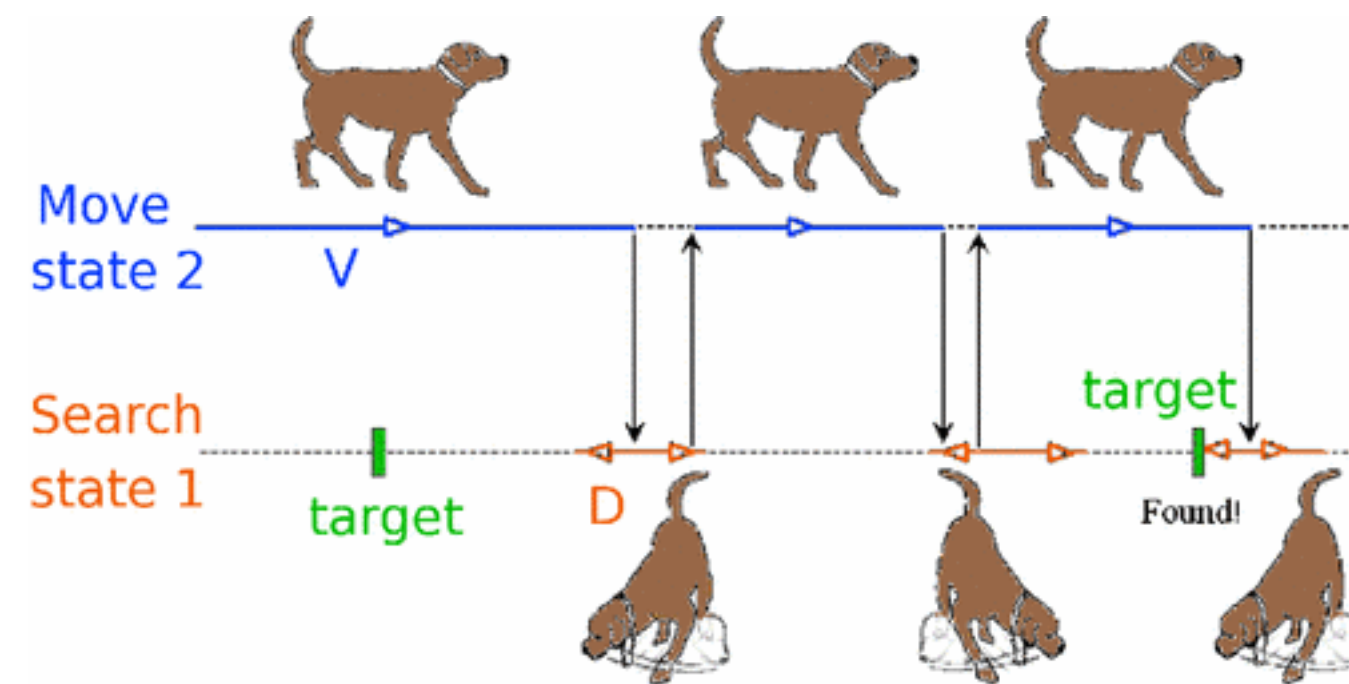
Navigation in noisy environments

Synthetic data



Switching behaviour in navigation

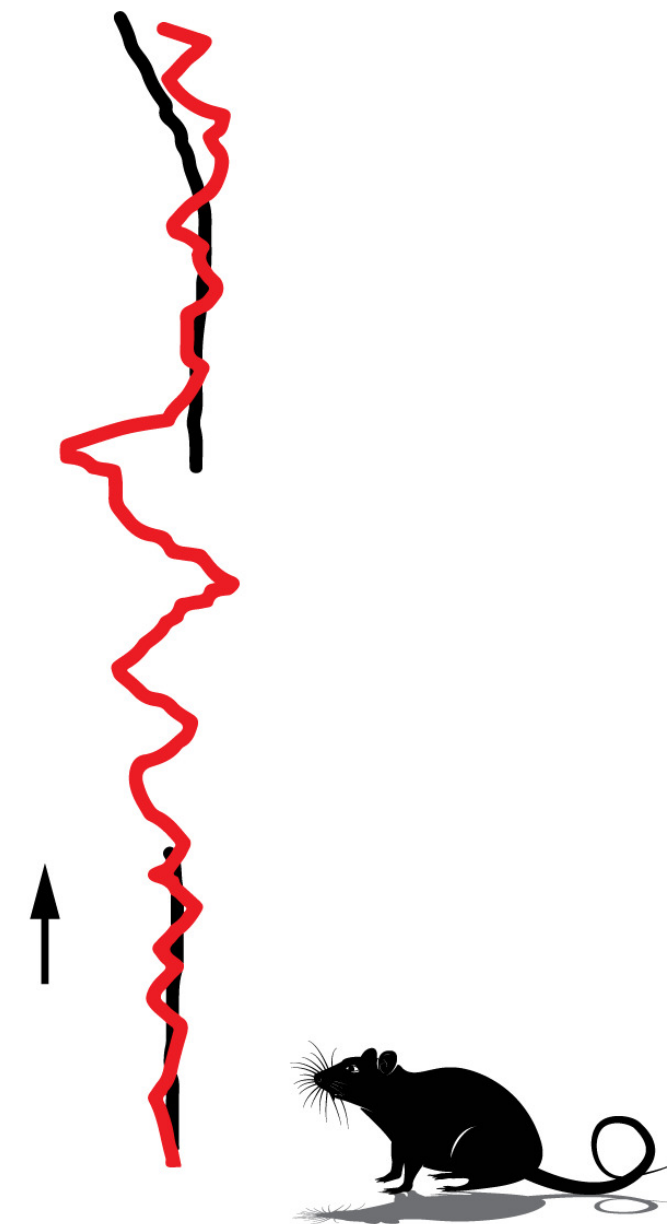
Intermittent search strategies



Thesen, Aud, Johan B. Steen, and Kjell B. Døving.
Journal of Experimental Biology 180.1 (1993):
247-251.

O. Bénichou, C. Loverdo, M. Moreau, and R.
Voituriez, *Rev. Mod. Phys.* **83**, 81

Trailing



A. G. Khan, M. Sarangi, and U. S. Bhalla, *Nature communications* **3**, 703 (2012).

G. Reddy, B. I. Shraiman, and M. Vergassola,
PNAS 119(1), e2107431118 (2022).

Zigzagging vs casting



E. Balkovsky, and B. I. Shraiman.
PNAS 99, 12589 (2002).

L.L. López, et al. IntechOpen, 2011.

G. Reddy, V. N. Murthy, and M.
Vergassola. *Ann. Rev. Cond. Matt. Phys.* 13 (2022): 191-213.

Egocentric vs geocentric strategies



Desert ant



Human with a
map

Egocentric



Geocentric



How to optimally switch?

Active Brownian motion

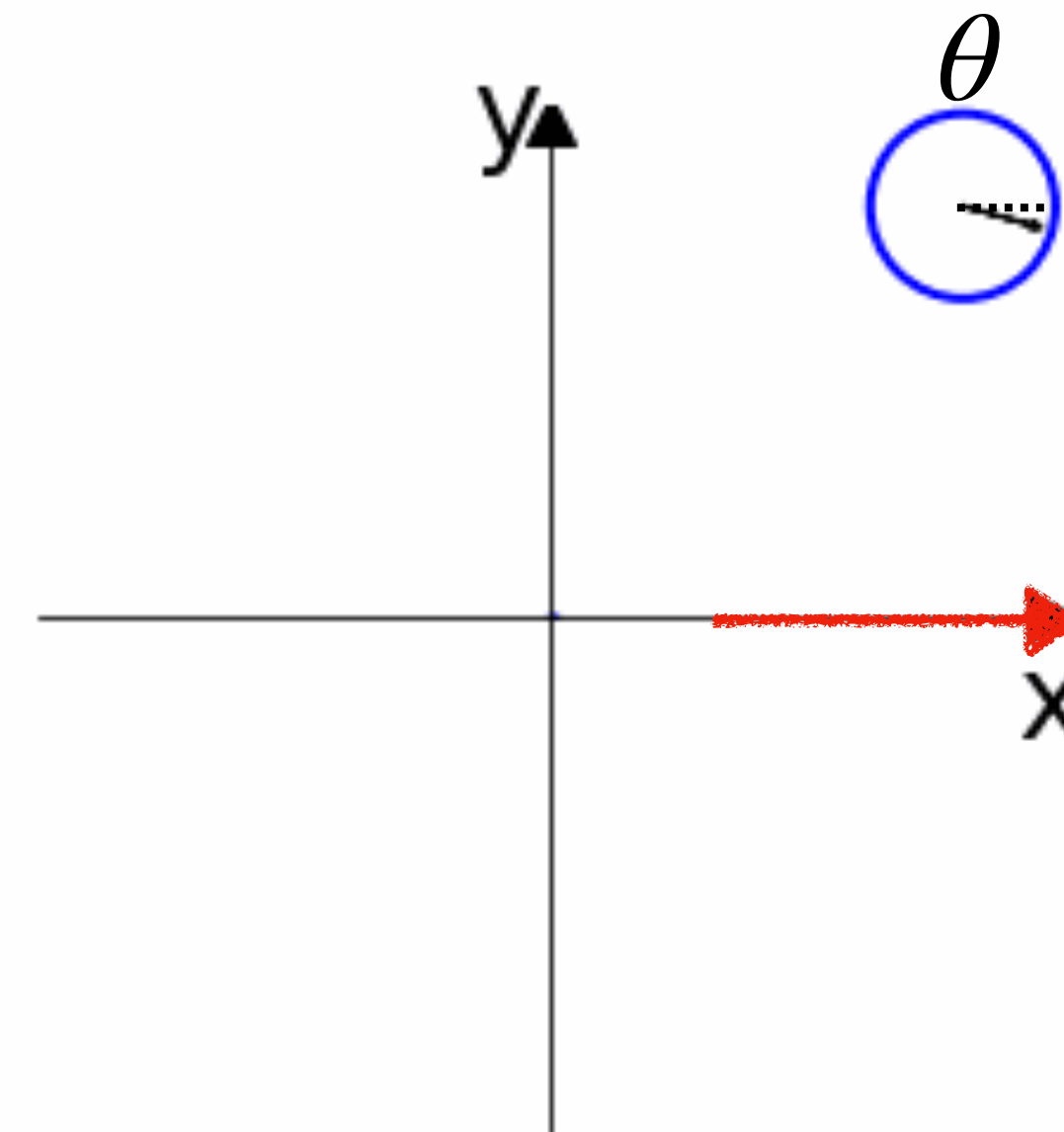
$$\begin{cases} \dot{x}(t) = v_0 \cos[\theta(t)], \\ \dot{y}(t) = v_0 \sin[\theta(t)], \\ \dot{\theta}(t) = \eta(t) \end{cases}$$



$$\begin{aligned} \langle \eta(t) \rangle &= 0 \\ \langle \eta(t) \eta(t') \rangle &= 2D \delta(t - t') \end{aligned}$$

GOAL

Maximise the displacement in the x direction (keep θ close to zero)



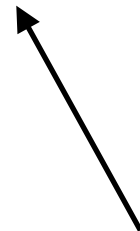
Active Brownian motion with corrections

GOAL

Maximise the displacement in the x direction (keep θ close to zero)

$$\theta(t + dt) = \begin{cases} \theta(t) + \sqrt{2D}dt\eta(t) & \text{if } s(\theta, t) = 0, \\ 0 & \text{if } s(\theta, t) = 1. \end{cases}$$

**Correction
policy**



Active Brownian motion with corrections

**Reward
function**

$$\mathcal{F}_{\theta_0, t_0}[s] = \left\langle \underbrace{v_0 \int_{t_0}^{t_f} \cos(\theta(\tau)) d\tau}_{\text{Reward}} - \underbrace{x_c (N(t_f) - N(t_0))}_{\text{Cost}} \right\rangle_{\theta_0}$$

t_f = Time horizon

$N(t)$ = number of reorientation events up to time t

x_c = Cost per reorientation

**Reward
function**

$$\mathcal{F}_{\theta_0, t_0}[s] = \langle v_0 \int_{t_0}^{t_f} d\tau \cos(\theta(\tau)) - x_c(N(t_f) - N(t_0)) \rangle$$

**Optimal
reward**

$$J(\theta, t) = \max_s \mathcal{F}_{\theta, t}[s]$$

Maximisation over all
possible strategies

**Optimal
strategy**

$$s^*(\theta, t) = \operatorname{argmax}_s \mathcal{F}_{\theta, t}[s]$$

ROYAL SOCIETY
OPEN SCIENCE

rsos.royalsocietypublishing.org

Research



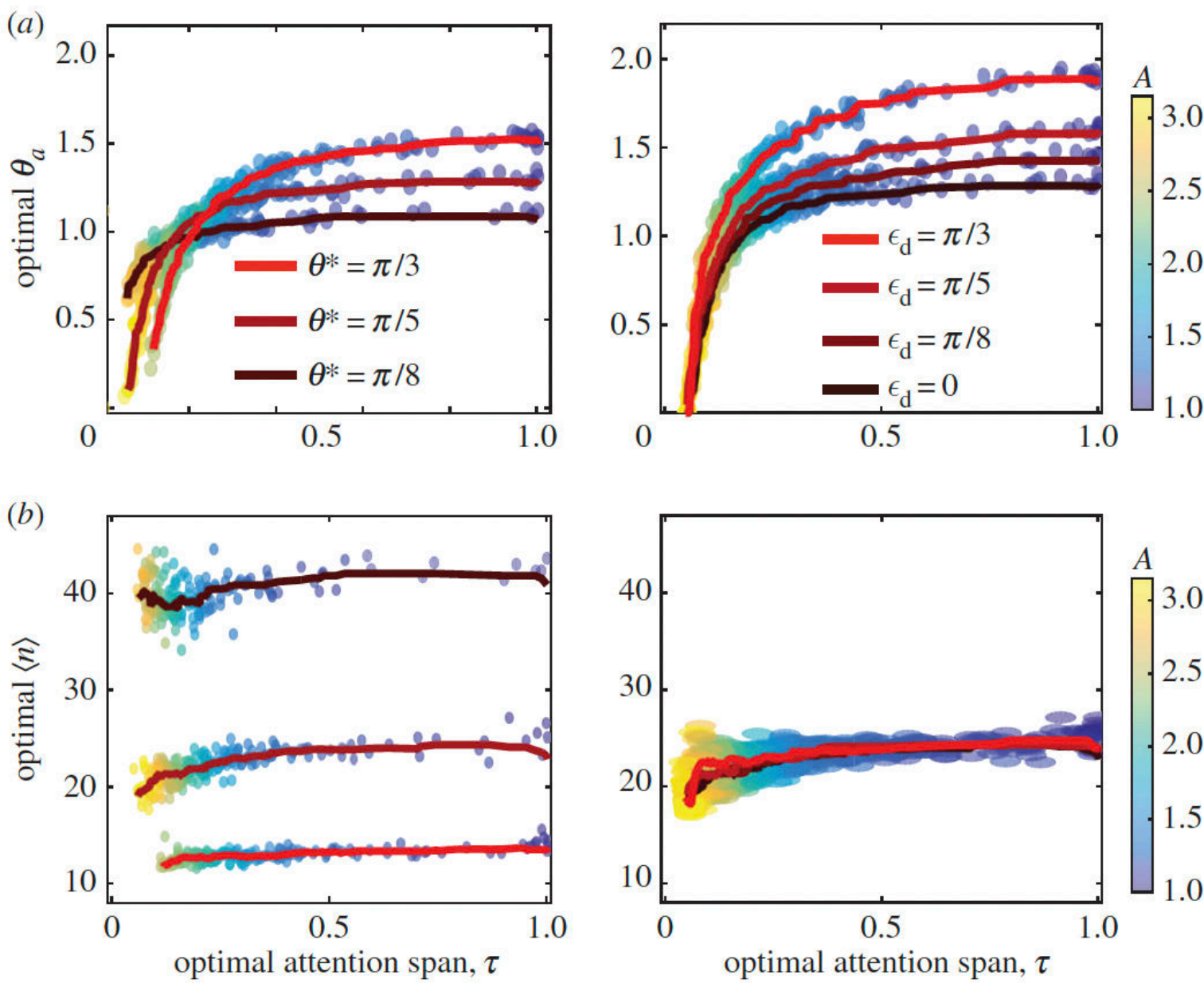
Cite this article: Peleg O, Mahadevan L. 2016
Optimal switching between geocentric and
egocentric strategies in navigation. *R. Soc.
open sci.* **3**: 160128.
<http://dx.doi.org/10.1098/rsos.160128>

Optimal switching between geocentric and egocentric strategies in navigation

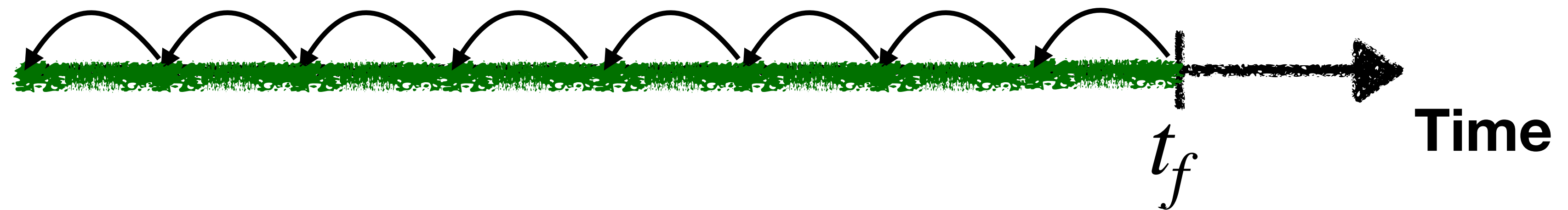
O. Peleg¹ and L. Mahadevan^{1,2,3}

¹Paulson School of Engineering and Applied Sciences, ²Department of Physics, and
³Department of Organismic and Evolutionary Biology, Kavli Institute for NanoBio
Science and Technology, Wyss Institute for Biologically Inspired Engineering, Harvard
University, Cambridge, MA 02138, USA

 OP, 0000-0001-9481-7967



Dynamic programming



R. Bellman, R. W. Kalaba, et al., Dynamic programming and modern control theory (1965)

Bellman equation



B. De Bruyne
(QRT)

**Optimal
reward**

$$J(\theta, t) = \max_s \mathcal{F}_{\theta, t}$$

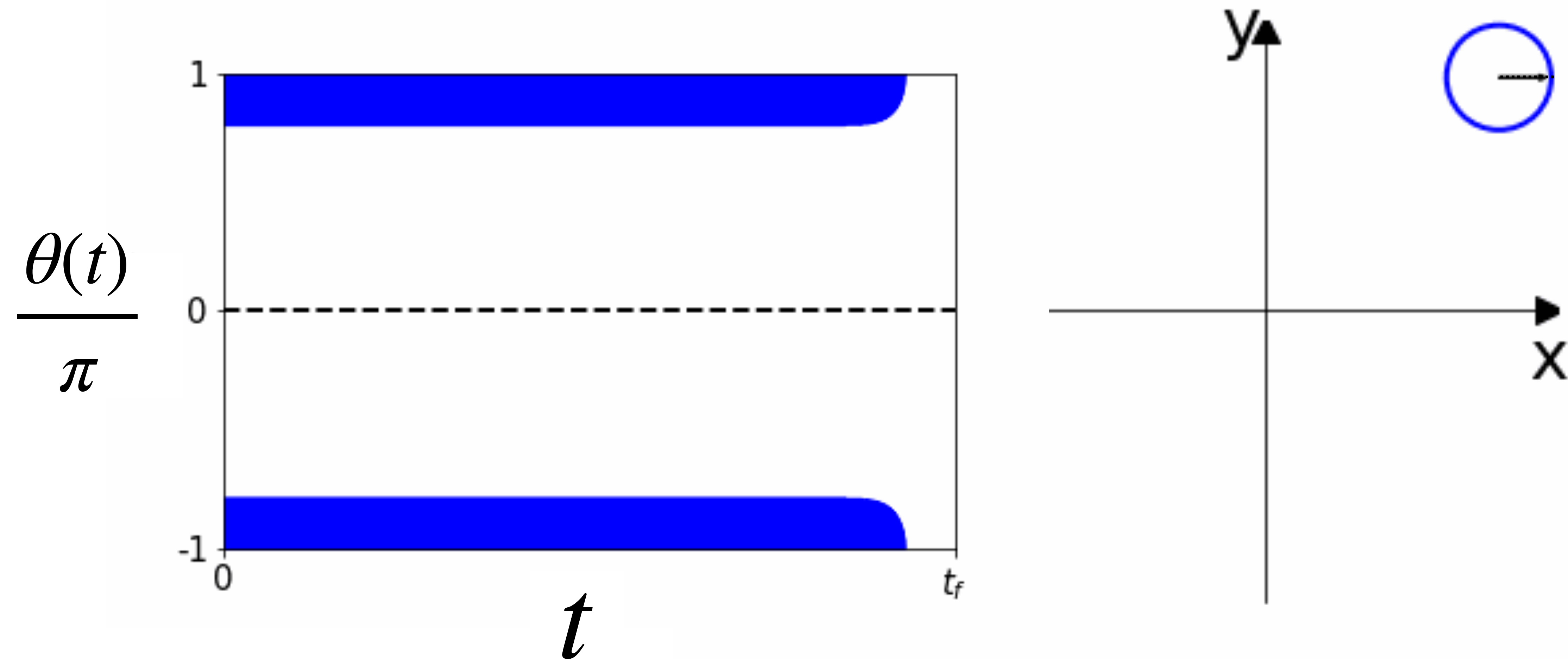
$$-\partial_t J(\theta, t) = D \partial_\theta^2 J(\theta, t) + v_0 \cos(\theta), \quad \theta \in \Omega(t),$$

$$\Omega(t) = \{\theta : J(\theta, t) \geq J(0, t) - x_c\}$$

**Optimal
strategy**

$$s^*(\theta, t) = \begin{cases} 1 & \text{if } \theta \notin \Omega(t), \\ 0 & \text{if } \theta \in \Omega(t), \end{cases}$$

Optimal reorientation policy



$$-\partial_t J(\theta, t) = D \partial_\theta^2 J(\theta, t) + v_0 \cos(\theta), \quad \theta \in \Omega(t),$$

$$\Omega(t) = \{\theta : J(\theta, t) \geq J(0, t) - x_c\}$$

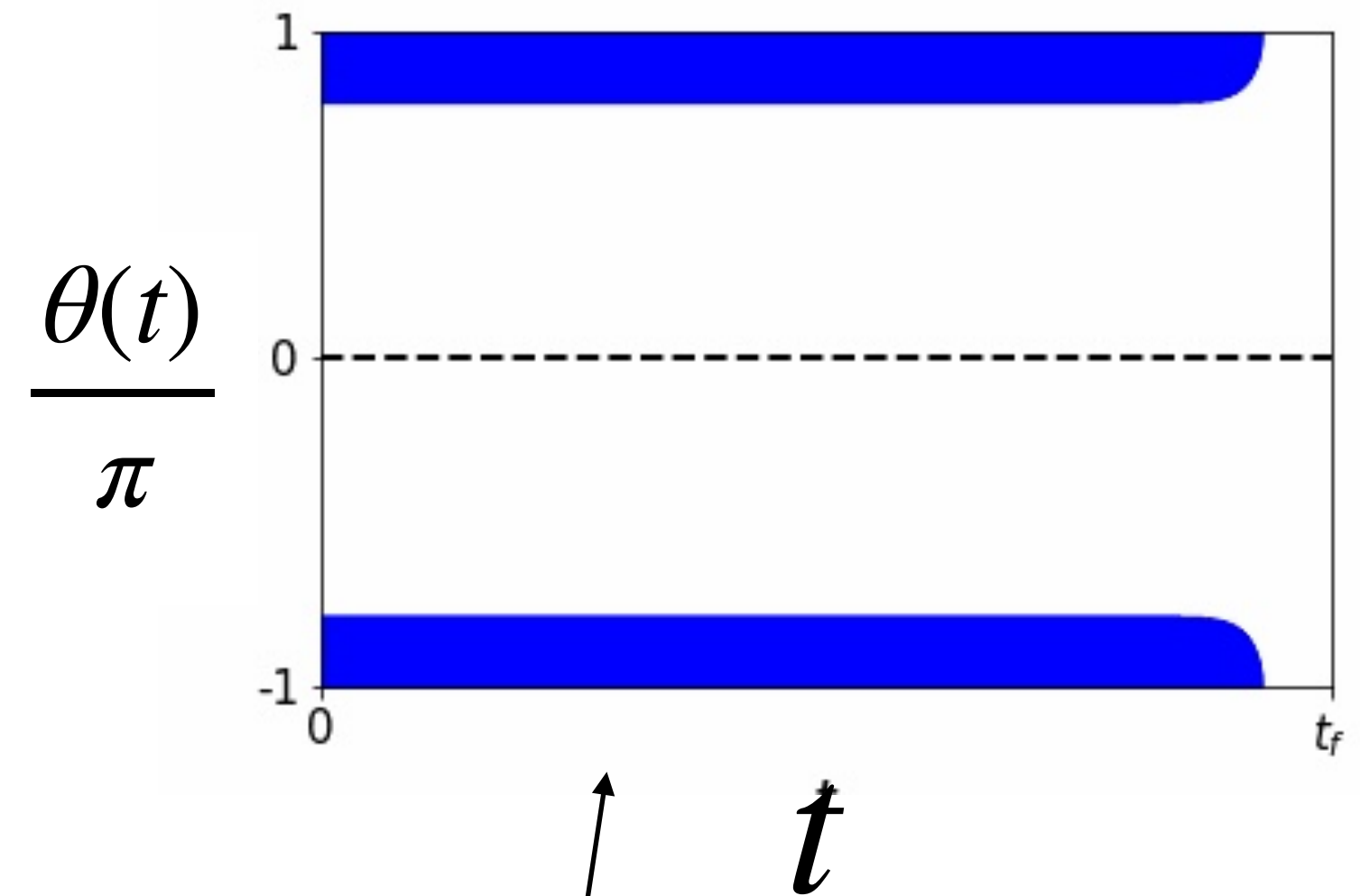
Ansatz

$$J(\theta, t) = j(\theta) + v^*(t_f - t)$$

$$\Omega(t) = [-\theta_a, \theta_a]$$

$$\frac{1}{2} \sin(\theta_a) \theta_a + \cos(\theta_a) = 1 - \frac{D x_c}{v_0}.$$

$$v^* = v_0 \frac{\sin(\theta_a)}{\theta_a}$$

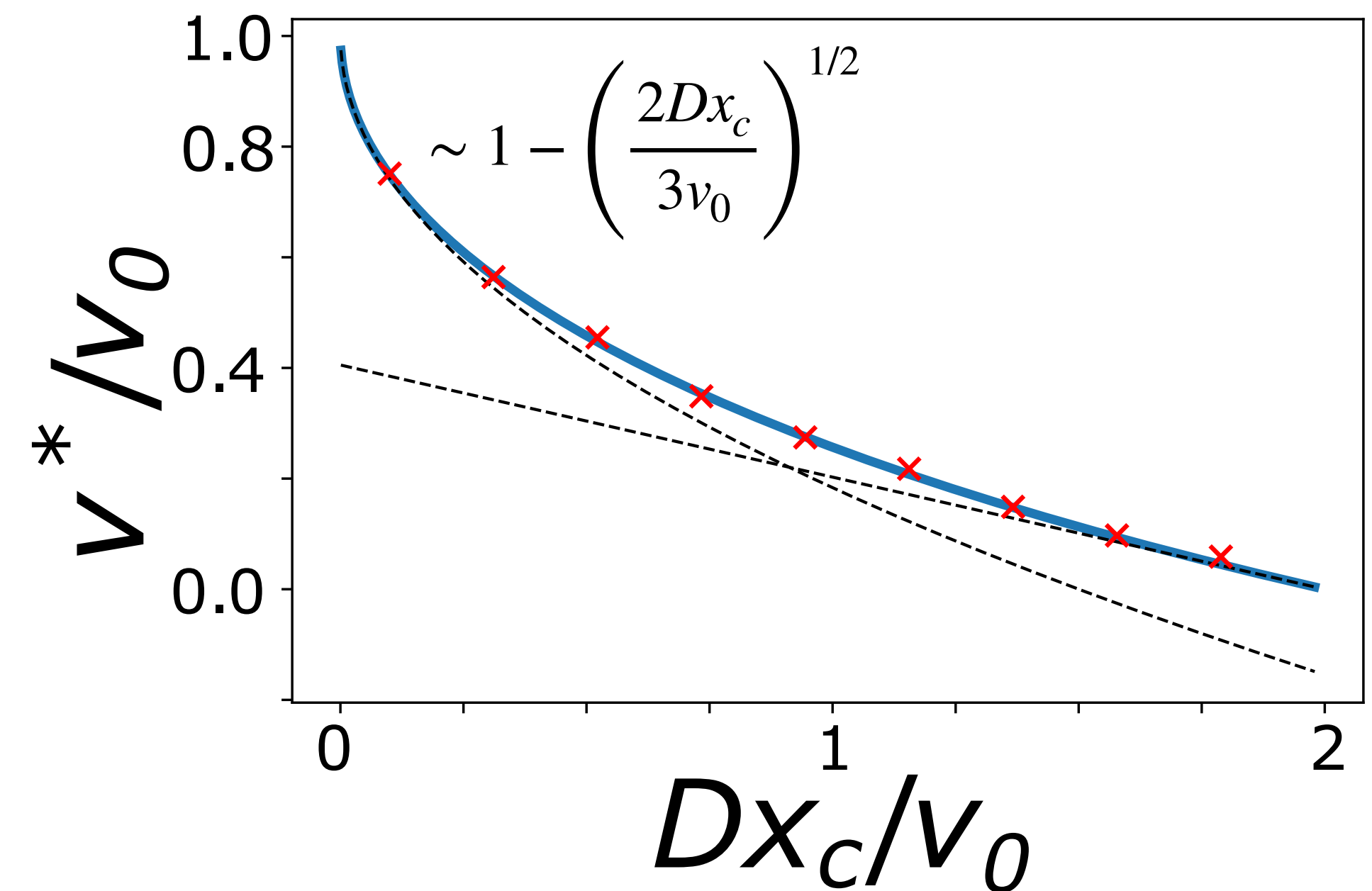
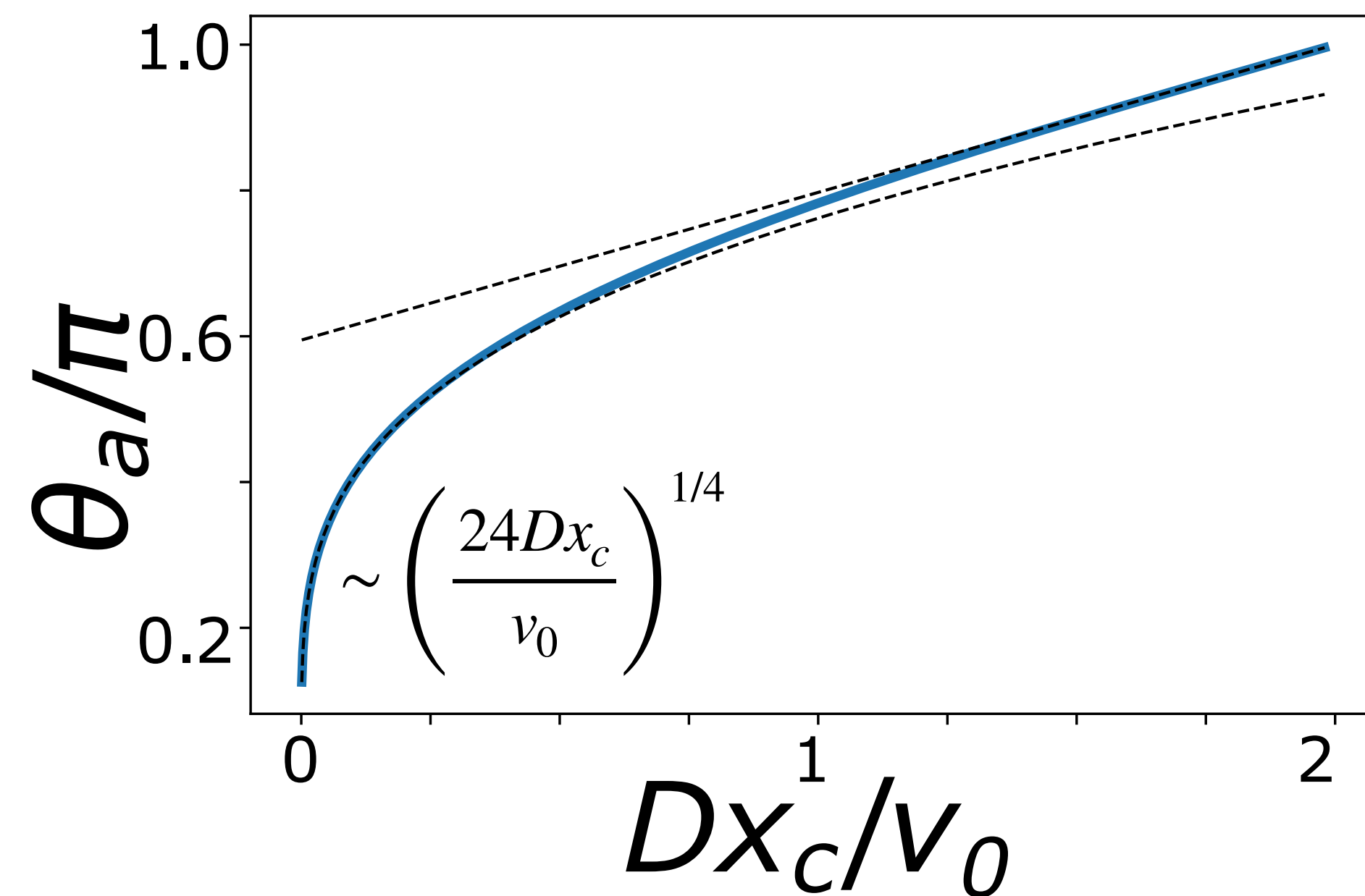


$$(t_f - t) \gg D^{-1}$$

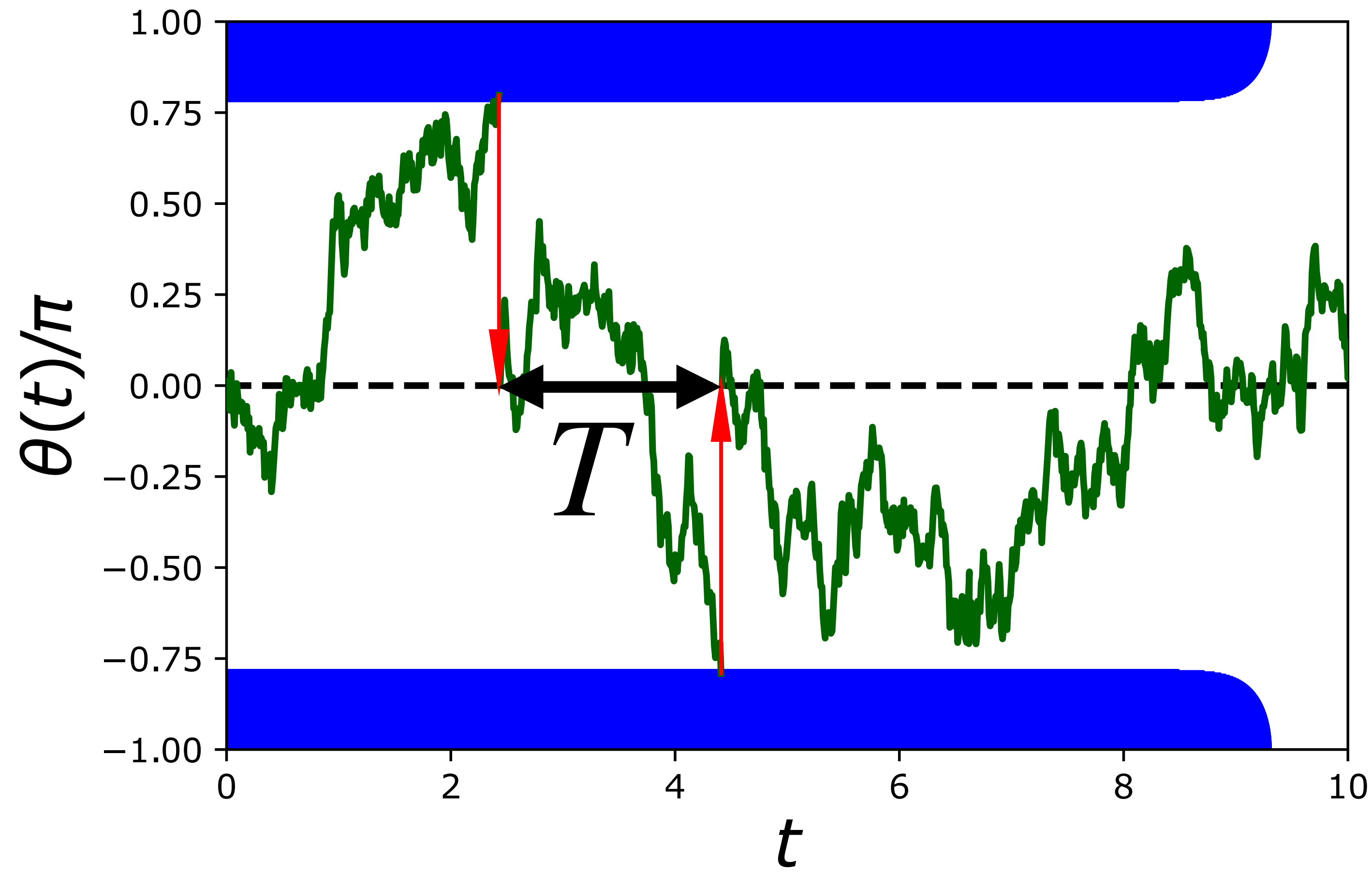
Infinite time horizon limit

$$\frac{1}{2} \sin(\theta_a) \theta_a + \cos(\theta_a) = 1 - \frac{Dx_c}{v_0}.$$

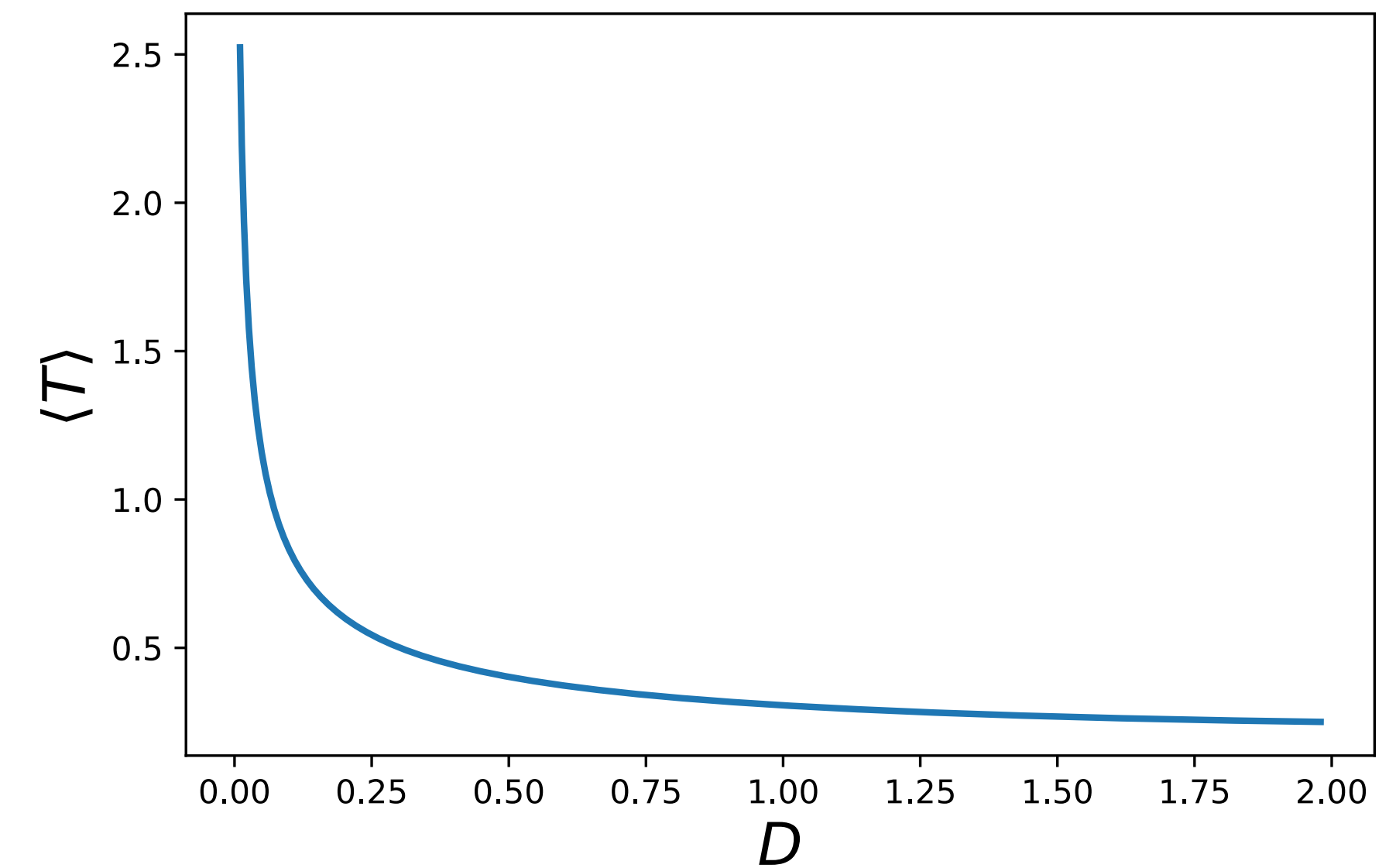
$$v^* = v_0 \frac{\sin(\theta_a)}{\theta_a}$$



Time between reorientations



$$\langle T \rangle = \frac{\theta_a^2}{2D} \sim \frac{1}{\sqrt{D}}$$



$$\begin{cases} \dot{x}(t) = v_0 \cos[\theta(t)], \\ \dot{y}(t) = v_0 \sin[\theta(t)], \\ \theta_1(t) = \theta(t) + \theta_n(t) \end{cases}$$

$\theta = \text{real angle}$	$\theta_n = \text{measurement noise}$
$\dot{\theta} = \eta(t)$	$\dot{\theta}_n = \eta_n(t)$
$\langle \eta(t)\eta(t') \rangle = 2D\delta(t - t')$	$\langle \eta_n(t)\eta_n(t') \rangle = 2D_n\delta(t - t')$

Strategy: reorient if $|\theta_1| > \theta_a$

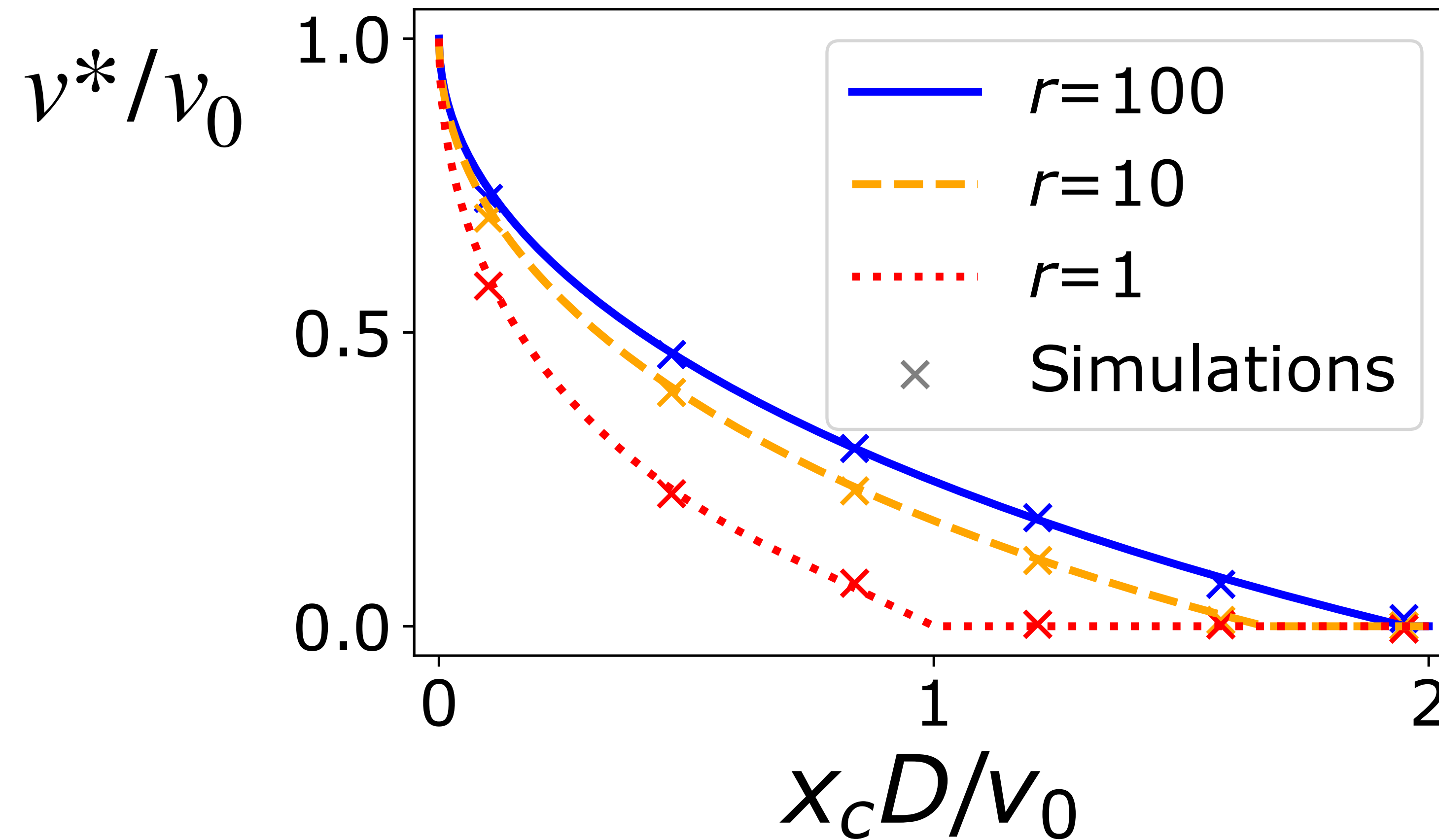
$$\begin{cases} \dot{x}(t) = v_0 \cos[\theta(t)], \\ \dot{y}(t) = v_0 \sin[\theta(t)], \\ \theta_1(t) = \theta(t) + \theta_n(t) \end{cases}$$

Strategy: reorient if $|\theta_1| > \theta_a$

$$v = \frac{2v_0}{\theta_a^2} (1 + 1/r) \left[1 - z - \frac{\cos\left(\frac{\theta_a}{1 + 1/r}\right)}{\cosh\left(\frac{\theta_a \sqrt{1/r}}{1 + 1/r}\right)} \right]$$

$$r = D/D_n \quad z = \frac{Dx_c}{v_0}$$

$$r = D/D_n$$



- Minimal model of animal navigation inspired by the dung beetle
- Applications other systems: olfactory navigation, search processes,...
- Experimental verification (**Lund Vision Group**)



DALL.E

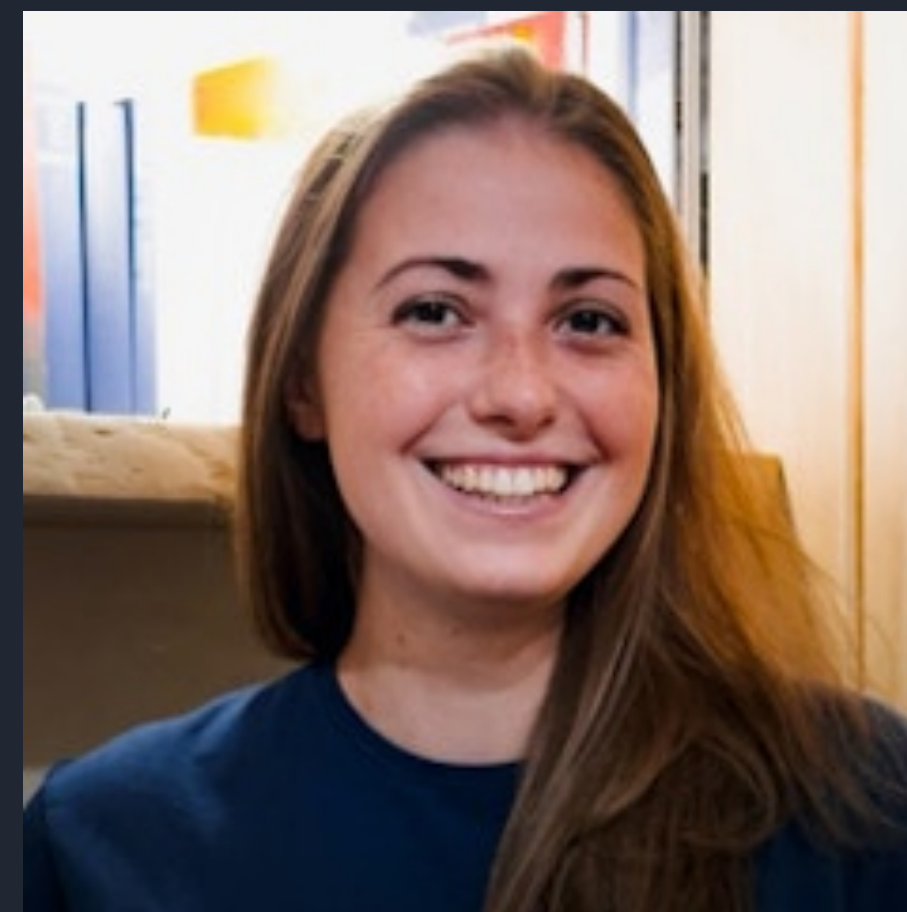
Part II:

Multi-task learning

“Optimal protocols for continual learning via statistical physics and control theory”

FM, Stefano Sarao Mannelli, Francesca Mignacco

arXiv preprint *arXiv:2409.18061*

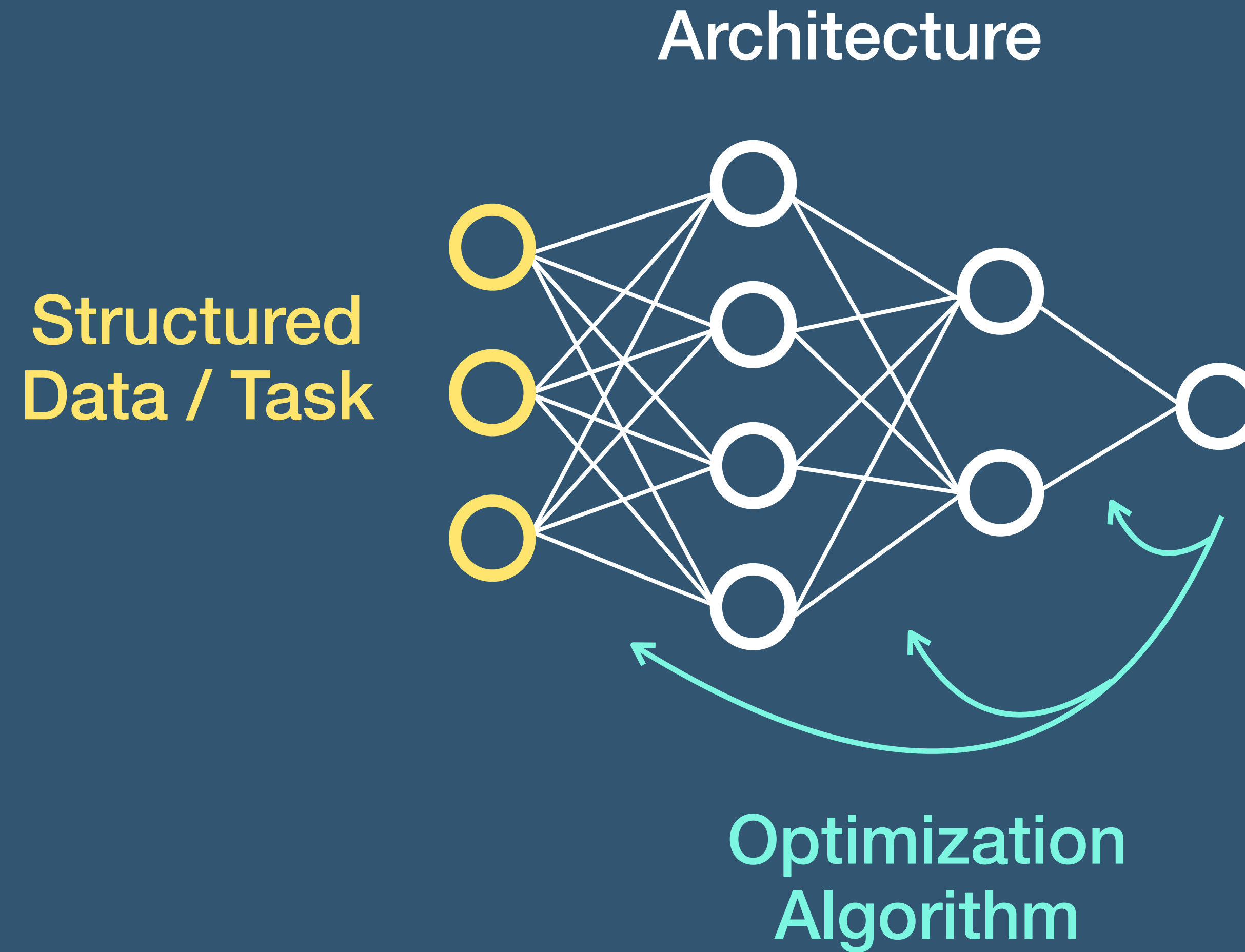


Francesca Mignacco
Princeton and CUNY

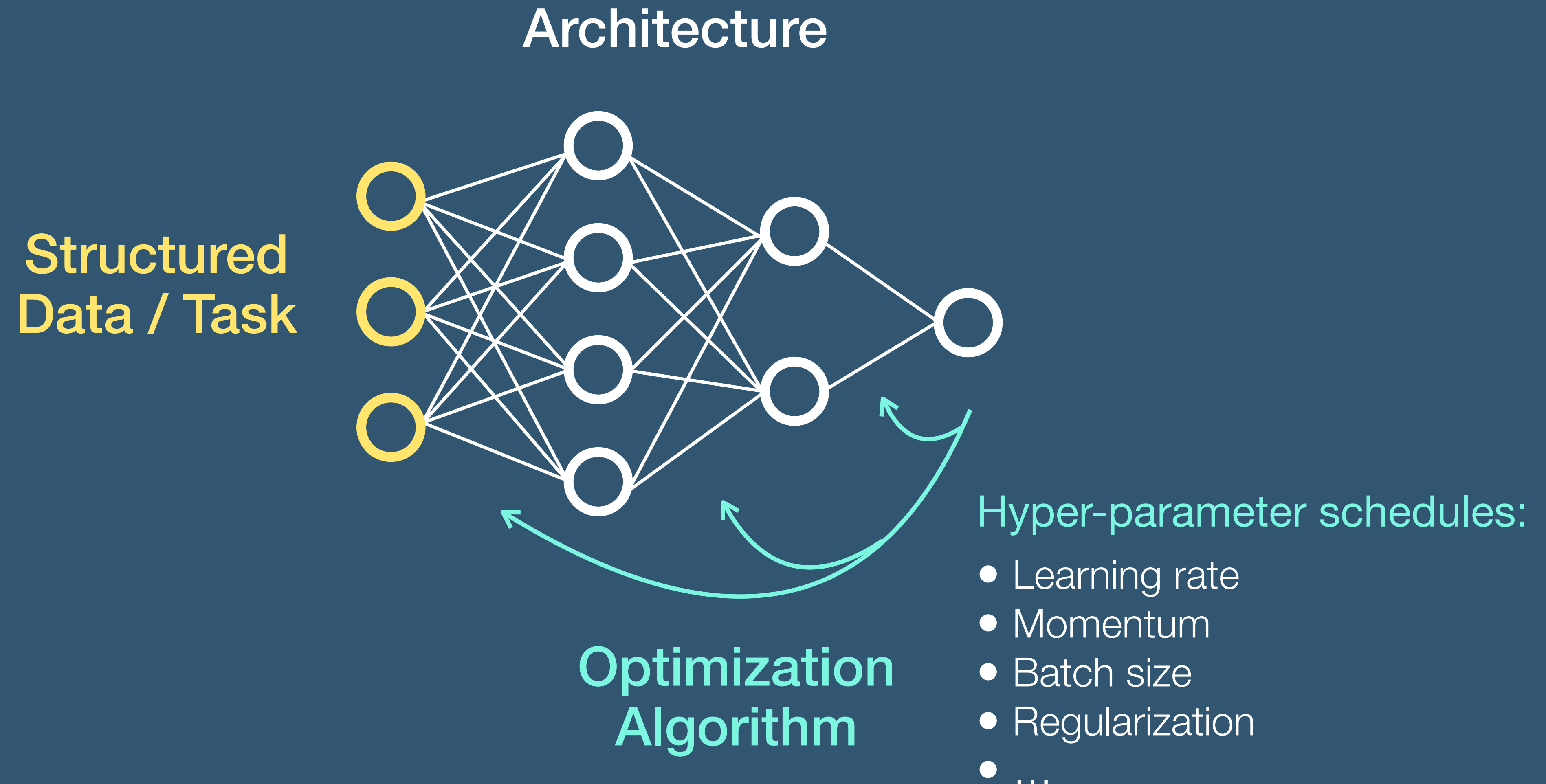


Stefano Sarao Mannelli
Chalmers University

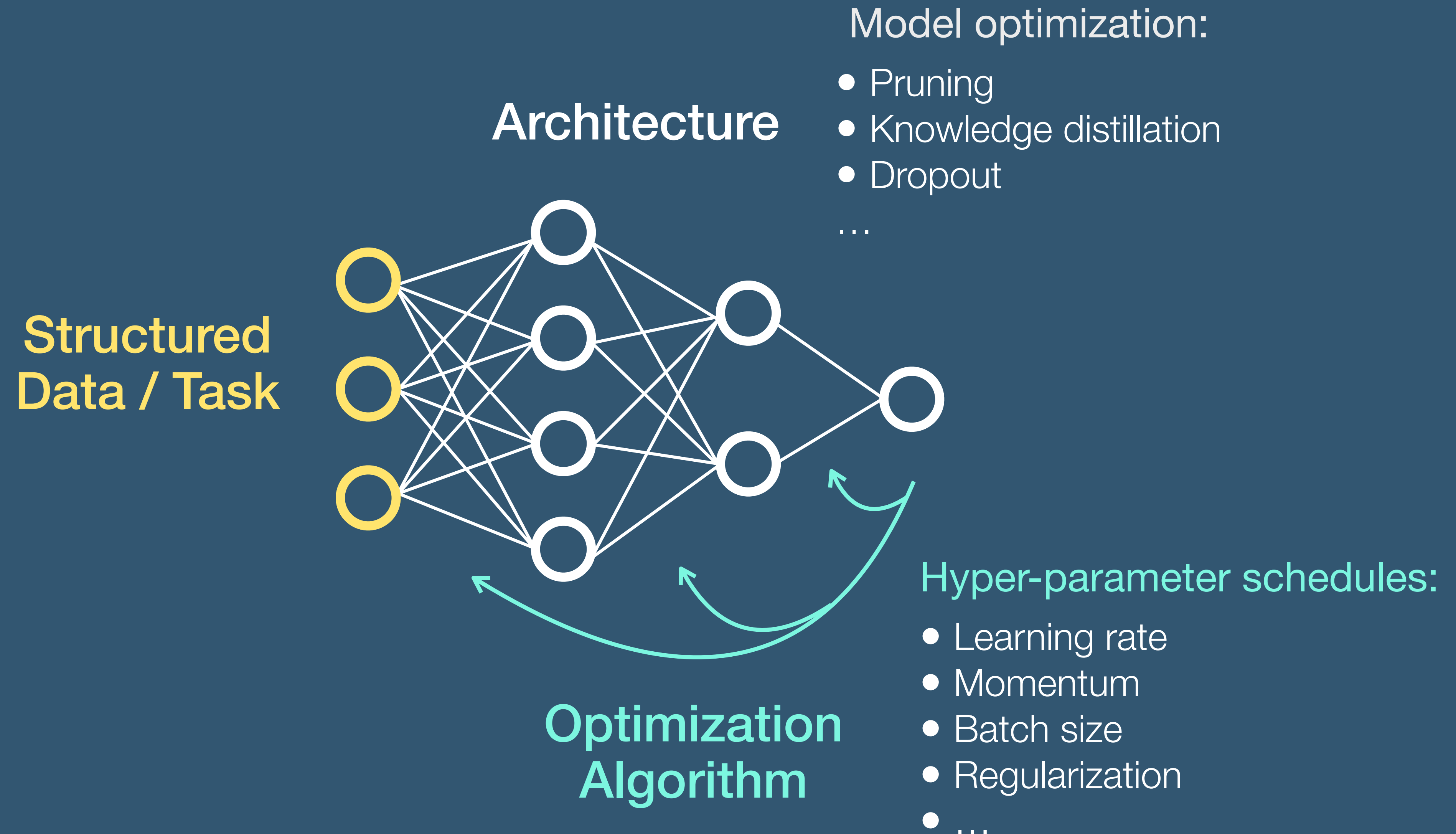
Learning protocols



Learning protocols



Learning protocols



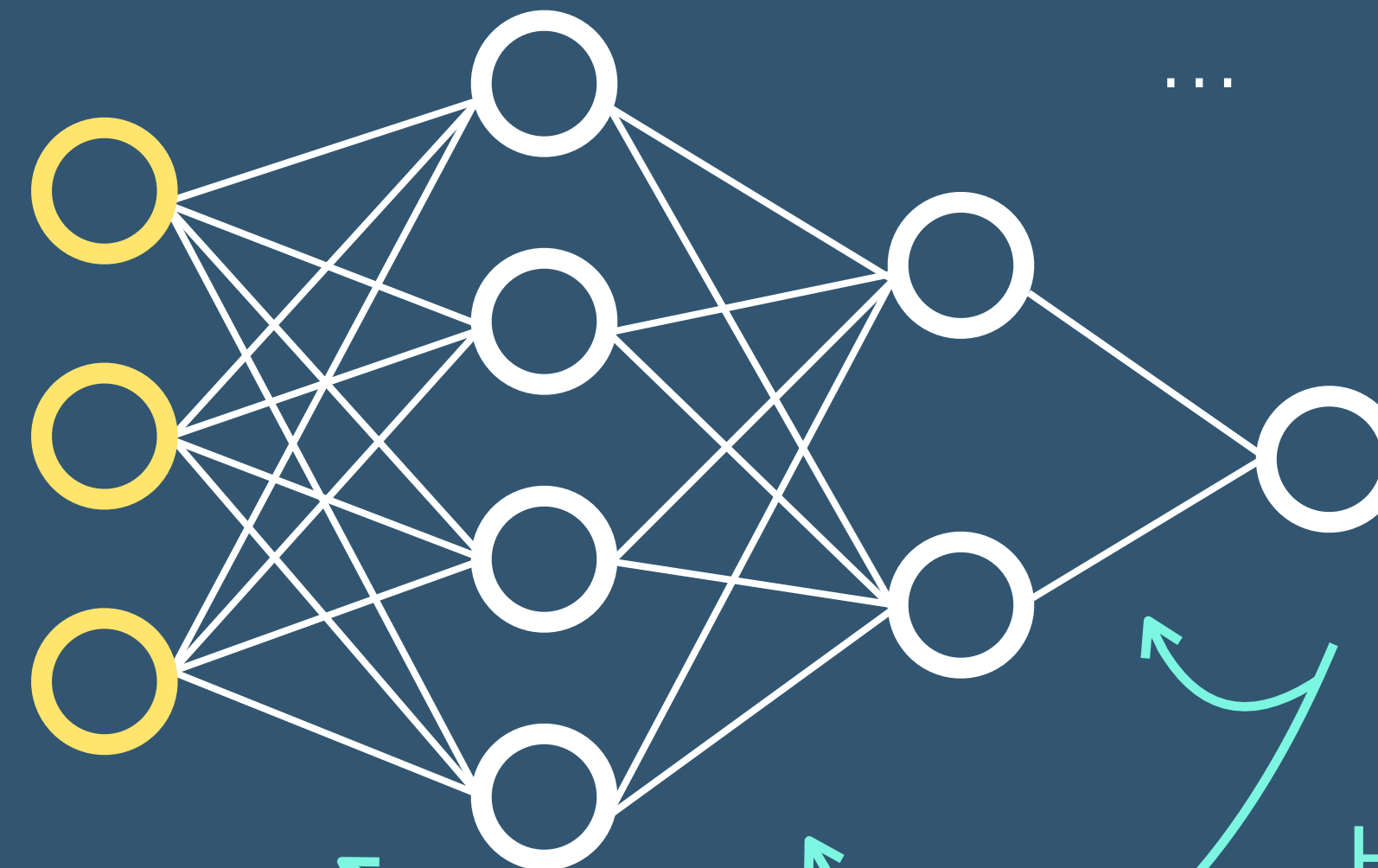
Learning protocols

Structured
Data / Task

Dynamic data / task selection:

- Active learning
- Curriculum learning
- Transfer learning
- Multi-task learning
- ...

Architecture



Optimization
Algorithm

Model optimization:

- Pruning
- Knowledge distillation
- Dropout
- ...

Hyper-parameter schedules:

- Learning rate
- Momentum
- Batch size
- Regularization
- ...

Learning protocols

Goals:

- **Speed-up** convergence
- **Guide** the training towards better regions of parameters space

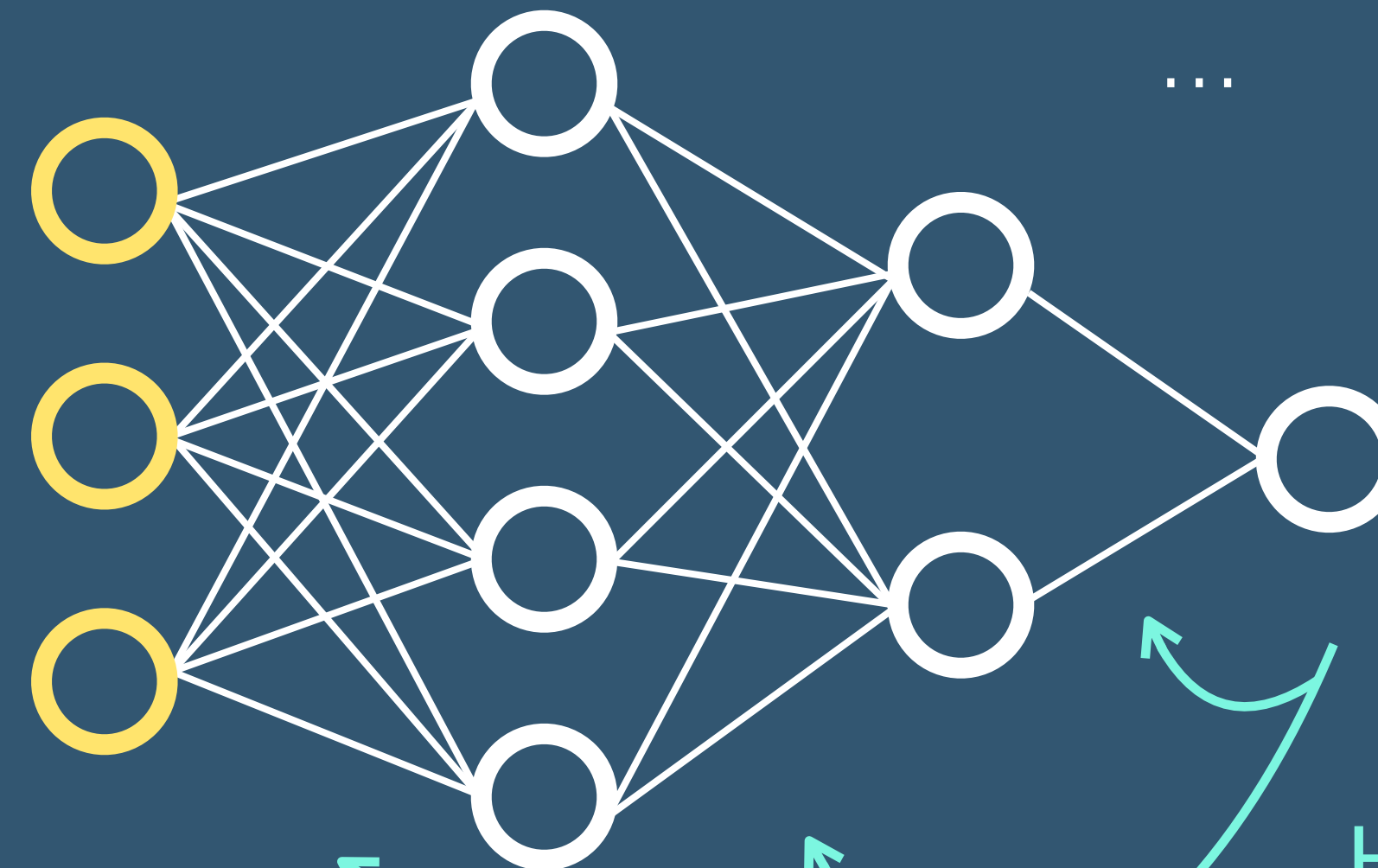
From smoother landscape to the target

Structured Data / Task

Dynamic data / task selection:

- Active learning
- Curriculum learning
- Transfer learning
- Multi-task learning
- ...

Architecture



Optimization Algorithm

Model optimization:

- Pruning
- Knowledge distillation
- Dropout
- ...

Hyper-parameter schedules:

- Learning rate
- Momentum
- Batch size
- Regularization
- ...

Learning protocols

Goals:

- **Speed-up** convergence
- **Guide** the training towards better regions of parameters space

From smoother landscape to the target

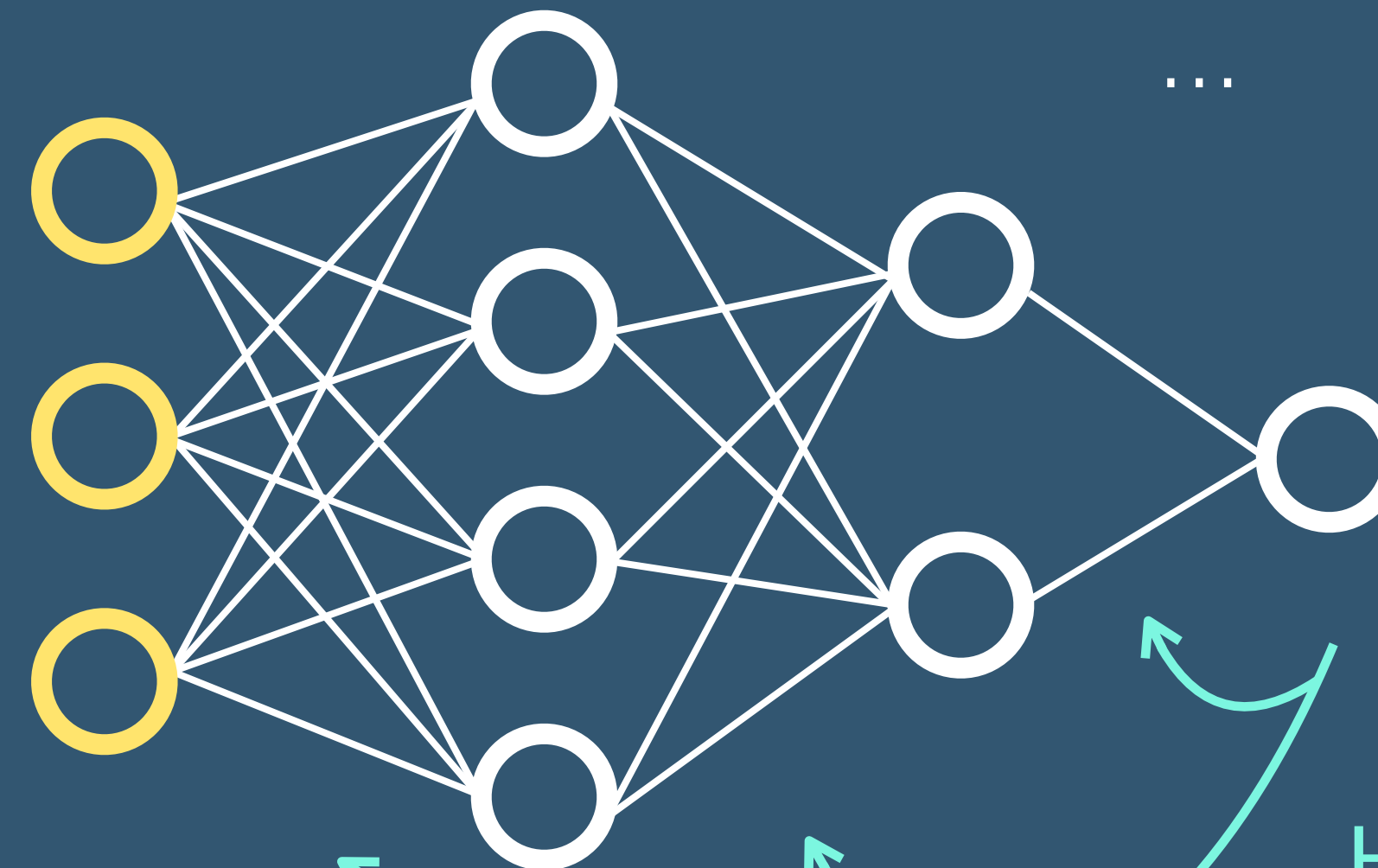
**Structured
Data / Task**

In this talk:

Dynamic data / task selection:

- Active learning
- Curriculum learning
- Transfer learning
- Multi-task learning
- ...

Architecture



Optimization Algorithm

Model optimization:

- Pruning
- Knowledge distillation
- Dropout
- ...

Hyper-parameter schedules:

- Learning rate
- Momentum
- Batch size
- Regularization
- ...

Model-based approaches to data / task selection

(Non-exhaustive list)

- Curriculum learning [Weinshall et al (2020); Saglietti et al. (2022); Abbe et al. (2023); Cornacchia et al (2023); Lee et al. (2024); Mannelli et al. (2024); ...]
- Transfer learning [Dhifallah & Lu (2021); Gerace et al. (2022, 2023, 2024); ...]
- Continual learning & catastrophic forgetting [Lee et al. (2021, 2022); Shan et al (2024); ...]
- Active learning [Cui et al. (2020); ...]
- ...

(Non-exhaustive list)

- Curriculum learning [Weinshall et al (2020); Saglietti et al. (2022); Abbe et al. (2023); Cornacchia et al (2023); Lee et al. (2024); Mannelli et al. (2024); ...]
- Transfer learning [Dhifallah & Lu (2021); Gerace et al. (2022, 2023, 2024); ...]
- Continual learning & catastrophic forgetting [Lee et al. (2021, 2022); Shan et al (2024); ...]
- Active learning [Cui et al. (2020); ...]
- ...

Can we compute the *optimal** strategy ?

* In terms of the final performance

Supervised learning - general setup

Dataset with labels

$$\mathcal{D} = \{x_i, y_i\}_{i=1}^P$$

Neural Network

$$\hat{y} = f_{\mathbf{w}}(x)$$

Example: $\hat{y} = \text{erf}(\mathbf{w}^T x)$

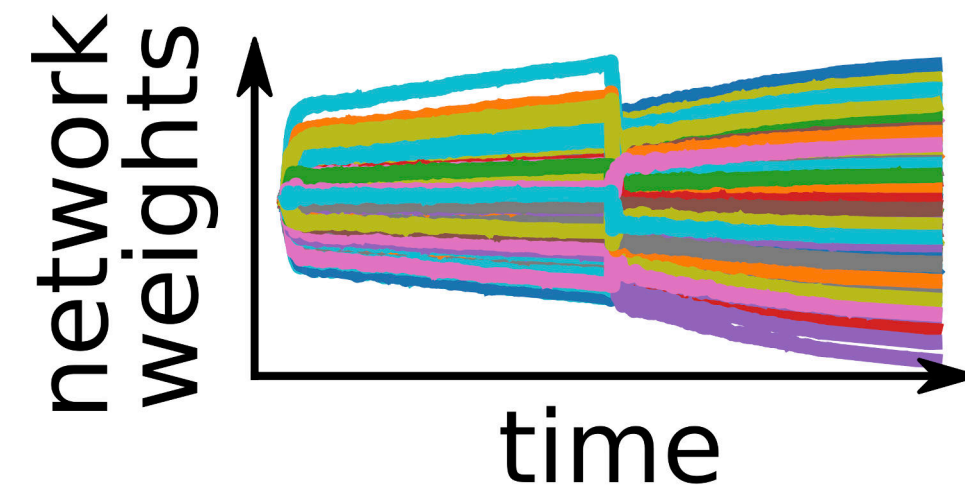
Error (aka loss)

$$\mathcal{L} = \frac{1}{2} (\hat{y} - y)^2$$

Stochastic gradient descent

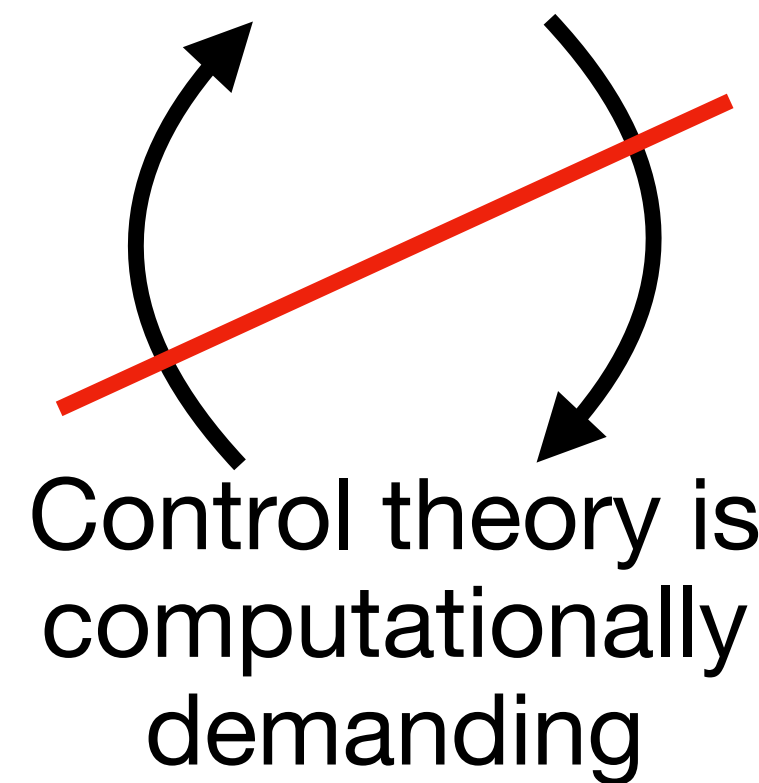
$$\mathbf{w}^{\mu+1} = \mathbf{w}^{\mu} - \eta \nabla \mathcal{L}^{\mu}$$

Dimensionality reduction + optimal control

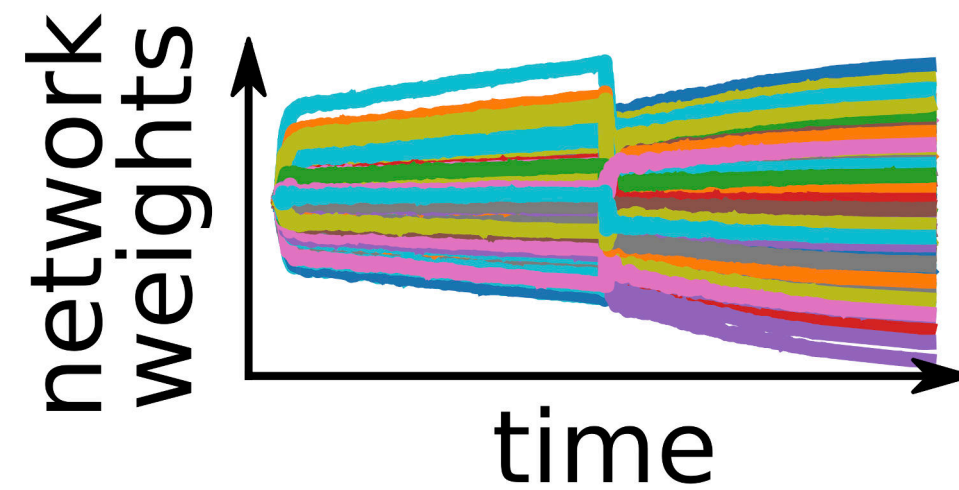


$$\mathbf{w}^{\mu+1} = \mathbf{w}^{\mu} - \eta \nabla \mathcal{L}^{\mu}$$

High-dimensional
complex dynamics

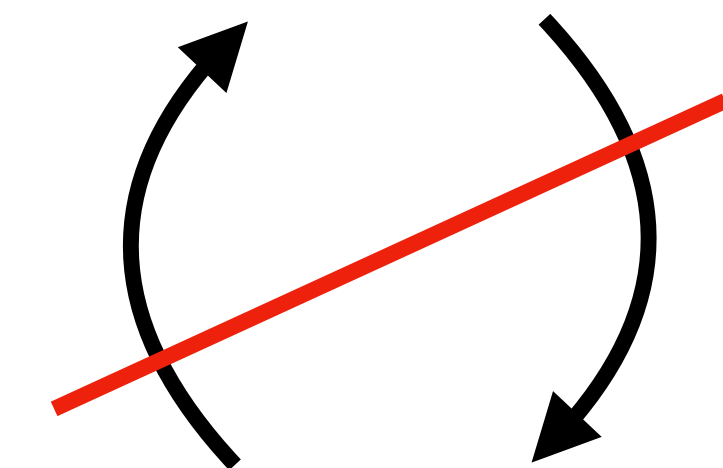


Dimensionality reduction + optimal control



$$\mathbf{w}^{\mu+1} = \mathbf{w}^{\mu} - \eta \nabla \mathcal{L}^{\mu}$$

High-dimensional
complex dynamics

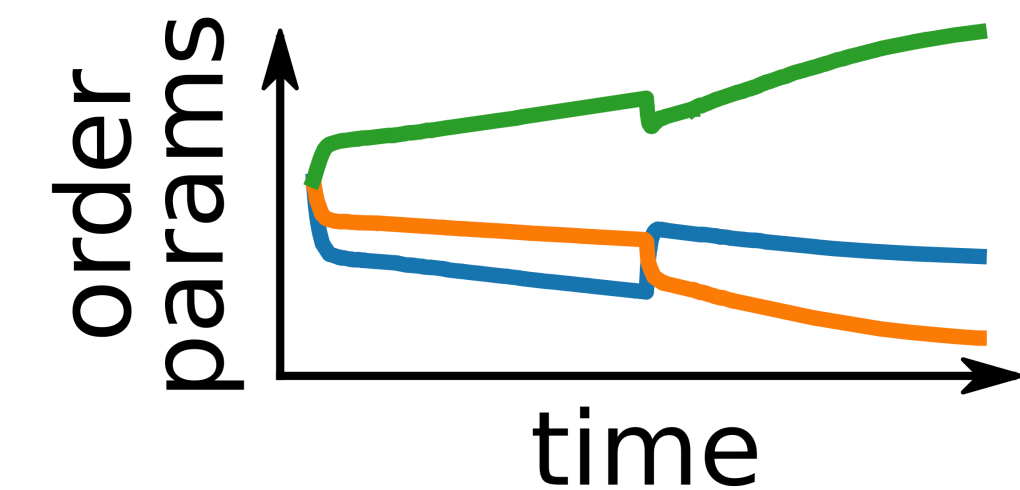


Control theory is
computationally
demanding

Statistical physics

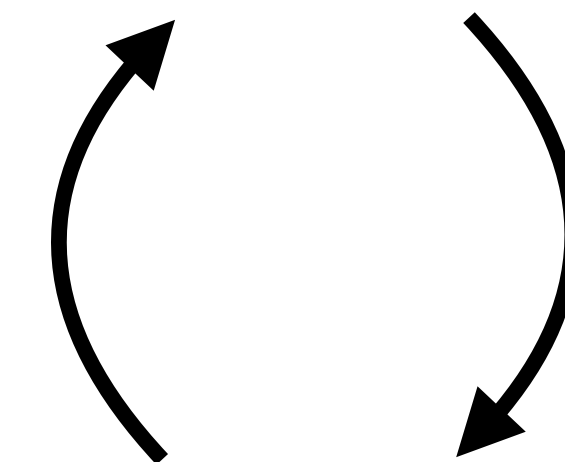


[Saad & Solla (1995); Biehl & Schwarze (1995); Riegler & Biehl (1995); Goldt et al (2019); Refinetti (2020); Veiga et al (2022); Arnaboldi et al (2023); ... Agoritsas et al (2018); Mignacco et al (2020); Mannelli et al (2021); Bonnaire et al (2023); ...]



$$\dot{\mathbf{Q}} = f(\mathbf{Q}, \mathbf{u})$$
 control variables

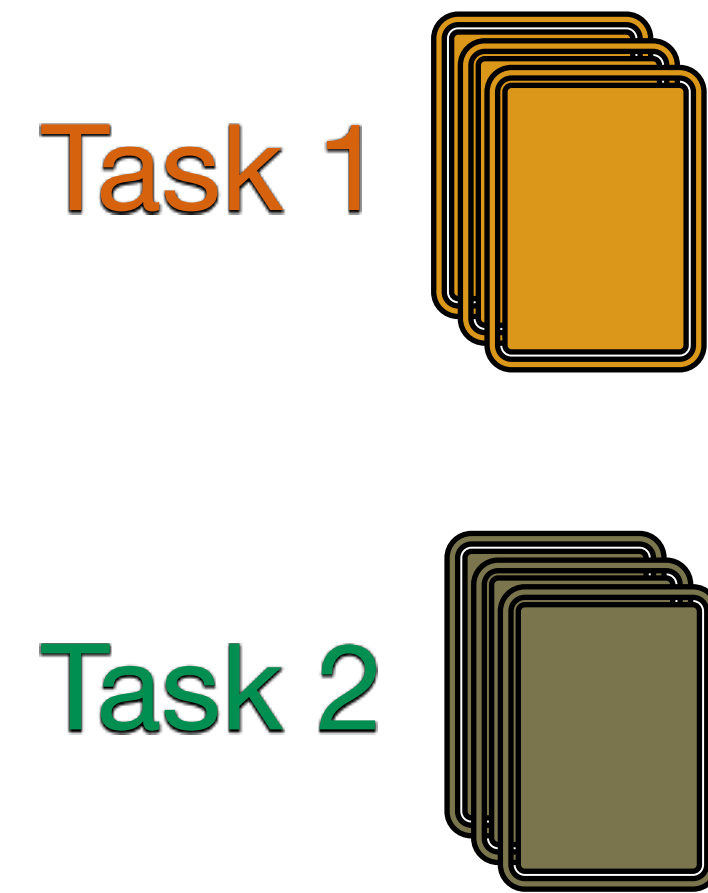
Low-dimensional
effective description



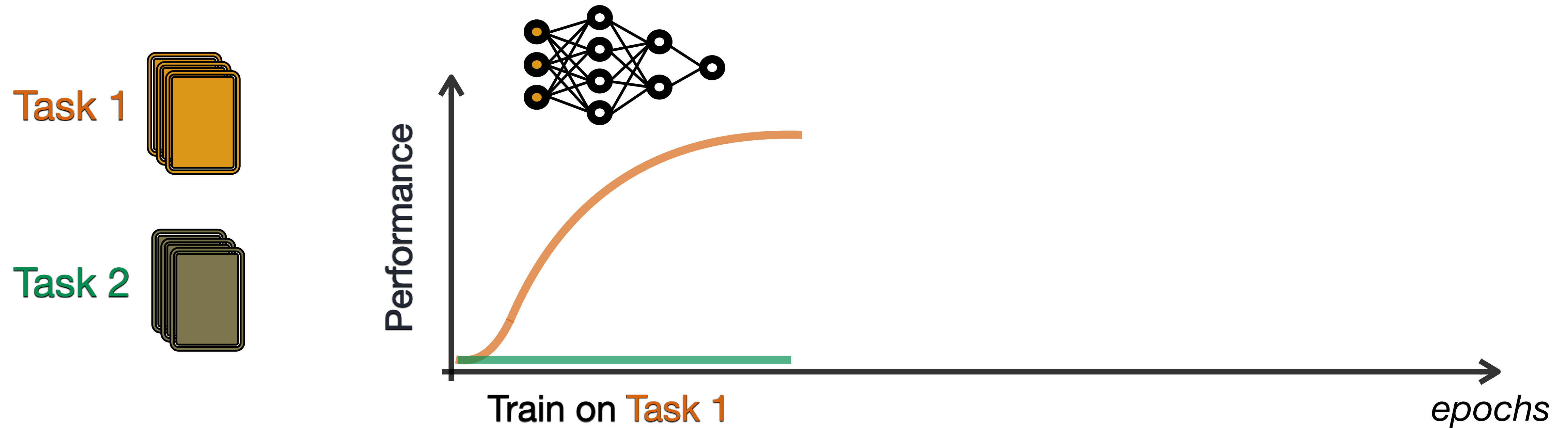
Control theory can be
applied

[Pontryagin (1962), Bellman (1965); Saad & Rattray (1997), Sivak & Crooks (2012); ...]

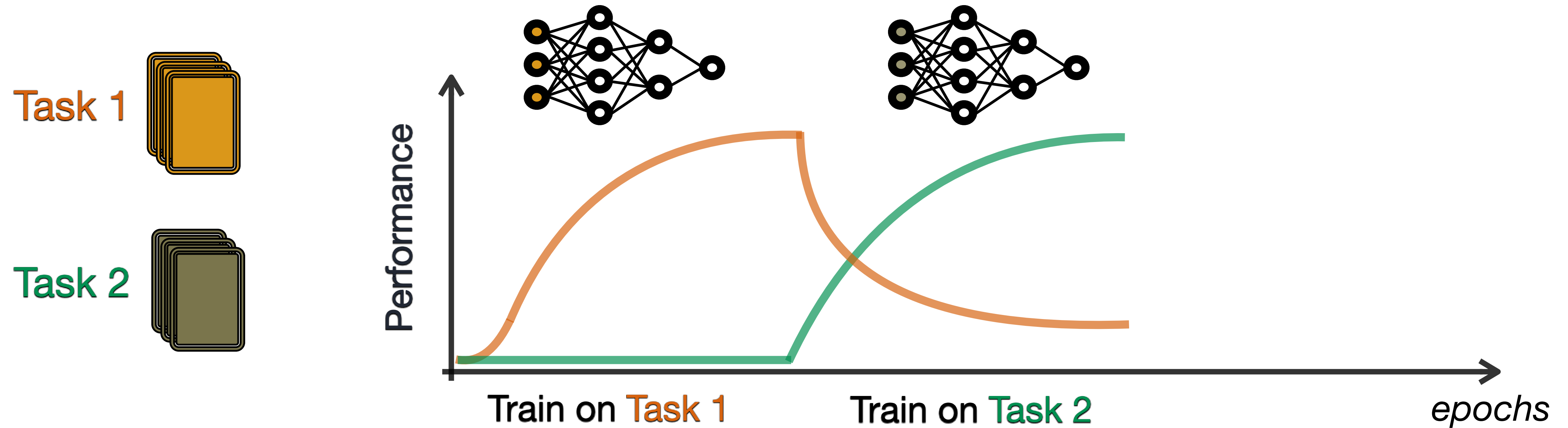
Continual learning & catastrophic forgetting



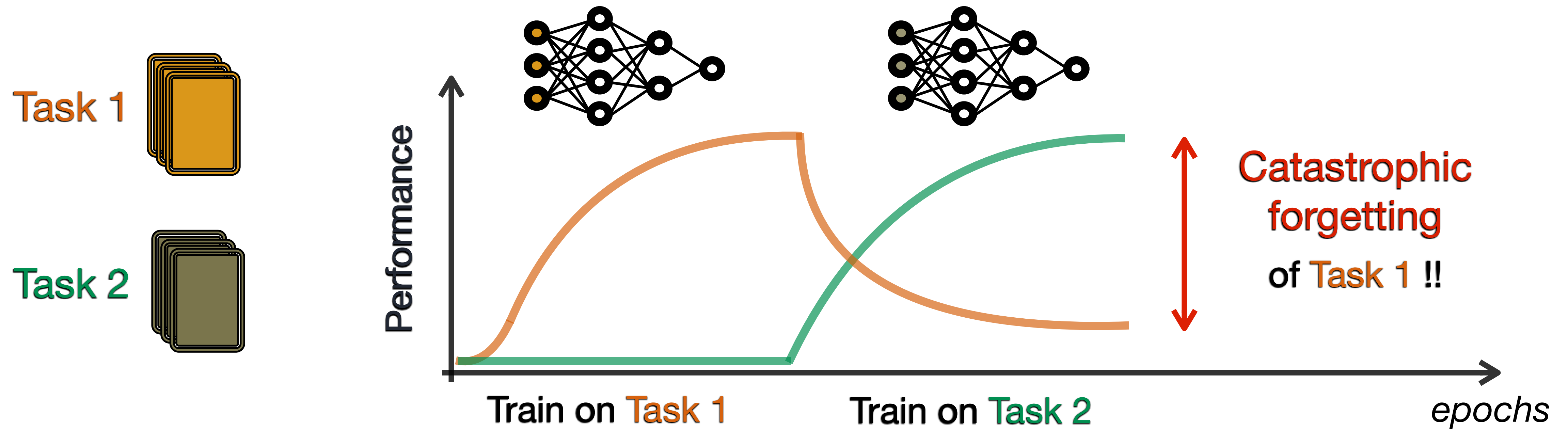
Continual learning & catastrophic forgetting



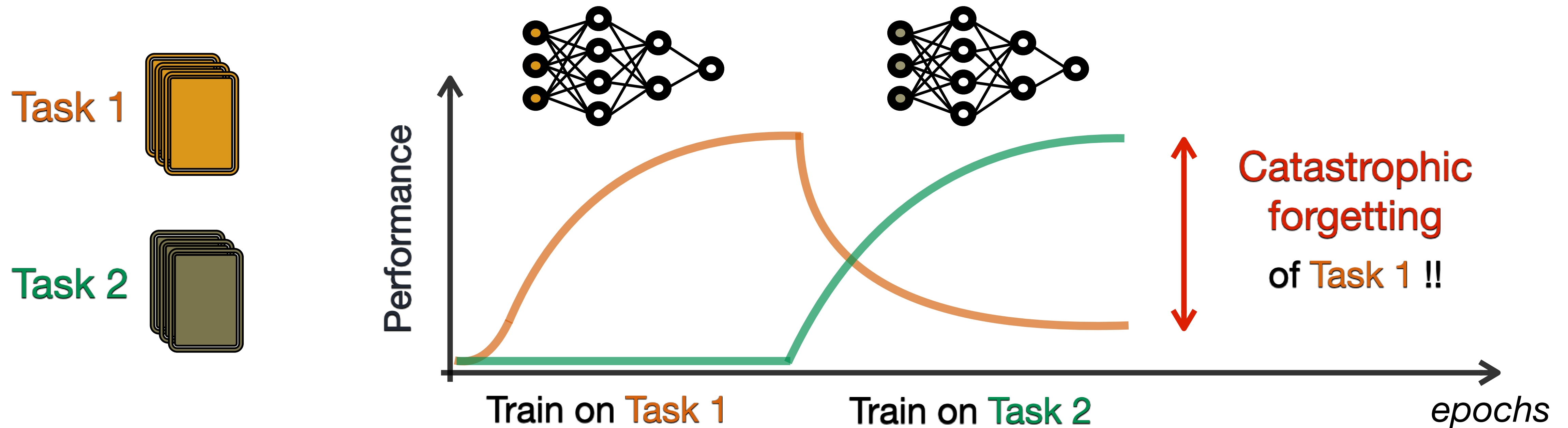
Continual learning & catastrophic forgetting



Continual learning & catastrophic forgetting



Continual learning & catastrophic forgetting



Humans & Animals

can learn sequentially without interference problems

[McClelland et al (1995); Barnett & Ceci (2002); Calvert et al., (2004); Mareschal et al (2007); Pallier et al (2003); ... Flesch et al (2018); Cichon & Gan (2015); Yang et al., 2014) ...]

ML (empirical):

Neural networks suffer from **catastrophic forgetting**

[Goodfellow et al (2014); Ruder & Planck (2014); Nguyen et al (2019); Parisi et al (2019); Mirzadeh et al (2020); Neyshabur et al (2020)]

ML (theory):

Key role of width, depth and **task similarity**

[Mirzadeh et al (2021); Lee et al. (2021, 2022); f Asanuma et al (2021); Doan et al (2021); Shan et al (2024)]

A teacher-student model of (**supervised**) continual learning

Introduced in: *Lee, Goldt, & Saxe (ICML 2021)*

Task 1

$$\mathcal{D}_1 = \{x_i^{(1)}, y_i^{(1)}\}$$

$$x \sim \mathcal{N}(0,1) \in \mathbf{R}^N$$

$$N \gg 1$$

Task 2

$$\mathcal{D}_2 = \{x_i^{(2)}, y_i^{(2)}\}$$

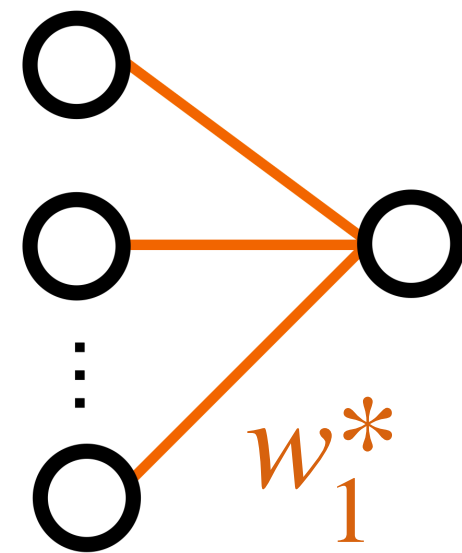
A teacher-student model of continual learning

Introduced in: *Lee, Goldt, & Saxe (ICML 2021)*

Task 1

$$y^{(1)} = g_* \left(w_1^* \cdot \frac{x}{\sqrt{N}} \right)$$

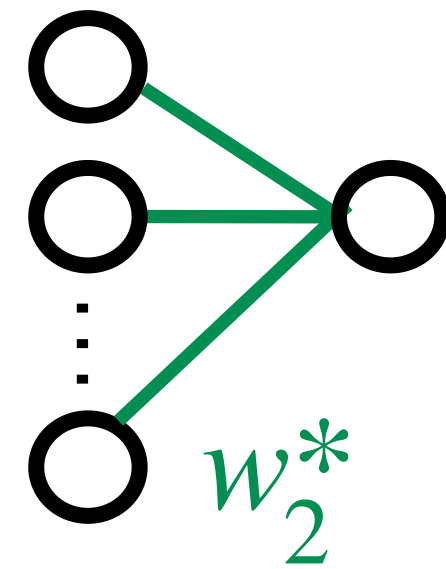
Teacher 1



Task 2

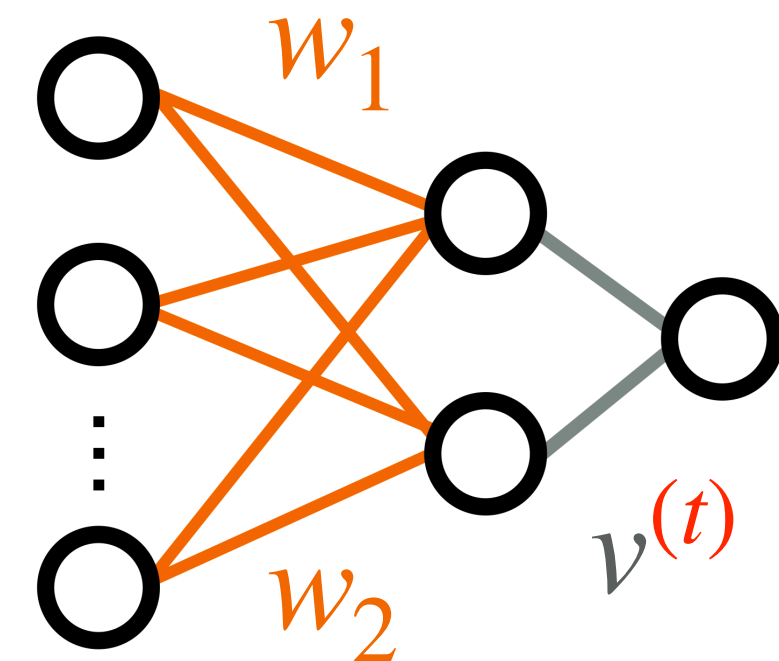
$$y^{(2)} = g_* \left(w_2^* \cdot \frac{x}{\sqrt{N}} \right)$$

Teacher 2



Student

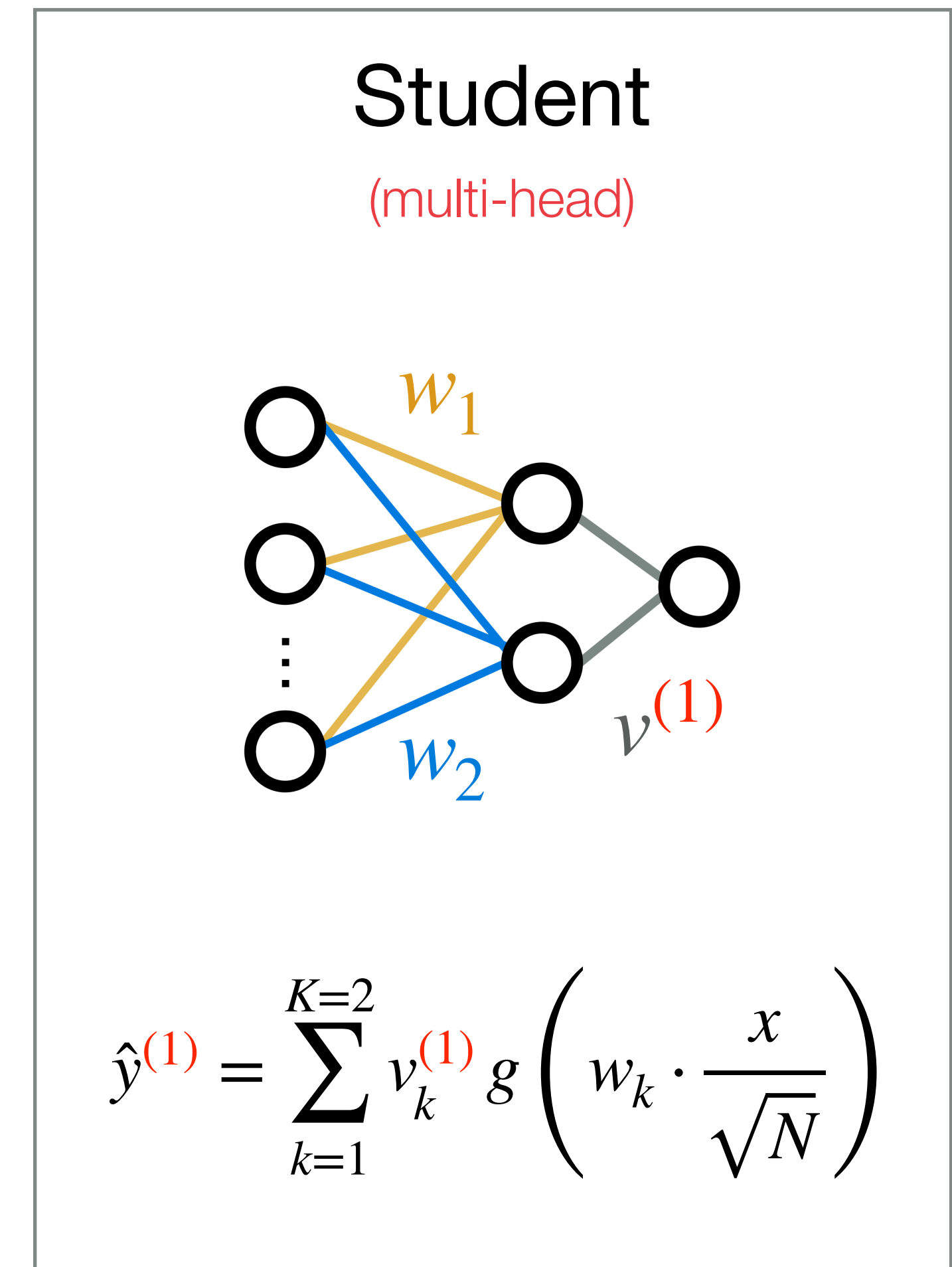
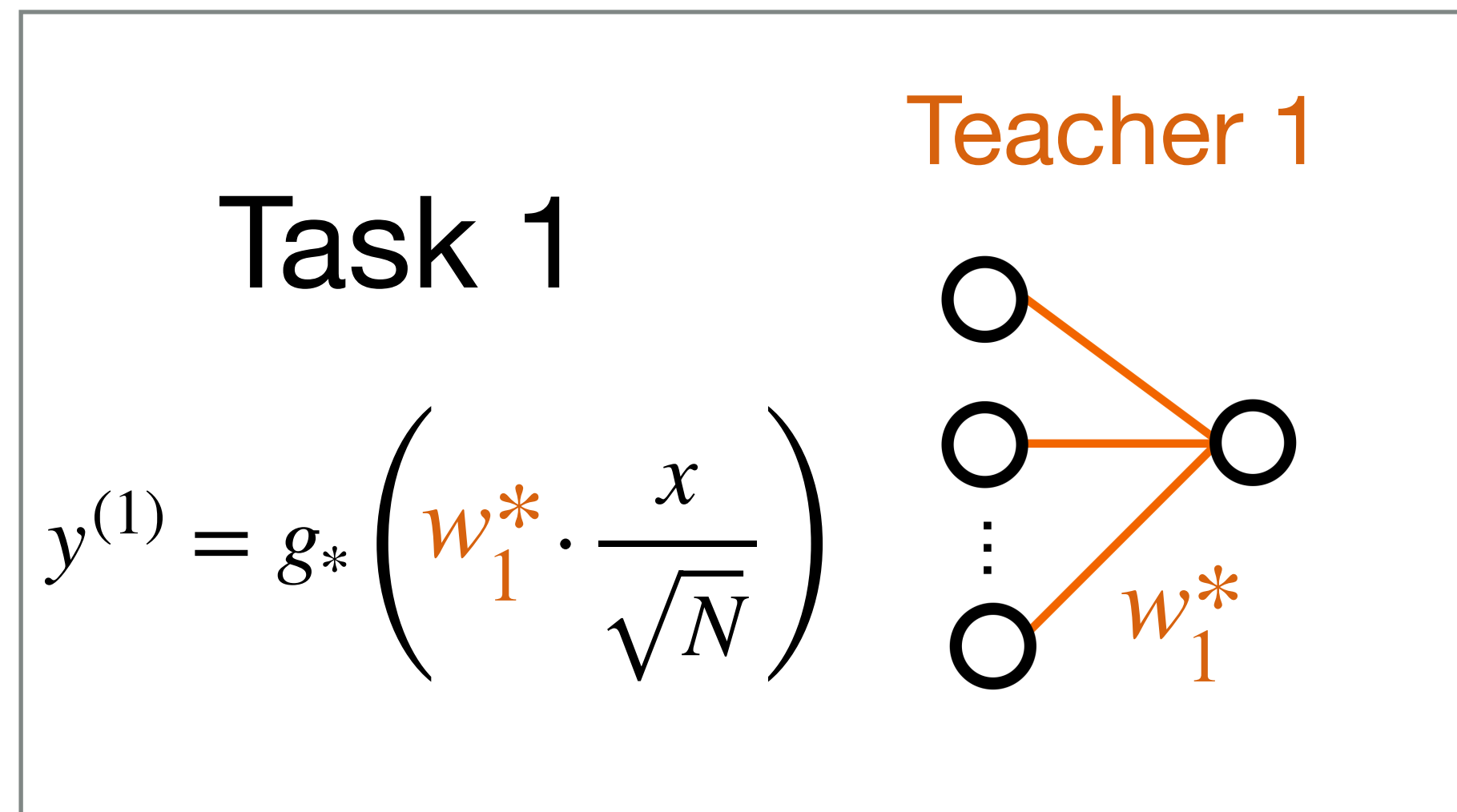
(multi-head)



$$\hat{y}^{(t)} = \sum_{k=1}^{K=2} v_k^{(t)} g \left(w_k \cdot \frac{x}{\sqrt{N}} \right)$$

A teacher-student model of continual learning

Introduced in: *Lee, Goldt, & Saxe (ICML 2021)*



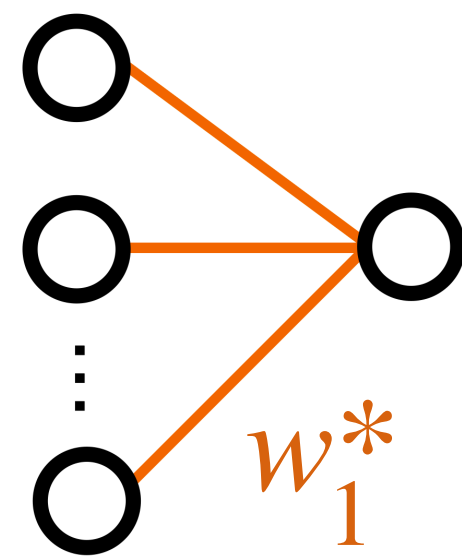
A teacher-student model of continual learning

Introduced in: *Lee, Goldt, & Saxe (ICML 2021)*

Task 1

$$y^{(1)} = g_* \left(w_1^* \cdot \frac{x}{\sqrt{N}} \right)$$

Teacher 1

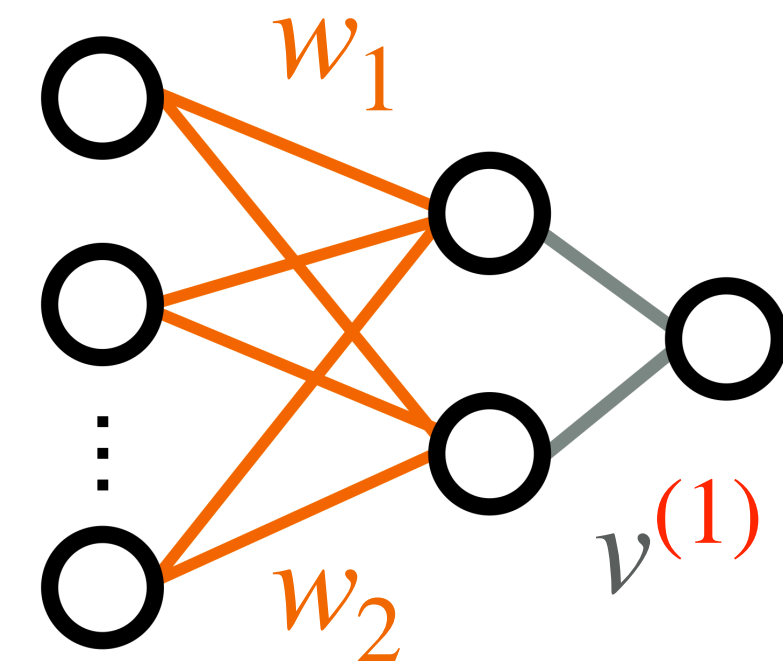


Training
phase 1

Learned ✓

Student

(multi-head)



$$\hat{y}^{(1)} = \sum_{k=1}^{K=2} v_k^{(1)} g \left(w_k \cdot \frac{x}{\sqrt{N}} \right)$$

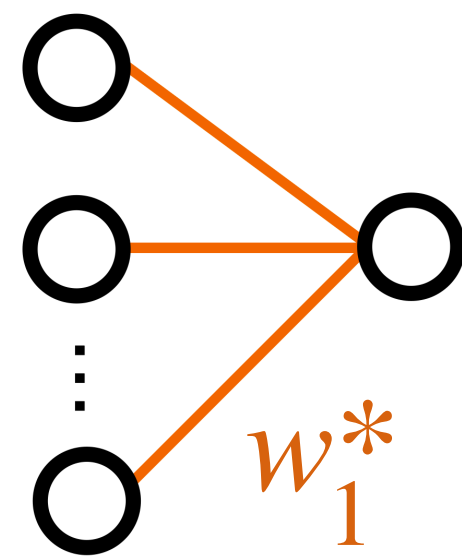
A teacher-student model of continual learning

Introduced in: *Lee, Goldt, & Saxe (ICML 2021)*

Task 1

$$y^{(1)} = g_* \left(w_1^* \cdot \frac{x}{\sqrt{N}} \right)$$

Teacher 1



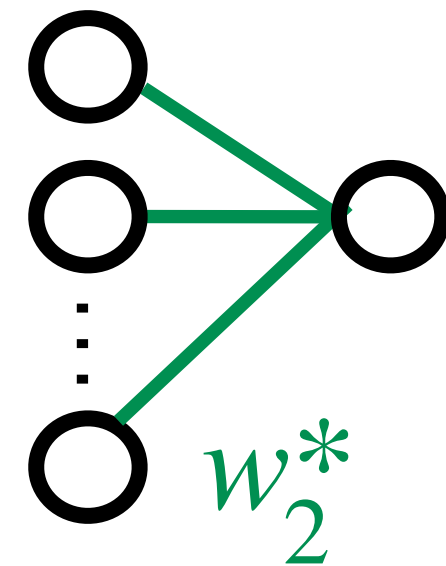
Training
phase 1

Learned ✓

Task 2

$$y^{(2)} = g_* \left(w_2^* \cdot \frac{x}{\sqrt{N}} \right)$$

Teacher 2

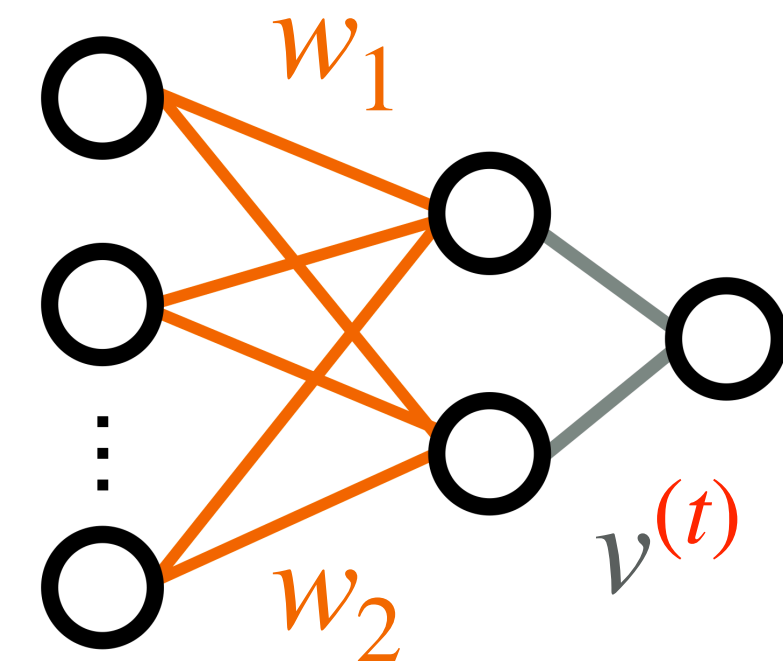


Training
phase 2

To be learned
without forgetting !!

Student

(multi-head)



$$\hat{y}^{(t)} = \sum_{k=1}^{K=2} v_k^{(t)} g \left(w_k \cdot \frac{x}{\sqrt{N}} \right)$$

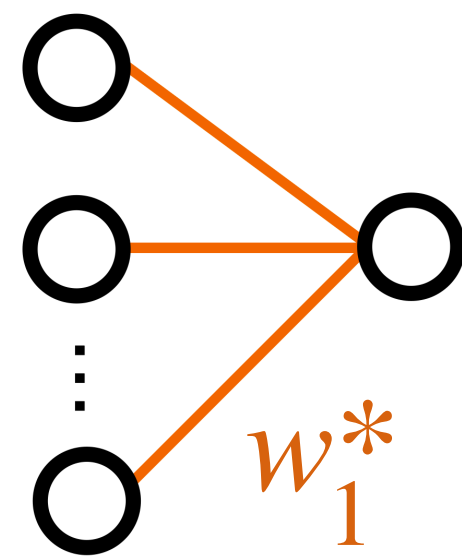
A teacher-student model of continual learning

Introduced in: *Lee, Goldt, & Saxe (ICML 2021)*

Task 1

$$y^{(1)} = g_* \left(w_1^* \cdot \frac{x}{\sqrt{N}} \right)$$

Teacher 1



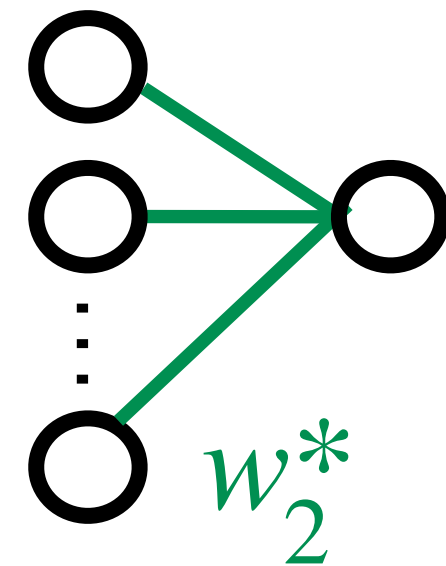
Training
phase 1

Learned ✓

Task 2

$$y^{(2)} = g_* \left(w_2^* \cdot \frac{x}{\sqrt{N}} \right)$$

Teacher 2

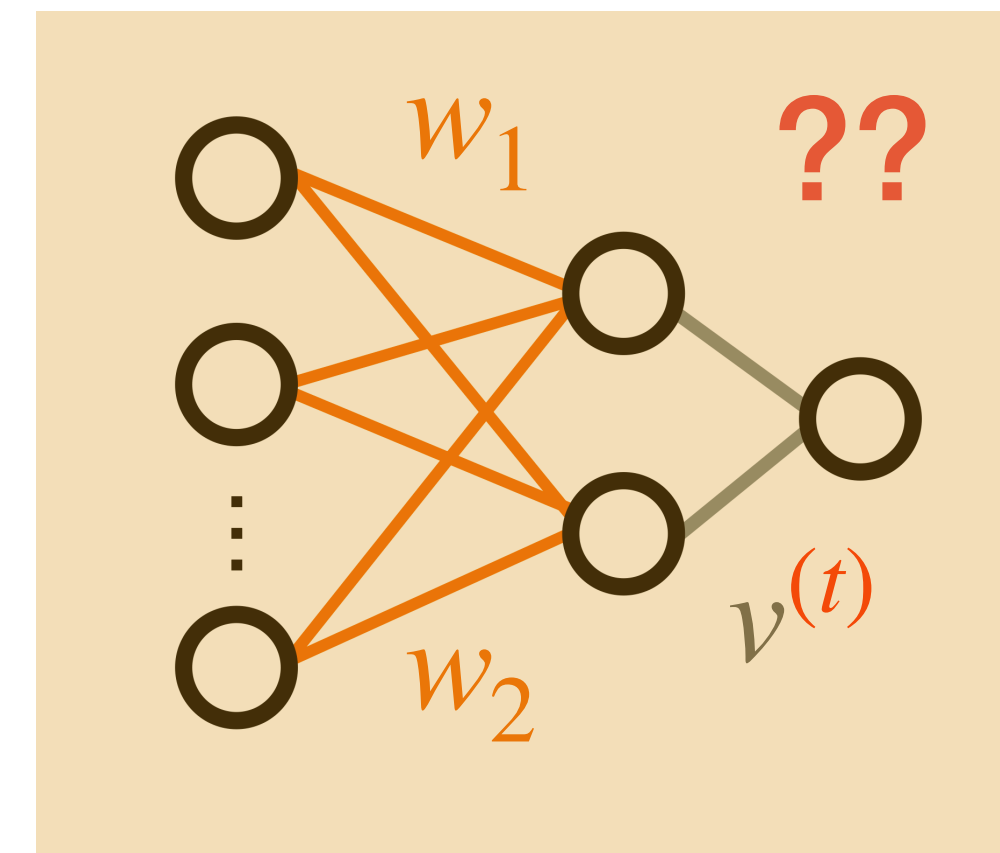


Training
phase 2

To be learned
without forgetting !!

Student

(multi-head)



$$\hat{y}^{(t)} = \sum_{k=1}^{K=2} v_k^{(t)} g \left(w_k \cdot \frac{x}{\sqrt{N}} \right)$$

1. The error depends only on the order params

Generalization error on task t :

$$\varepsilon_t(\mathbf{W}, \mathbf{V}, \mathbf{W}_*) = \frac{1}{2} \mathbb{E}_x \left[\left(y^{(t)} - \hat{y}^{(t)} \right)^2 \right]$$

$$y^{(t)} = g_* \left(w_t^* \cdot \frac{x}{\sqrt{N}} \right)$$

$$\hat{y}^{(t)} = \sum_{k=1}^{K=2} v_k^{(t)} g \left(w_k \cdot \frac{x}{\sqrt{N}} \right)$$

1. The error depends only on the order params

Generalization error on task t :

$$\begin{aligned}
 \varepsilon_t(\mathbf{W}, \mathbf{V}, \mathbf{W}_*) &= \frac{1}{2} \mathbb{E}_x \left[\left(y^{(t)} - \hat{y}^{(t)} \right)^2 \right] \\
 &= \frac{1}{2} \sum_{k,h} v_k^{(t)} v_h^{(t)} \mathbb{E}_{\boldsymbol{\lambda}, \boldsymbol{\lambda}_*} [g(\lambda_k) g(\lambda_h)] + \frac{1}{2} \mathbb{E}_{\boldsymbol{\lambda}, \boldsymbol{\lambda}_*} \left[g_*(\lambda_*^{(t)})^2 \right] \\
 &\quad - \sum_k v_k^{(t)} \mathbb{E}_{\boldsymbol{\lambda}, \boldsymbol{\lambda}_*} \left[g(\lambda_k) g_*(\lambda_*^{(t)}) \right]
 \end{aligned}$$

Pre-activations: $\lambda_k = \frac{w_k \cdot x}{\sqrt{N}}$ and $\lambda_*^{(t)} = \frac{w_*^{(t)} \cdot x}{\sqrt{N}}$

1. The error depends only on the order params

Generalization error on task t :

$$\varepsilon_t(\mathbf{W}, \mathbf{V}, \mathbf{W}_*) = \mathbb{E}_x \left[f(\lambda, \lambda_*) \right]$$

Pre-activations: $\lambda_k = \frac{w_k \cdot x}{\sqrt{N}}$ and $\lambda_*^{(t)} = \frac{w_*^{(t)} \cdot x}{\sqrt{N}}$ are *jointly Gaussian*.

1. The error depends only on the order params

Generalization error on task t :

$$\varepsilon_t(\mathbf{W}, \mathbf{V}, \mathbf{W}_*) = \mathbb{E}_{\mathbf{x}} \left[f(\lambda, \lambda_*) \right]$$

$$P(\boldsymbol{\lambda}, \boldsymbol{\lambda}_*) = \frac{1}{\sqrt{(2\pi)^{K+T} |\mathbf{C}|}} \exp \left(-\frac{1}{2} (\boldsymbol{\lambda}, \boldsymbol{\lambda}_*)^\top \mathbf{C}^{-1} (\boldsymbol{\lambda}, \boldsymbol{\lambda}_*) \right)$$

$$\mathbf{C} = \begin{pmatrix} \mathbf{Q} & \mathbf{M} \\ \mathbf{M}^\top & \mathbf{S} \end{pmatrix}.$$

Pre-activations: $\lambda_k = \frac{w_k \cdot x}{\sqrt{N}}$ and $\lambda_*^{(t)} = \frac{w_*^{(t)} \cdot x}{\sqrt{N}}$ are *jointly Gaussian*.

The second moments are given by the *order parameters* :

$$M_{kt} := \mathbb{E}_{\mathbf{x}} \left[\lambda_k \lambda_*^{(t)} \right] = \frac{\mathbf{w}_k \cdot \mathbf{w}_*^{(t)}}{N},$$

$$Q_{kh} := \mathbb{E}_{\mathbf{x}} \left[\lambda_k \lambda_h \right] = \frac{\mathbf{w}_k \cdot \mathbf{w}_h}{N},$$

$$S_{tt'} := \mathbb{E}_{\mathbf{x}} \left[\lambda_*^{(t)} \lambda_*^{(t')} \right] = \frac{\mathbf{w}_*^{(t')} \cdot \mathbf{w}_*^{(t)}}{N}.$$

2. Exact ODEs for the order parameters

$\mu = 1, \dots, P =$ training epoch

Online SGD: $\mathbf{w}^{\mu+1} = \mathbf{w}^{\mu} - \eta \nabla_{\mathbf{w}} \mathcal{L}^{\mu}$

$$\mathcal{L} = \frac{1}{2} \left(y^{(t)} - \hat{y}^{(t)} \right)^2$$

2. Exact ODEs for the order parameters

$\mu = 1, \dots, P =$ training epoch

Online SGD:

$$\mathbf{w}_k^{\mu+1} = \mathbf{w}_k^\mu - \eta^\mu \Delta^{(t_c)\mu} v_k^{(t_c)\mu} g'(\lambda_k^\mu) \frac{\mathbf{x}^\mu}{\sqrt{N}},$$

$$\Delta^{(t)\mu} := \hat{y}^{(t)\mu} - y^{(t)\mu} = \sum_{k=1}^K v_k^{(t)} g(\lambda_k^\mu) - g_*(\lambda_*^{(t)\mu})$$

Pre-activations: $\lambda_k^\mu := \frac{\mathbf{x}^\mu \cdot \mathbf{w}_k^\mu}{\sqrt{N}}, \quad \lambda_*^{(t)\mu} := \frac{\mathbf{x}^\mu \cdot \mathbf{w}_*^{(t)}}{\sqrt{N}}.$

2. Exact ODEs for the order parameters

$\mu = 1, \dots, P =$ training epoch

Online SGD:

$$\mathbf{w}_k^{\mu+1} = \mathbf{w}_k^\mu - \eta^\mu \Delta^{(t_c)\mu} v_k^{(t_c)\mu} g'(\lambda_k^\mu) \frac{\mathbf{x}^\mu}{\sqrt{N}},$$

Example derivation: ODE for the “magnetization”

$$M_{kt} := \mathbb{E}_{\mathbf{x}} \left[\lambda_k \lambda_*^{(t)} \right] = \frac{\mathbf{w}_k \cdot \mathbf{w}_*^{(t)}}{N}$$

$$\frac{\mathbf{w}_k^{\mu+1} \cdot \mathbf{w}_*^{(t)}}{N} - \frac{\mathbf{w}_k^\mu \cdot \mathbf{w}_*^{(t)}}{N} = -\frac{\eta^\mu}{N} \Delta^{(t_c)\mu} v_k^{(t_c)\mu} g'(\lambda_k^\mu) \lambda_*^{(t)\mu}$$

2. Exact ODEs for the order parameters

$\mu = 1, \dots, P =$ training epoch

Online SGD:

$$\mathbf{w}_k^{\mu+1} = \mathbf{w}_k^\mu - \eta^\mu \Delta^{(t_c)\mu} v_k^{(t_c)\mu} g'(\lambda_k^\mu) \frac{\mathbf{x}^\mu}{\sqrt{N}},$$

Example derivation: ODE for the “magnetization” $M_{kt} := \mathbb{E}_{\mathbf{x}} [\lambda_k \lambda_*^{(t)}] = \frac{\mathbf{w}_k \cdot \mathbf{w}_*^{(t)}}{N}$

$$\frac{\mathbf{w}_k^{\mu+1} \cdot \mathbf{w}_*^{(t)}}{N} - \frac{\mathbf{w}_k^\mu \cdot \mathbf{w}_*^{(t)}}{N} = -\frac{\eta^\mu}{N} \Delta^{(t_c)\mu} v_k^{(t_c)\mu} g'(\lambda_k^\mu) \lambda_*^{(t)\mu}$$

$\alpha = \mu/N =$ training “time”, high-dimensional limit with $P/N = \mathcal{O}_N(1)$:

$$\frac{dM_{kt}}{d\alpha} = -\eta v_k^{(t_c)} \mathbb{E}_{\lambda, \lambda_*} \left[\Delta^{(t_c)} g'(\lambda_k) \lambda_*^{(t)} \right] := f_{\mathbf{M}, kt}$$

Saad & Solla (1995); Biehl & Schwarze (1995);

2. Exact ODEs for the order parameters

ODEs for the order parameters:

$$\mathbb{Q} = (\text{vec}(\mathbf{Q}), \text{vec}(\mathbf{M}), \text{vec}(\mathbf{V}))^\top$$

$$\frac{d\mathbb{Q}(\alpha)}{d\alpha} = f_{\mathbb{Q}}(\mathbb{Q}(\alpha), \mathbf{u}(\alpha))$$

$$\alpha \in (0, \alpha_F]$$

Control

- Task
- Learning rate
- ...

3. Meta-optimization

ODEs for the order parameters:

$$\mathbb{Q} = (\text{vec}(\mathbf{Q}), \text{vec}(\mathbf{M}), \text{vec}(\mathbf{V}))^\top$$

$$\frac{d\mathbb{Q}(\alpha)}{d\alpha} = f_{\mathbb{Q}}(\mathbb{Q}(\alpha), \mathbf{u}(\alpha))$$

$$\alpha \in (0, \alpha_F]$$

Forward dynamics

Optimization objective:

$$h(\mathbb{Q}(\alpha_F)) = \sum_{t=1}^T c_t \varepsilon_t(\mathbb{Q}(\alpha_F))$$

with $c_t \geq 0$ and $\sum_{t=1}^T c_t = 1$

Cost functional: $\mathcal{F}[\mathbb{Q}, \hat{\mathbb{Q}}, \mathbf{u}] = h(\mathbb{Q}(\alpha_F)) + \int_0^{\alpha_F} d\alpha \hat{\mathbb{Q}}(\alpha)^\top \left[-\frac{d\mathbb{Q}(\alpha)}{d\alpha} + f_{\mathbb{Q}}(\mathbb{Q}(\alpha), \mathbf{u}(\alpha)) \right]$

[Pontryagin (1962)]

3. Meta-optimization

ODEs for the order parameters:

$$\mathbb{Q} = (\text{vec}(\mathbf{Q}), \text{vec}(\mathbf{M}), \text{vec}(\mathbf{V}))^\top$$

$$\frac{d\mathbb{Q}(\alpha)}{d\alpha} = f_{\mathbb{Q}}(\mathbb{Q}(\alpha), \mathbf{u}(\alpha))$$

$$\alpha \in (0, \alpha_F]$$

Forward dynamics

Optimization objective:

$$h(\mathbb{Q}(\alpha_F)) = \sum_{t=1}^T c_t \varepsilon_t(\mathbb{Q}(\alpha_F))$$

with $c_t \geq 0$ and $\sum_{t=1}^T c_t = 1$

Cost functional: $\mathcal{F}[\mathbb{Q}, \hat{\mathbb{Q}}, \mathbf{u}] = h(\mathbb{Q}(\alpha_F)) + \int_0^{\alpha_F} d\alpha \hat{\mathbb{Q}}(\alpha)^\top \left[-\frac{d\mathbb{Q}(\alpha)}{d\alpha} + f_{\mathbb{Q}}(\mathbb{Q}(\alpha), \mathbf{u}(\alpha)) \right]$

[Pontryagin (1962)]

$$-\frac{d\hat{\mathbb{Q}}(\alpha)^\top}{d\alpha} = \hat{\mathbb{Q}}(\alpha)^\top \nabla_{\mathbb{Q}} f_{\mathbb{Q}}(\mathbb{Q}(\alpha), \mathbf{u}(\alpha))$$

$$\hat{\mathbb{Q}}(\alpha_F) = \nabla_{\mathbb{Q}} h(\mathbb{Q}(\alpha_F))$$

Backwards dynamics

3. Meta-optimization

ODEs for the order parameters:

$$\mathbb{Q} = (\text{vec}(\mathbf{Q}), \text{vec}(\mathbf{M}), \text{vec}(\mathbf{V}))^\top$$

$$\frac{d\mathbb{Q}(\alpha)}{d\alpha} = f_{\mathbb{Q}}(\mathbb{Q}(\alpha), \mathbf{u}(\alpha))$$

$$\alpha \in (0, \alpha_F]$$

Forward dynamics

Optimization objective:

$$h(\mathbb{Q}(\alpha_F)) = \sum_{t=1}^T c_t \varepsilon_t(\mathbb{Q}(\alpha_F))$$

with $c_t \geq 0$ and $\sum_{t=1}^T c_t = 1$

Cost functional: $\mathcal{F}[\mathbb{Q}, \hat{\mathbb{Q}}, \mathbf{u}] = h(\mathbb{Q}(\alpha_F)) + \int_0^{\alpha_F} d\alpha \hat{\mathbb{Q}}(\alpha)^\top \left[-\frac{d\mathbb{Q}(\alpha)}{d\alpha} + f_{\mathbb{Q}}(\mathbb{Q}(\alpha), \mathbf{u}(\alpha)) \right]$

[Pontryagin (1962)]

$$-\frac{d\hat{\mathbb{Q}}(\alpha)^\top}{d\alpha} = \hat{\mathbb{Q}}(\alpha)^\top \nabla_{\mathbb{Q}} f_{\mathbb{Q}}(\mathbb{Q}(\alpha), \mathbf{u}(\alpha))$$

$$\hat{\mathbb{Q}}(\alpha_F) = \nabla_{\mathbb{Q}} h(\mathbb{Q}(\alpha_F))$$

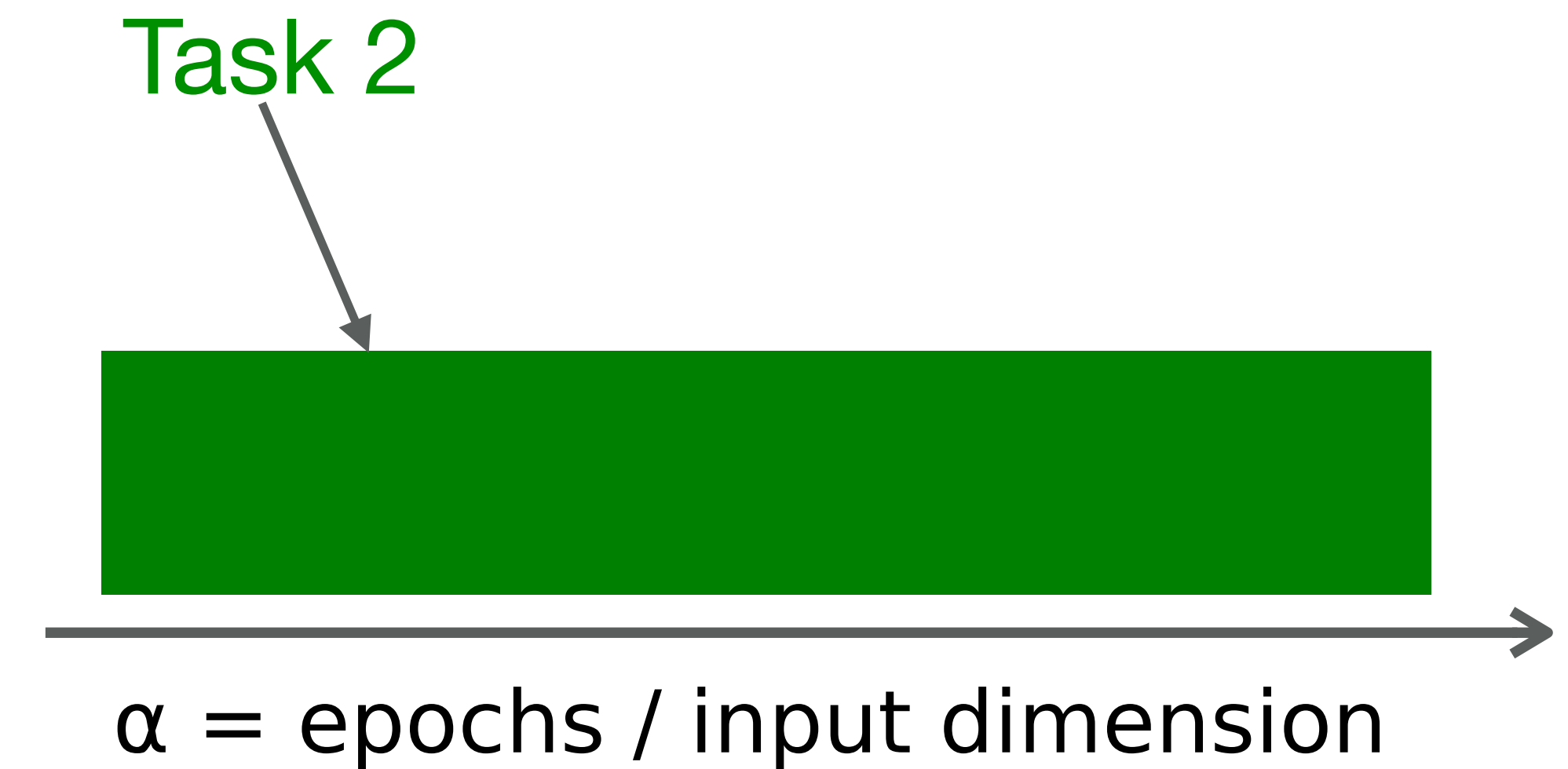
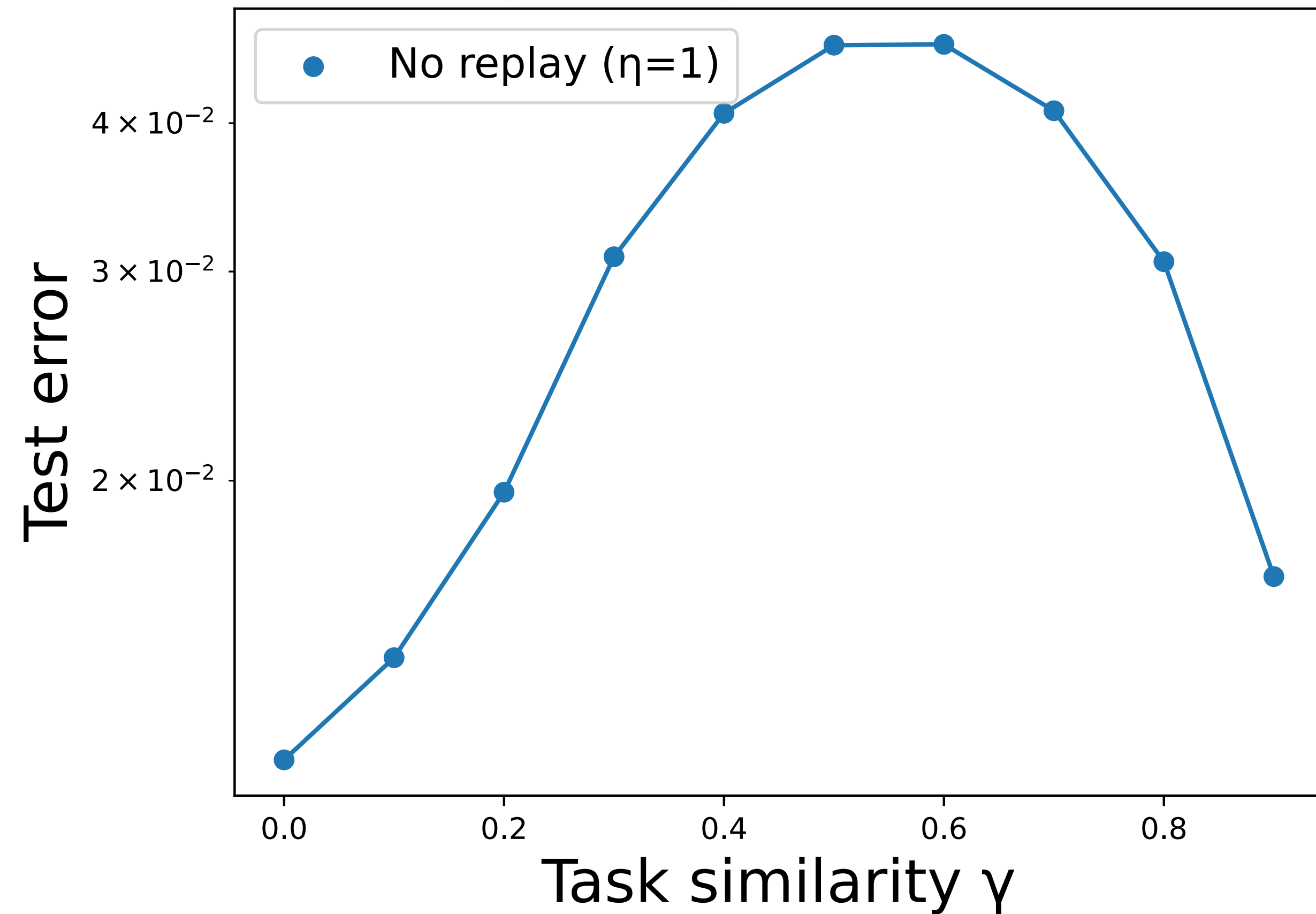
Backwards dynamics

Optimal protocol: $\mathbf{u}^*(\alpha) = \text{argmin}_{\mathbf{u}} \left\{ \hat{\mathbb{Q}}(\alpha)^\top f_{\mathbb{Q}}(\mathbb{Q}(\alpha), \mathbf{u}) \right\}$

Training only on the new task

Control: $\mathbf{u} = \text{task}$

$\alpha = \text{epochs} / \text{input dimension} = 25.0$



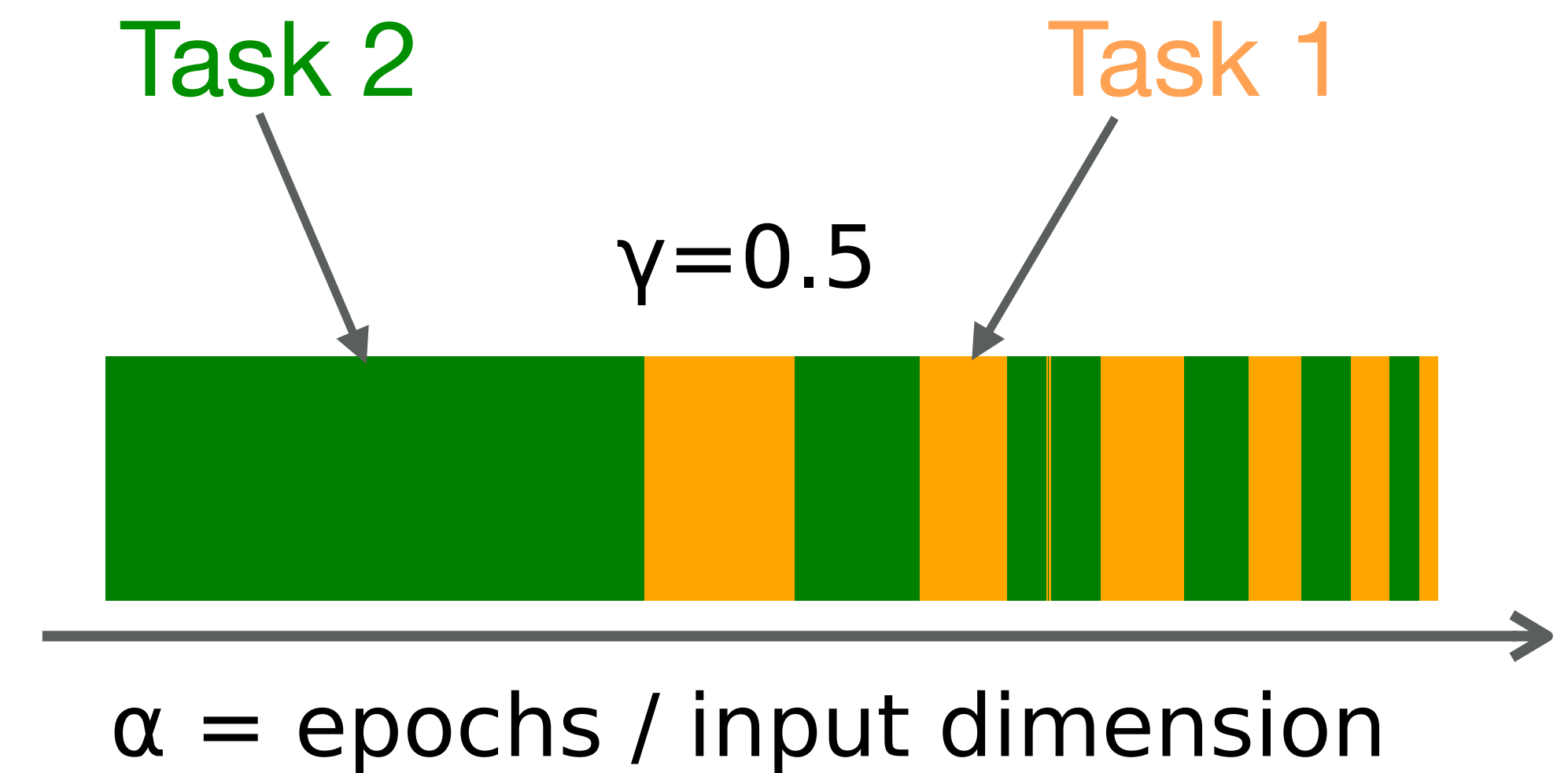
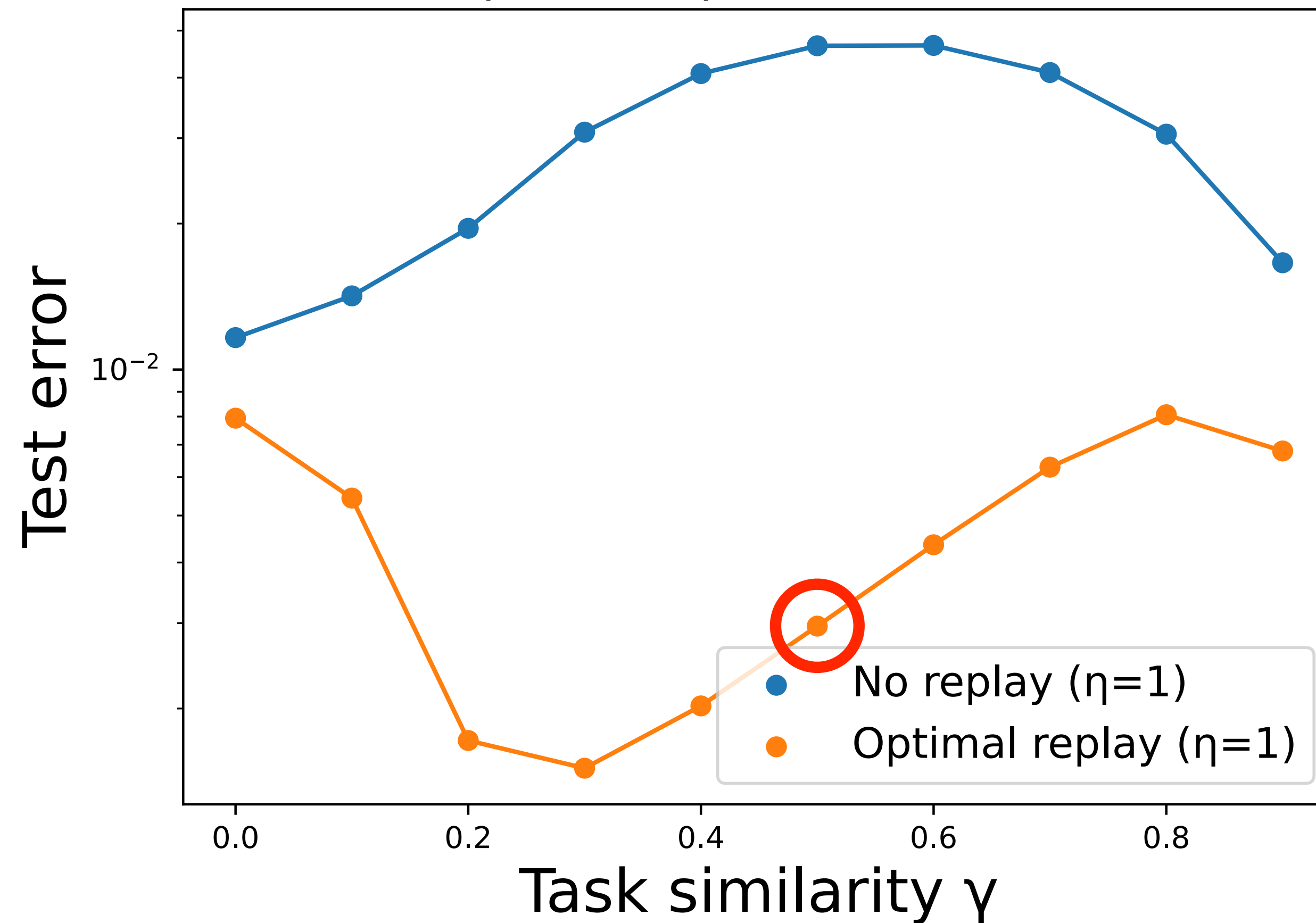
Intermediate task similarity leads to the worst forgetting.

[Lee et al (2021); Ramesh et al (2020); Doan et al (2020); Nguyen et al (2019)]

The impact of replay on the performance

Control: $\mathbf{u} = \text{task}$

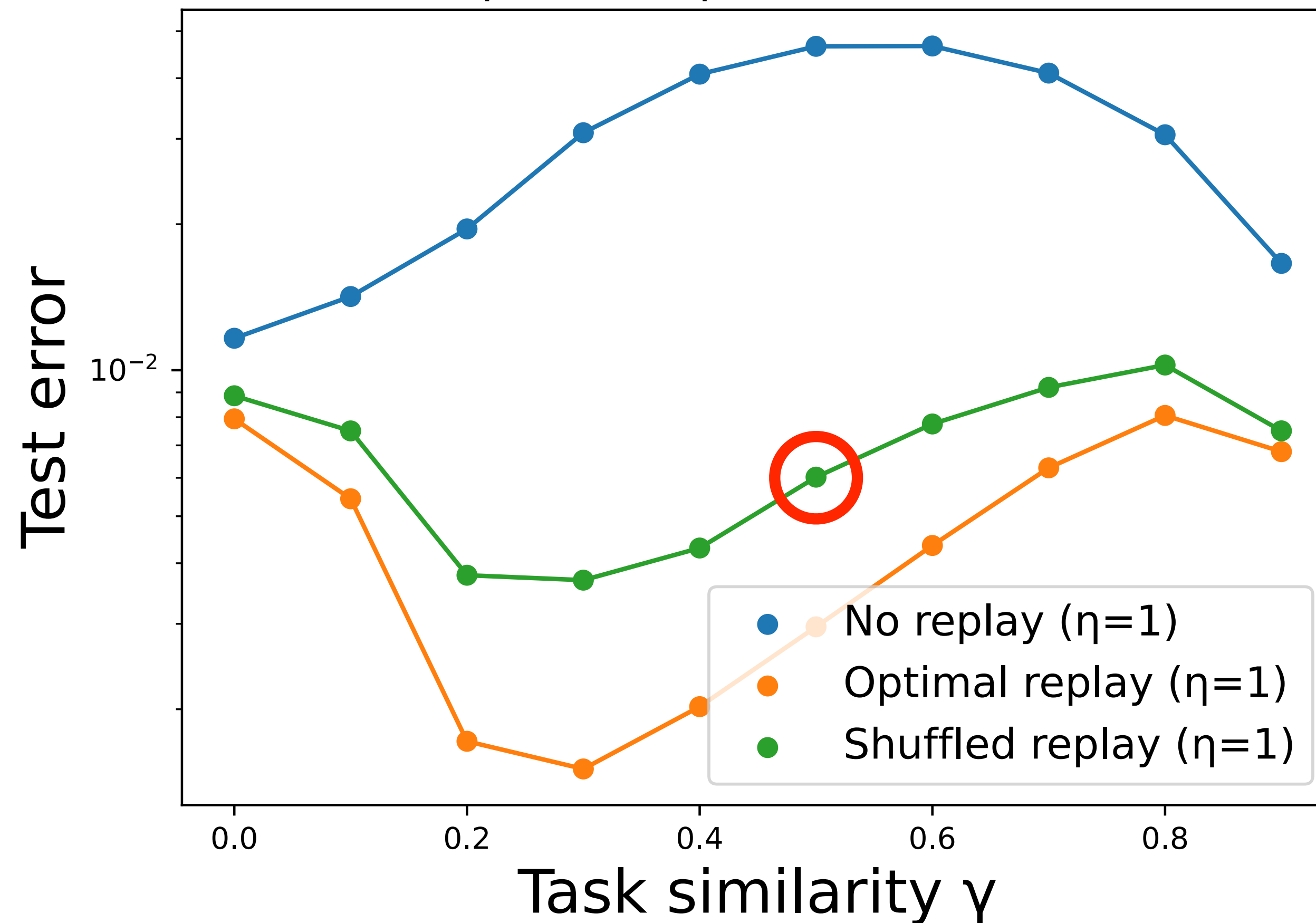
$\alpha = \text{epochs} / \text{input dimension} = 25.0$



The impact of replay on the performance

Control: $\mathbf{u} = \text{task}$

$\alpha = \text{epochs} / \text{input dimension} = 25.0$



$\gamma = 0.5$



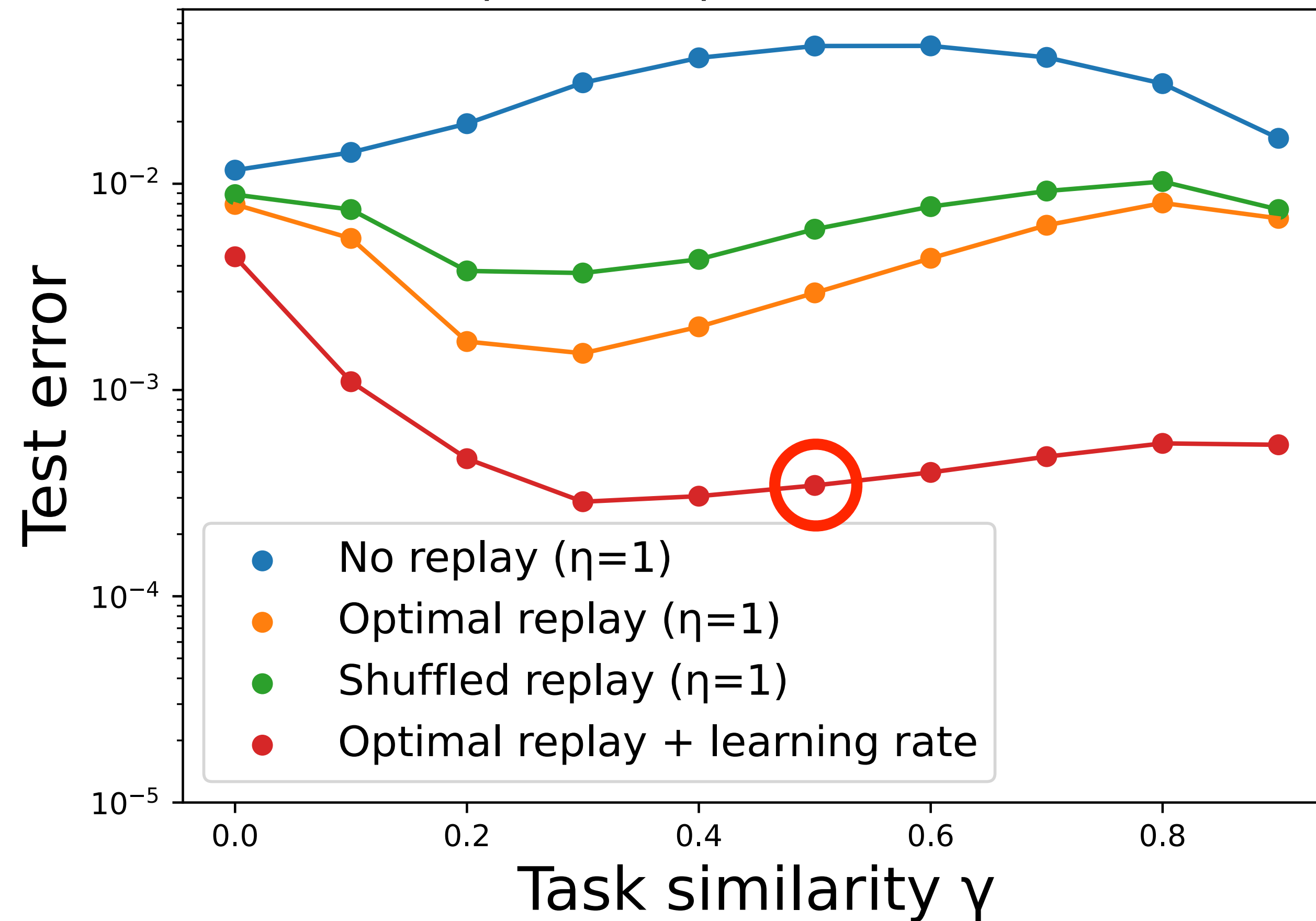
$\alpha = \text{epochs} / \text{input dimension}$

Order matters!

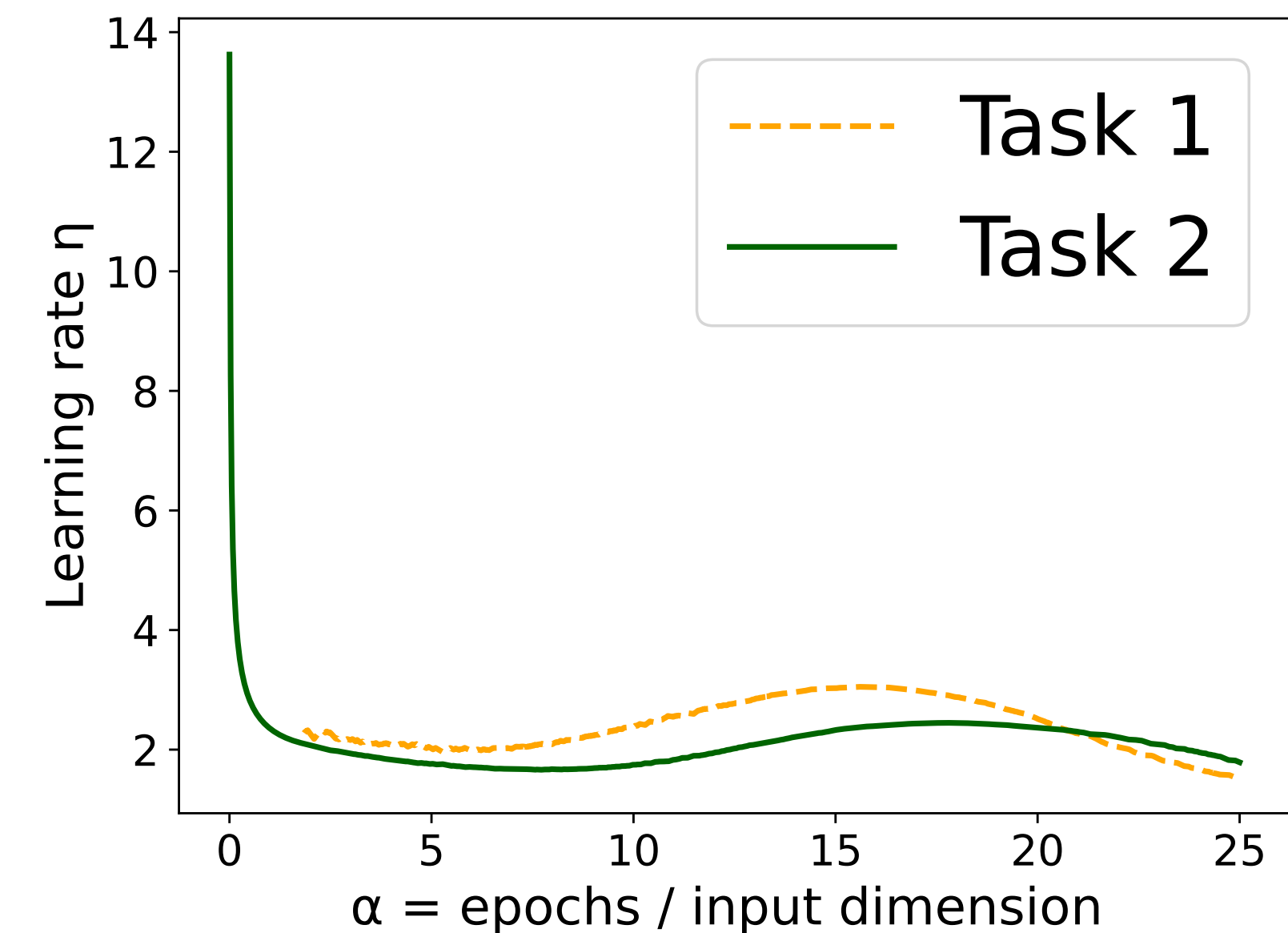
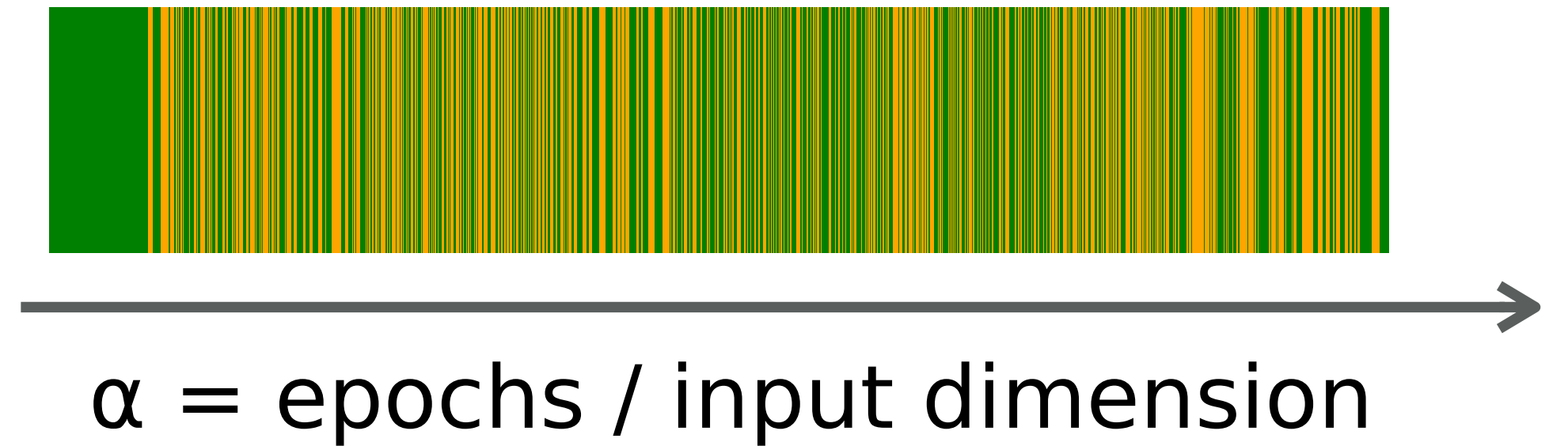
Optimizing replay + learning rate schedule

Control: $\mathbf{u} = \text{task}$

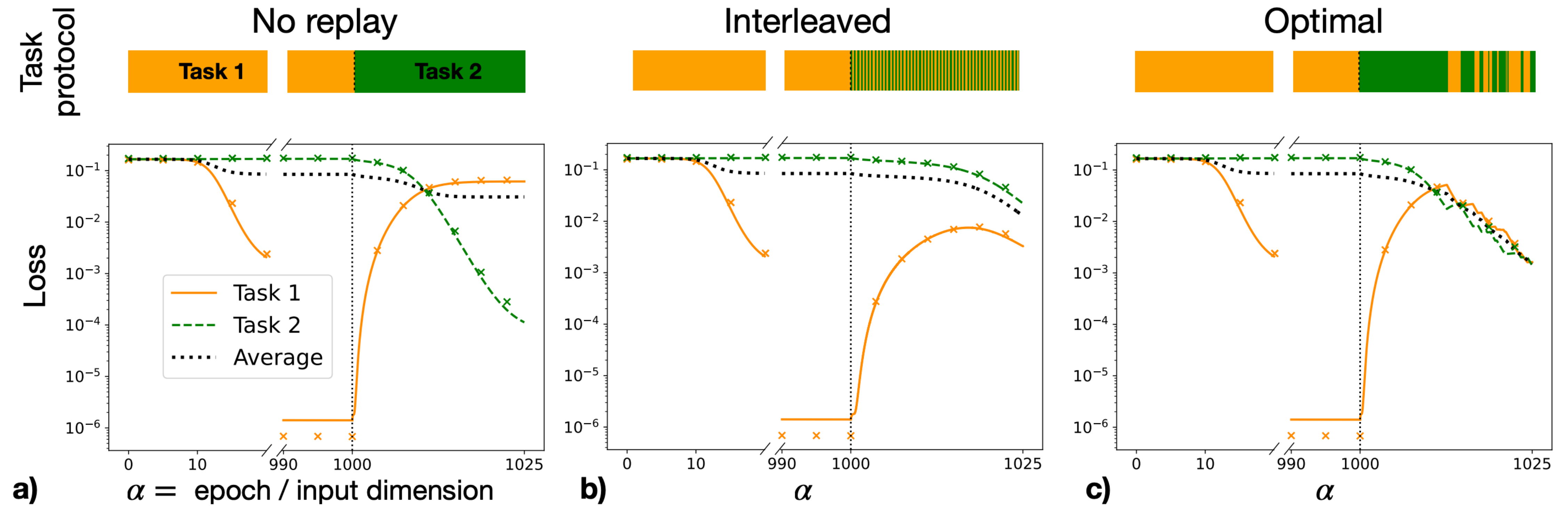
$\alpha = \text{epochs} / \text{input dimension} = 25.0$



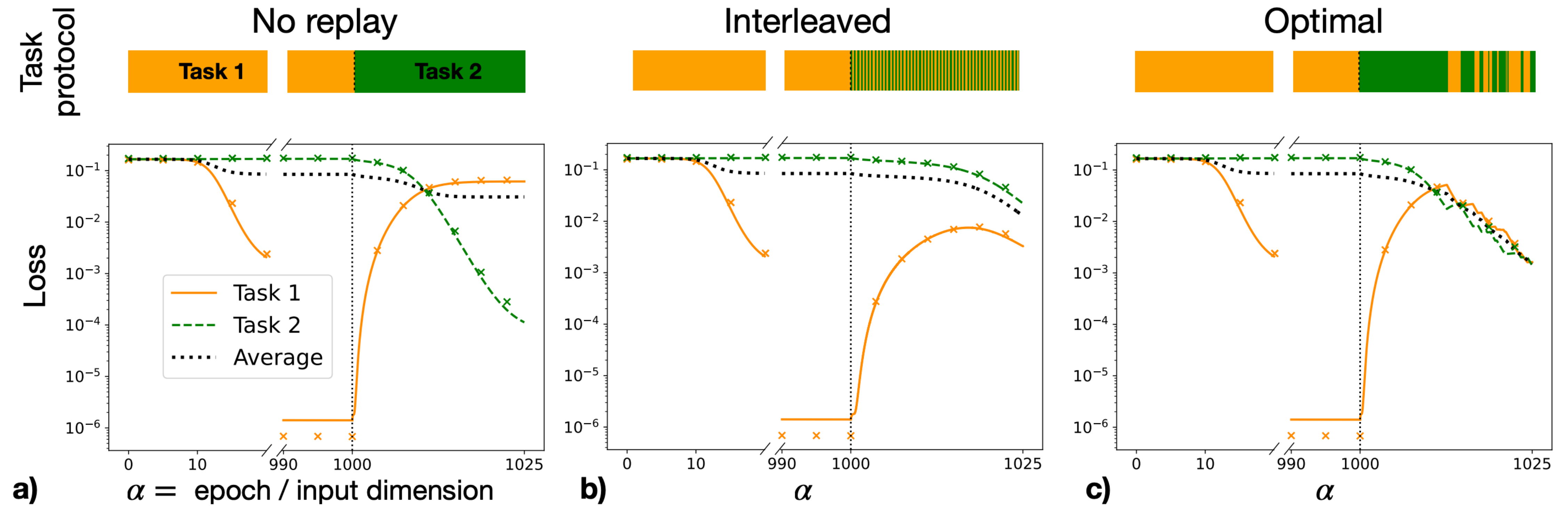
$\gamma=0.5$



Results: optimal strategy vs benchmarks

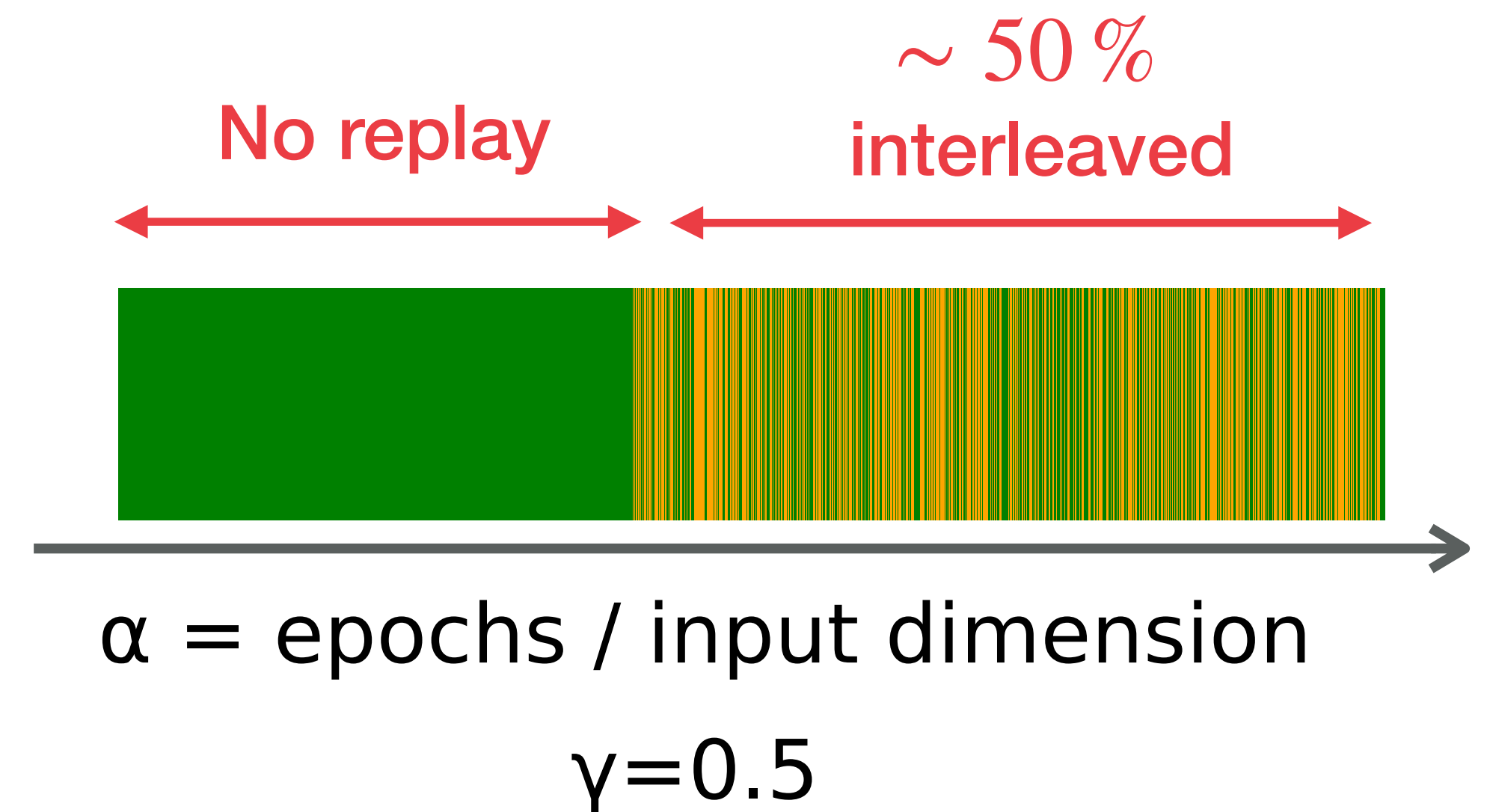
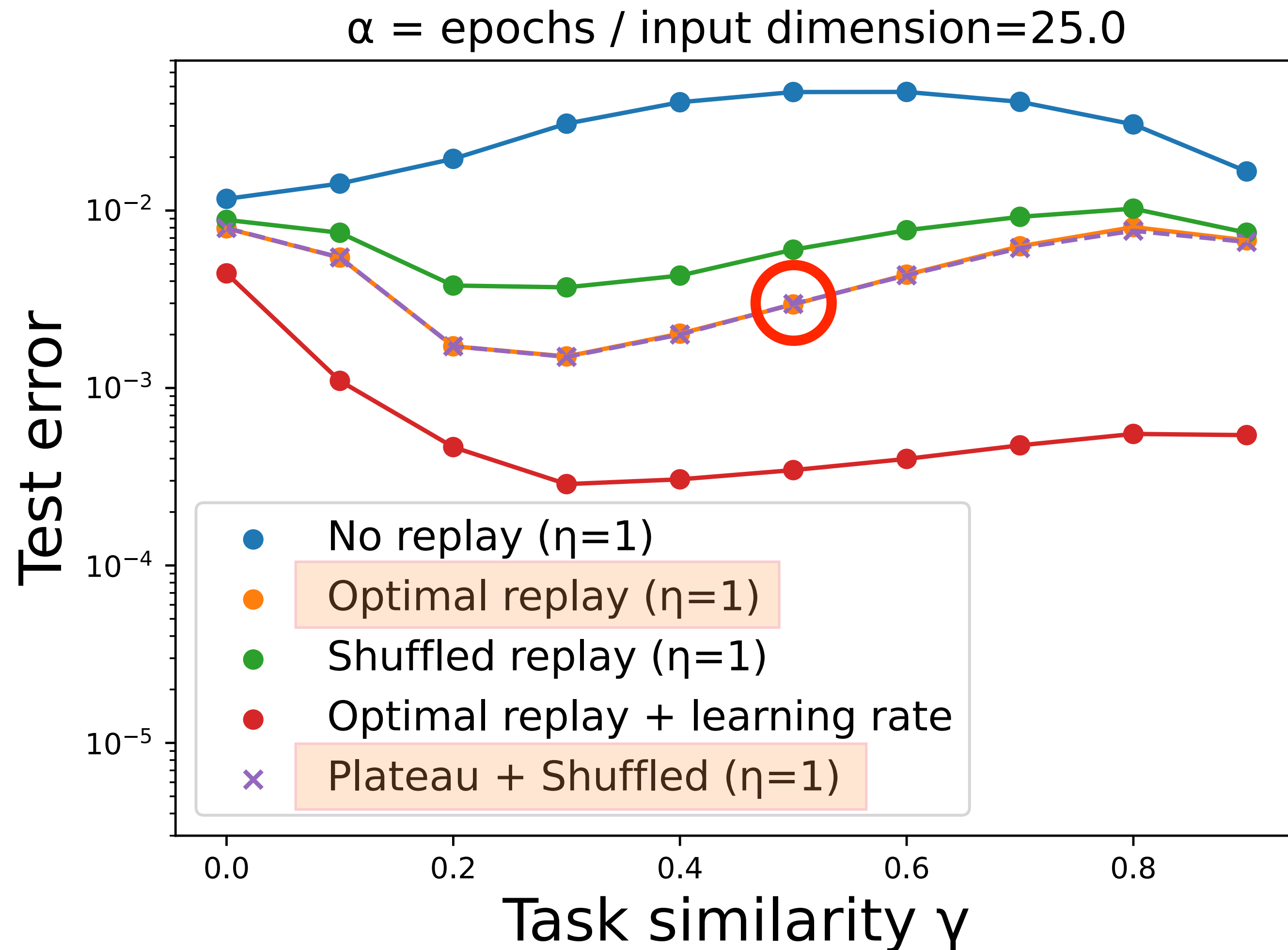


Results: optimal strategy vs benchmarks



Does the order of replayed episodes matter?

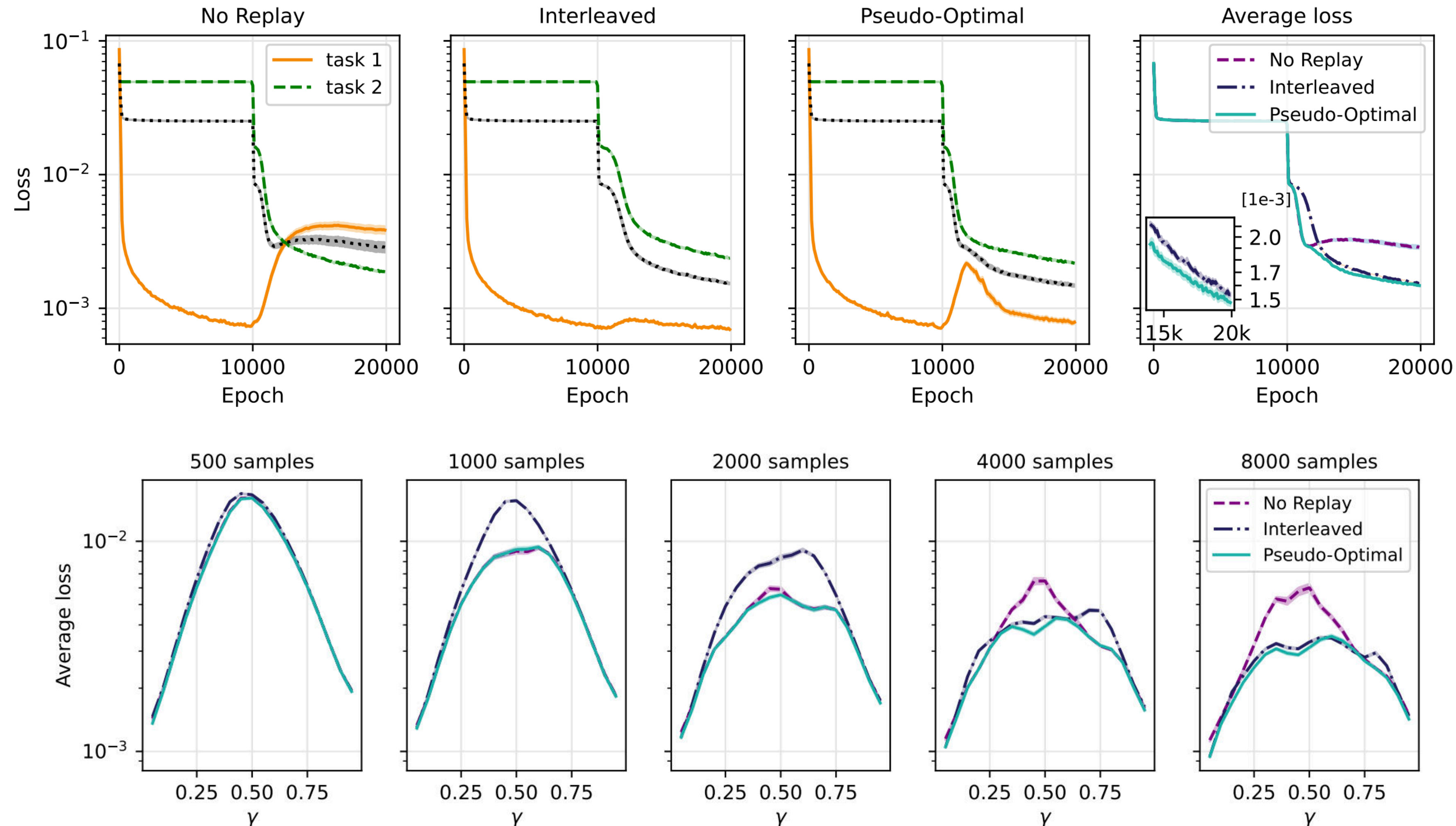
Learning the new vs Replaying the old



**Order does not matter
after the no-replay phase!**

Experiments on Fashion MNIST

$$\mathcal{D}_1 = \{\mathbf{x}_i^{(1)}, y_i^{(1)}\}_i \quad \mathcal{D}_2 = \{\mathbf{x}_i^{(2)}, y_i^{(2)}\}_i = \{\gamma \mathbf{x}_i^{(1)} + (1 - \gamma) \tilde{\mathbf{x}}_i, \gamma y_i^{(1)} + (1 - \gamma) \tilde{y}_i\}_i$$



In summary:

- Optimal control of effective learning dynamics reveals **nontrivial training protocols**.
- Continual learning: **non-homogeneous replay** avoids forgetting.

Many open directions!

- Experiments beyond toy models: incorporate models of structured data;
- Batch learning;
- Other learning paradigms : shaping, transfer learning, active learning, ...

Thank you!

Bonus Slides

Curriculum learning

(in progress)

Curriculum learning

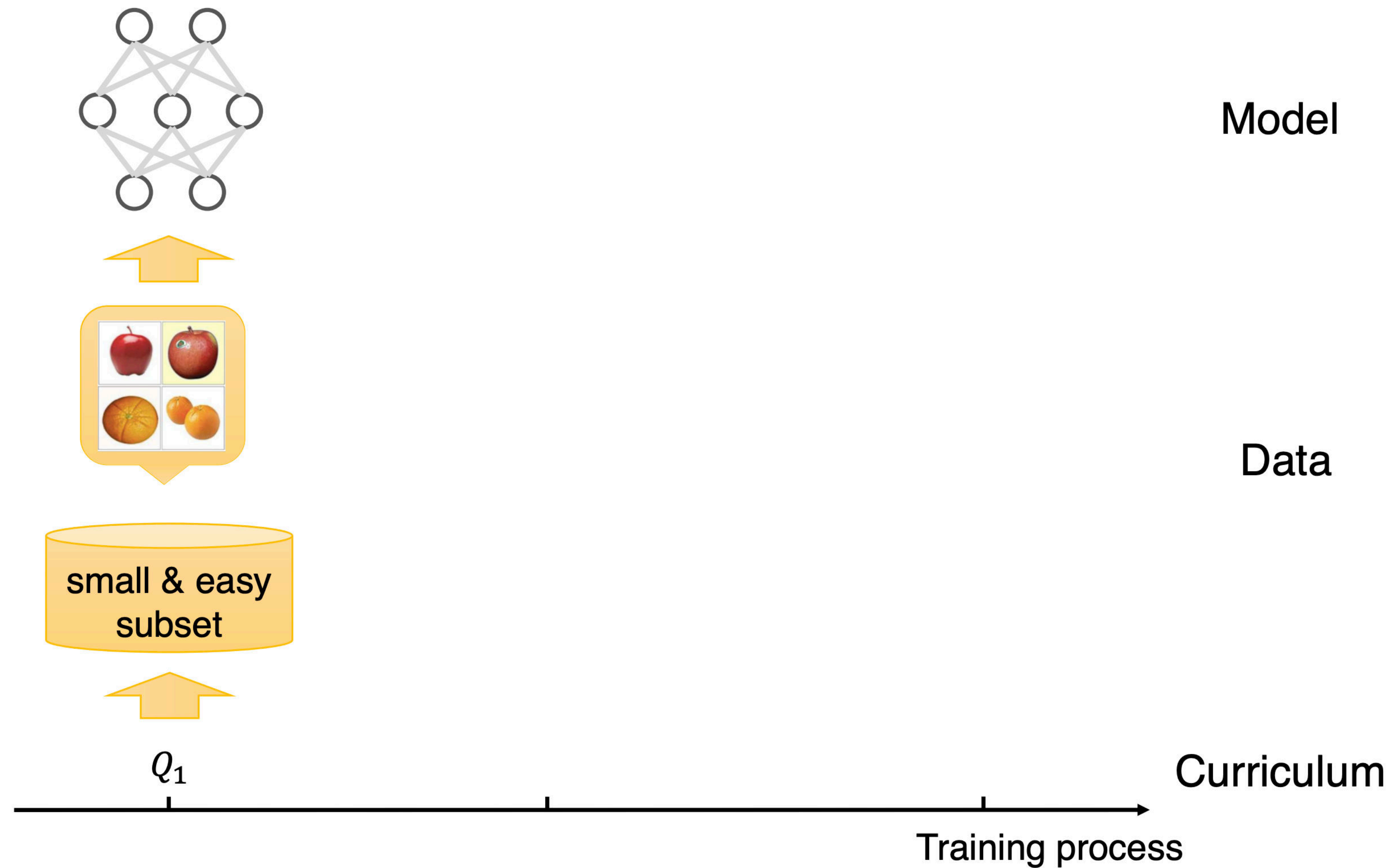


Image from: Wang, Xin, Yudong Chen, and Wenwu Zhu. *IEEE transactions on pattern analysis and machine intelligence* 44.9 (2021):4555-4576.

Curriculum learning

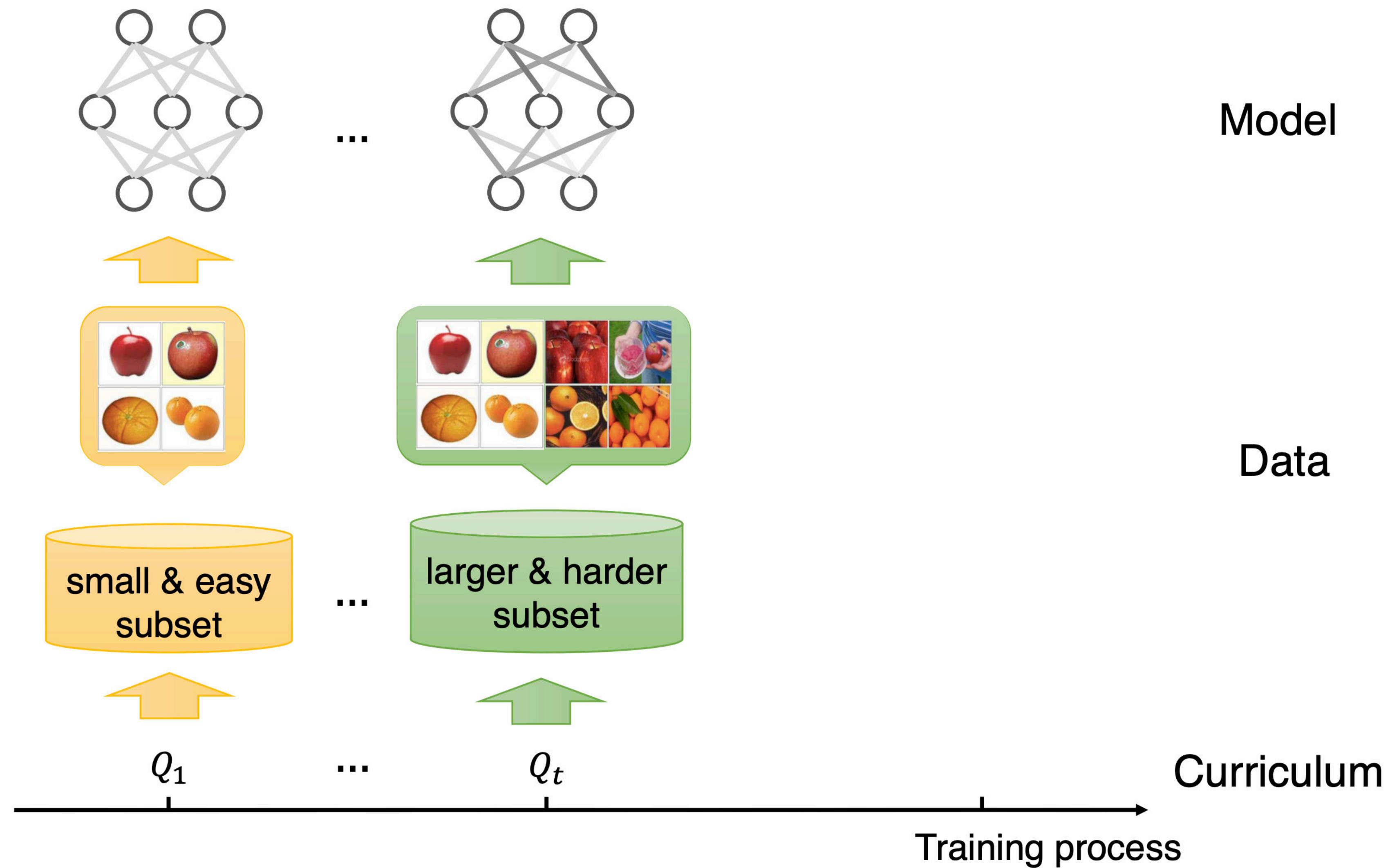


Image from: Wang, Xin, Yudong Chen, and Wenwu Zhu. *IEEE transactions on pattern analysis and machine intelligence* 44.9 (2021):4555-4576.

Curriculum learning

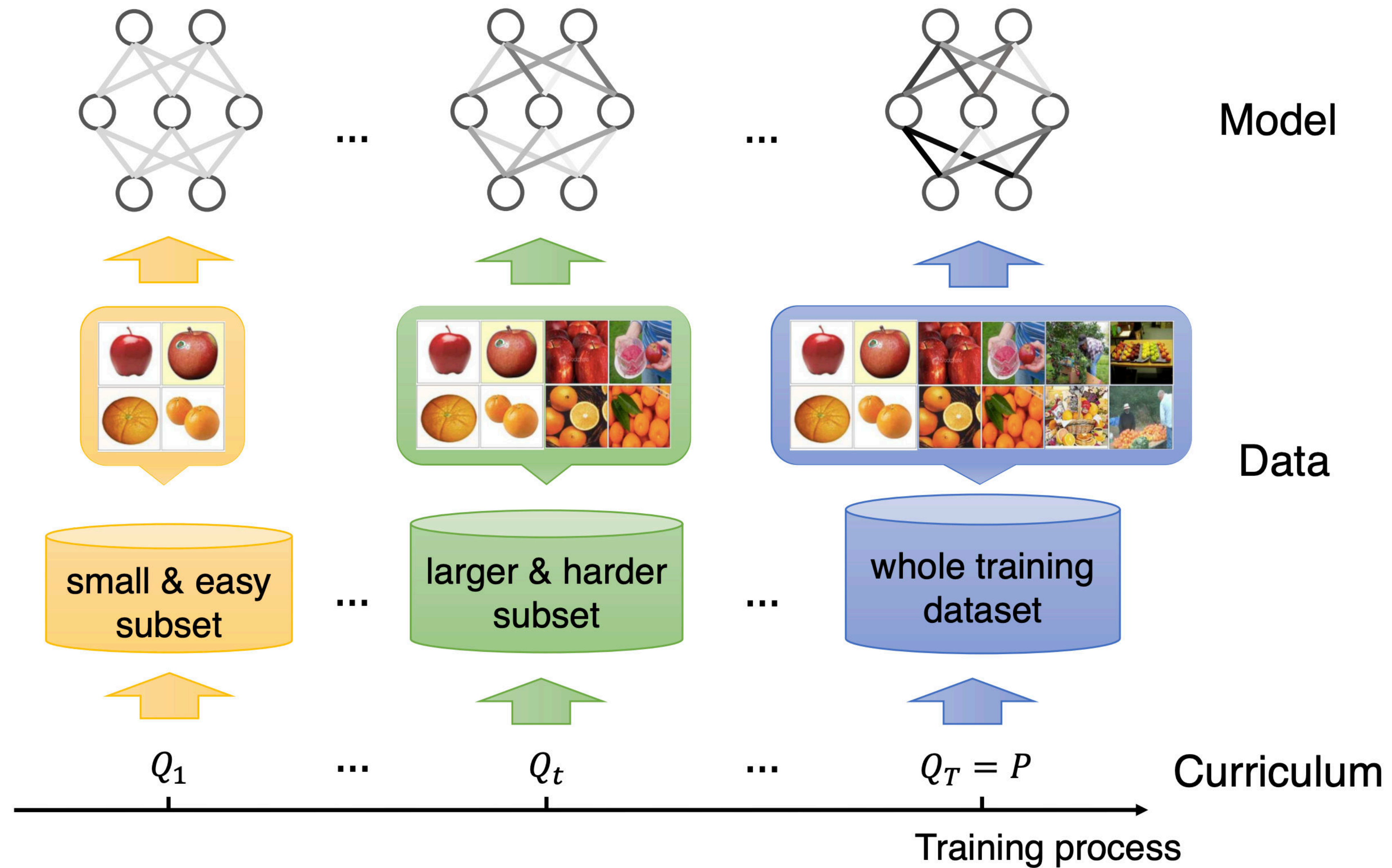


Image from: Wang, Xin, Yudong Chen, and Wenwu Zhu. *IEEE transactions on pattern analysis and machine intelligence* 44.9 (2021):4555-4576.

Animals:

- Conditional reflexes (dogs) [Pavlov (1927)]
- Shaping (rats, pigeons) [Skinner (1938)]
- Discrimination along a continuum (rats) [Lawrence (1952)]
- Cross-species auditory identification (rats, humans) [Liu et al. (2008)]

Humans:

- Discrimination along a continuum [Baker, Stanley (1954)]
- Past tense [Plunkett et al (1990; 1991)]
- Fading with auditory and visual stimuli [Pashler, Mozer (2013)]
- Eureka effect [Ahissar, Hochstein (1997)]

Animals:

- Conditional reflexes (dogs) [Pavlov (1927)]
- Shaping (rats, pigeons) [Skinner (1938)]
- Discrimination along a continuum (rats) [Lawrence (1952)]
- Cross-species auditory identification (rats, humans) [Liu et al. (2008)]

ML (empirical):

- Easy-to-hard training [Bengio et al (2009)]
- Anti-curriculum [Zhang et al (2019); Hachohen & Weinshall (2019)]
- No effect of CL in vision benchmarks [Wu et al (2020)]
- Convincing results for LLMs and RL [Brown et al (2020); Narvekar et al (2020)]

Humans:

- Discrimination along a continuum [Baker, Stanley (1954)]
- Past tense [Plunkett et al (1990; 1991)]
- Fading with auditory and visual stimuli [Pashler, Mozer (2013)]
- Eureka effect [Ahissar, Hochstein (1997)]

ML (theory):

- Single-update advantage of easy samples [Weinshall et al (2020)]
- Speed benefit but limited performance upgrade in convex problems [Saglietti et al (2021); Lee et al. (2024)]
- Computational benefit in parity machines [Abbe et al. (2023); Cornacchia et al (2023)]
- Asymptotic benefit in non-convex models [Mannelli et al (2024)]

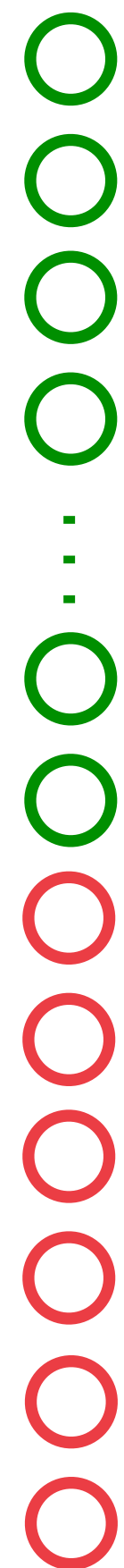
A teacher-student model of curriculum learning

Introduced in: *Bengio, et al. (ICML 2009)*, *Saglietti, et al. (NeurIPS 2022)*

Input: $\mathbf{x} = (\mathbf{x}_r, \mathbf{x}_i) \in \mathbb{R}^N$

Relevant

Irrelevant



$$\mathbf{x}_r \in \mathbb{R}^{\rho N}$$

Unit variance

$$\mathbf{x}_i \in \mathbb{R}^{(1-\rho)N}$$

Variance Δ

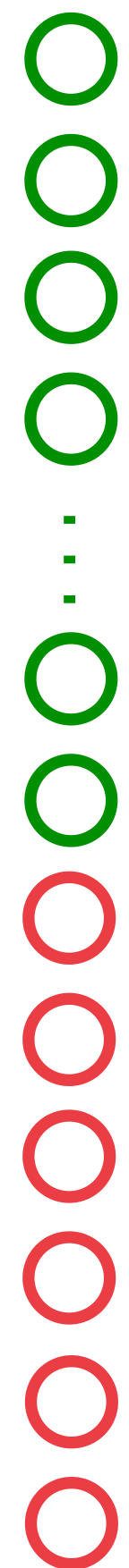
A teacher-student model of curriculum learning

Introduced in: *Bengio, et al. (ICML 2009)*, *Saglietti, et al. (NeurIPS 2022)*

Input: $\mathbf{x} = (\mathbf{x}_r, \mathbf{x}_i) \in \mathbb{R}^N$

Relevant

Irrelevant



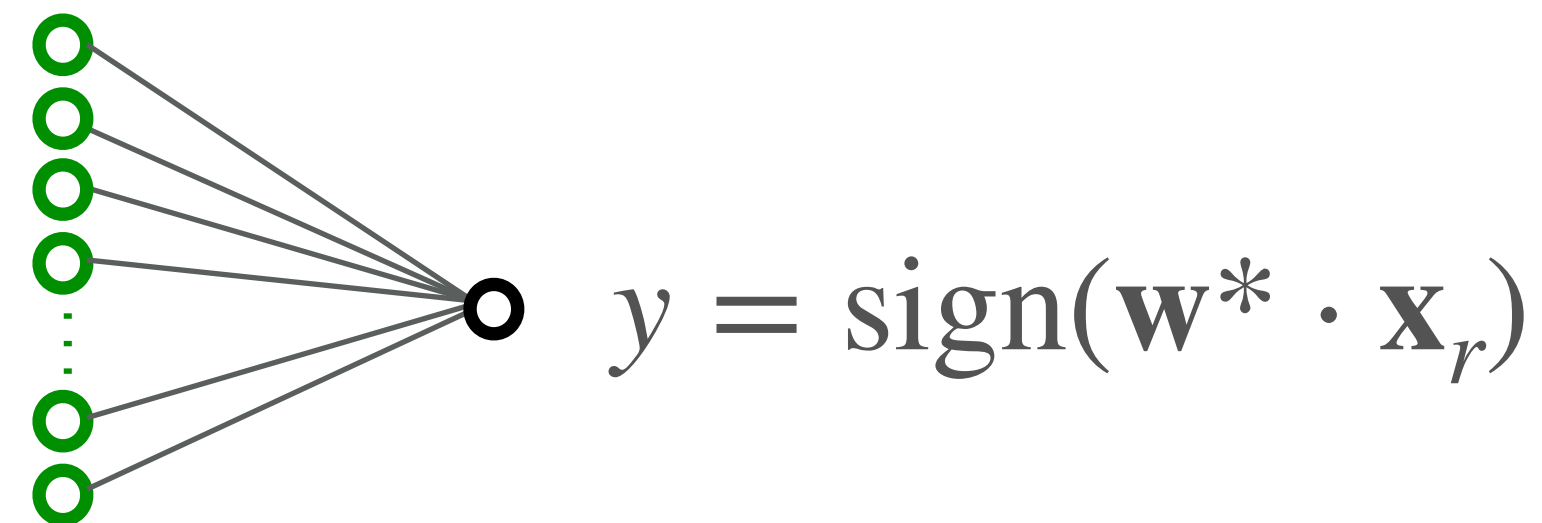
$$\mathbf{x}_r \in \mathbb{R}^{\rho N}$$

Unit variance

$$\mathbf{x}_i \in \mathbb{R}^{(1-\rho)N}$$

Variance Δ

Teacher



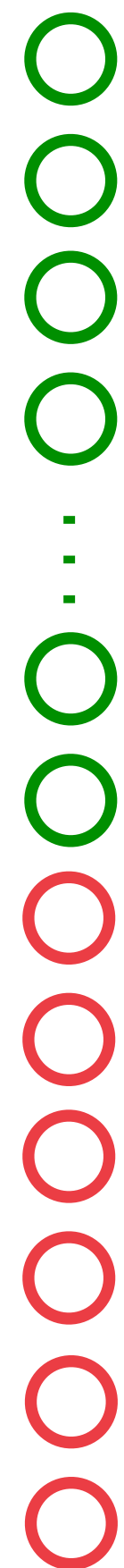
A teacher-student model of curriculum learning

Introduced in: *Bengio, et al. (ICML 2009)*, *Saglietti, et al. (NeurIPS 2022)*

Input: $\mathbf{x} = (\mathbf{x}_r, \mathbf{x}_i) \in \mathbb{R}^N$

Relevant

Irrelevant



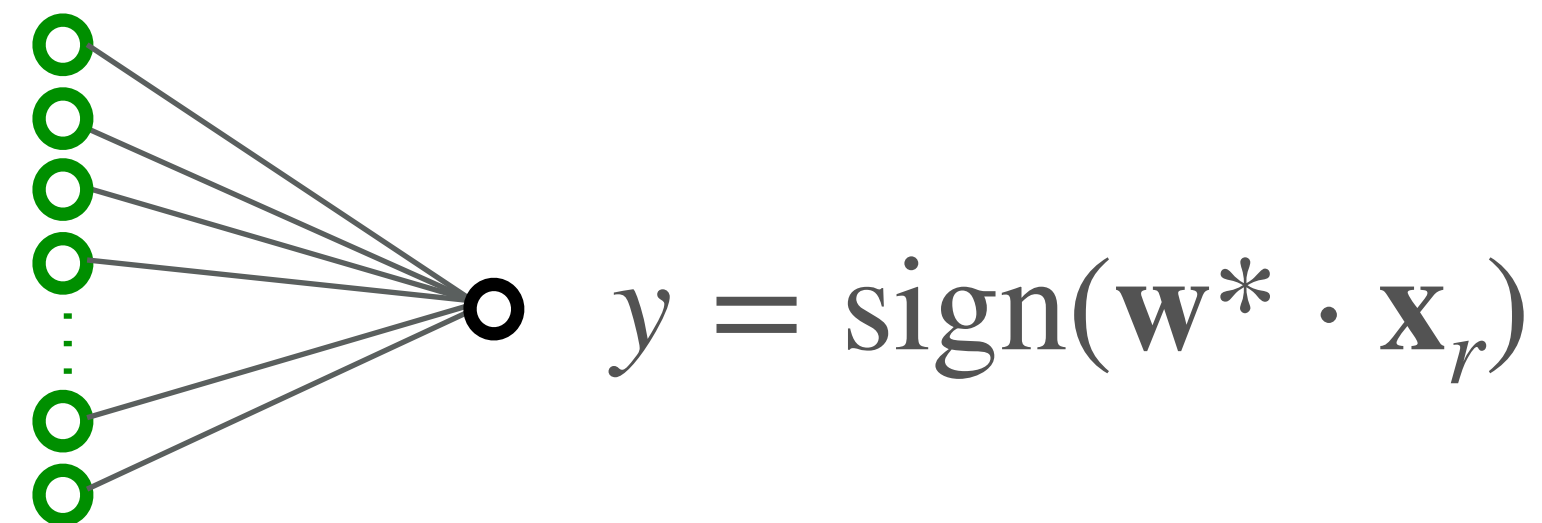
$$\mathbf{x}_r \in \mathbb{R}^{\rho N}$$

Unit variance

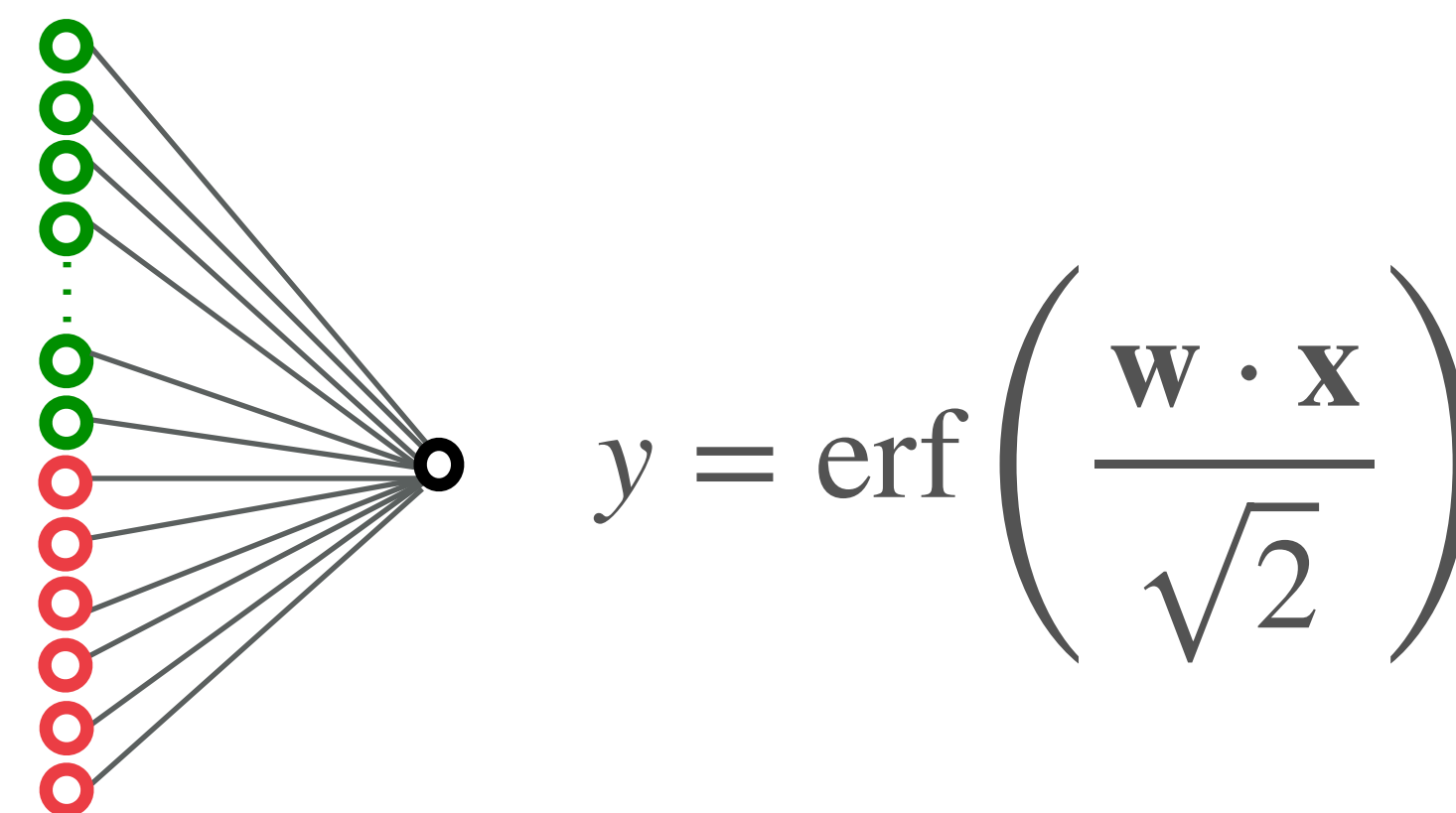
$$\mathbf{x}_i \in \mathbb{R}^{(1-\rho)N}$$

Variance Δ

Teacher



Student



Ridge-regularized MSE loss:

$$\mathcal{L} = \frac{1}{2}(y - \hat{y})^2 + \lambda \|\mathbf{w}\|_2^2$$

An Analytical Theory of Curriculum Learning in Teacher-Student Networks

Luca Saglietti^{†,*}, Stefano Sarao Mannelli^{‡,*}, and Andrew Saxe^{‡,§}

The evolution of the dynamics can be tracked using four order parameters:

$$Q_r = \frac{1}{N} \mathbf{W}_r \cdot \mathbf{W}_r,$$

$$R = \frac{1}{N} \mathbf{W}_r \cdot \mathbf{W}_T,$$

$$Q_i = \frac{1}{N} \mathbf{W}_i \cdot \mathbf{W}_i,$$

$$T = \frac{1}{N} \mathbf{W}_T \cdot \mathbf{W}_T ;$$

[Saglietti, et al (*NeurIPS* 2022)]

An Analytical Theory of Curriculum Learning in Teacher-Student Networks

Luca Saglietti^{†,*}, Stefano Sarao Mannelli^{‡,*}, and Andrew Saxe^{‡,§}

The evolution of the dynamics can be tracked using four order parameters:

$$\begin{aligned} Q_r &= \frac{1}{N} \mathbf{W}_r \cdot \mathbf{W}_r, & R &= \frac{1}{N} \mathbf{W}_r \cdot \mathbf{W}_T, \\ Q_i &= \frac{1}{N} \mathbf{W}_i \cdot \mathbf{W}_i, & T &= \frac{1}{N} \mathbf{W}_T \cdot \mathbf{W}_T; \end{aligned}$$

Online learning: [Biehl & Schwarze (1995); Saad & Solla (1995); ...]

$$Q_r \leftarrow f_{Q_r}(Q_r, Q_i, R, T)$$

$$Q_i \leftarrow f_{Q_i}(Q_r, Q_i, R, T)$$

$$R \leftarrow f_R(Q_r, Q_i, R, T)$$

[Saglietti, et al (*NeurIPS* 2022)]

A teacher-student model of curriculum learning

An Analytical Theory of Curriculum Learning in Teacher-Student Networks

Luca Saglietti^{†,*}, Stefano Sarao Mannelli^{‡,*}, and Andrew Saxe^{‡,§}

The evolution of the dynamics can be tracked using four order parameters:

$$Q_r = \frac{1}{N} \mathbf{W}_r \cdot \mathbf{W}_r,$$

$$R = \frac{1}{N} \mathbf{W}_r \cdot \mathbf{W}_T,$$

$$Q_i = \frac{1}{N} \mathbf{W}_i \cdot \mathbf{W}_i,$$

$$T = \frac{1}{N} \mathbf{W}_T \cdot \mathbf{W}_T;$$

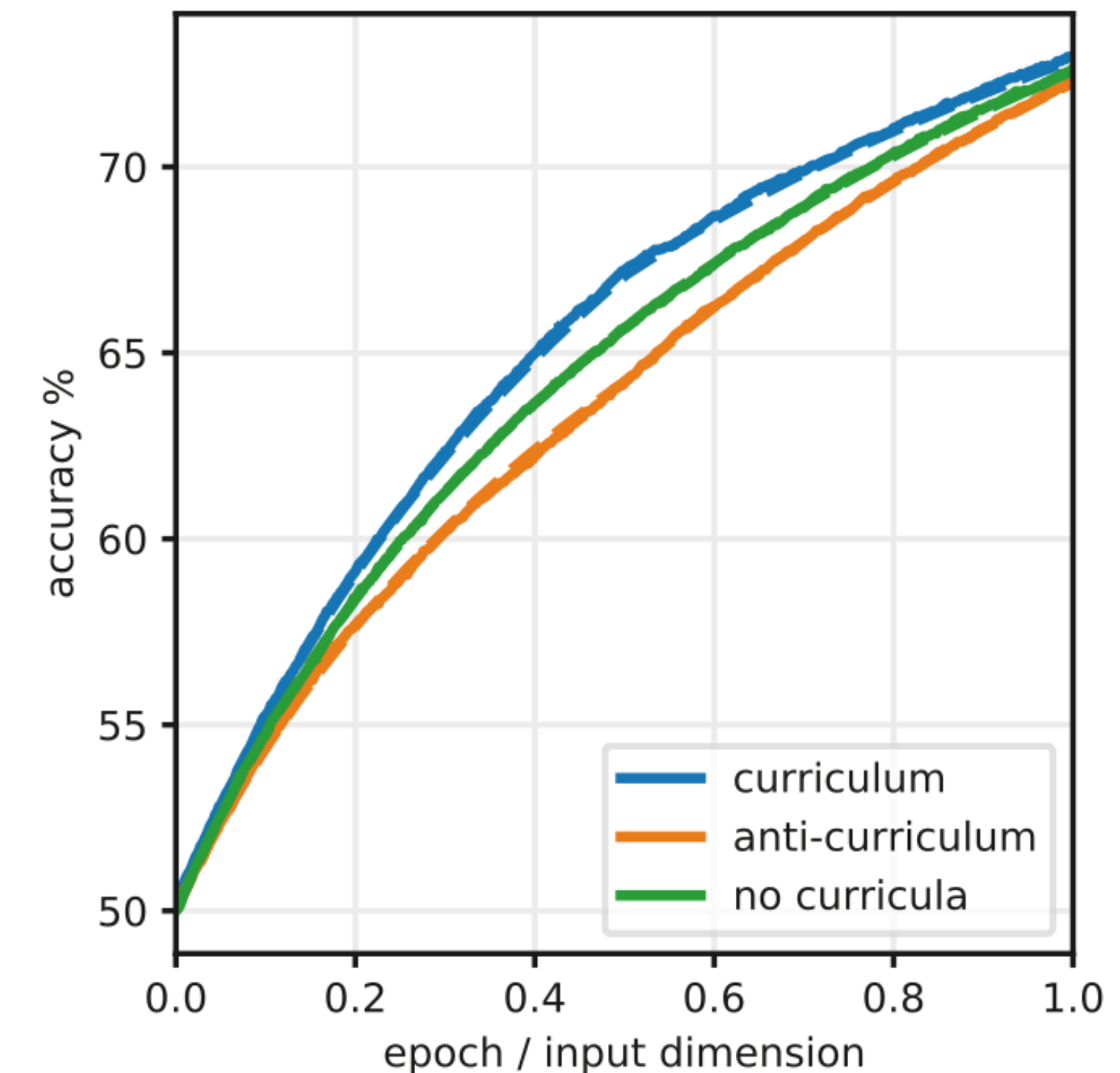
Online learning: [Biehl & Schwarze (1995); Saad & Solla (1995); ...]

$$Q_r \leftarrow f_{Q_r}(Q_r, Q_i, R, T)$$

$$Q_i \leftarrow f_{Q_i}(Q_r, Q_i, R, T)$$

$$R \leftarrow f_R(Q_r, Q_i, R, T)$$

[Saglietti, et al (*NeurIPS* 2022)]



Curriculum:

Easy

Hard

$\alpha = \text{epochs} / \text{input dimension}$

Anti-curriculum:

Hard

Easy

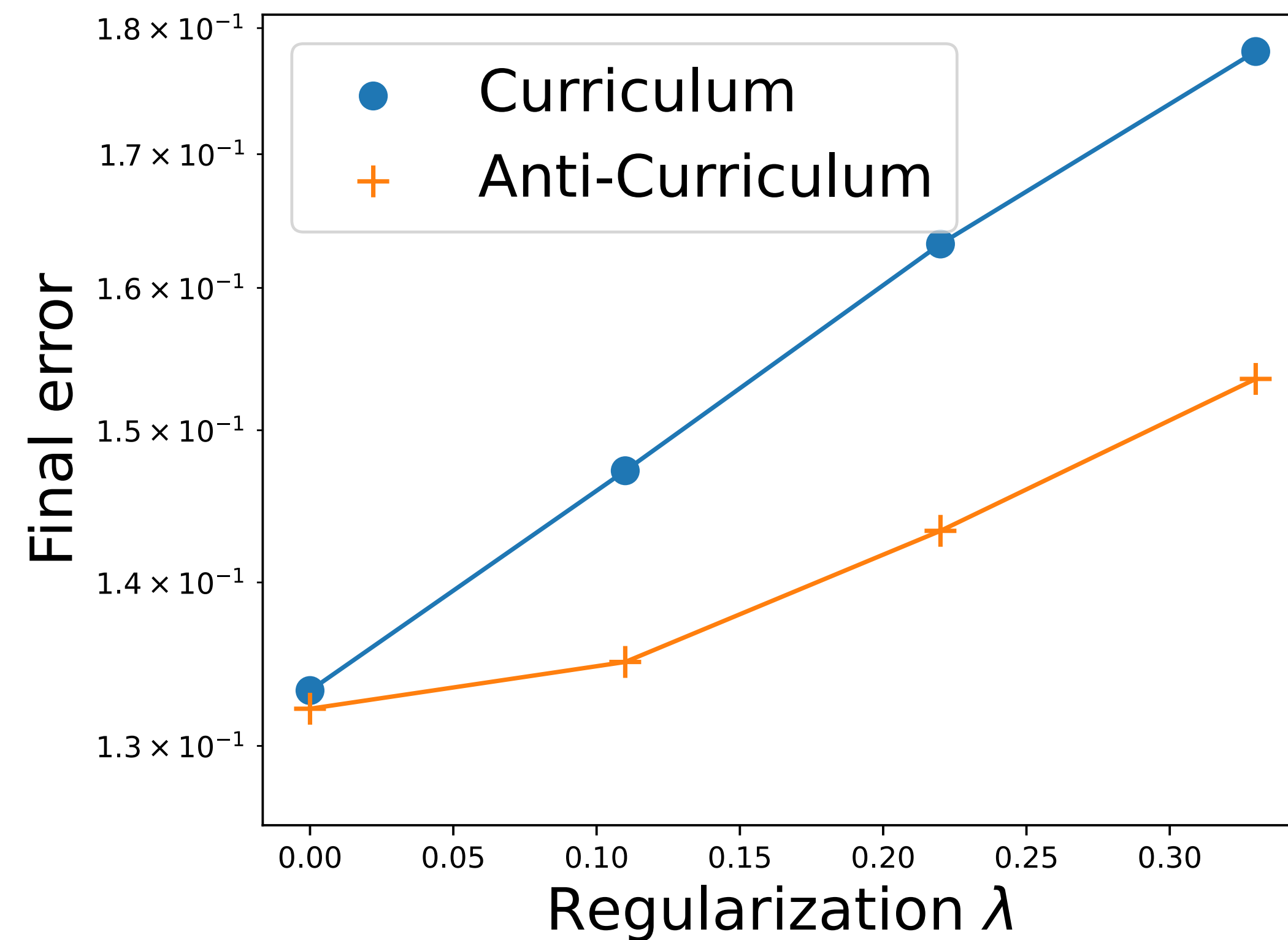
$\alpha = \text{epochs} / \text{input dimension}$

Optimal curriculum protocol

- Easy-to-hard curriculum is often *suboptimal*.

$$\rho = 0.55, \eta = 2.58$$

Control: $\mathbf{u} = \Delta$

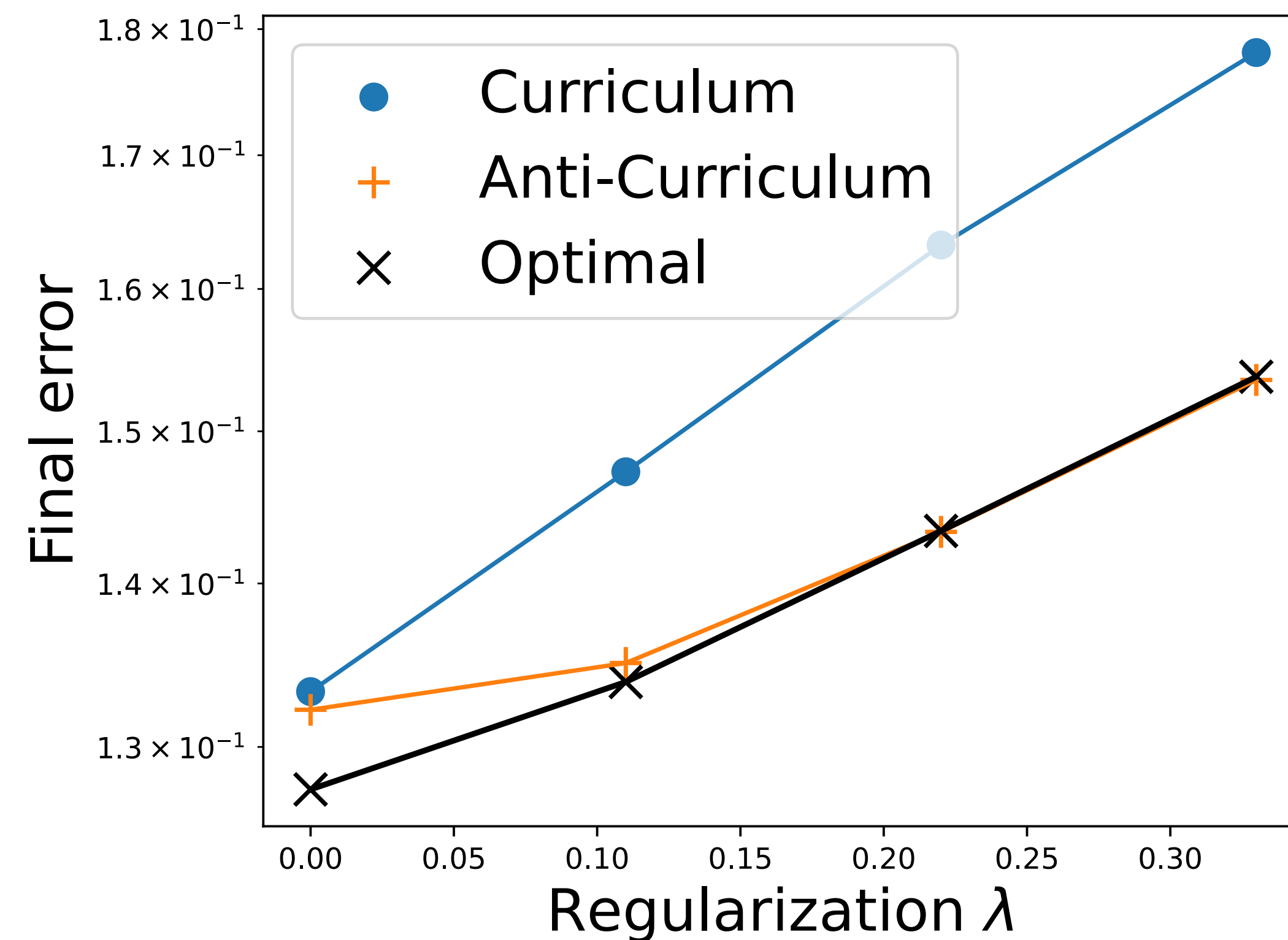


Optimal curriculum protocol

- Easy-to-hard curriculum is often *suboptimal*.

$$\rho = 0.55, \eta = 2.58$$

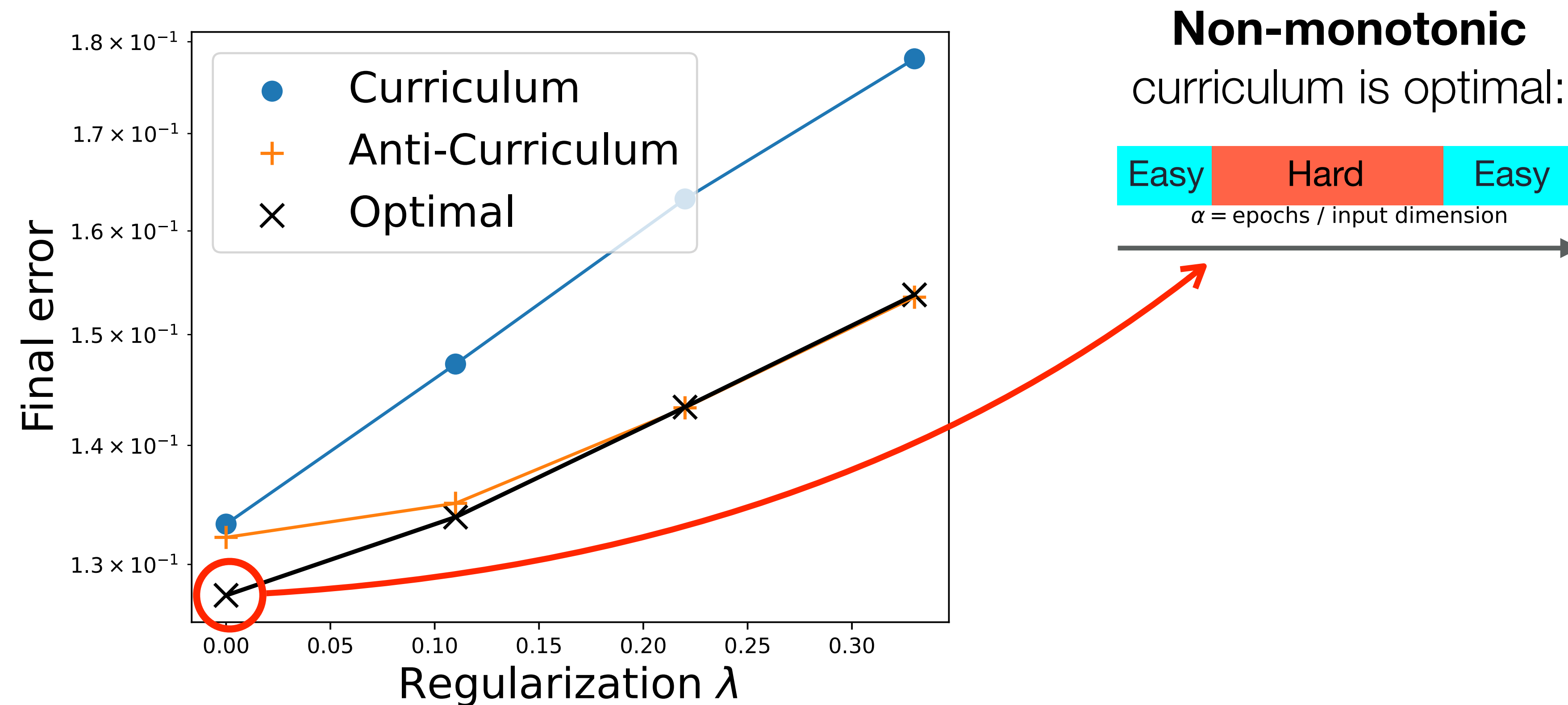
Control: $\mathbf{u} = \Delta$



Optimal curriculum protocol

Control: $\mathbf{u} = \Delta$

- Easy-to-hard curriculum is often *suboptimal*. $\rho = 0.55, \eta = 2.58$

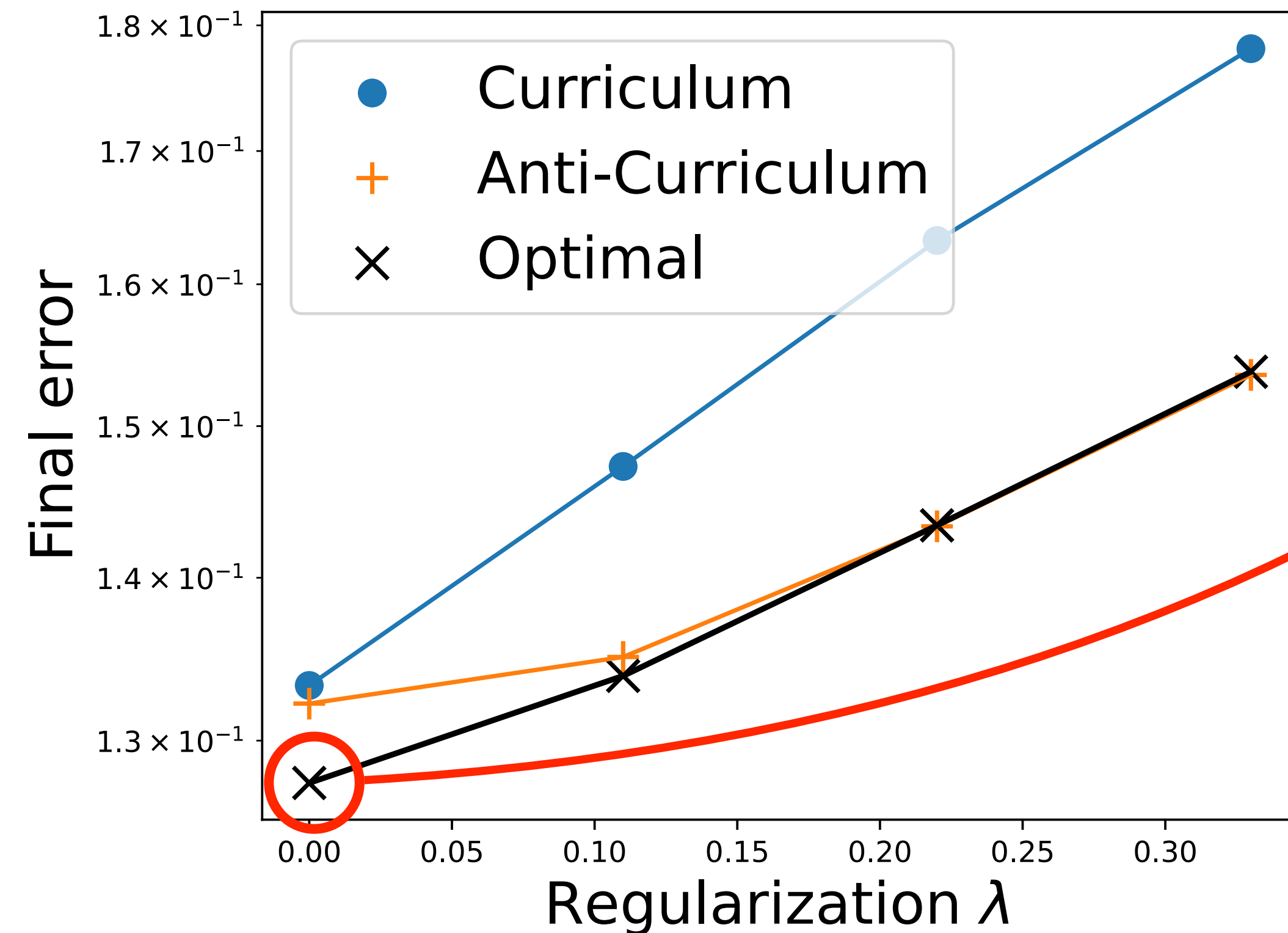


Optimal curriculum protocol

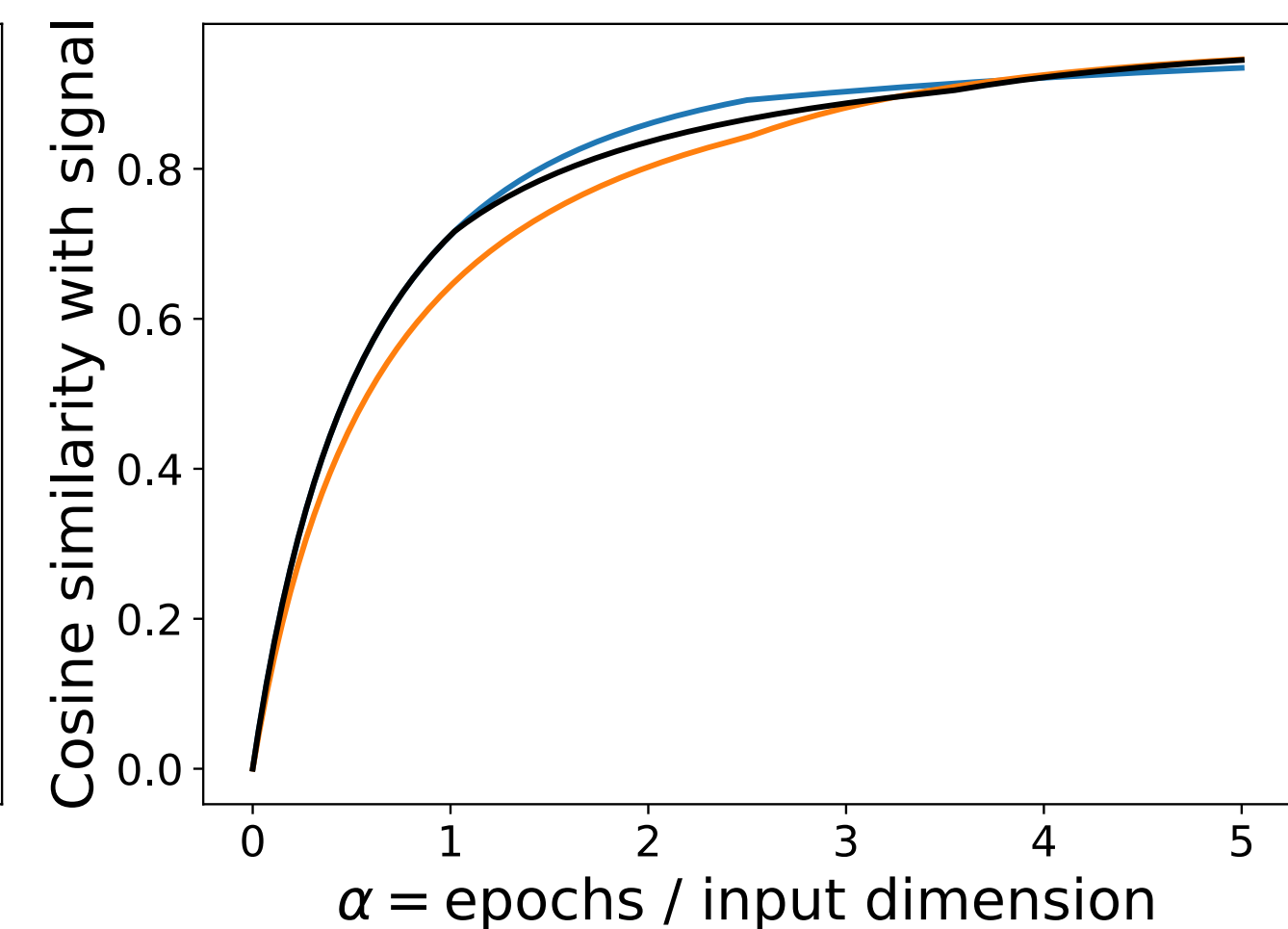
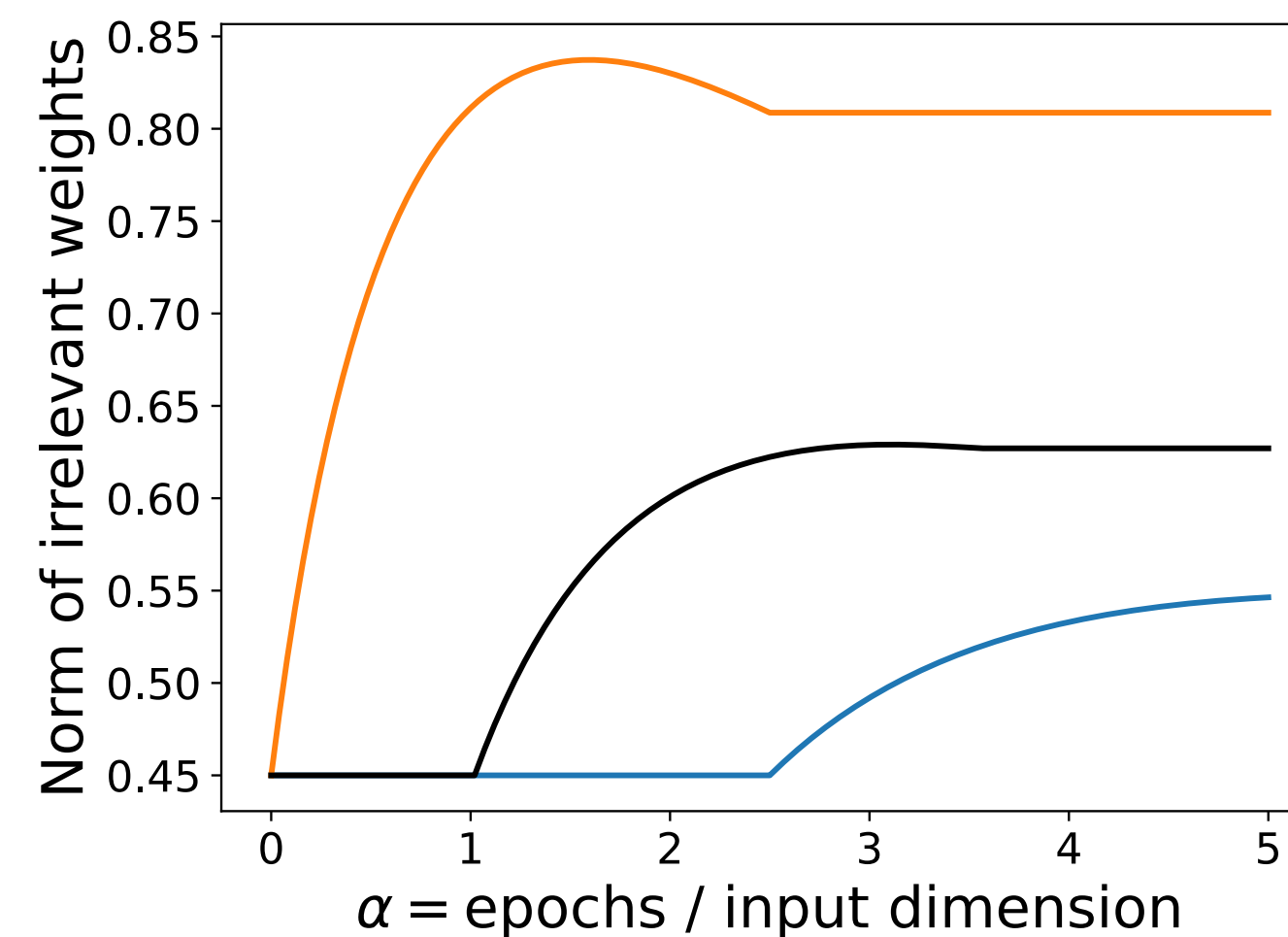
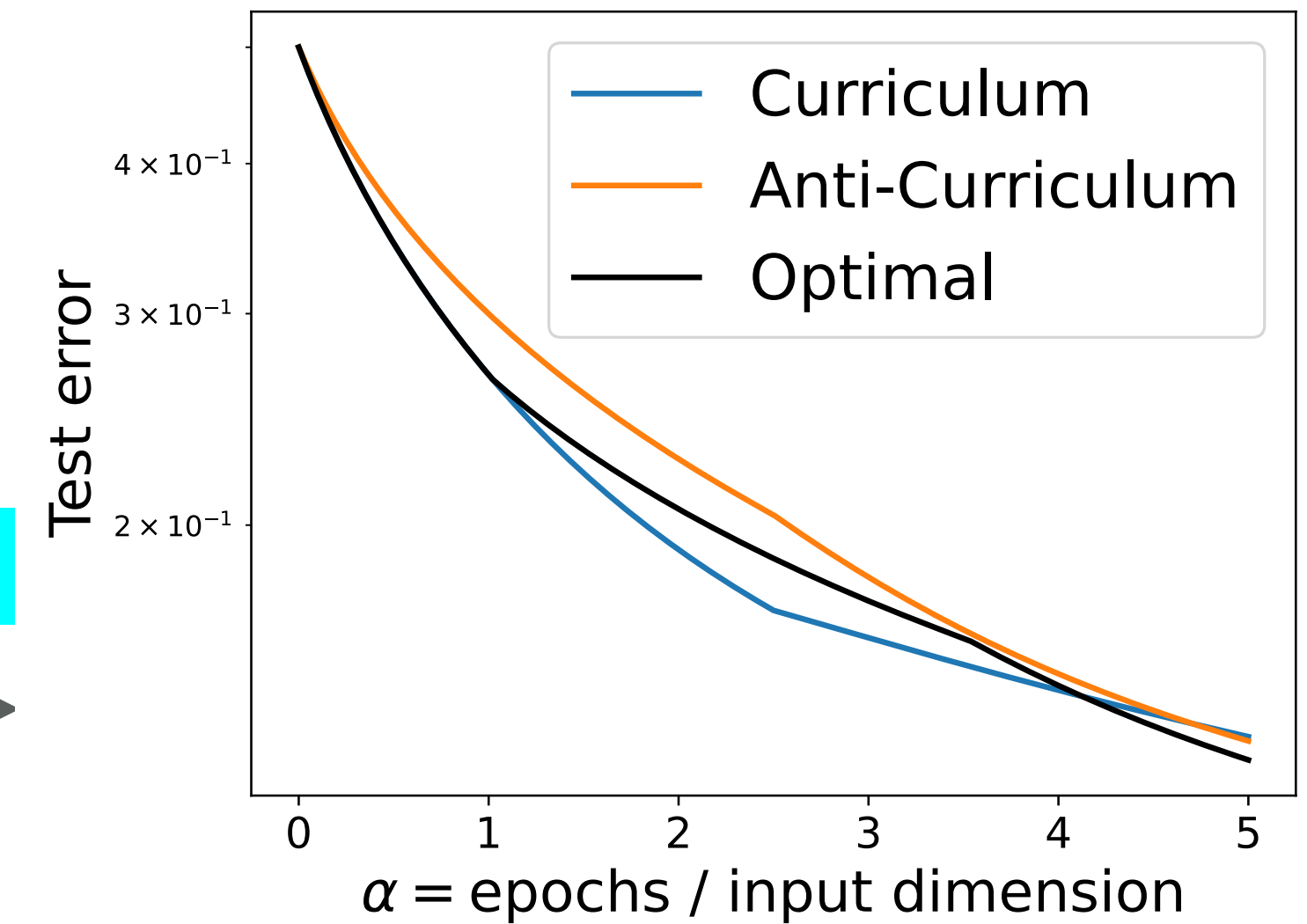
- Easy-to-hard curriculum is often *suboptimal*.

$$\rho = 0.55, \eta = 2.58$$

Control: $\mathbf{u} = \Delta$



Non-monotonic
curriculum is optimal:

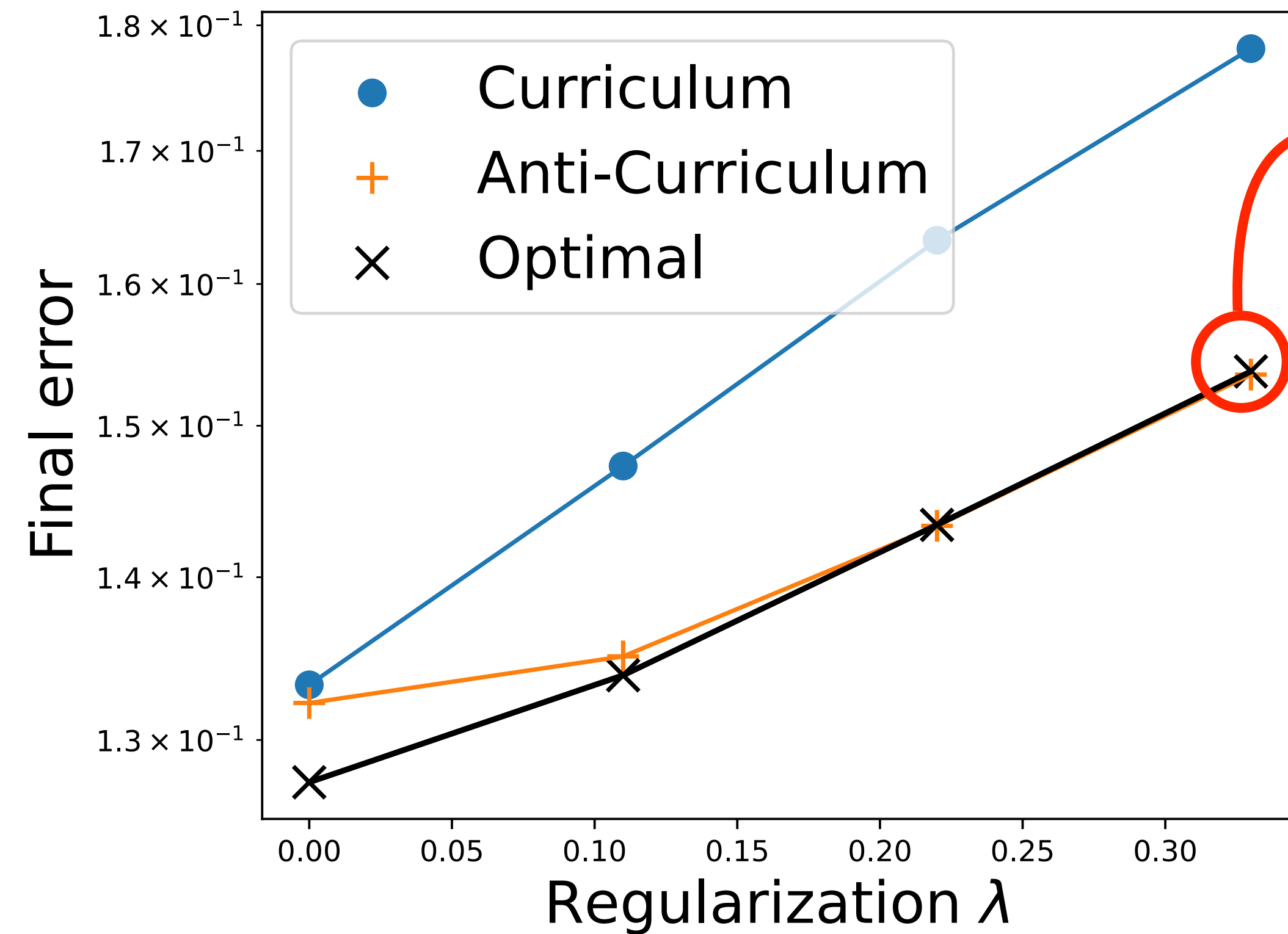


Optimal curriculum protocol

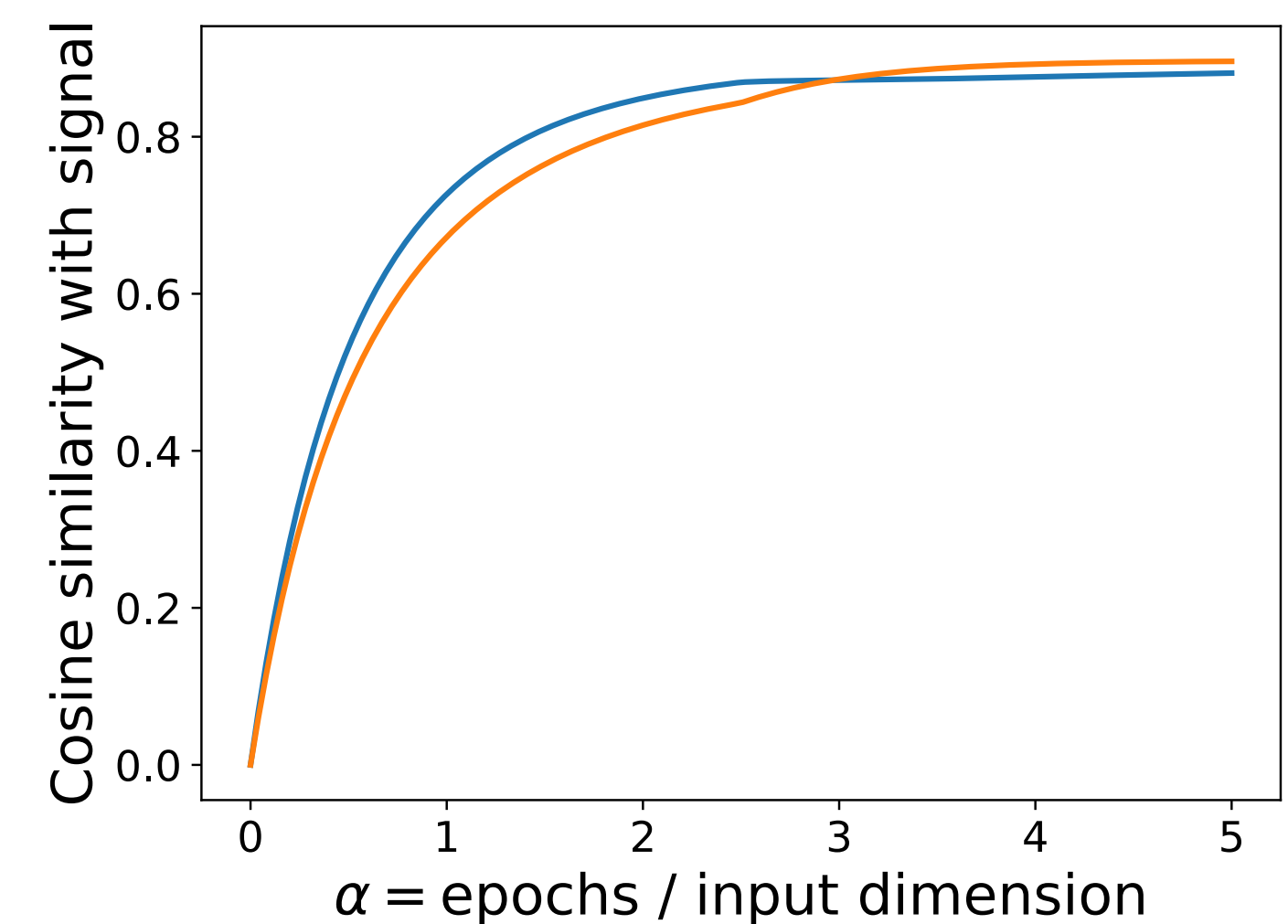
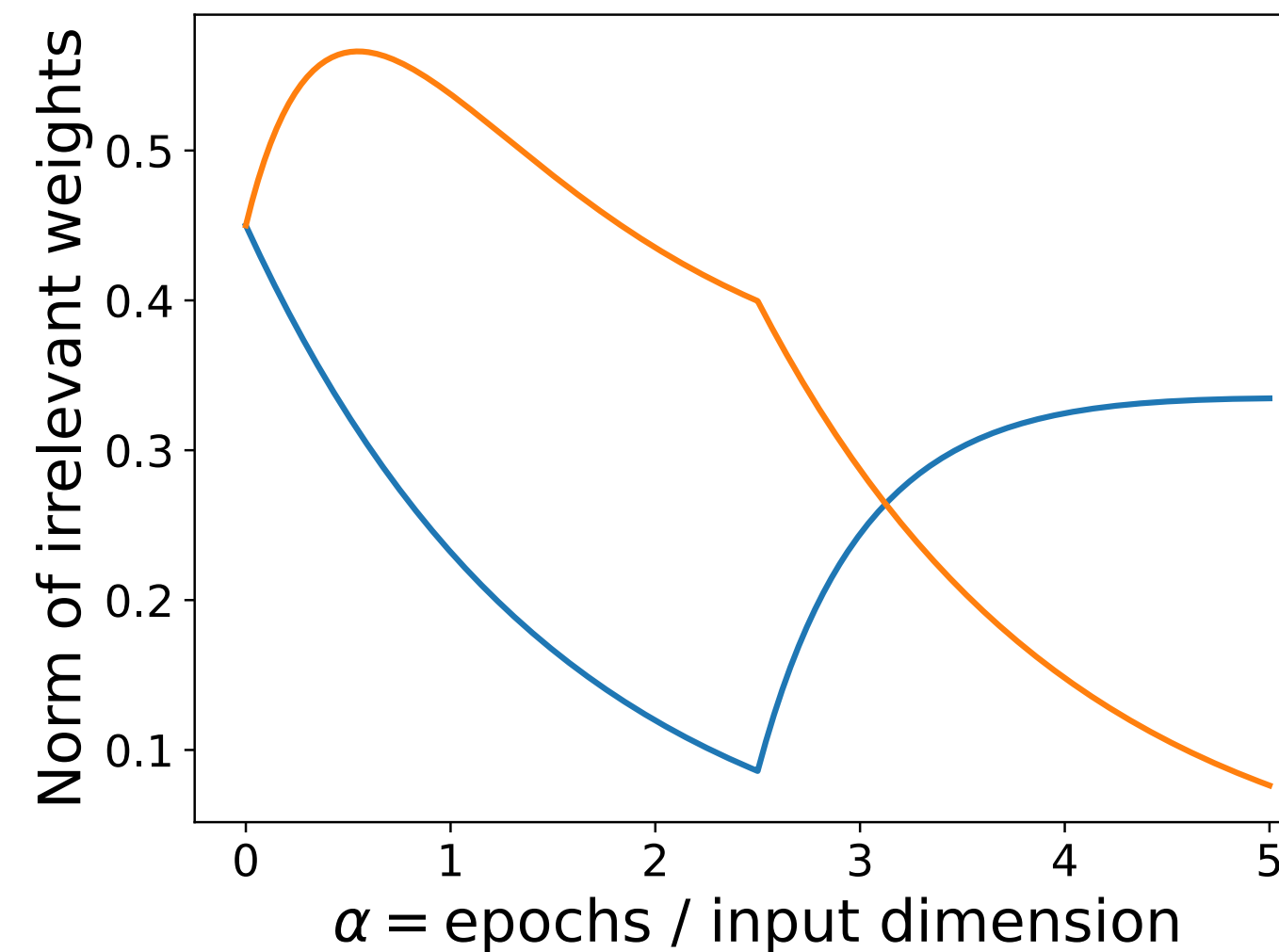
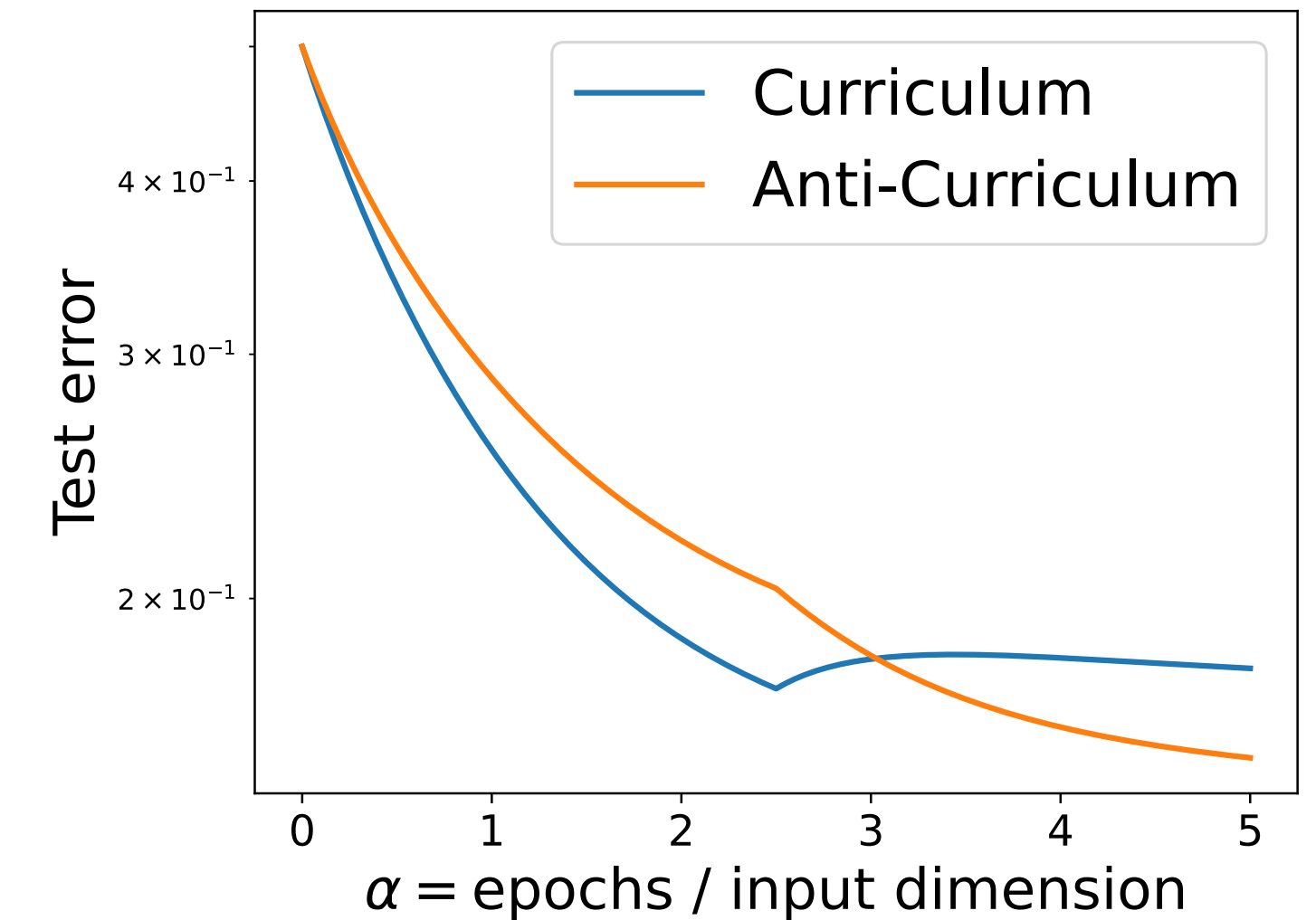
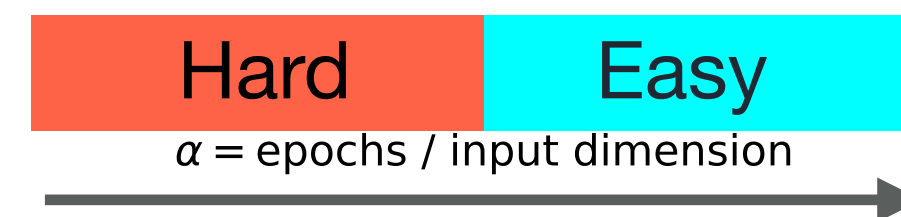
- Easy-to-hard curriculum is often *suboptimal*.

$$\rho = 0.55, \eta = 2.58$$

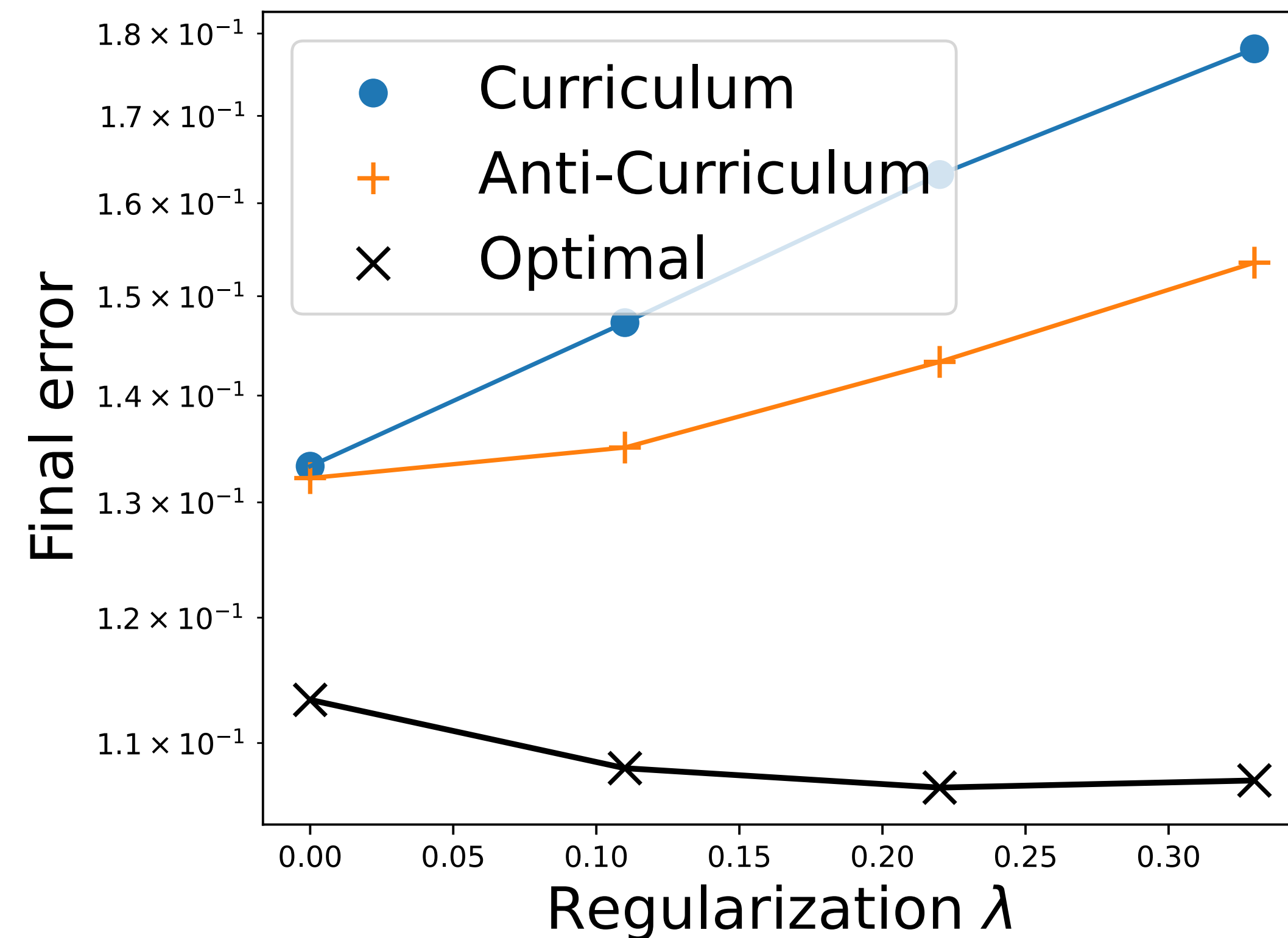
Control: $\mathbf{u} = \Delta$



Anti-curriculum
is optimal:

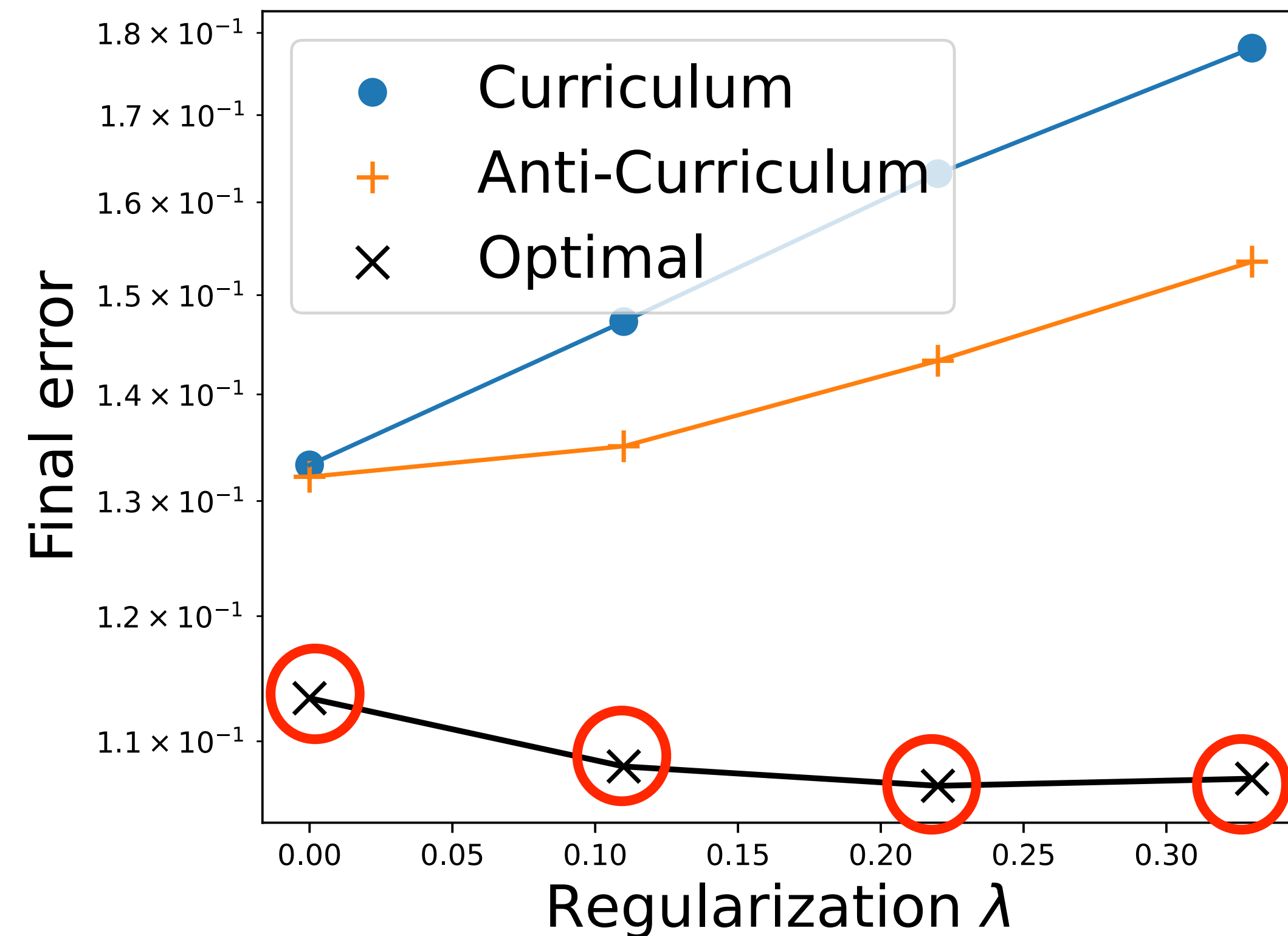


- Easy-to-hard curriculum becomes *optimal* if we also optimize over the learning rate schedule.

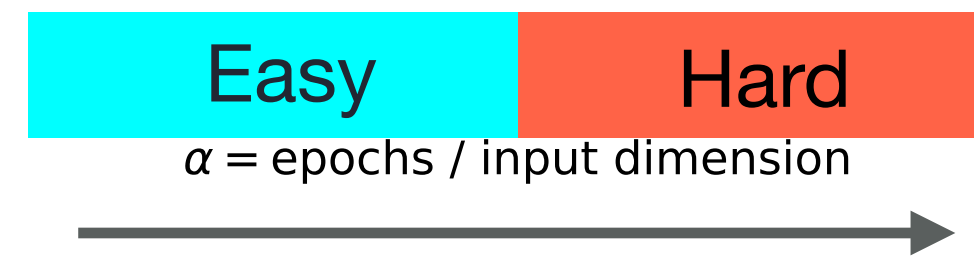


Optimal curriculum protocol

- Easy-to-hard curriculum becomes *optimal* if we also optimize over the learning rate schedule.

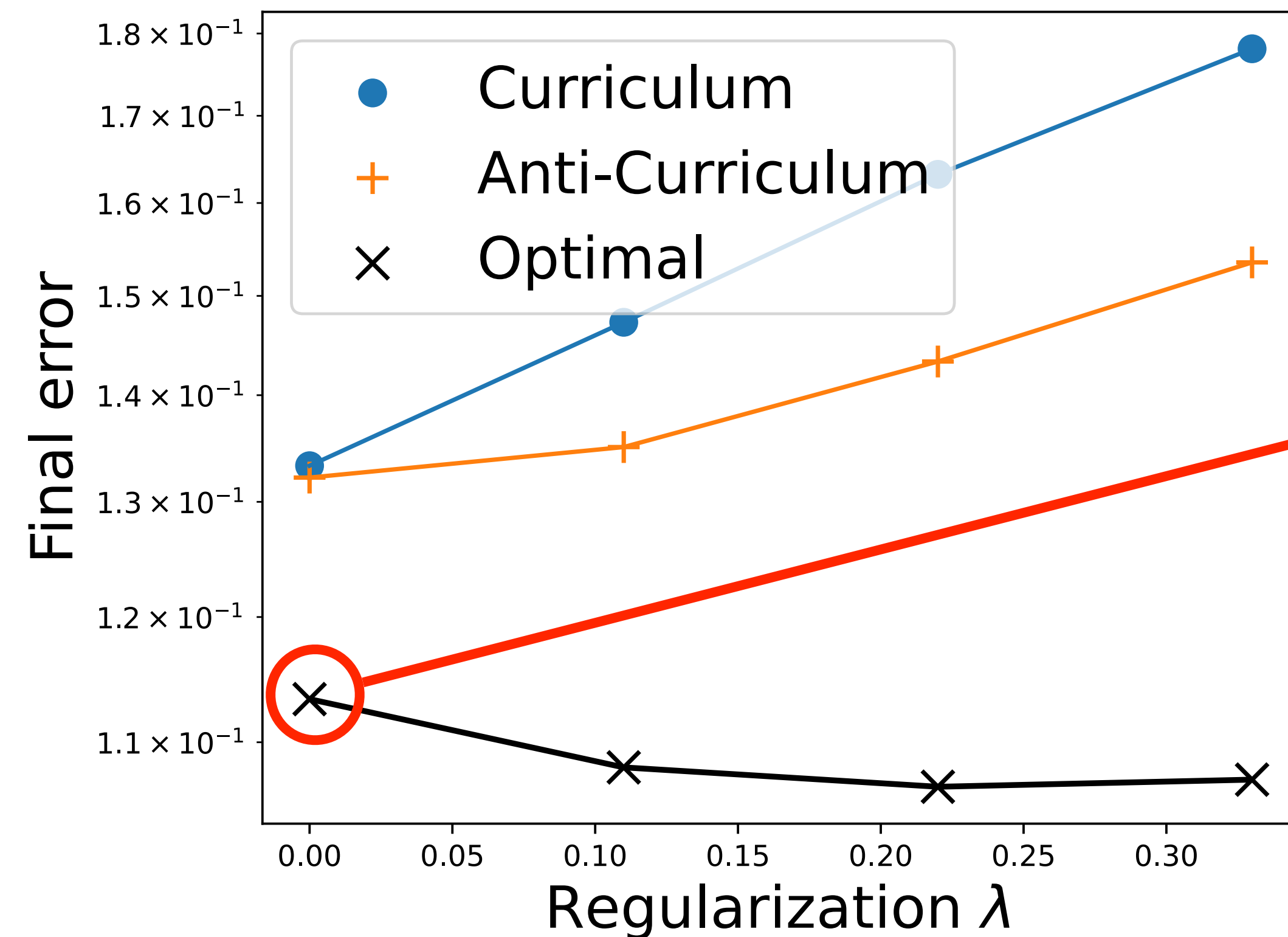


Easy-to-hard
curriculum is optimal:



Optimal curriculum protocol

- Easy-to-hard curriculum becomes *optimal* if we also optimize over the learning rate schedule.



Easy-to-hard
curriculum is optimal:

