

TRAINING NEURAL NETWORKS

WHERE TRAINING GOES AND ITS STABILITY

Lecture 18

PIERFRANCESCO BENEVENTANO
PIERB@MIT.EDU



Part 0

Recap from last classes and motivations

THE GOAL WAS

- We have data coming in real life: $X \sim P$
- We need to find the neural network $f_\theta(X)$ that has the right output $\varphi(X)$.
- A measure of the Error
 $distance(\phi(.), f_\theta(.))$.

They exist.

Can we find them?

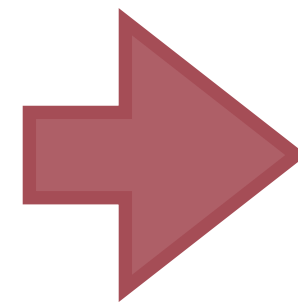
We minimize the measure of error

Optimization

LEARNABILITY *VS* OPTIMIZATION

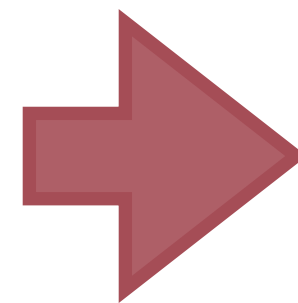
We have:

- A dataset.
- Some compute
(Memory + number of computations constraint).



Do we have enough samples?

Learnability



In case, does the training routine find it within the compute?

Optimization

COMPUTE / MEMORY

- The dimensionality N of the problem ($\theta \in \mathbb{R}^N$) is very high, e

- We have many n loss is an

$$\frac{1}{n} \sum_{x \in D} L(\theta, x)$$

We work with:
(for both learnability and compute)

Mini-batch Gradient Methods

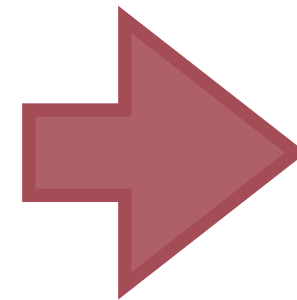
Gradient-based

$$\theta_{k+1} = \theta_k - \eta \nabla L(\theta_k)$$

Mini-batch methods

Gradient Based Methods

Designing:
***Gradient-Based
Optimizers***



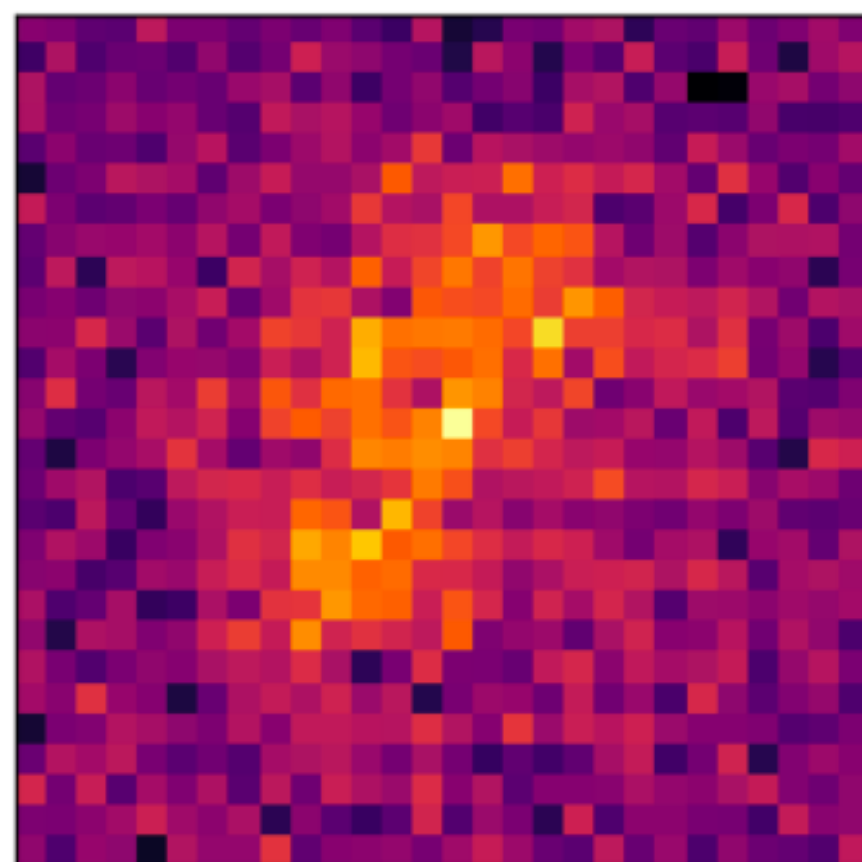
How to go:
from infinitesimal to local

-
- Step size = Norm + Trust.
 - About second derivative.

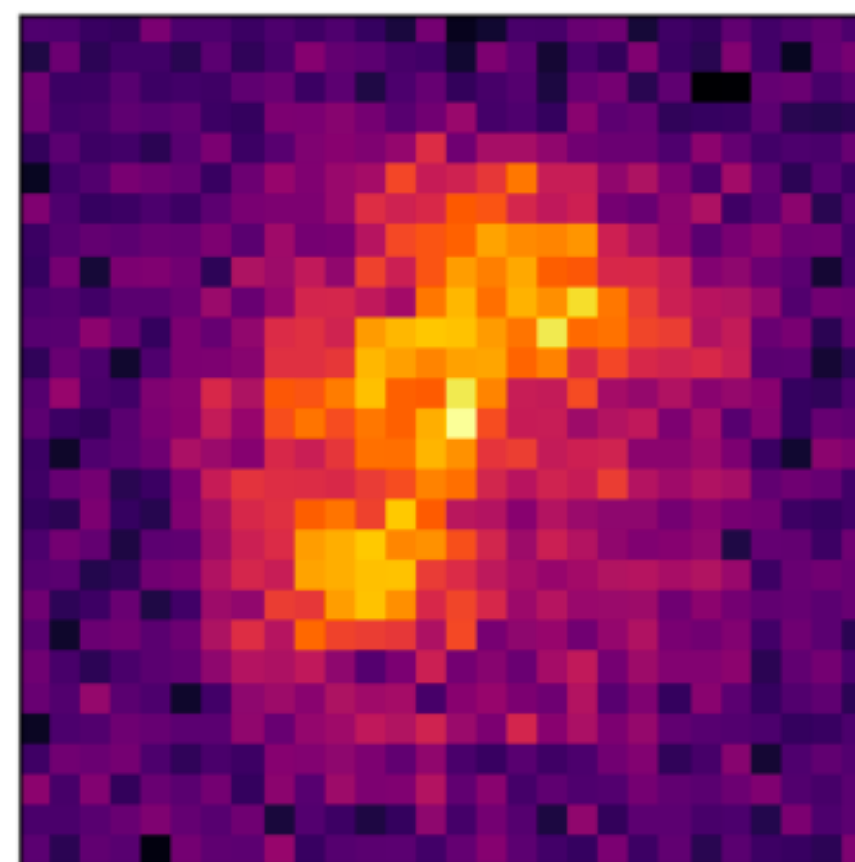
Using history

What is this Picture Missing?

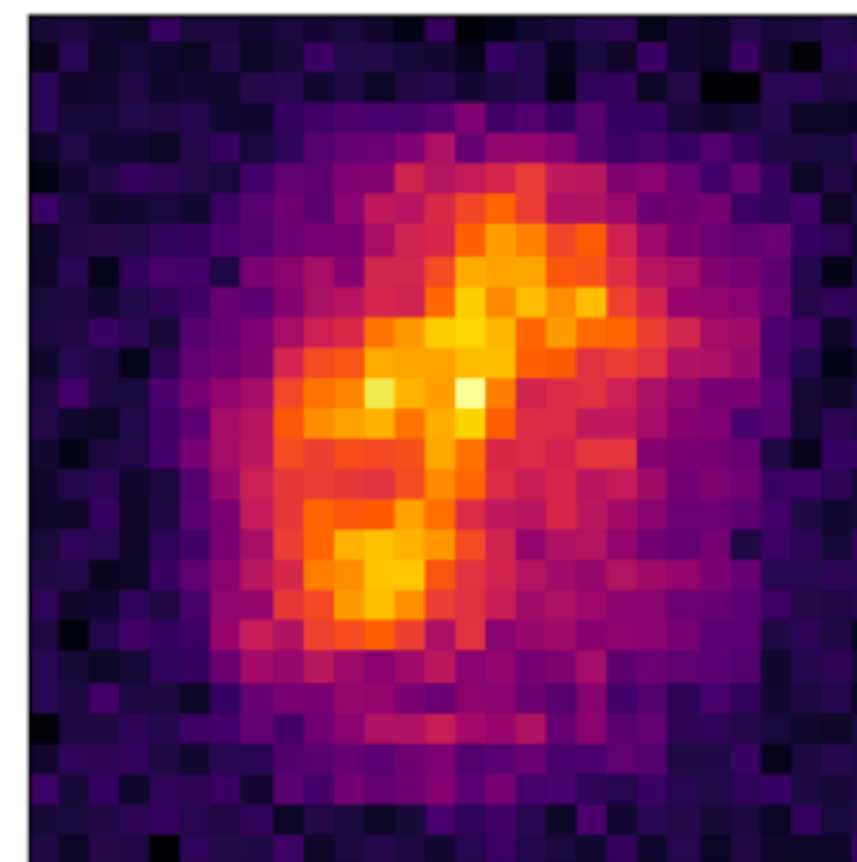
SGD-2048



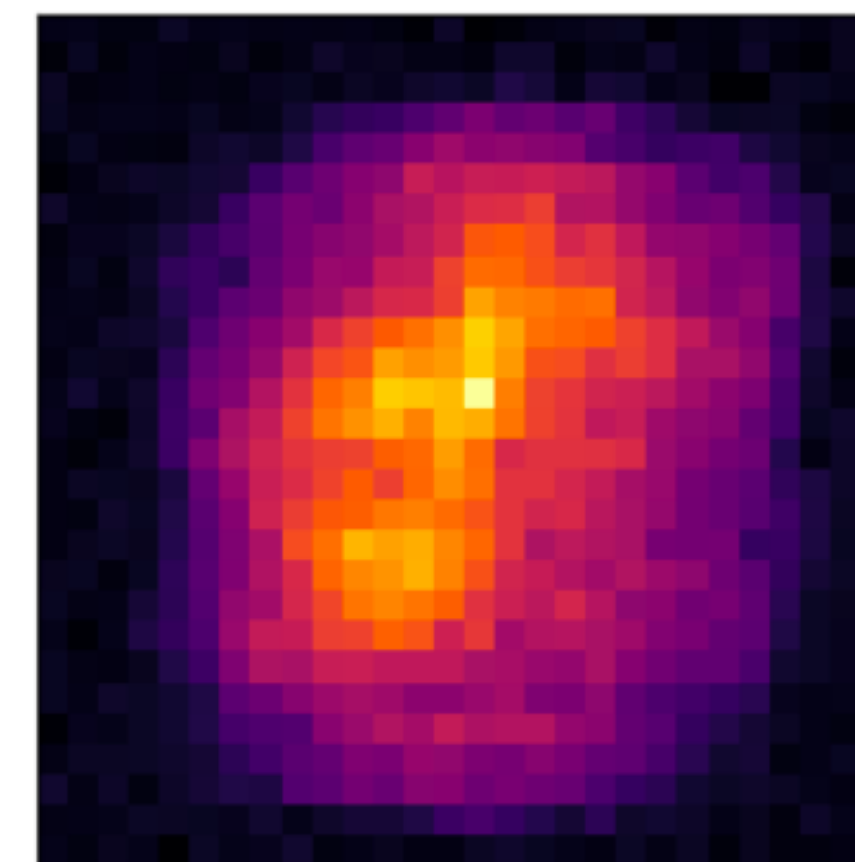
SGD-512



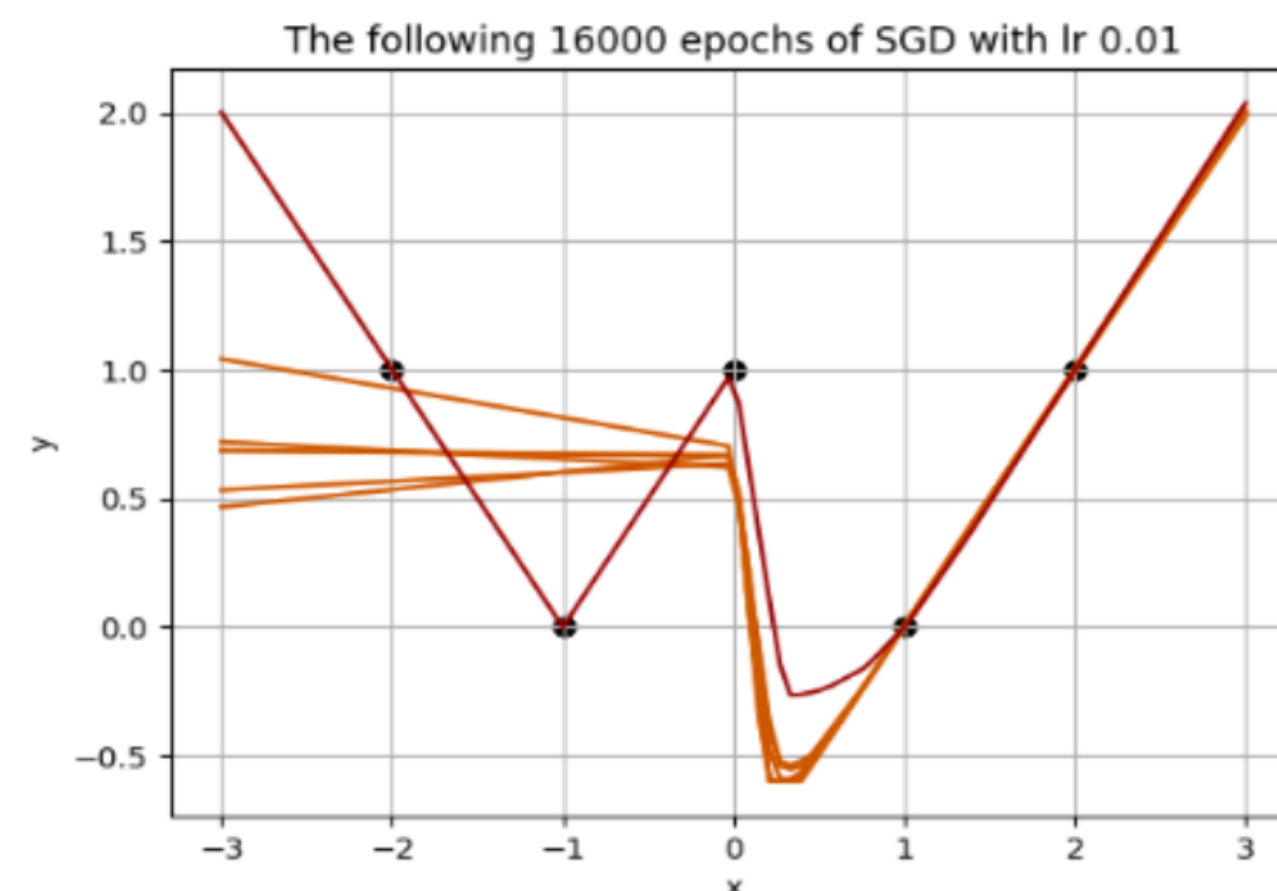
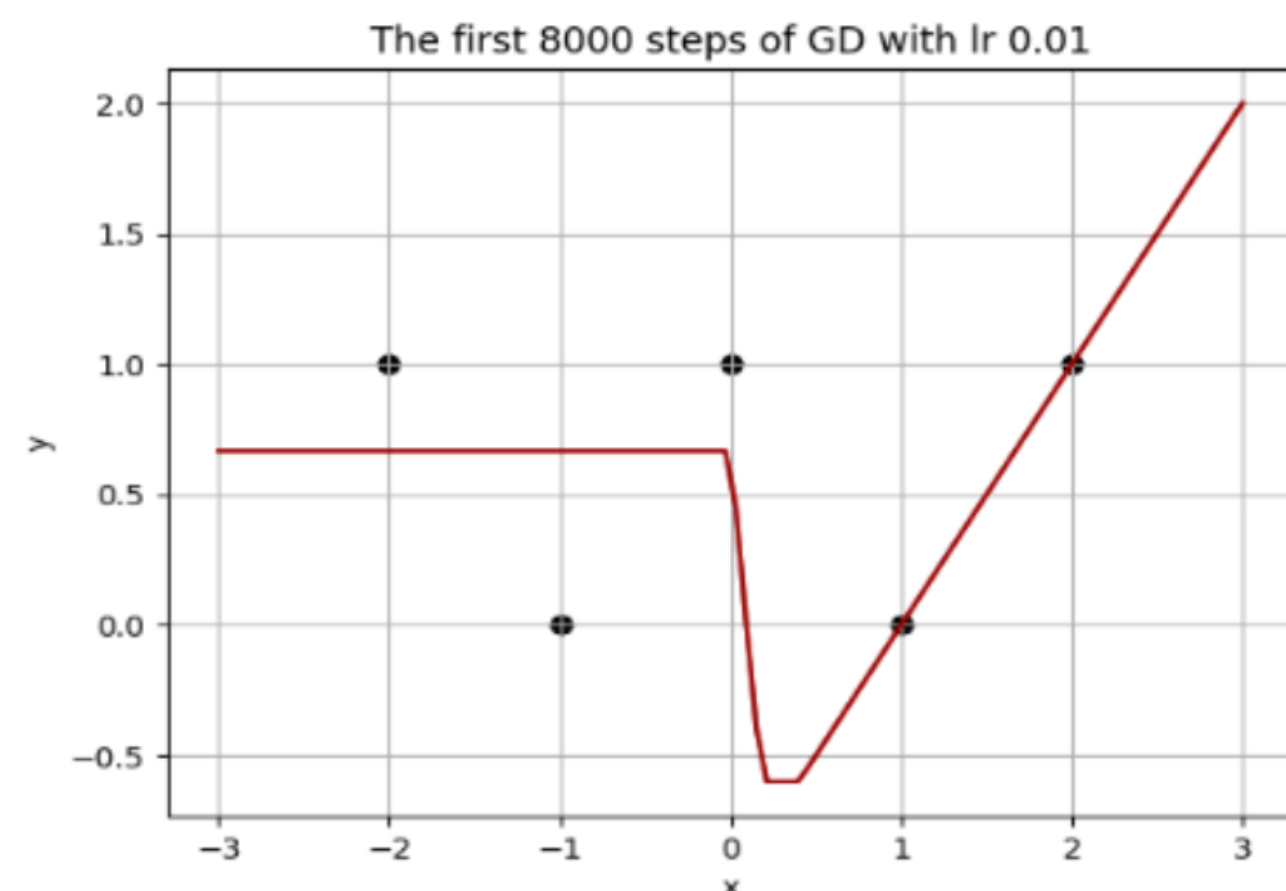
SGD-128



SGD-32



SGD smaller batch size →



Location of Convergence

It is not just about if you go down within compute:

- It's also about where you go within compute!
- We discussed how to do your best in the step:

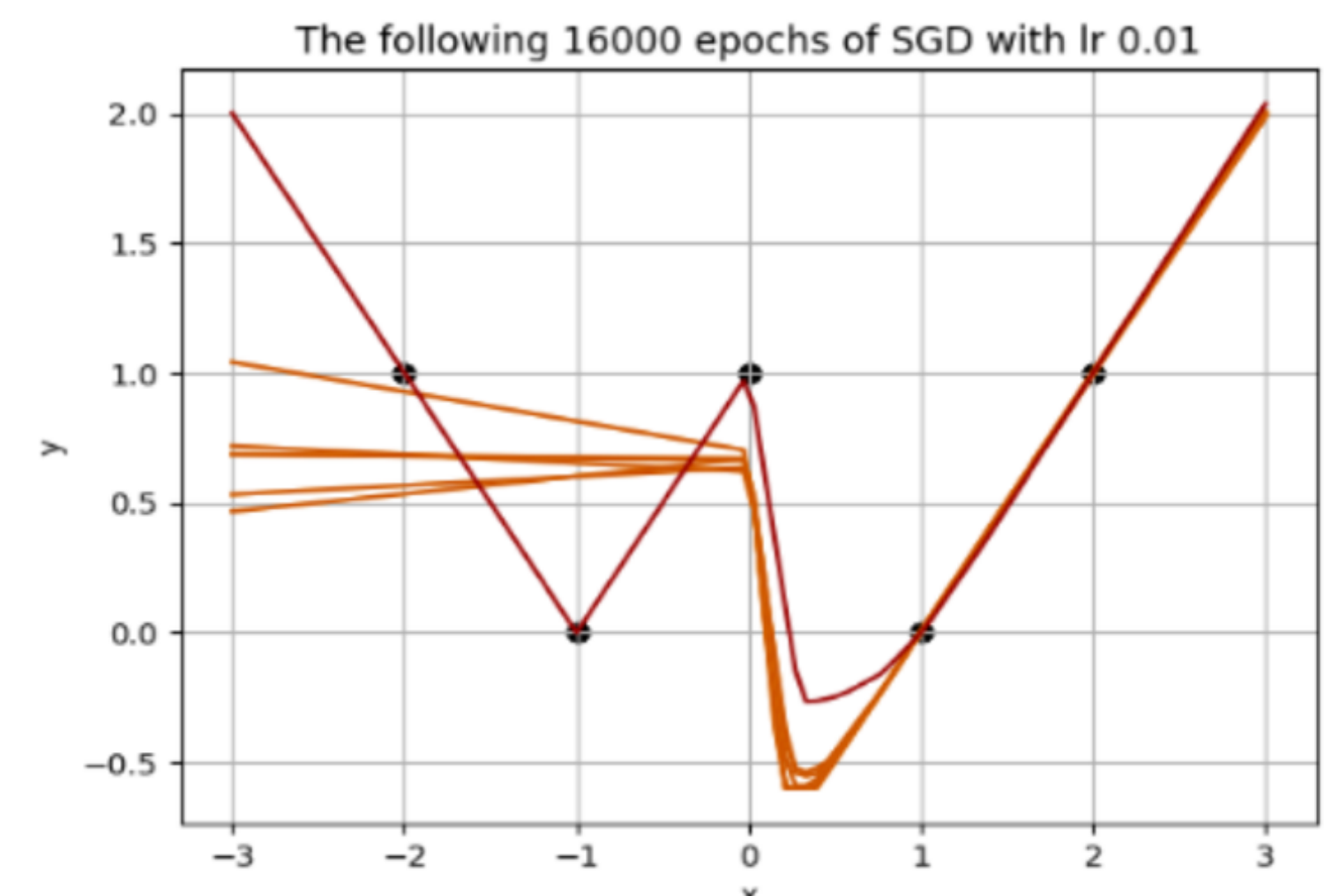
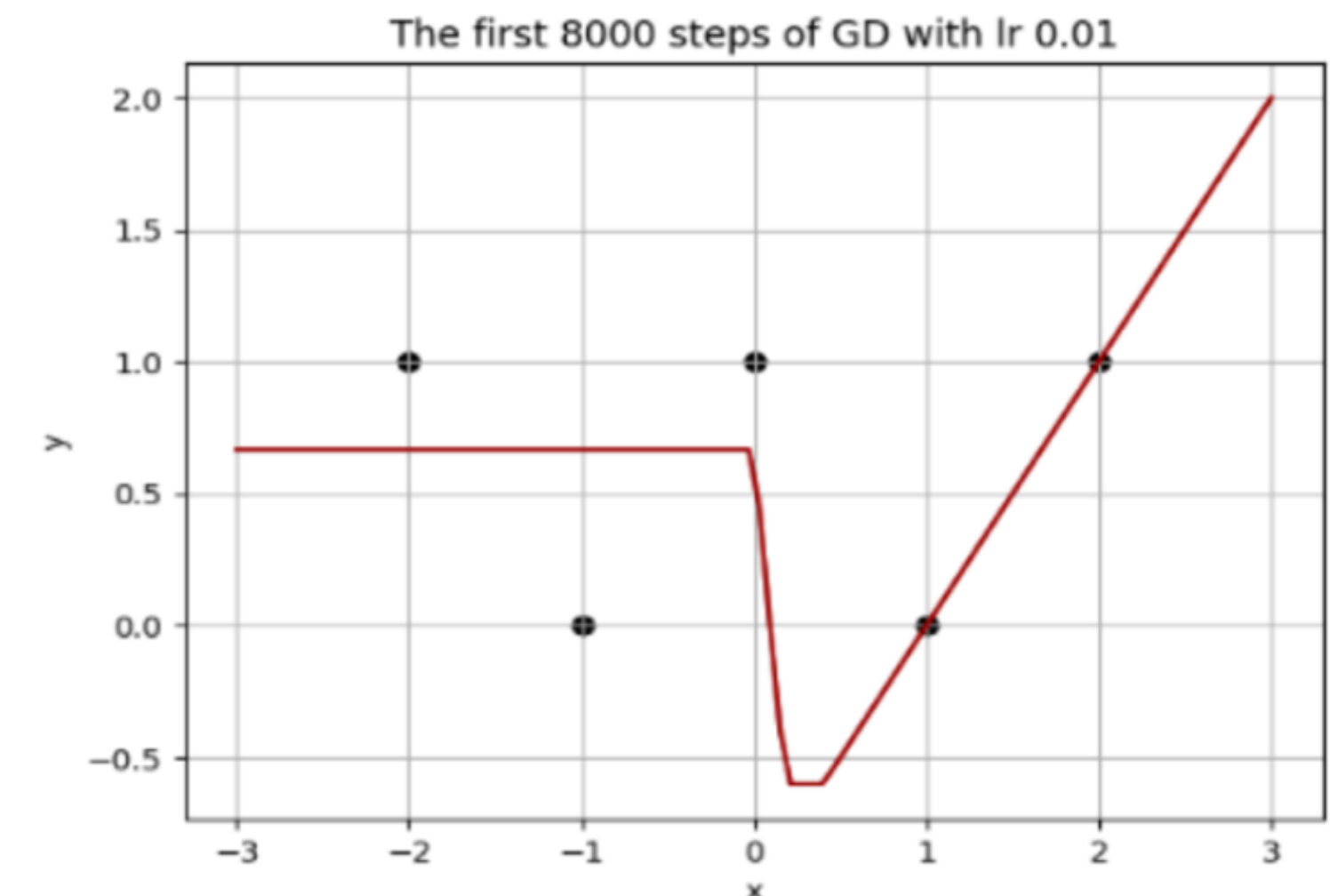
Infinitesimal to *local*

- We did not say anything about:

Local to *Global*

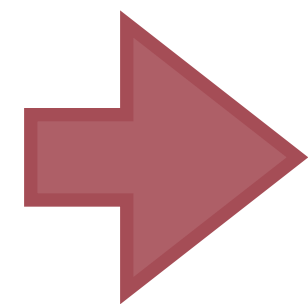
Next This class is about local to global:

*On the location of convergence
and the phases of training*



Message 1, Class 17

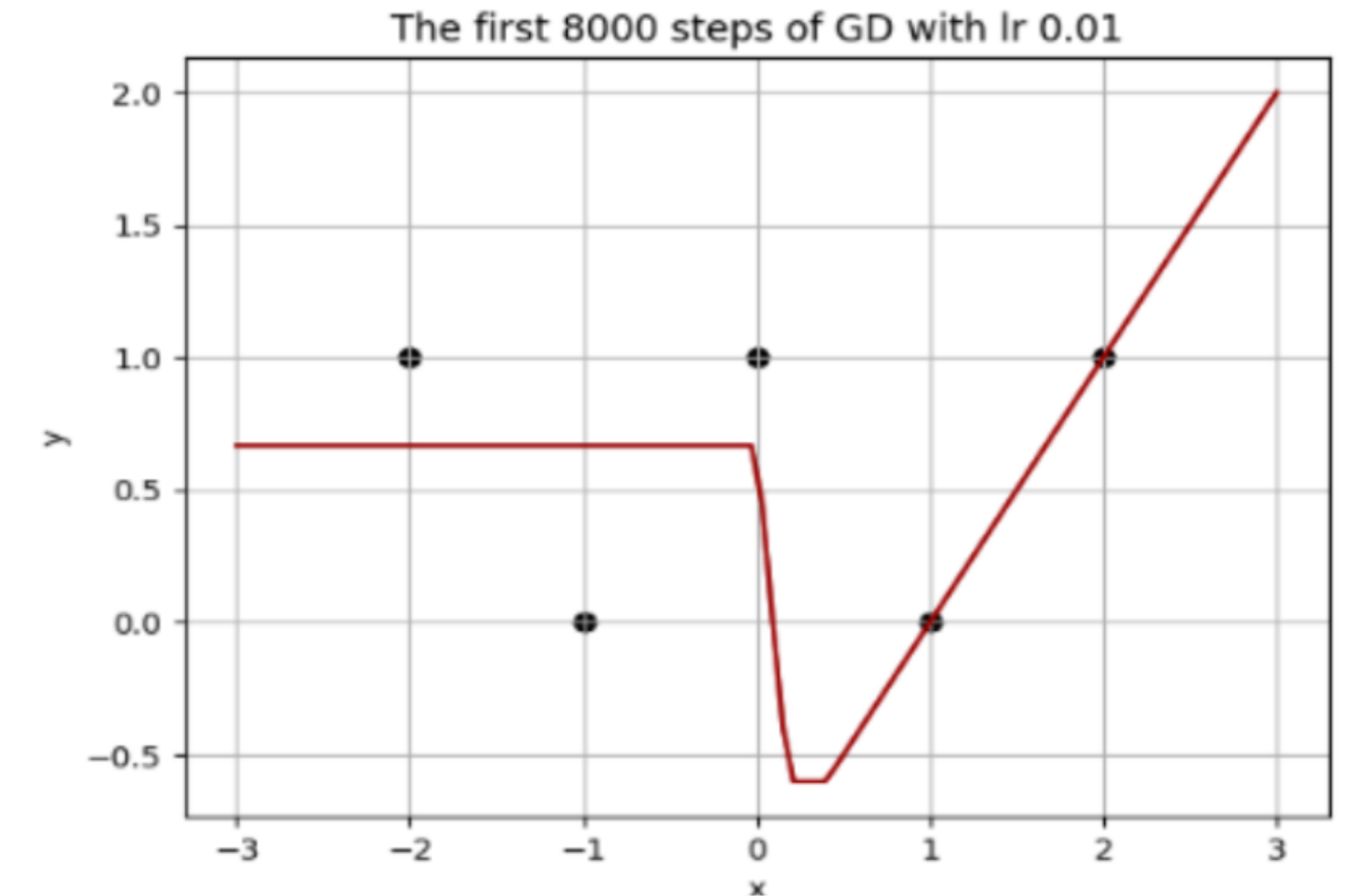
Designing Gradient
Based Optimizers



How to go:
from *infinitesimal* to *local*
(TO GLOBAL)

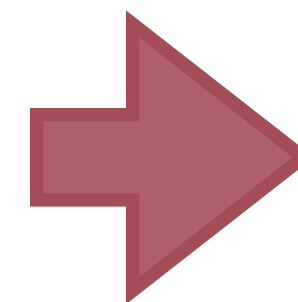
REGIME 2

- Not enough exact data!
- Enough compute



We want a model that:

- Can **fit them all**.
- It's the **simplest** such model.



A training pipeline which goes to the best global minima.

CONTENTS

- Location of Convergence
 - Some examples with NN
 - What gradient methods were
- (Stable) Implicit Regularization
- Edge of Stability

THE QUESTION

How to go optimally
from local to global

Where does our training algorithm
go in the case of neural networks?

A preliminary

SPEAKING MATH

WHERE DO ALGORITHMS GO?

This is not an *optimization* question:
It is a dynamical systems question!

Questions are:

- What are the basins of attraction? Do trajectory recur? (**Long-term behavior**)
- Is the invariant set stable (**Stability in the sense of Lyapunov, ...**)
- How does the behavior changes with the (hyper)parameters (**Bifurcation Analysis**)

CLASS 17

Infinitesimally, to go down *optimally* we follow the gradient:

$$\begin{cases} \theta_0 = \text{initialization} \\ \frac{d\theta_t}{dt} = -\nabla L(\theta_t) \end{cases}$$

How to pick initialization?

What is the optimal discrete step size version of this?

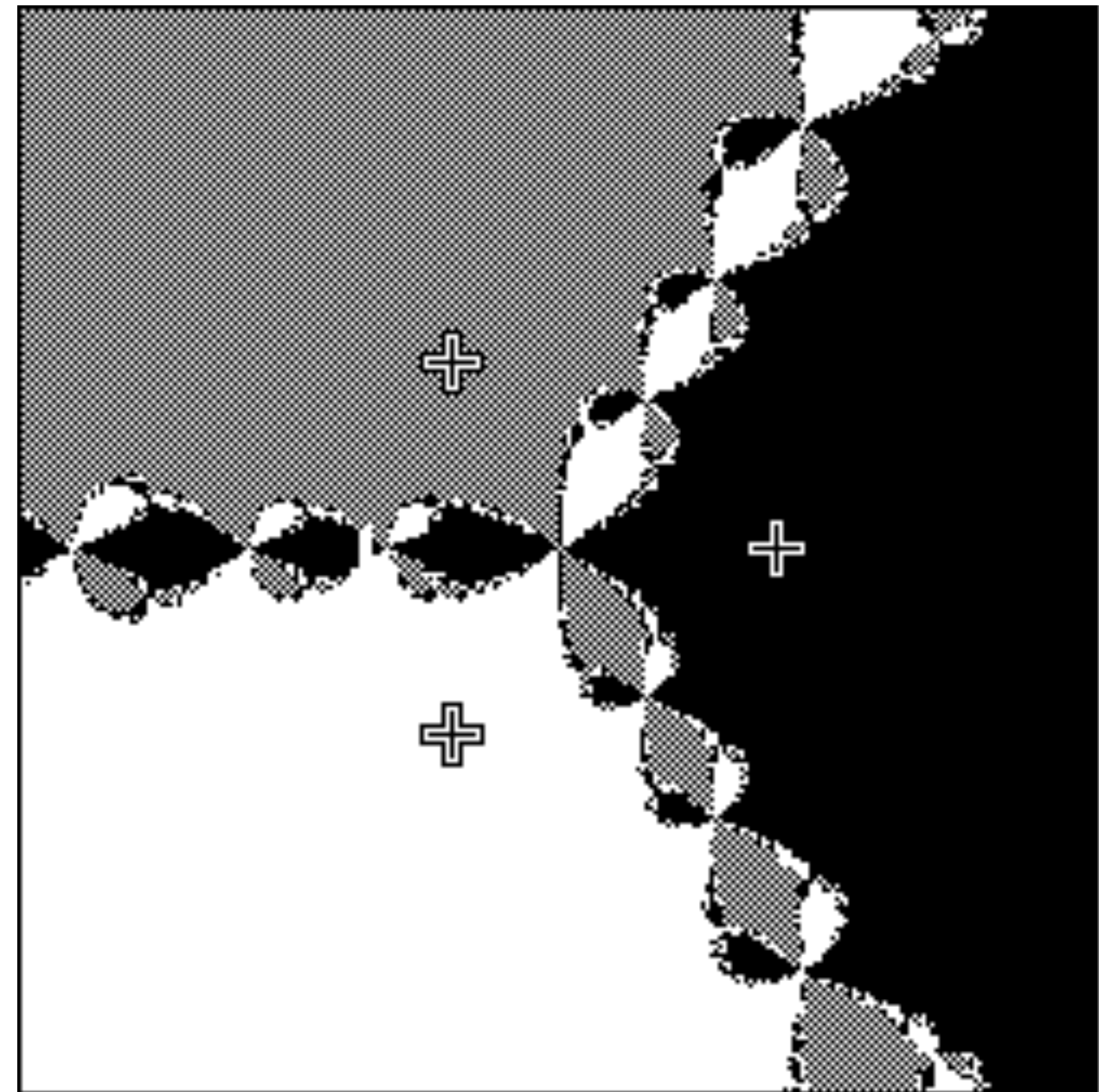
CLASSICAL PoV

Classically, initialization is about:

$$\begin{cases} \theta_0 = \text{initialization} \\ \frac{d\theta_t}{dt} = -\nabla L(\theta_t) \end{cases}$$

In what basin of attraction you are:

It is about:
which minimum



Basins for Newton on $L(\theta) = \theta^3 - 1$

Message 3, Class 17

- We want to pick the **architecture** and the **initialization** to enhance **trainability**.
- Not really for purposes of **approximation** or location of **convergence**!

So let's study the *optimizer*
instead of the initialization

Questions

- We would like to design a **globally optimal** gradient based optimizer.
- We will see that they have been designed for something else:

Convex functions

- We do not have the tools to discuss which one is local to global optimal, yet.
- So we will just explore:

Where they go in NN

What tools can we figure out for the problem

Part 1

What can we expect empirically
&
Refreshing the gradient based methods

LOCATION OF CONVERGENCE

- **Location of Convergence**
 - **Some examples with NN**
 - What gradient methods were
- (Stable) Implicit Regularization
- Edge of Stability

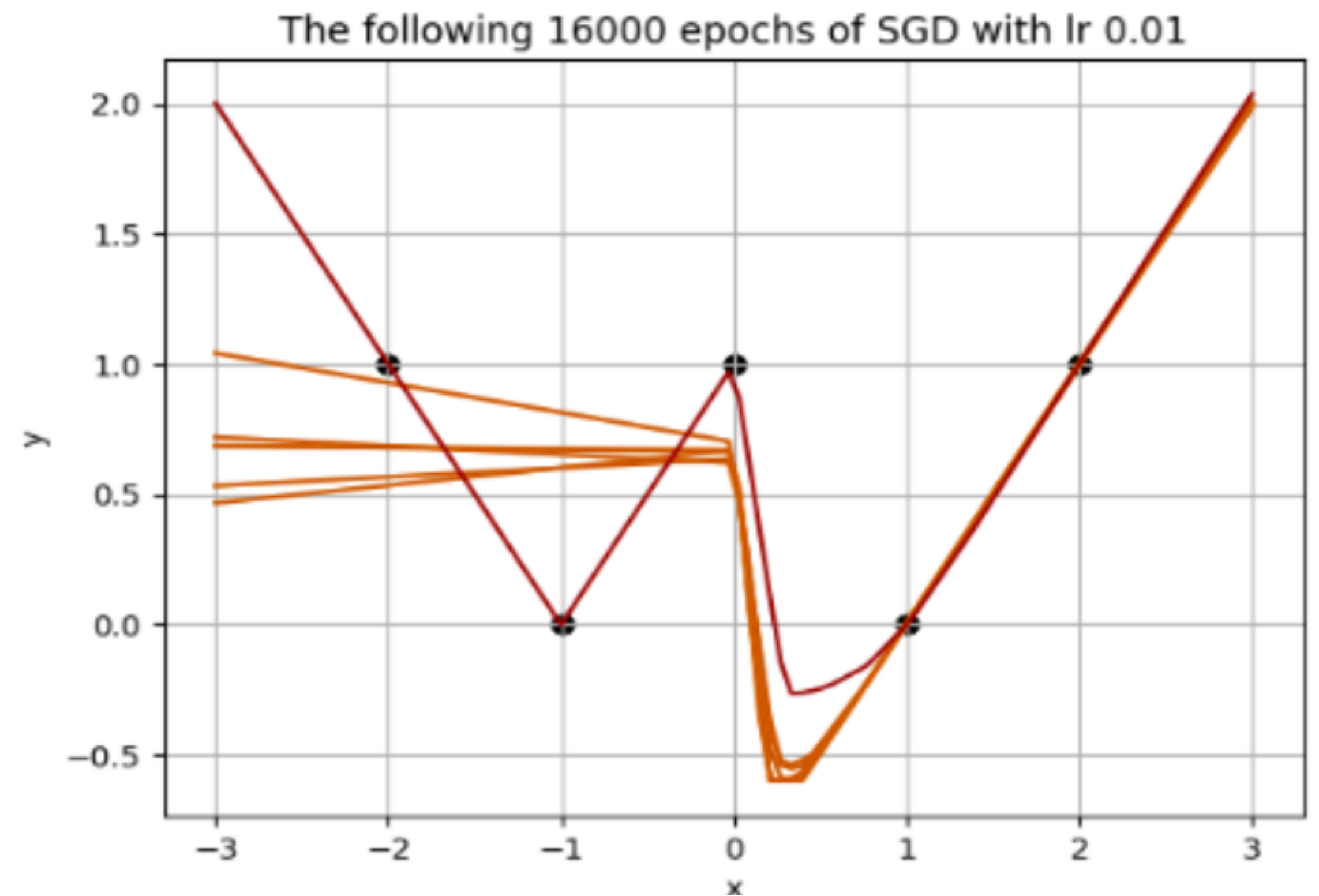
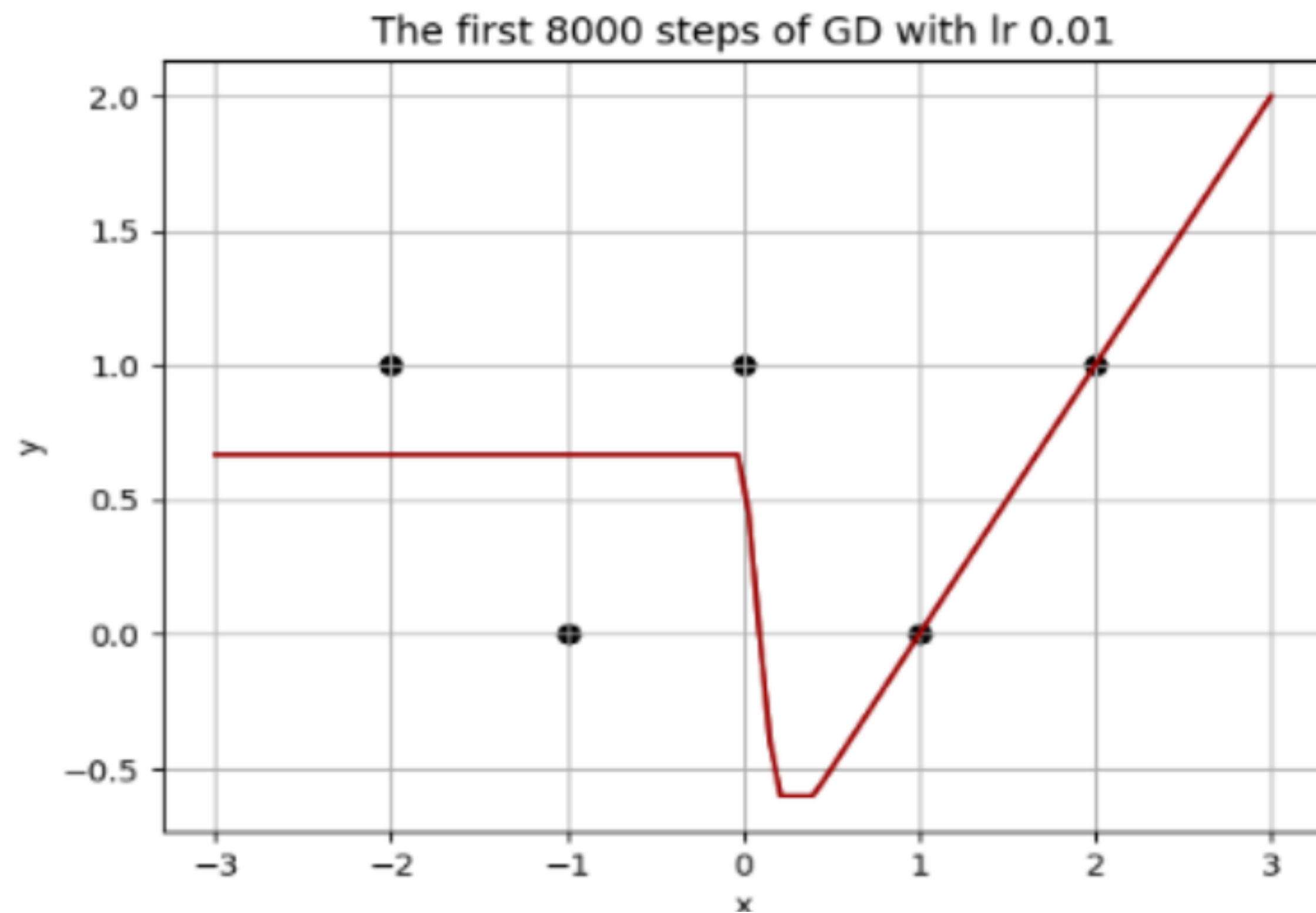
Location of Convergence

Example 1:

Different algorithms go to minima with different training loss.

SGD vs GD

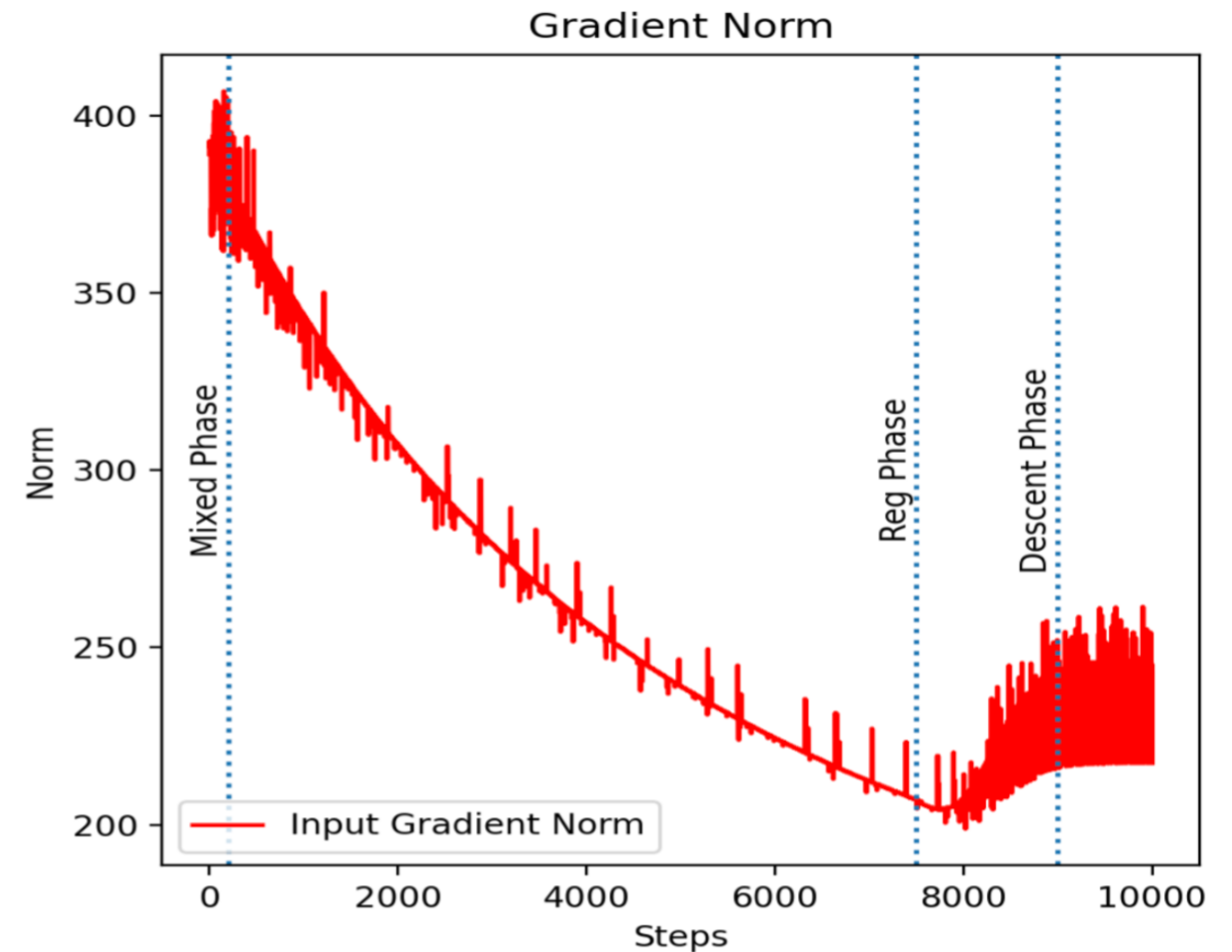
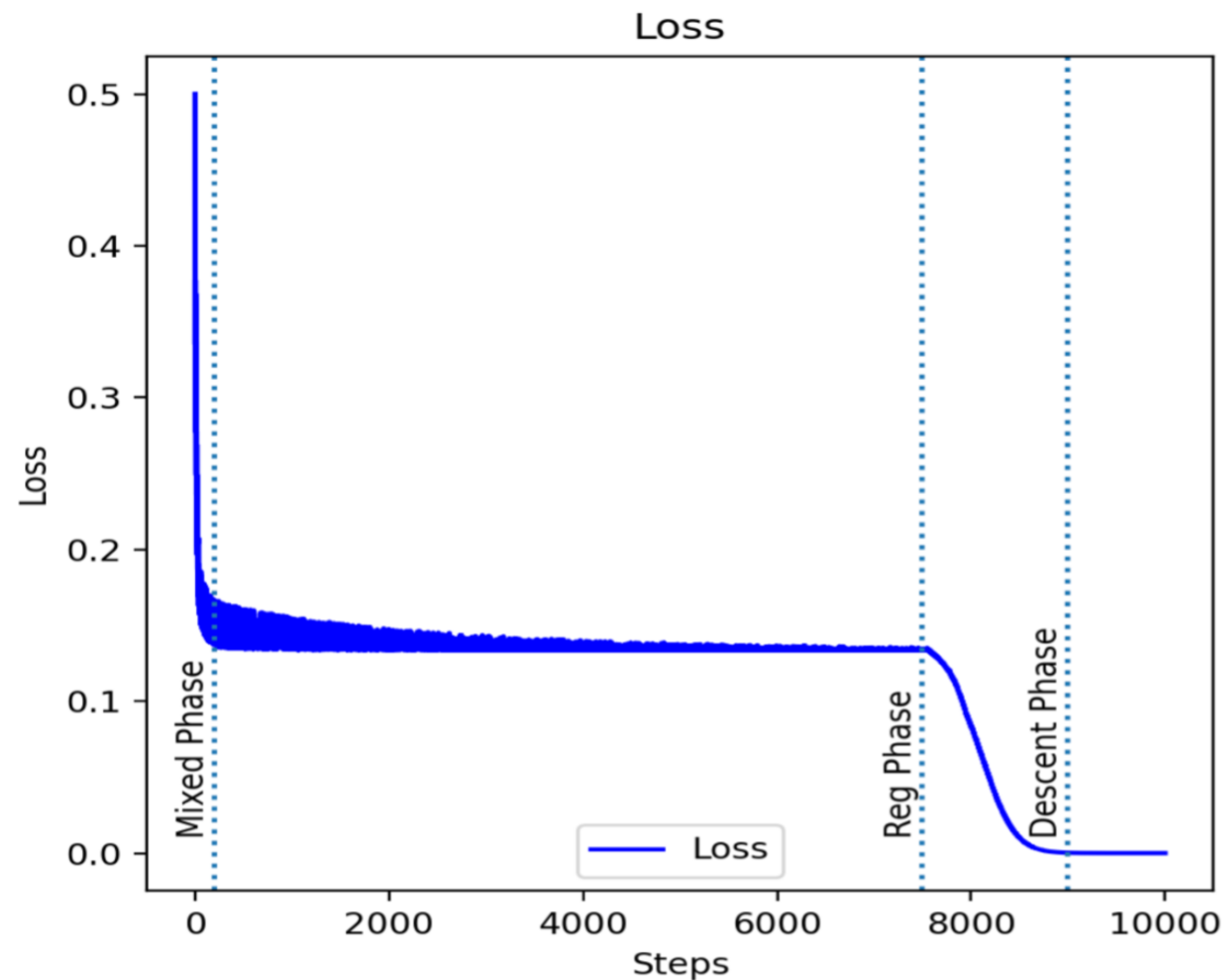
Practical example 1: ReLU network on W-data.



MSE on ReLU Shallow Network with 20 neurons

Implicit Regularization

SGD is cooking something down there



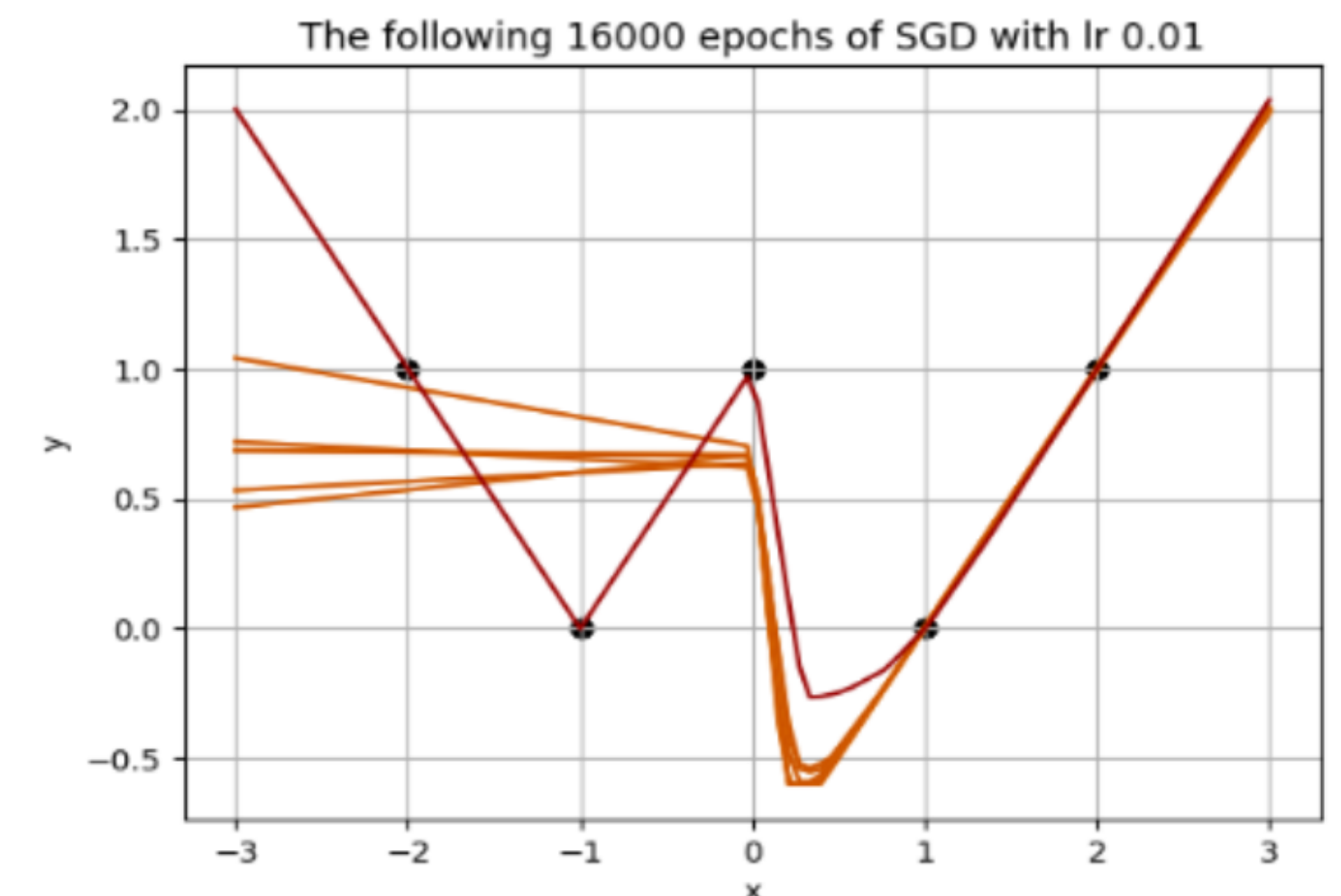
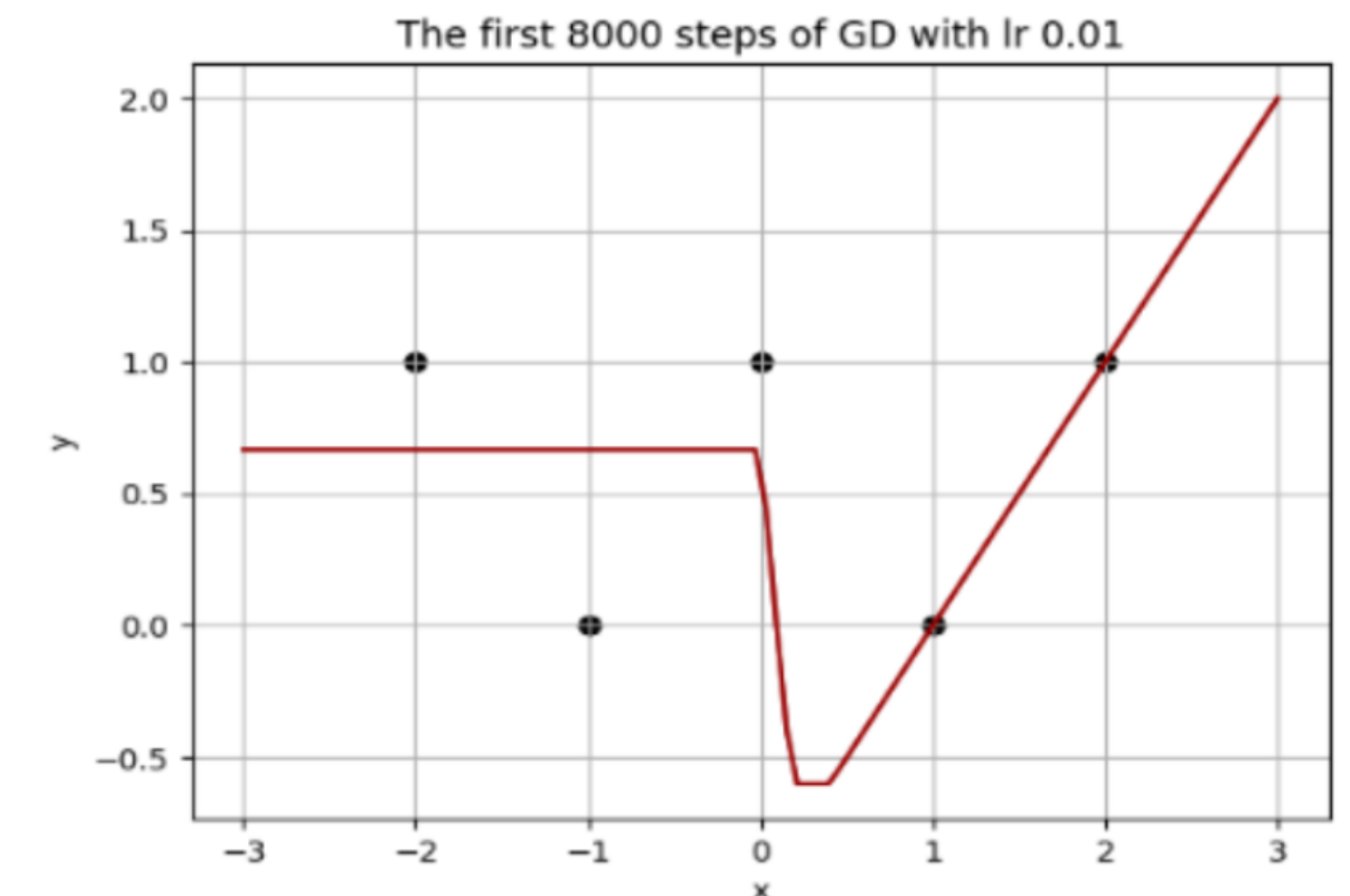
SGD vs GD

Practical example 1: ReLU network on W-data.

- There are many minima.
- Which one you find matters.



- GD stops in a bad one (for the training set).
- SGD escapes to a good one—through flat areas



Training Networks

What **qualitatively** matters empirically:

- Deterministic **vs** Stochastic

Sampling the Batches

- **With Replacement:** Sample independently at every step b data points.
- **Without Replacement:** Reshuffle the data and partition it in n/b steps.

The diagram shows the parameter update equation $\theta_{i+1} = \theta_i - \frac{\eta}{b} \sum_{z \in B_{i+1}} \nabla_{\theta} L(\theta_i, z)$ with four callout boxes: 'Parameters at i+1-th step' pointing to θ_{i+1} , 'step size batch size' pointing to $\frac{\eta}{b}$, 'i+1-th batch of training data' pointing to B_{i+1} , and 'Gradient of the loss' pointing to $\nabla_{\theta} L(\theta_i, z)$.

$$\theta_{i+1} = \theta_i - \frac{\eta}{b} \sum_{z \in B_{i+1}} \nabla_{\theta} L(\theta_i, z), \quad i = 0, 1, \dots,$$

Parameters at i+1-th step

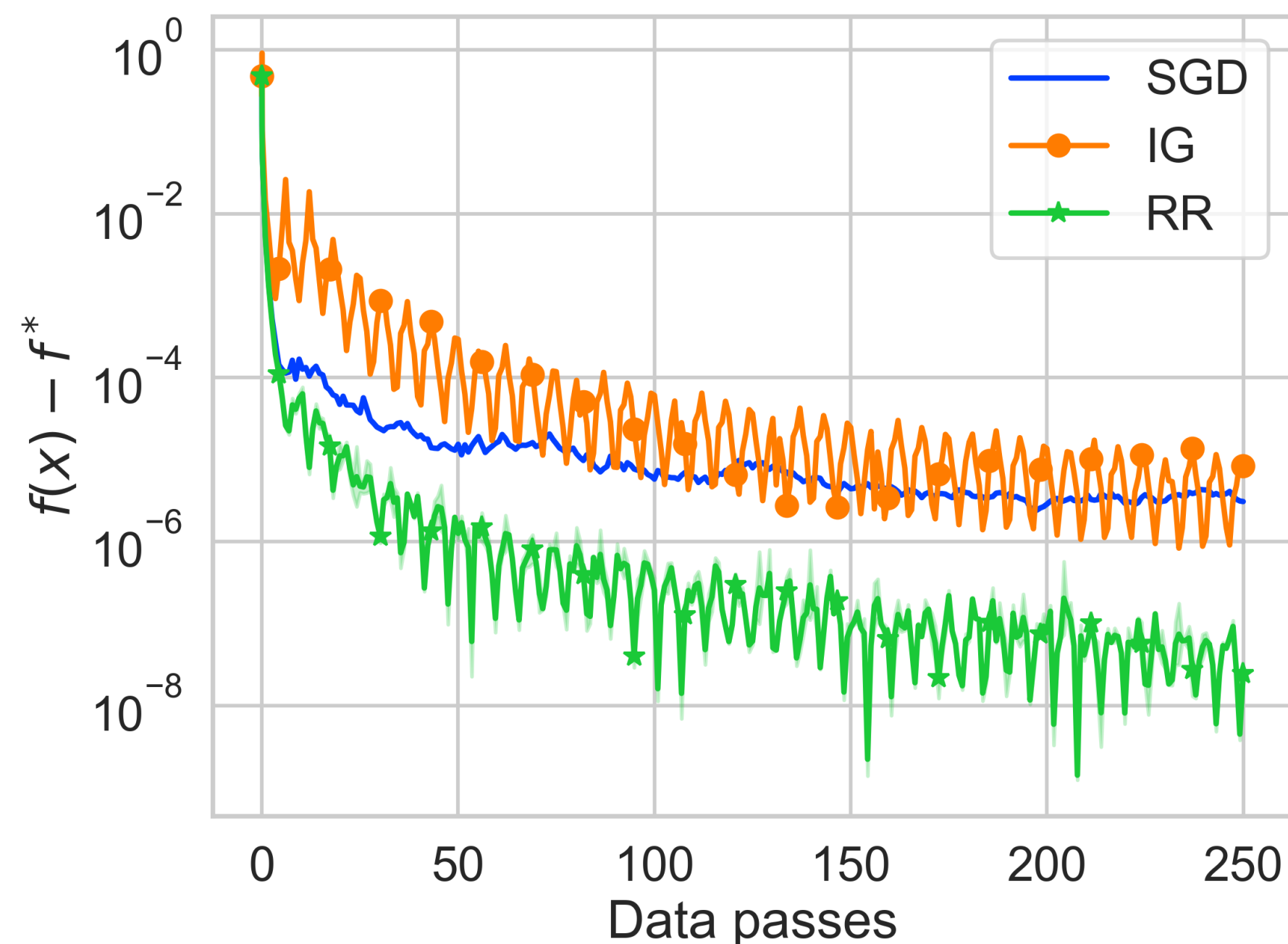
step size batch size

i+1-th batch of training data

Gradient of the loss

SGD *is* Without Replacement

SGD without replacement is the standard in Pytorch and TensorFlow



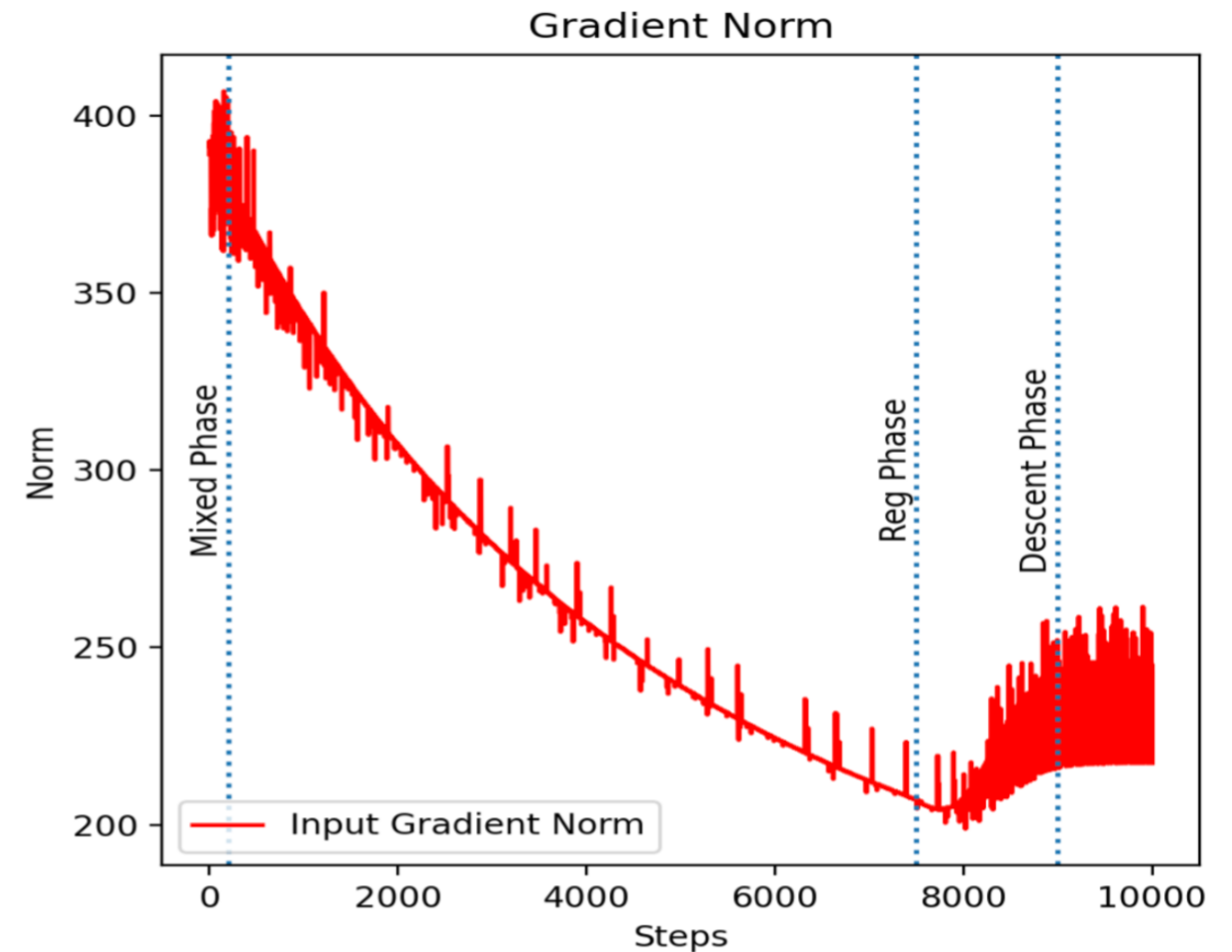
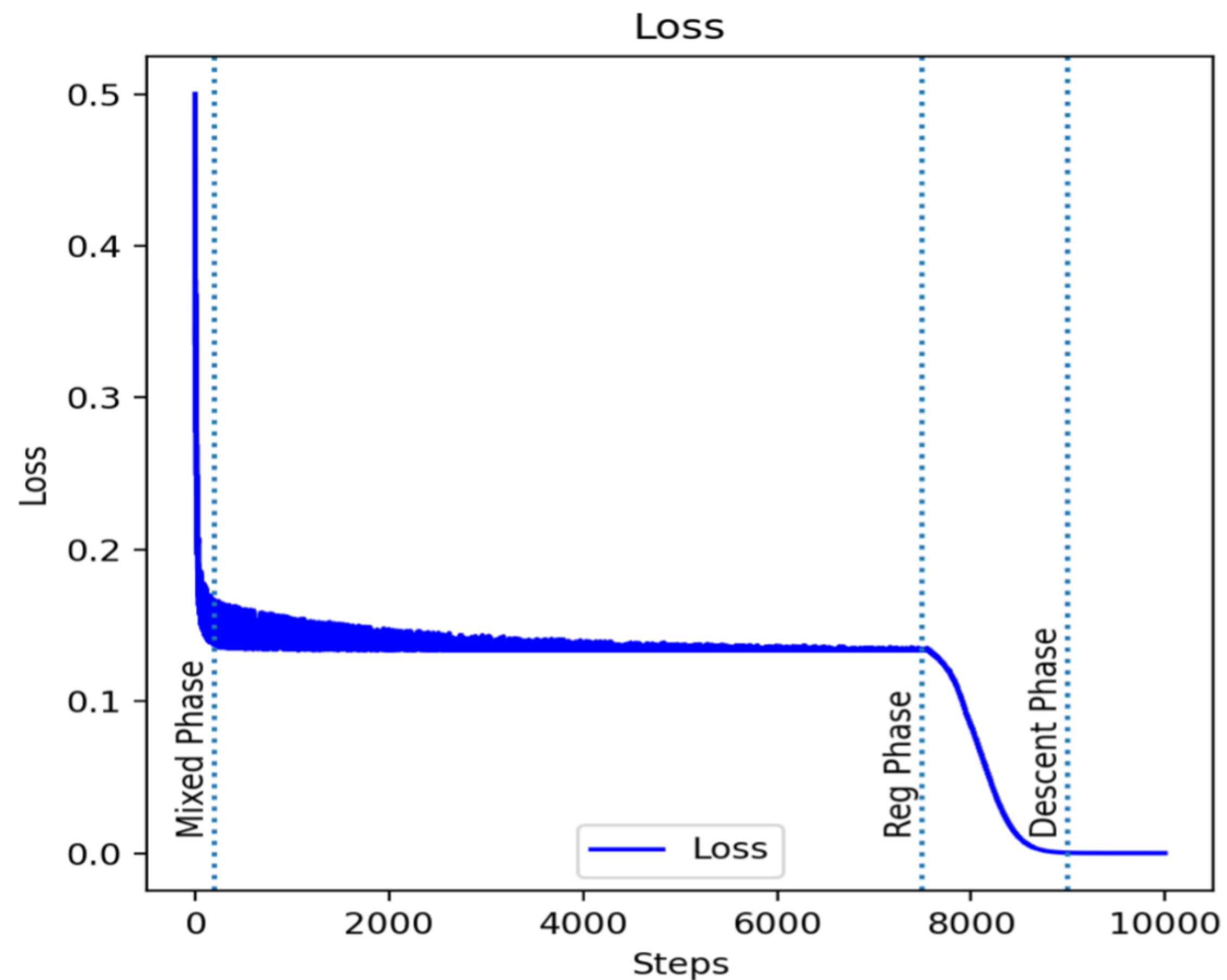
Generated with code based
on [\(Mishchenko et al., 2020\)](#)

1. It is the one that better generalizes in practice ([Bottou, 2009](#); [Recht and Ré, 2013](#)).
2. It is *faster* (step and convergence) ([Bottou, 2009](#); [Bengio, 2012](#); [Mishchenko et al., 2020](#)).

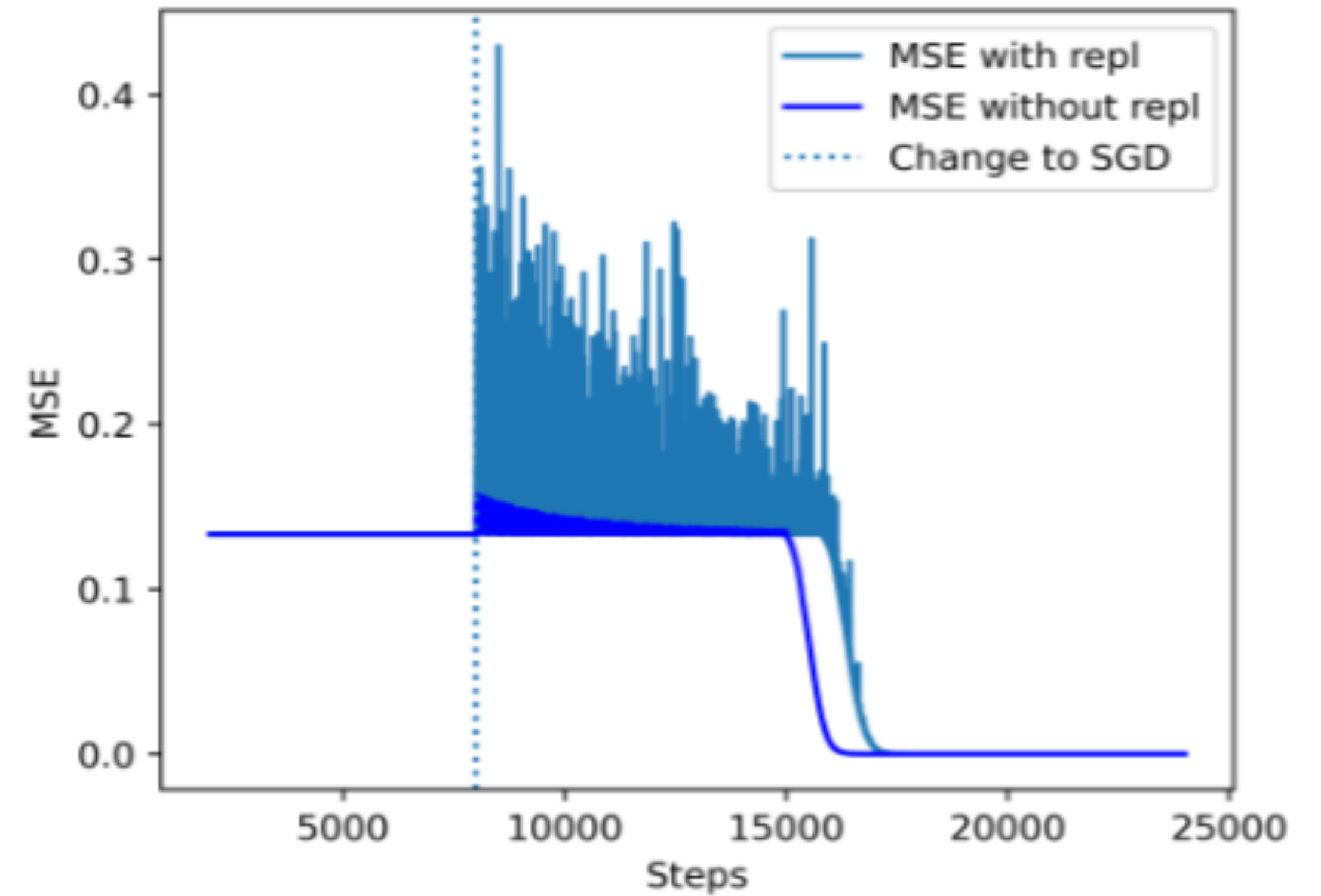
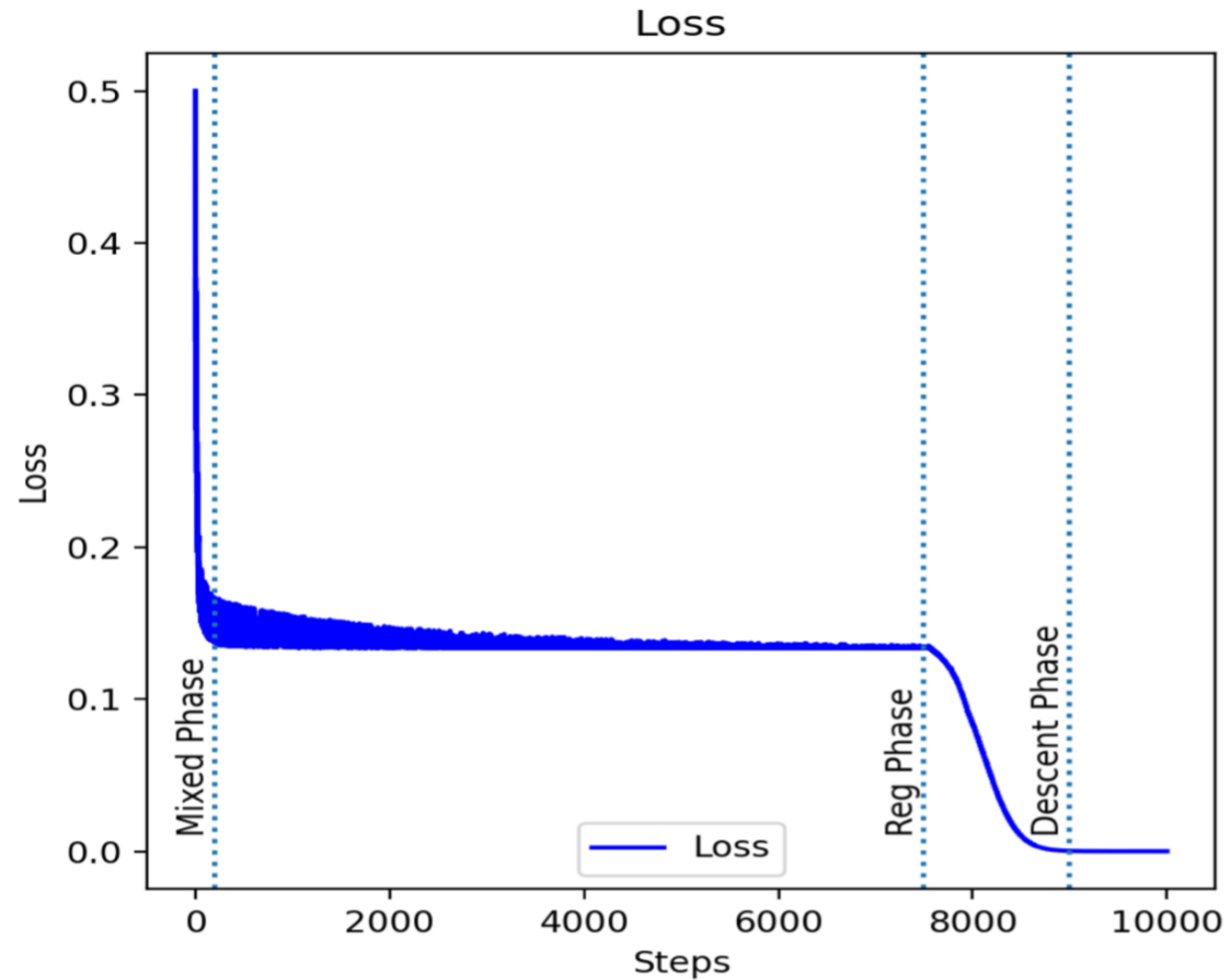
can not assume i.i.d. steps

Implicit Regularization

SGD is cooking something down there



Implicit Regularization



Same hyperparameters

Training Networks

What **qualitatively** matters empirically:

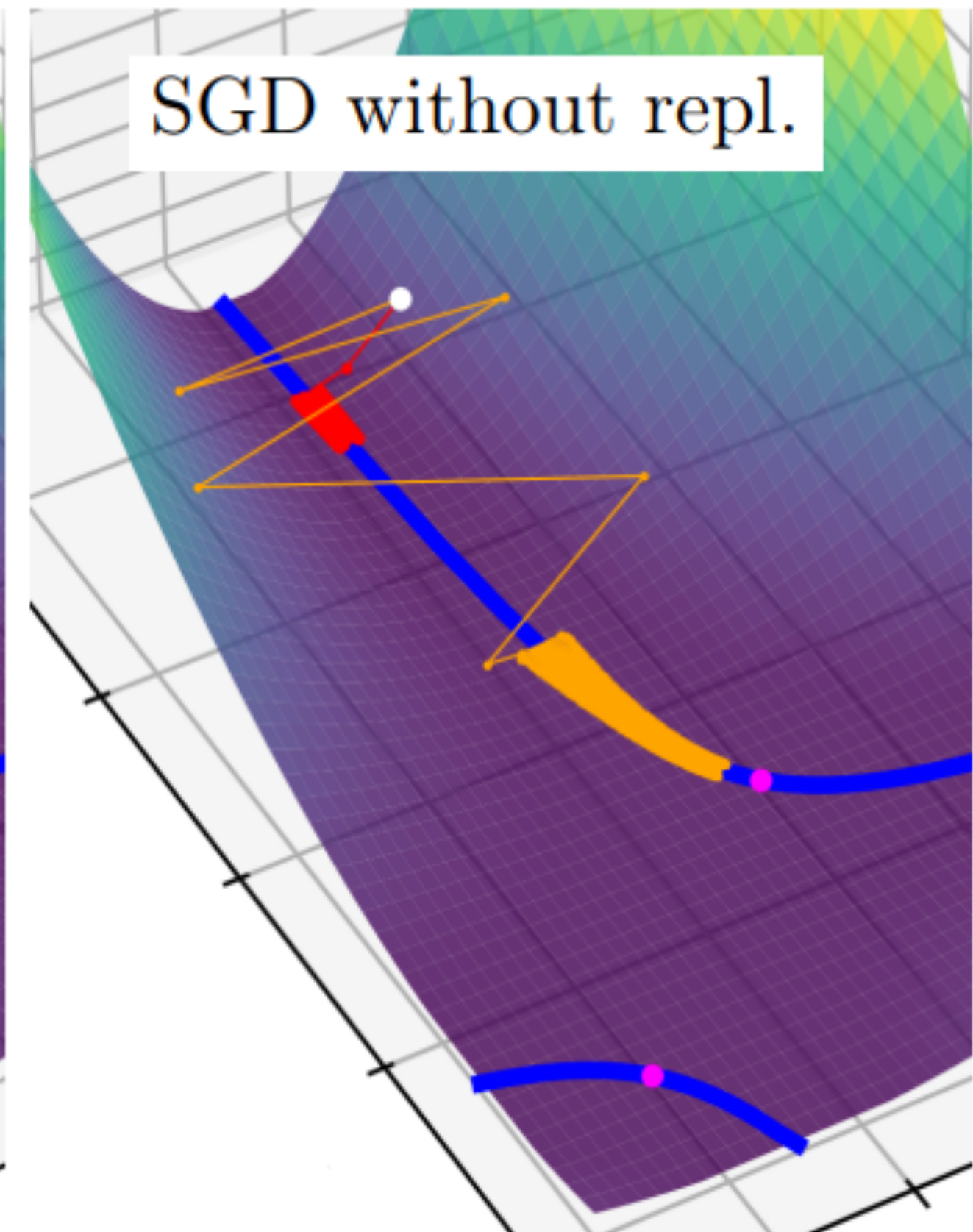
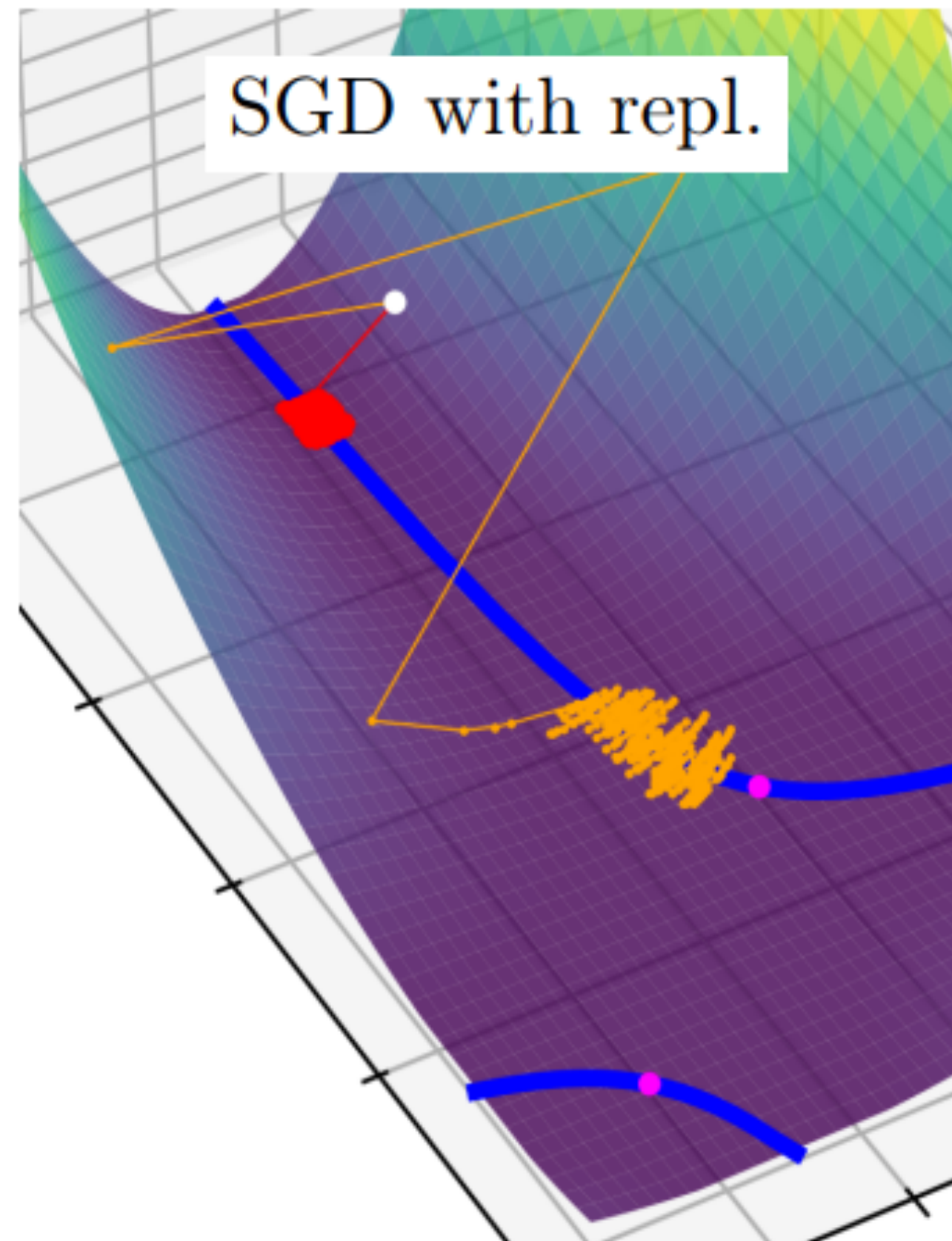
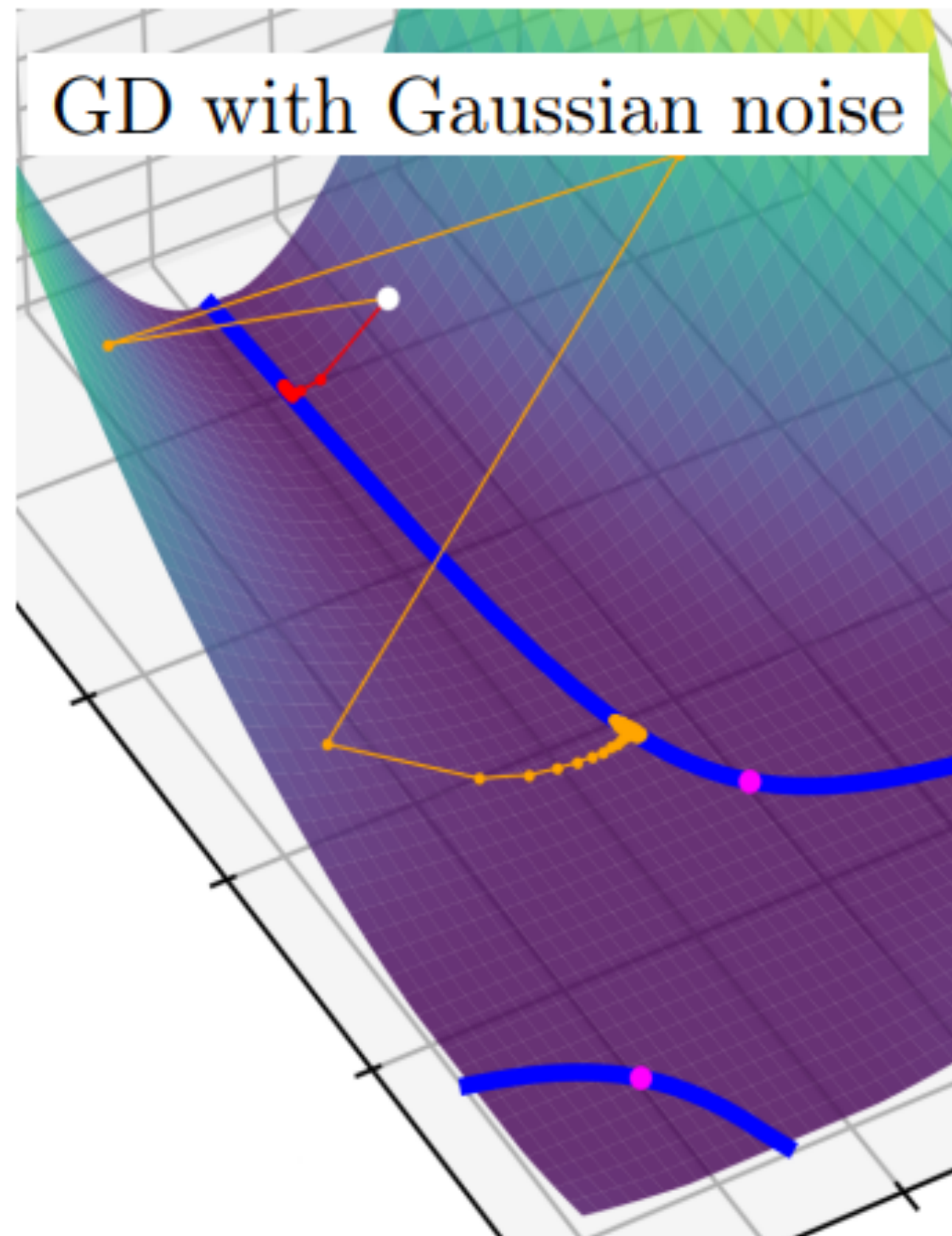
- Deterministic **vs** Stochastic
- What **kind of** stochasticity

Location of Convergence

Example 2:

Same function represented but lower norm representation.

Comparing SGDs



$$\frac{1}{2n} \sum_{i=1}^n (\theta_1 \cdot \theta_2 \cdot x_i - y_i)^2 \quad \text{with dataset } \{(x_i, y_i)\}_{i=1}^n \subseteq \mathbb{R} \times \mathbb{R} \text{ and parameters } (\theta_1, \theta_2) \in \mathbb{R}^2$$

Location of Convergence

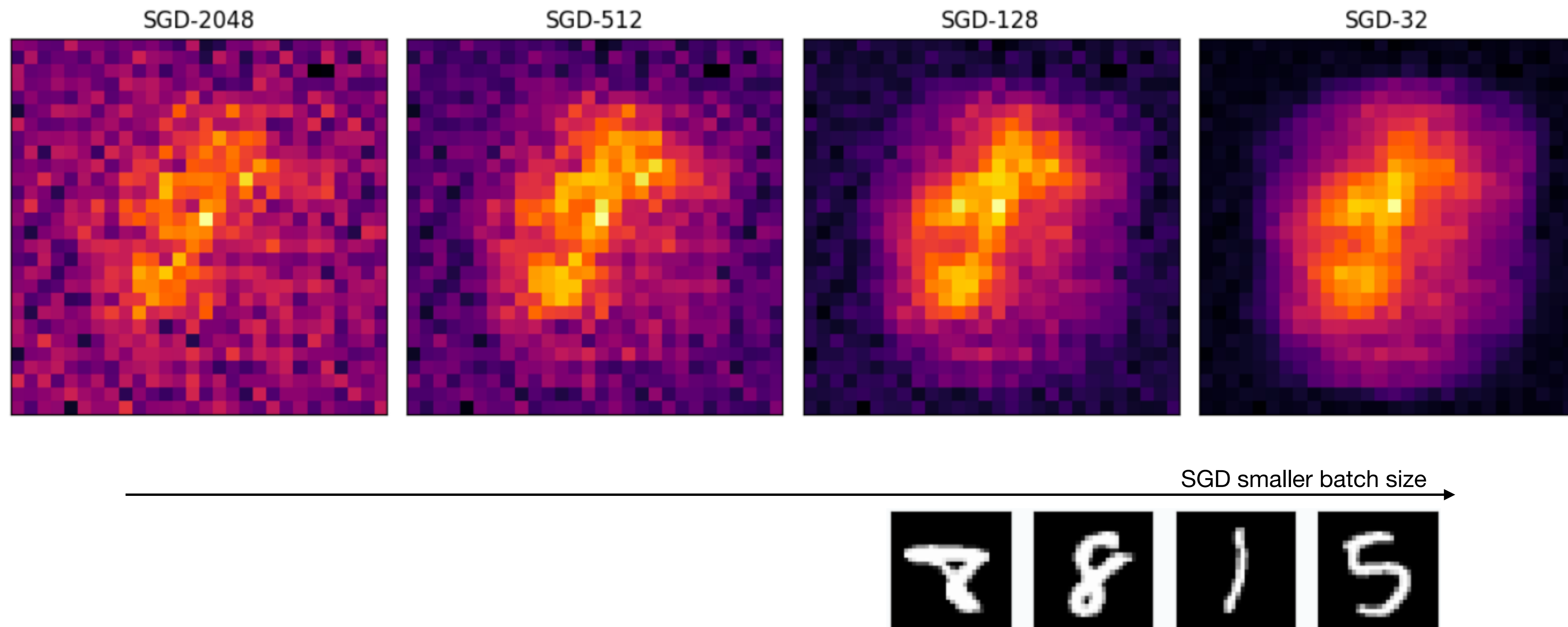
Example 3:

Different algorithms go to minima with the same training loss
but different test loss.

Support Selection

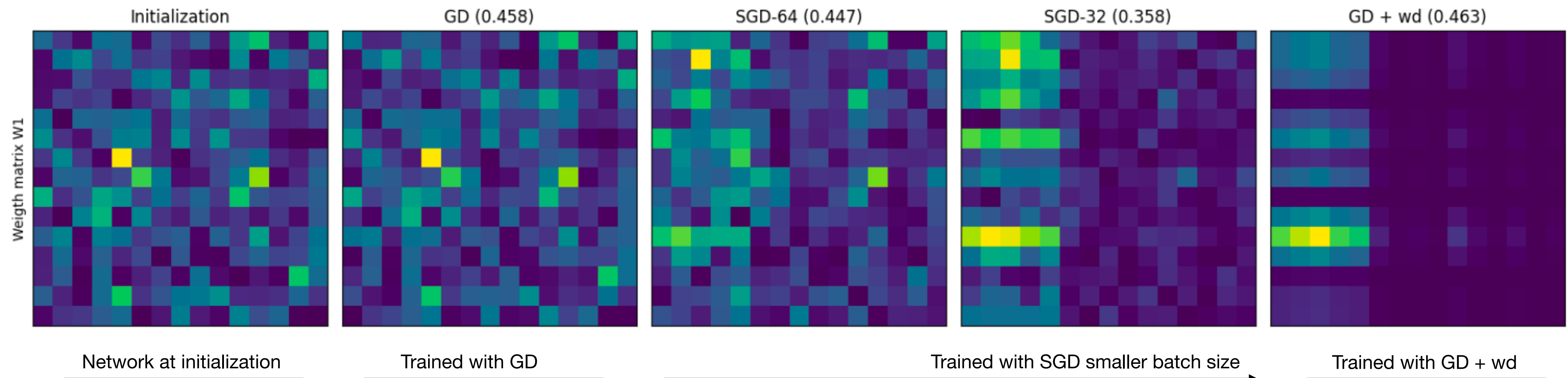
SGD identifies support in the first layer, GD does not

- Size of the first layer weights of an MLP learning MNIST with different optimizers



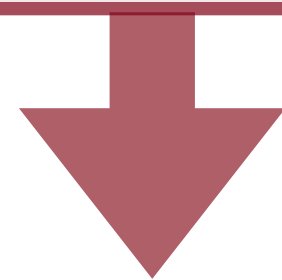
Support Selection

- How neural networks learn features from the data depends on the optimizer.
- Lower batch size yields to more compressible, interpretable, networks.
- Transfer learning may work better $\implies$ better generalization.

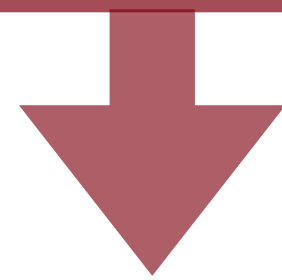


THE MOTIVATION

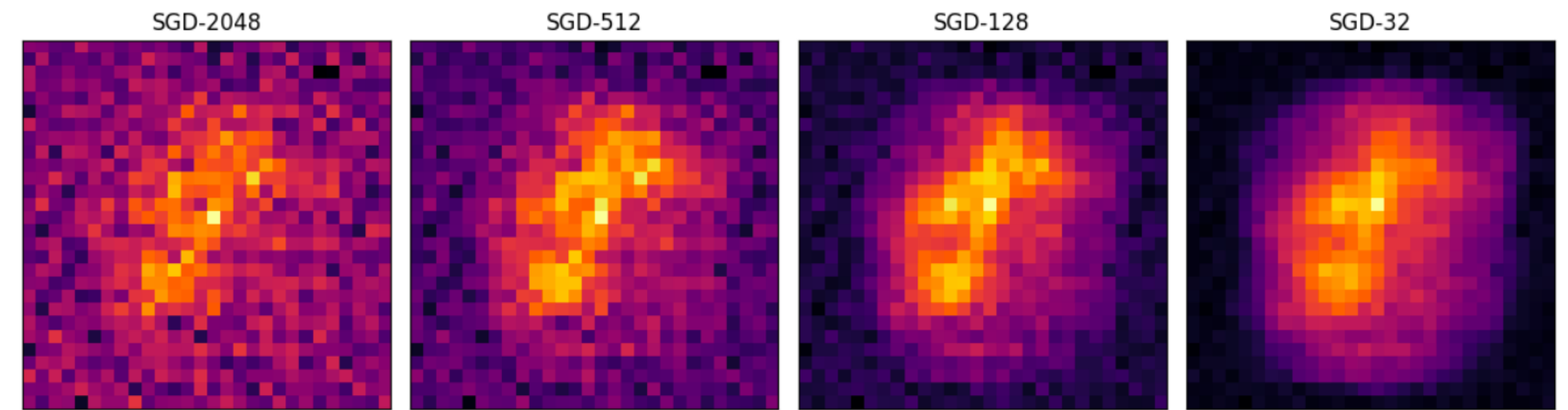
Goal is having best
“performance”



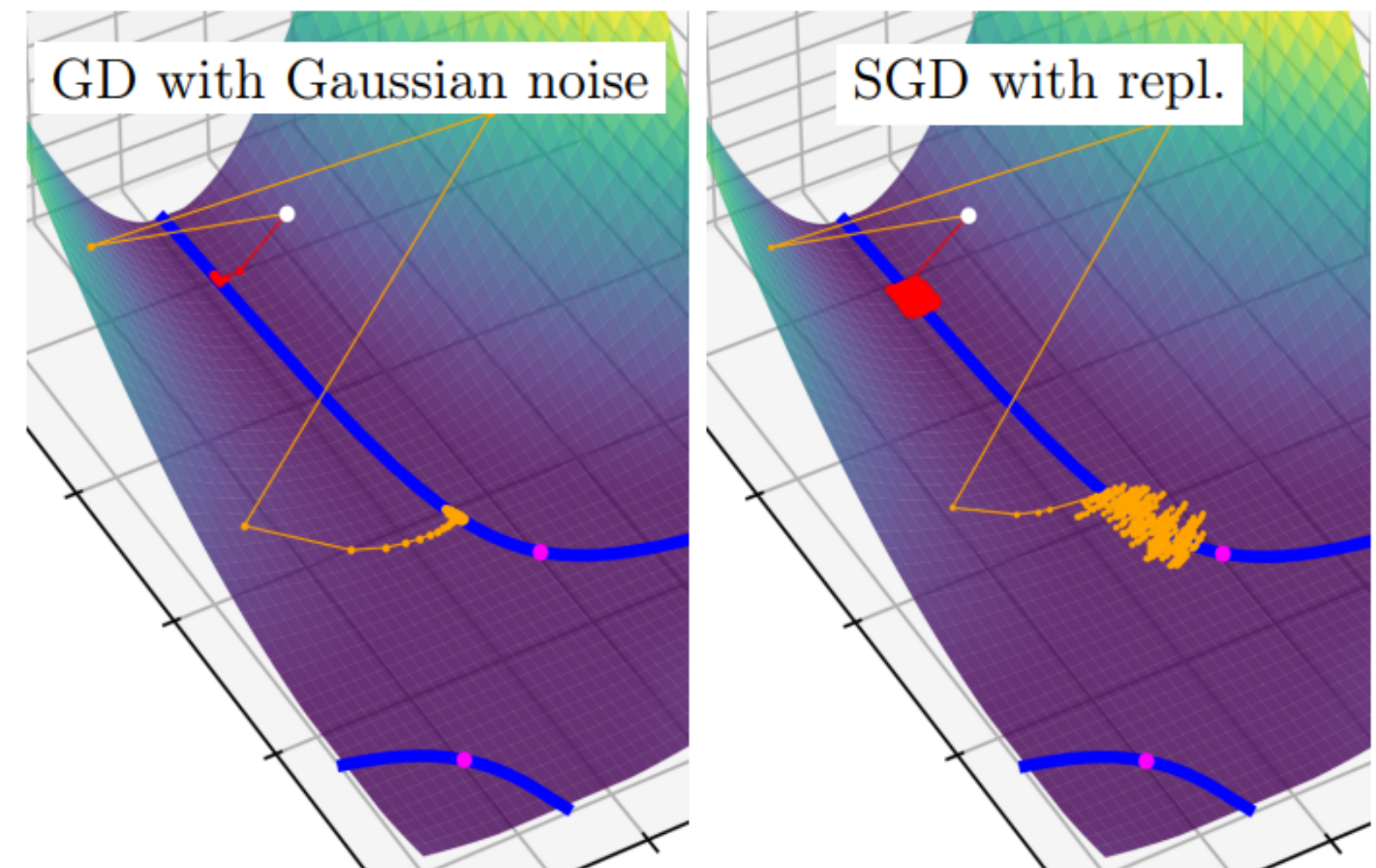
There are many different minima
with different performance



Different algorithms find
different minima, *consistently*



SGD smaller batch size



Training Networks

What **qualitatively** matters empirically *(for performance in Regime 2)*:

- Deterministic **vs** Stochastic
- What **kind of** stochasticity
- The choice of **hyperparameters**

Consistently

LOCATION OF CONVERGENCE

- **Location of Convergence**
 - Some NN examples
 - **What gradient methods were**
- (Stable) Implicit Regularization
- Edge of Stability

WHAT IS GD

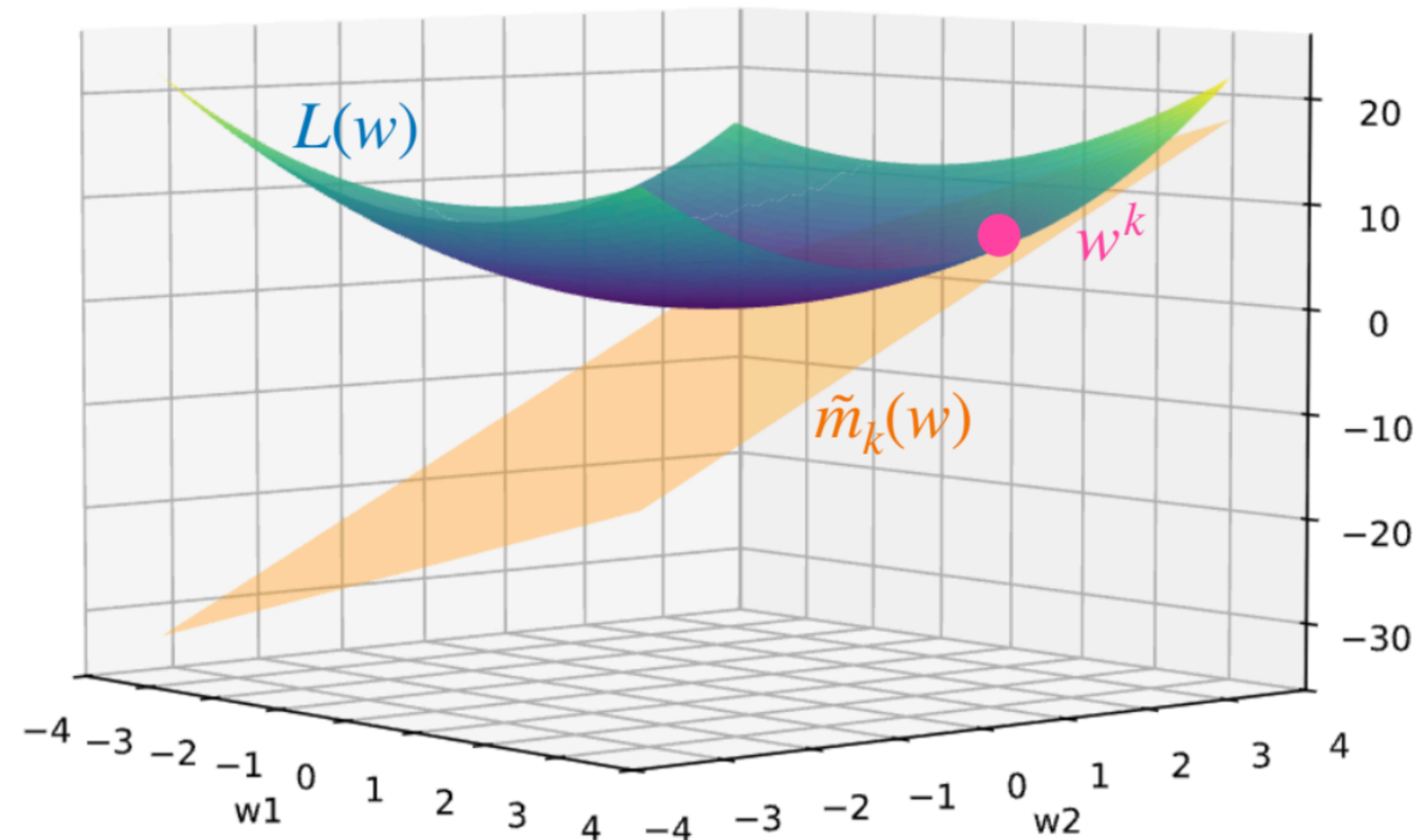
- Locally around θ_t we have the following approx (order 1):

$$\tilde{L}(\vartheta) := L(\theta_t) + \nabla L(\theta_t)^\top (\vartheta - \theta_t).$$

- What is the min in ϑ of $\tilde{L}(\vartheta)$?

Minus infinity.

This is not an accurate approximation.



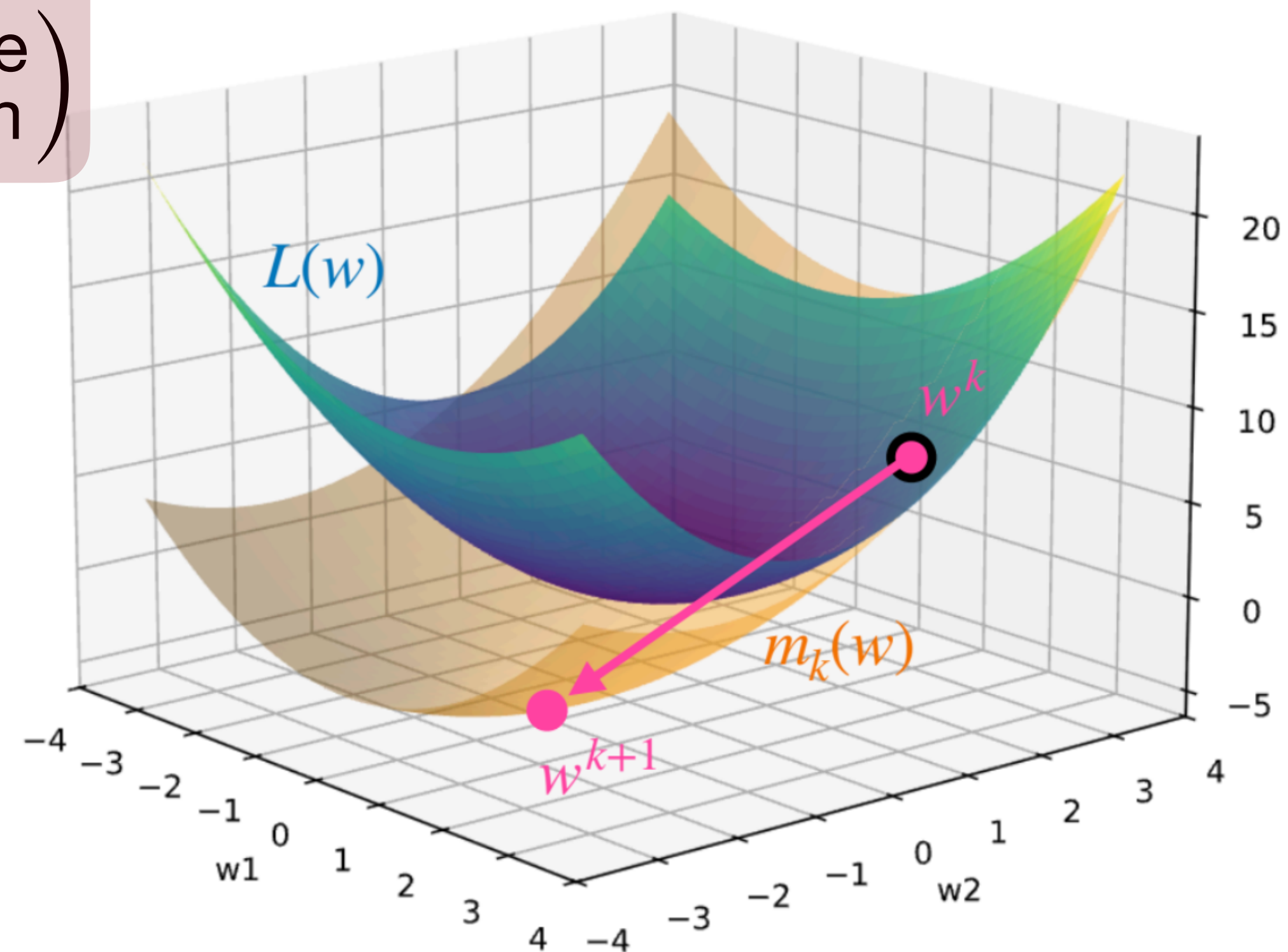
PROXIMAL METHODS

- Let say we trust it “ $1/\lambda$ ”, we want to stay closer:

$$\tilde{L}(\vartheta) := L(\theta_t) + \nabla L(\theta_t)^\top (\vartheta - \theta_t) + \lambda \|\vartheta - \theta_t\| \left(\begin{array}{c} \text{some} \\ \text{norm} \end{array} \right)$$

- Minimize $\tilde{L}(\theta_t)$ to obtain θ_{t+1} .

- Any norm.
- Any “trust” $1/\lambda$.



GD *IS* A PROXIMAL METHOD

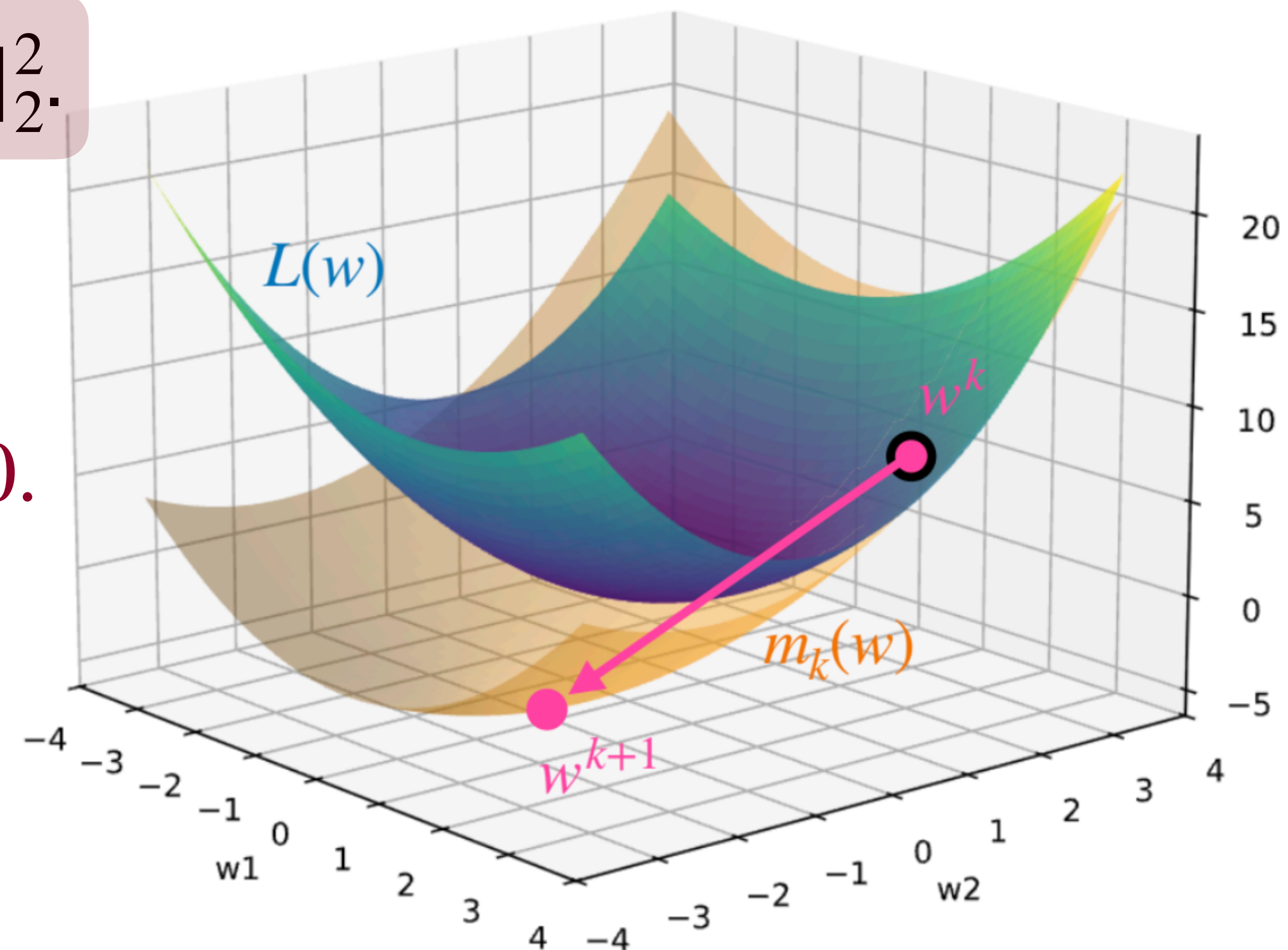
- If we trust it “ $1/\lambda$ ” with Euclidean norm:

$$\tilde{L}(\vartheta) := L(\theta_t) + \nabla L(\theta_t)^\top (\vartheta - \theta_t) + \lambda \|\vartheta - \theta_t\|_2^2.$$

- What is the minimum in ϑ of that?

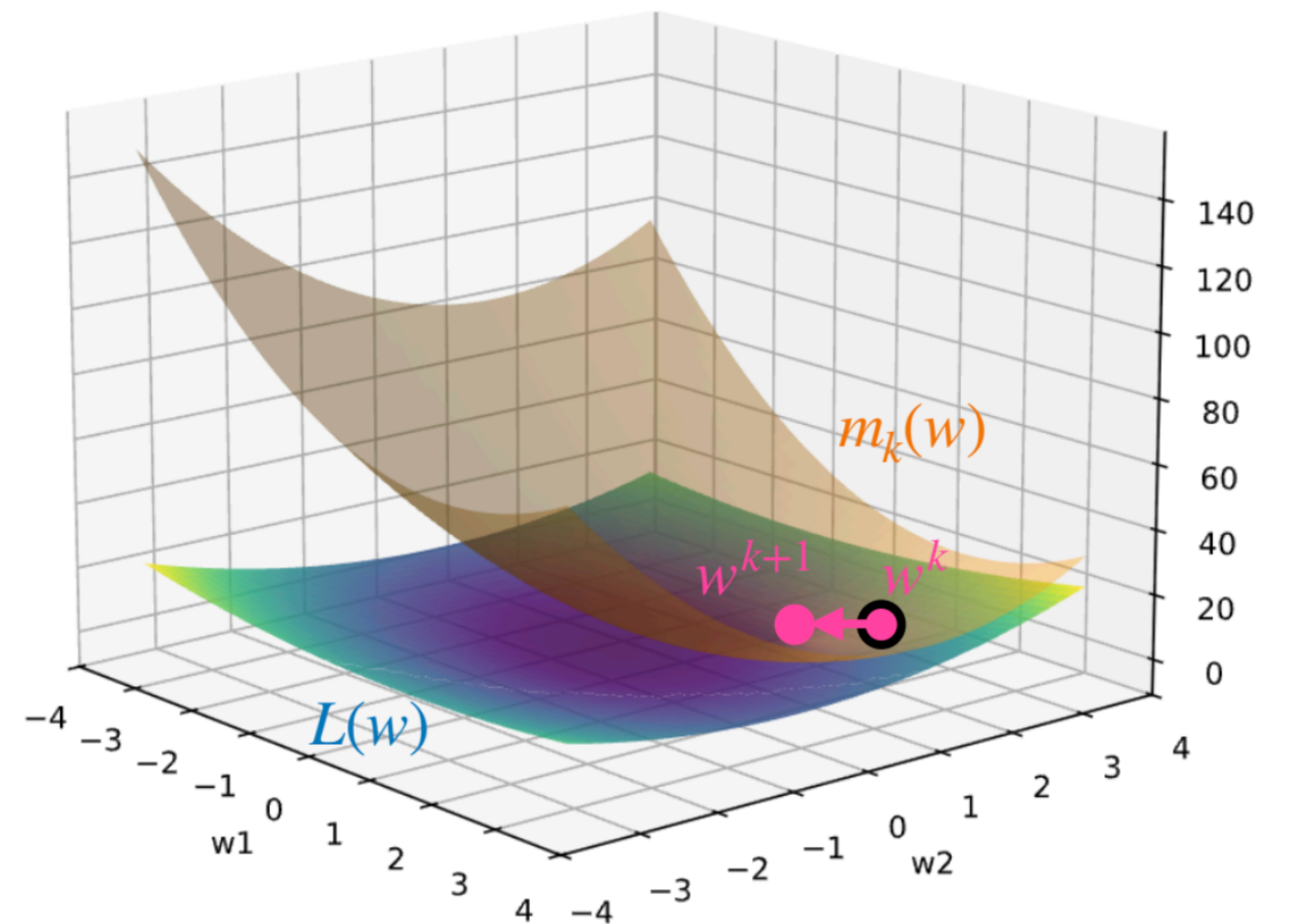
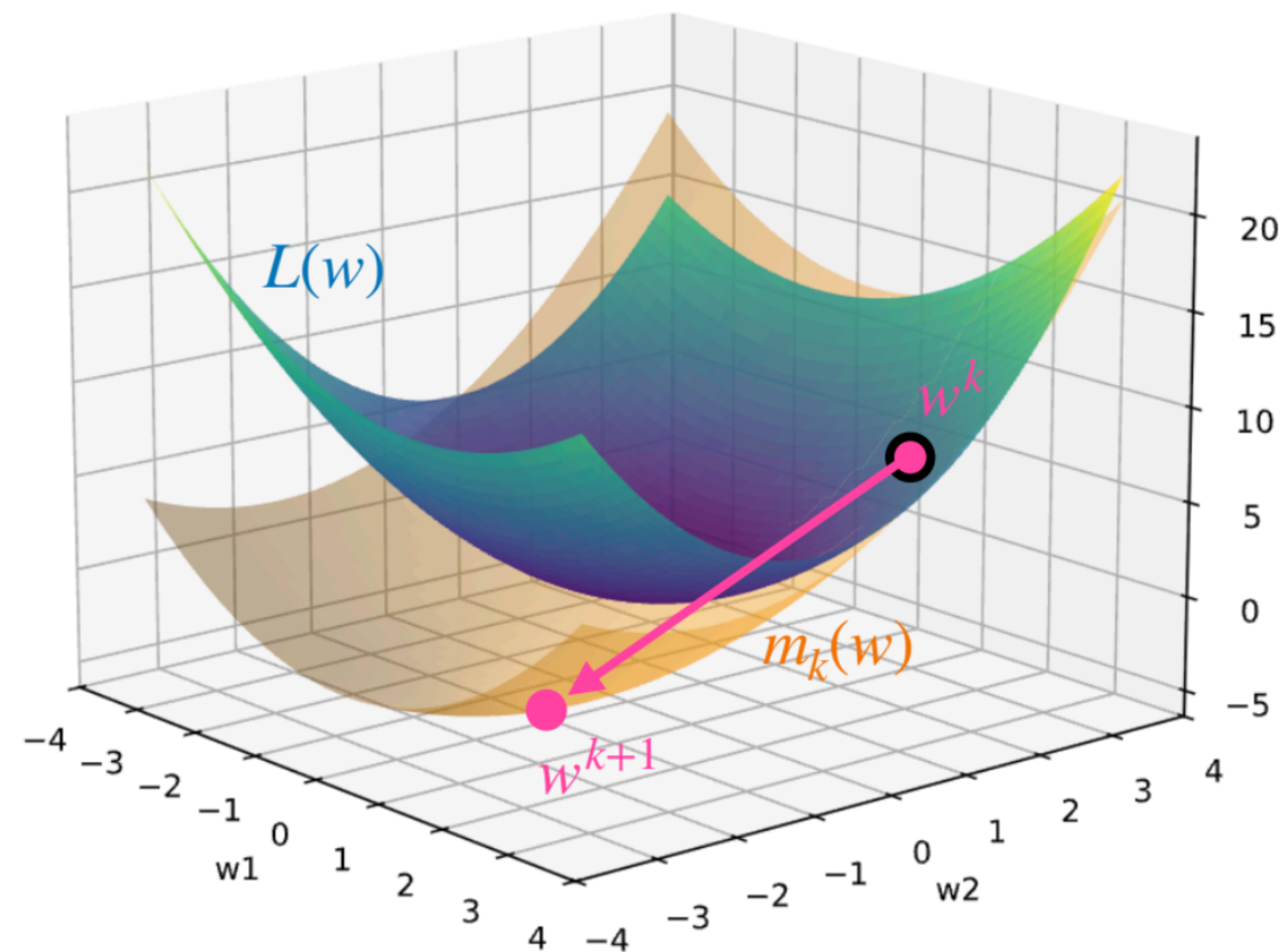
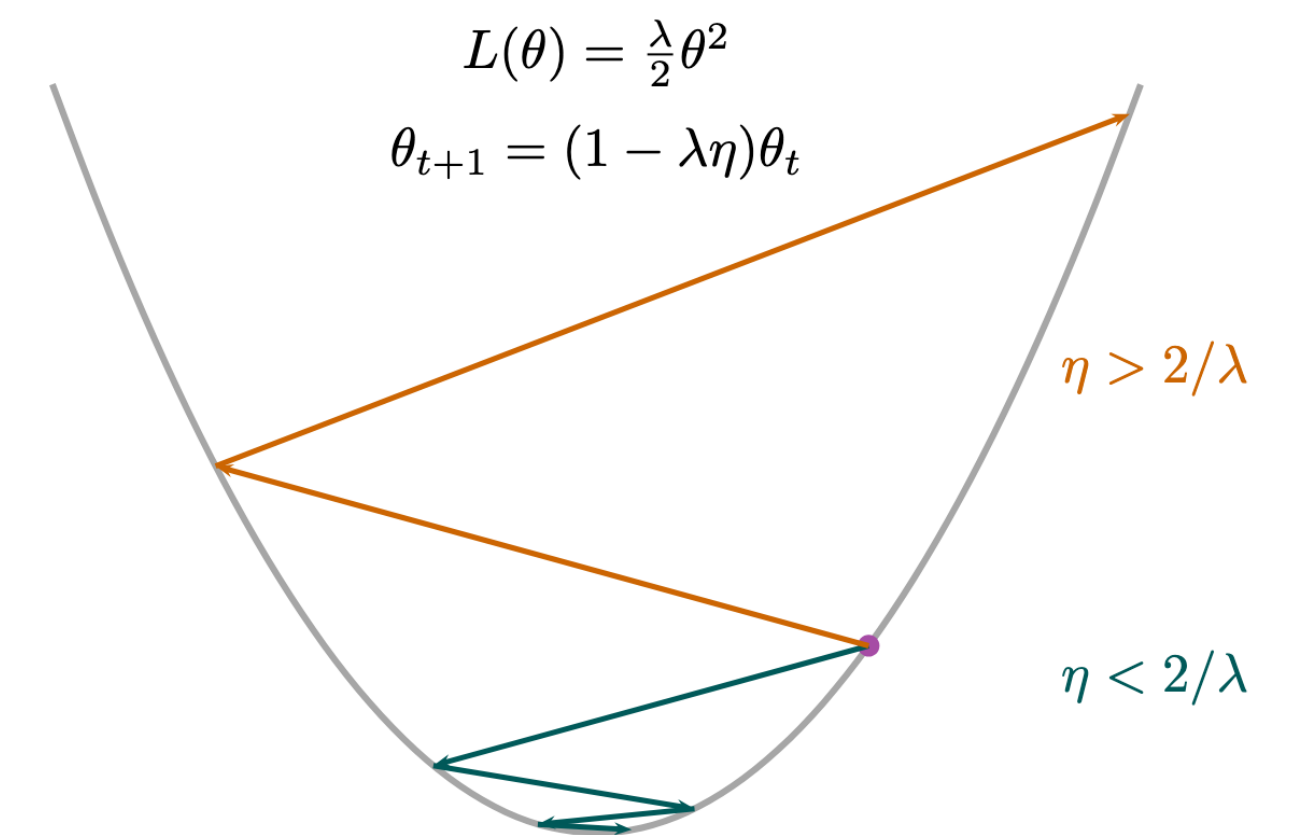
$$\operatorname{argmin}_{\vartheta} \left(\tilde{L}(\vartheta) \right) \iff \nabla L(\theta_t) + 2\lambda \cdot (\vartheta - \theta_t) = 0.$$

$$\vartheta - \theta_t = -\frac{1}{2\lambda} \nabla L(\theta_t).$$



WHAT IS GD

Increasing my *trust* $1/\lambda$:



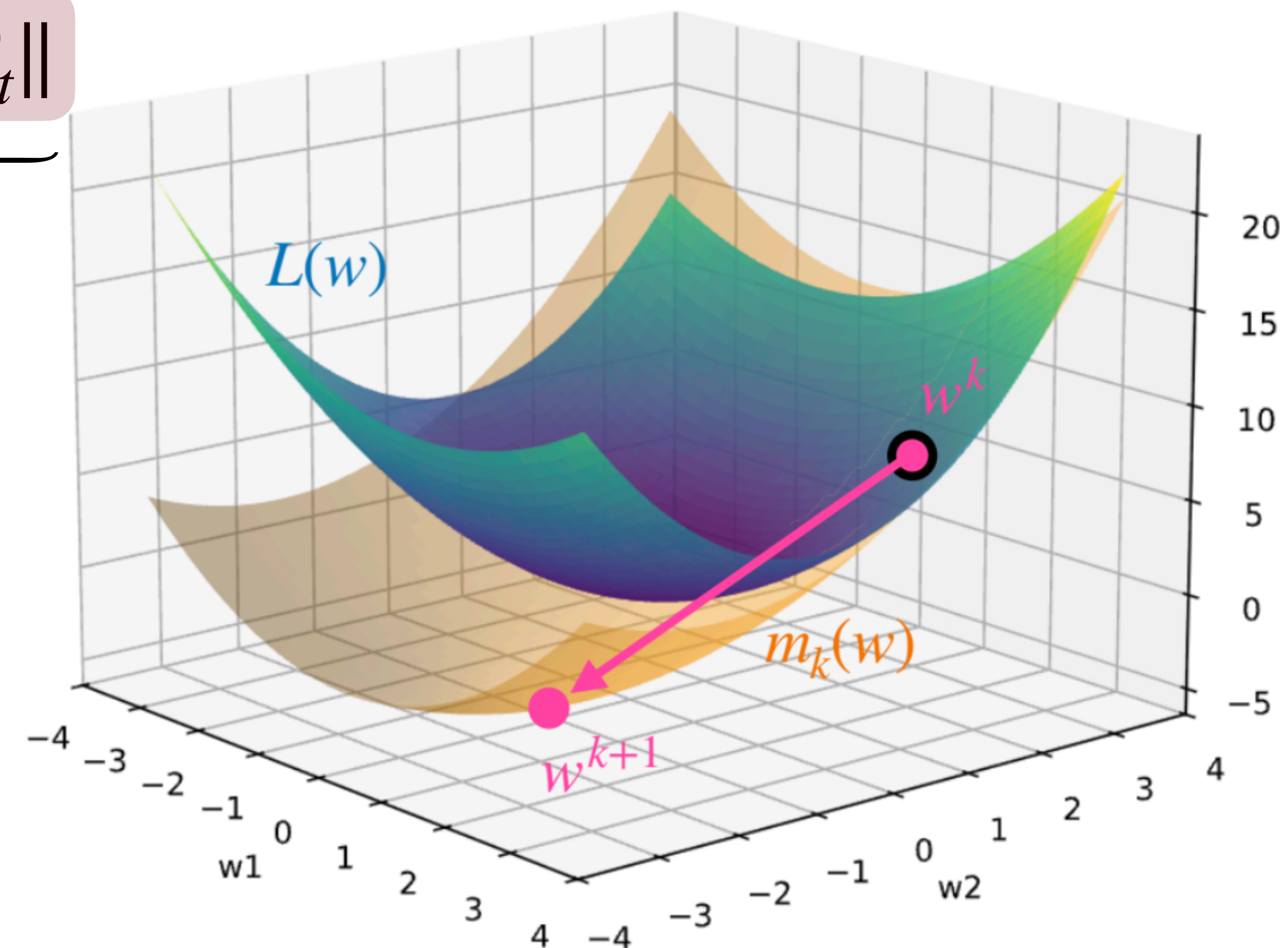
PROXIMAL METHODS

- Let say we trust it “ $1/\lambda$ ”, we want to stay closer:

$$\tilde{L}(\vartheta) := \underbrace{L(\theta_t) + \nabla L(\theta_t)^\top (\vartheta - \theta_t)}_{\text{Convex Function}} + \lambda \|\vartheta - \theta_t\|$$

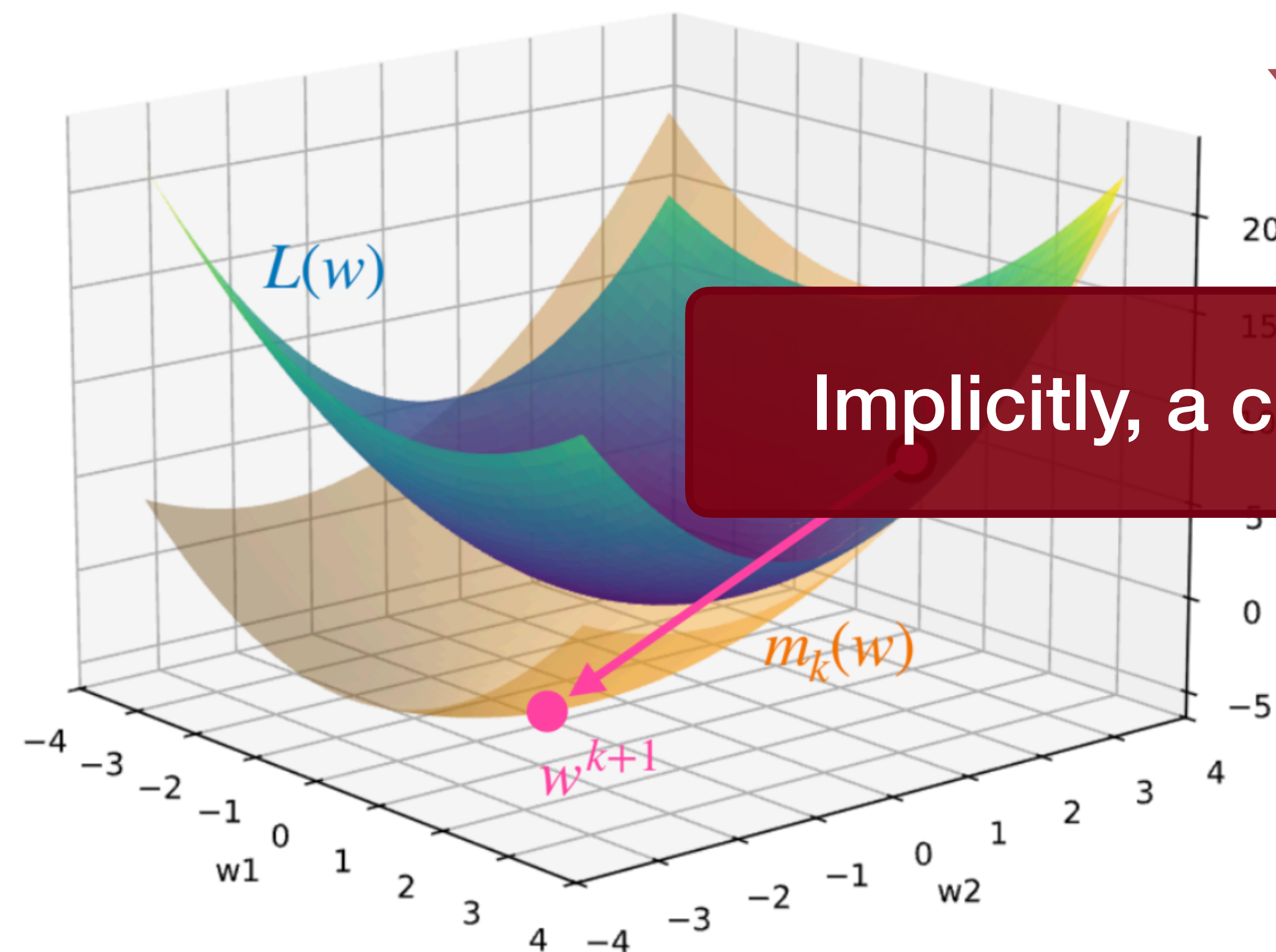
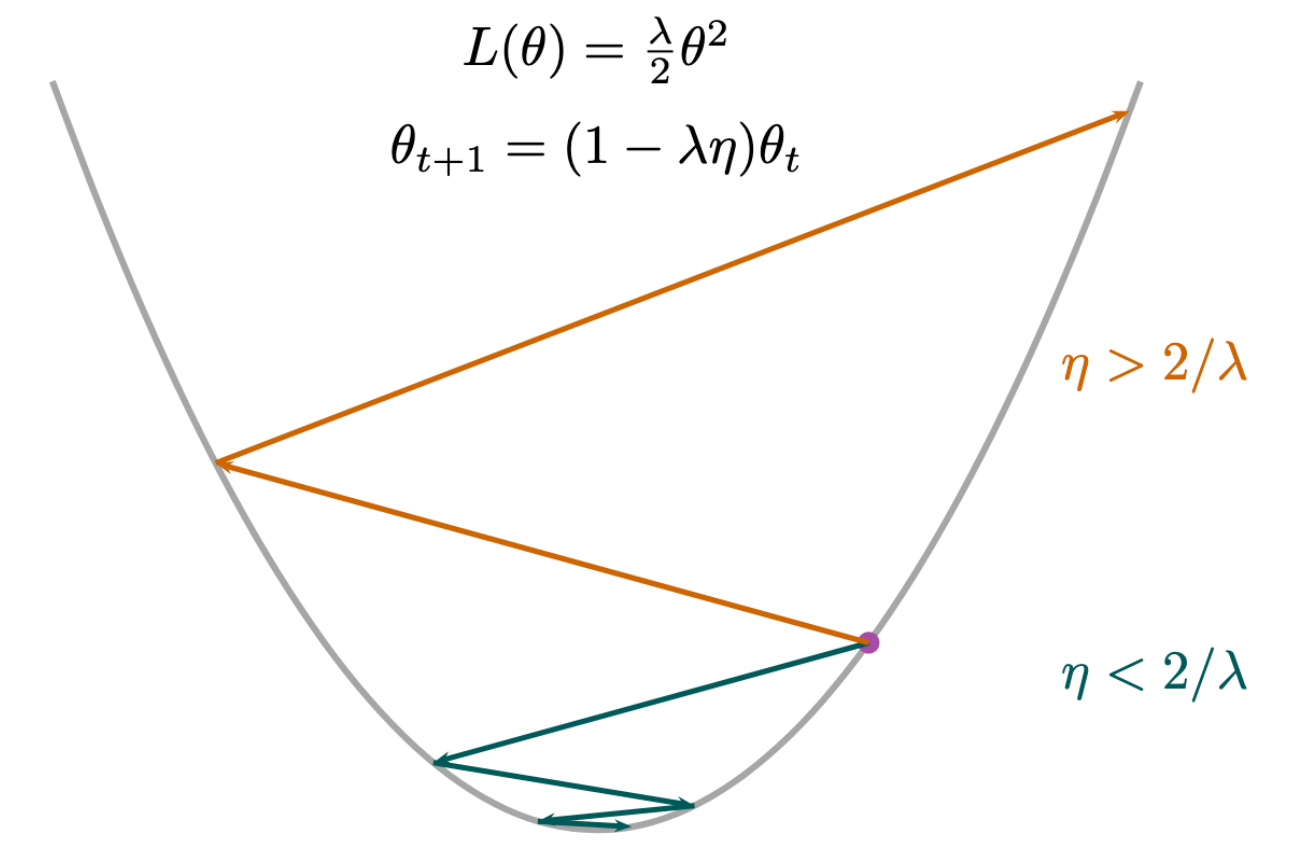
- Minimize $\tilde{L}(\theta_t)$ to obtain θ_{t+1} .

This is a convex relaxation!!

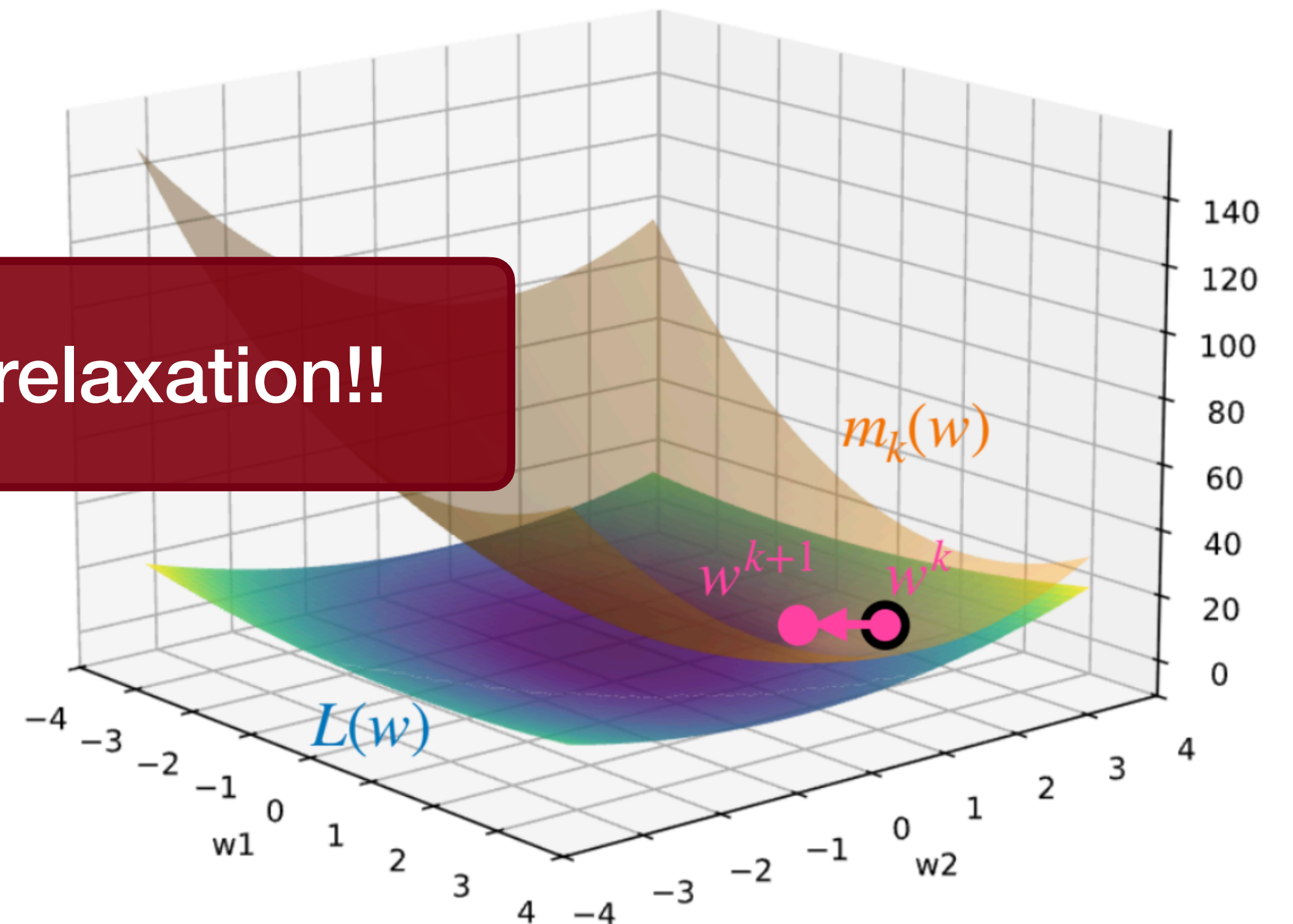


WHAT IS GD

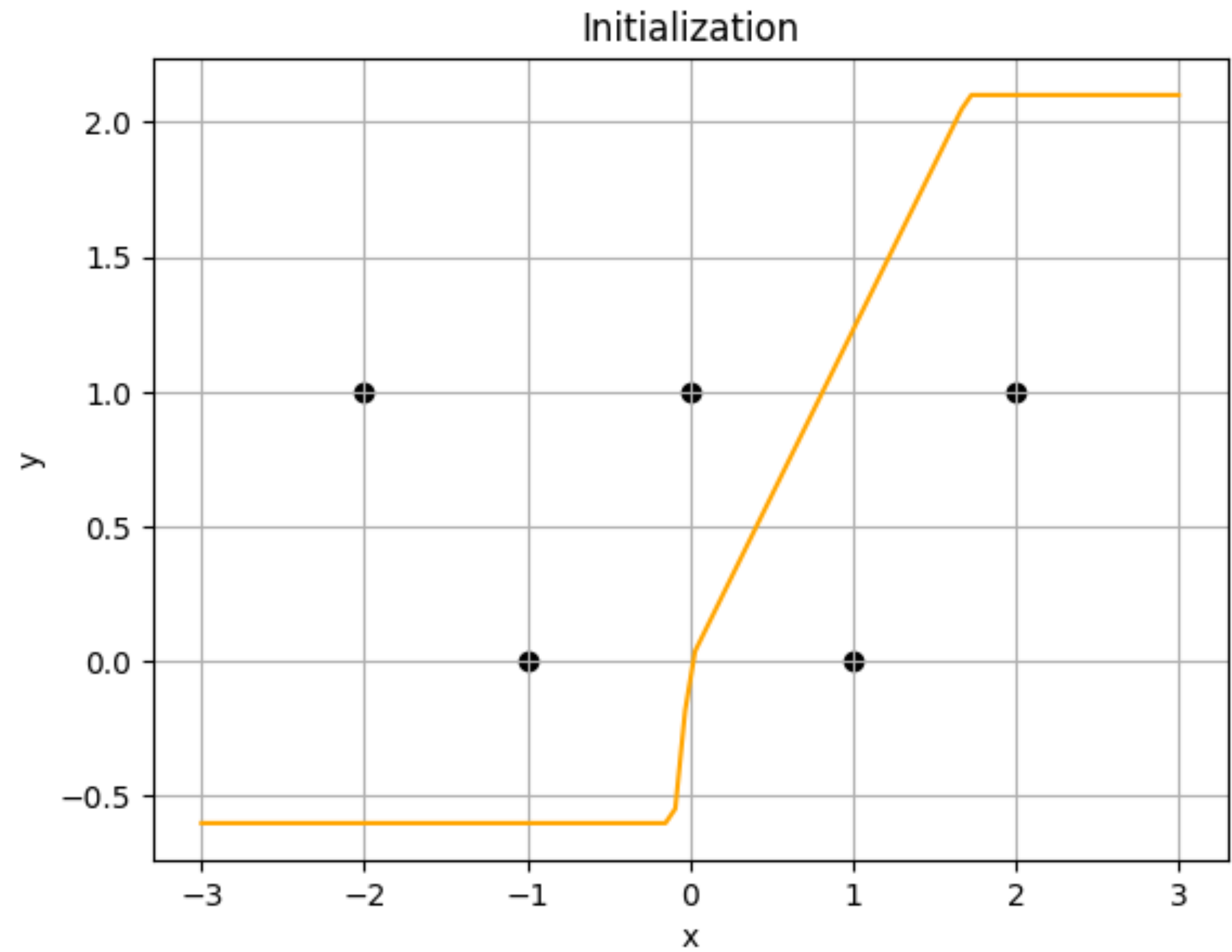
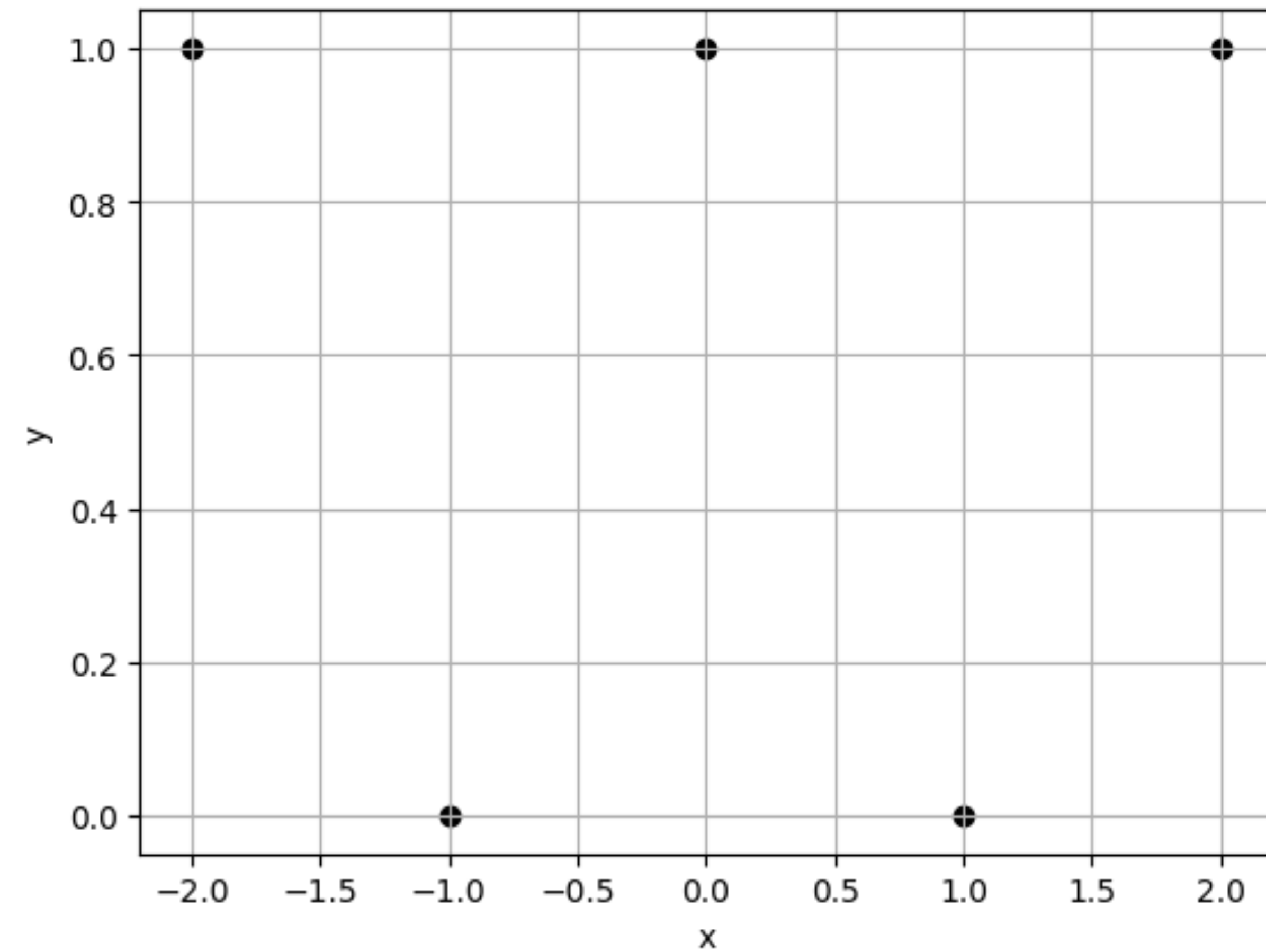
What are we saying here?



Implicitly, a convex relaxation!!

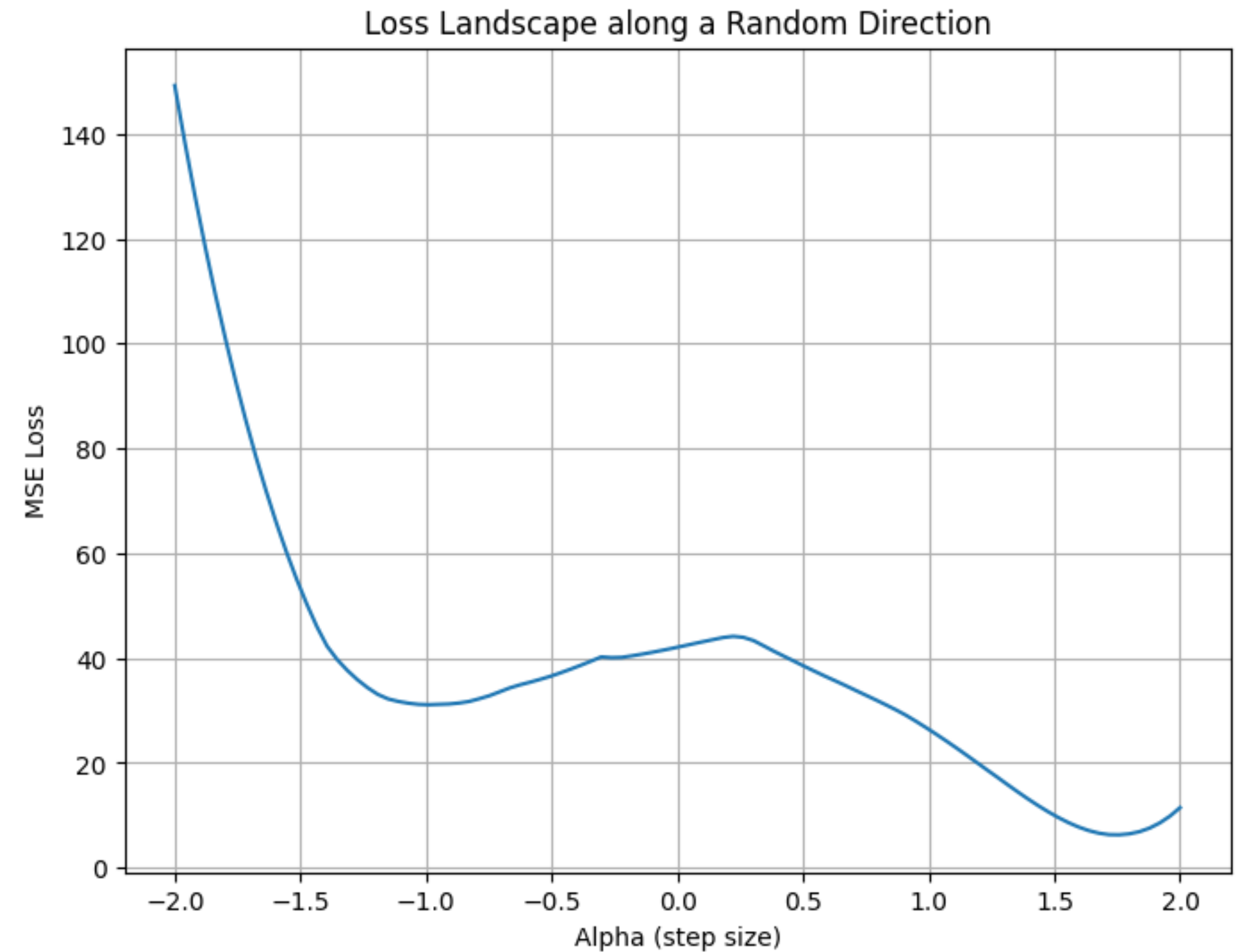
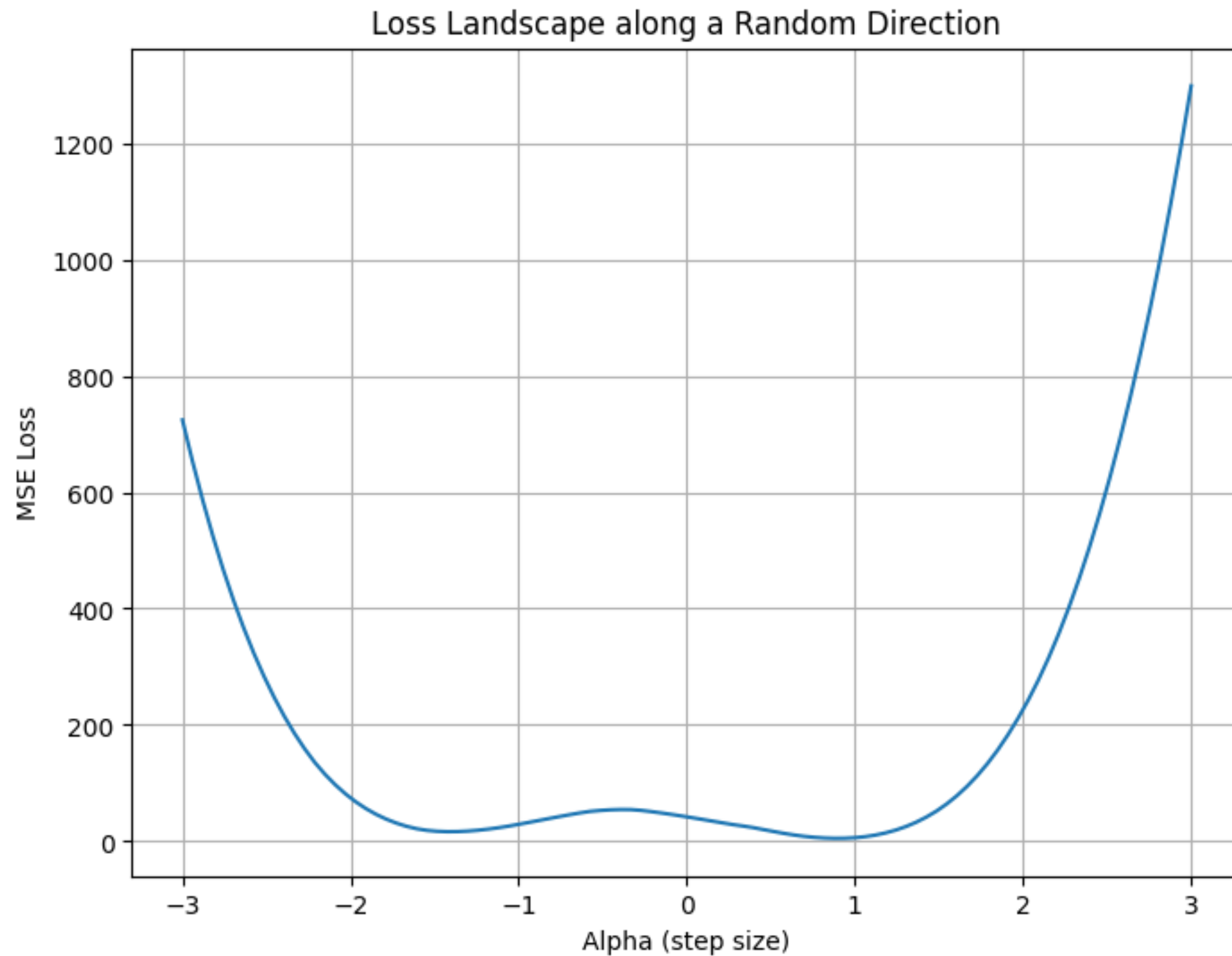


Data on a W



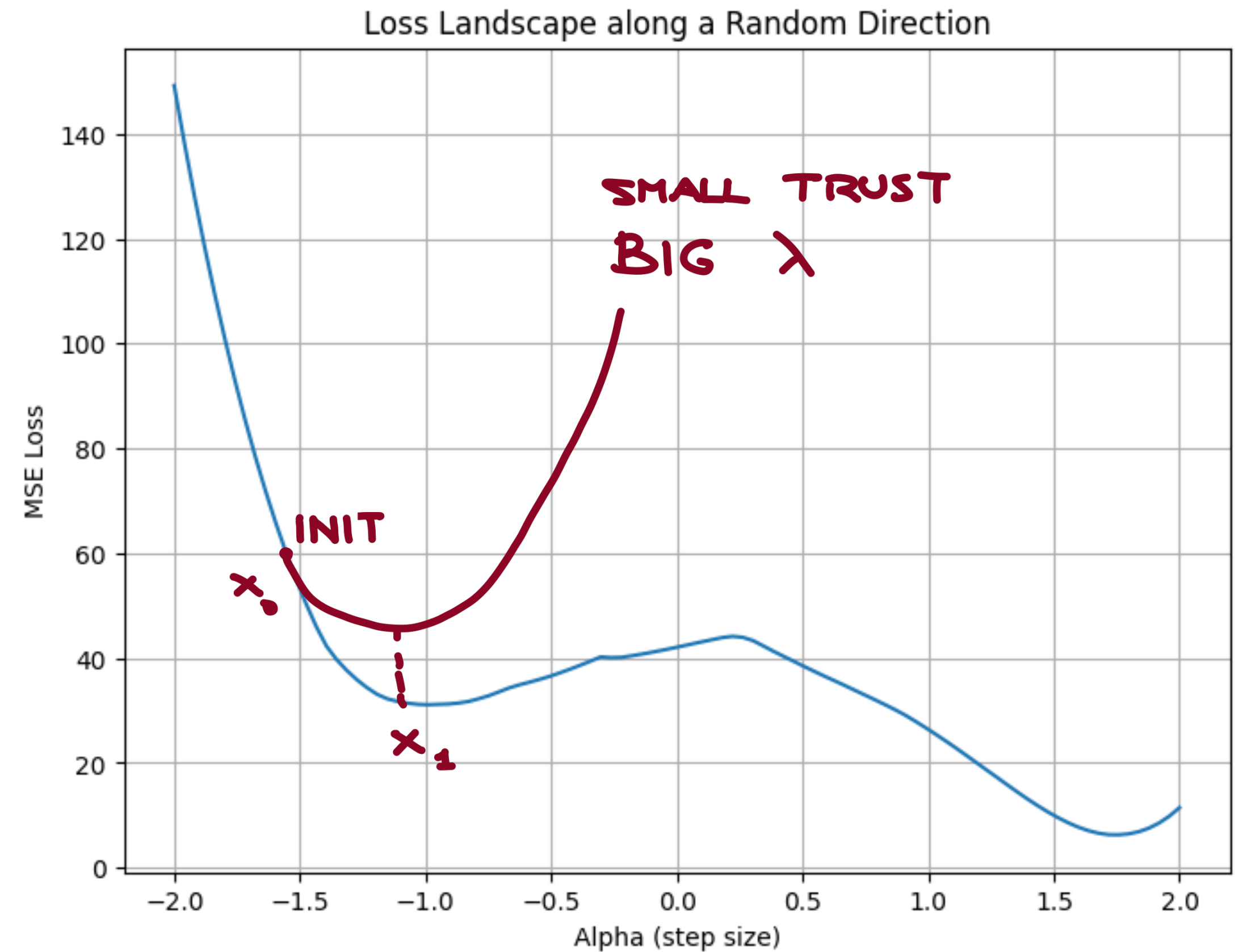
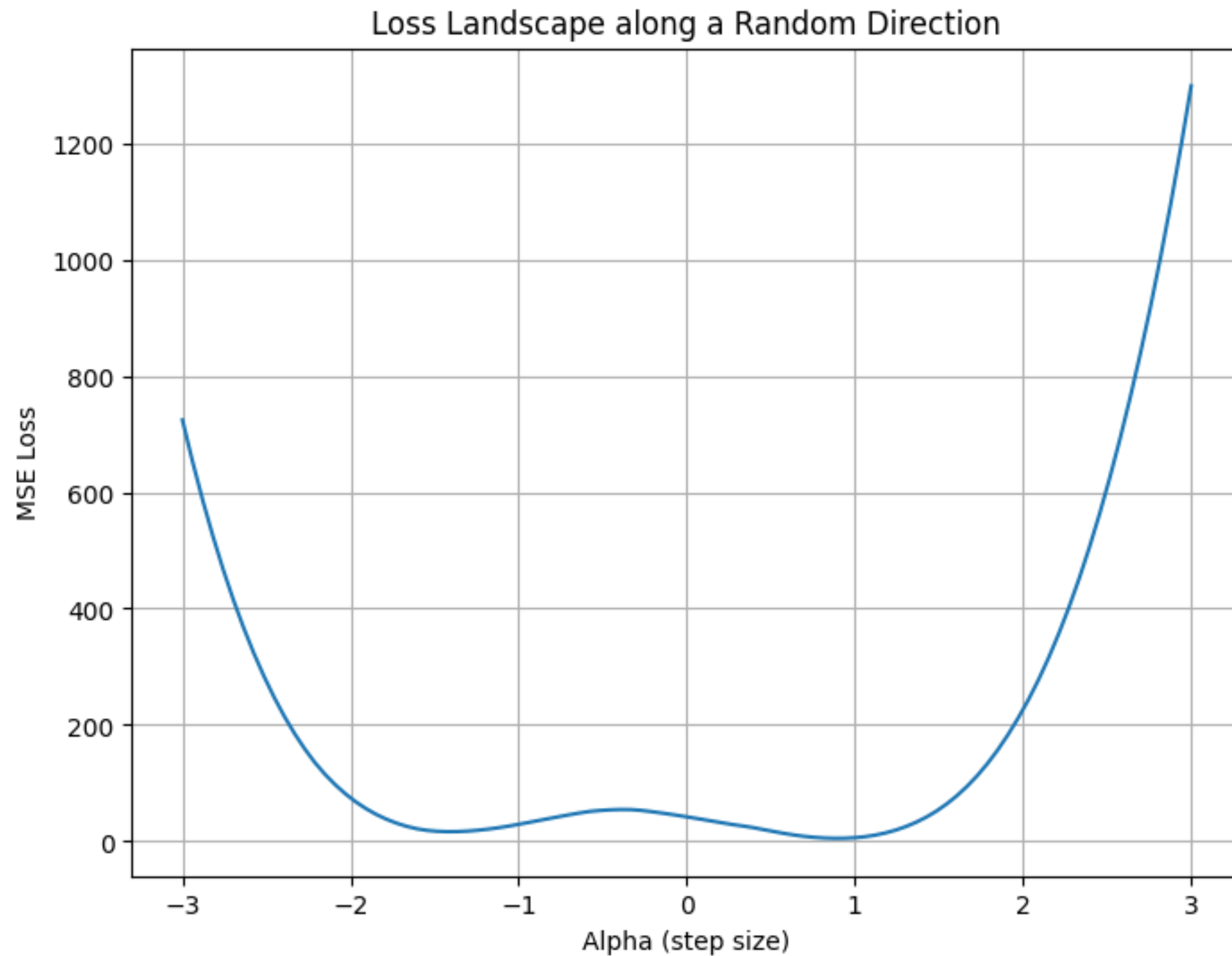
Model
$$f_{a,b,c}(x) = \sum_{i=1}^{20} a_i \cdot \text{ReLU}(b_i x + c_i)$$

We are training NN this way:



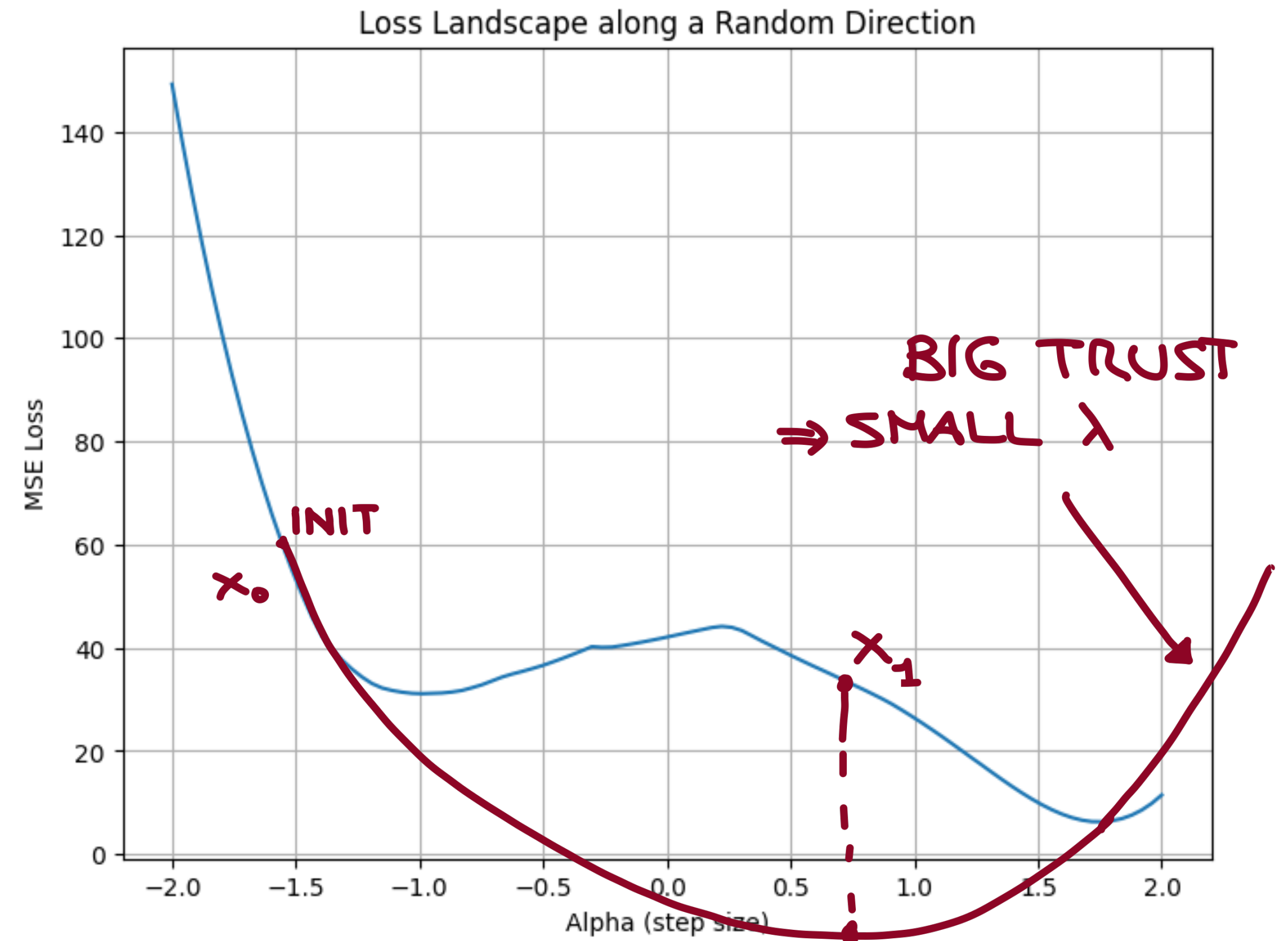
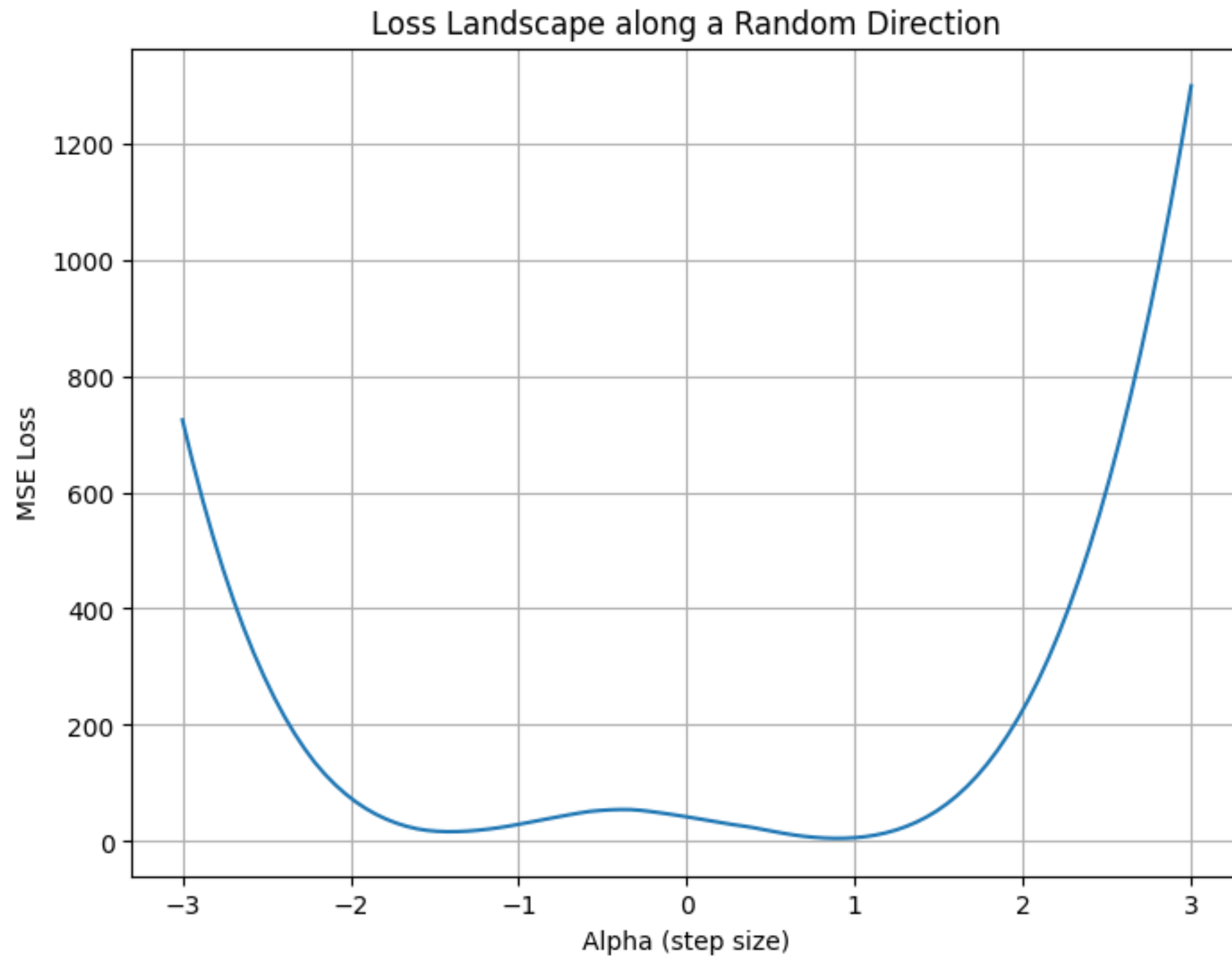
Plot of the loss landscape along a line

It Should Not Work



Convex Relaxation

It Should Not Work



Convex Relaxation

Questions

- How well can Gradient Based Methods do?

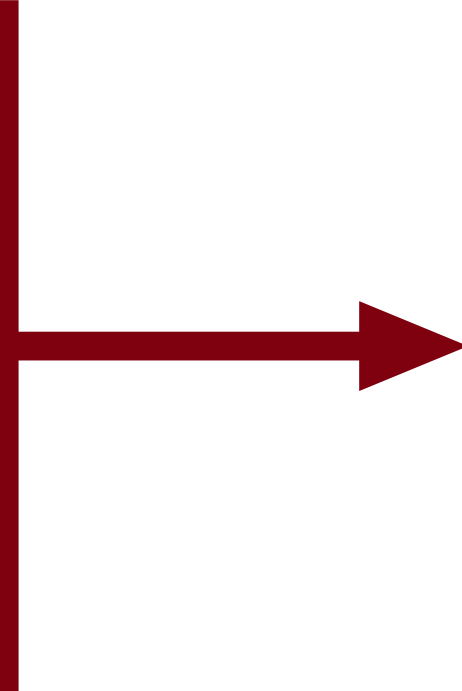
Very well on convex problems...

- How well can these *Iterative Convex Relaxations* do on non-convex problems?

Depends on the geometry of the problem...?

Questions

- How do they **escape local minima**?
- How do they **escape saddles**?
- How do they **pick** the minimum in the high dimensional manifold?



These two are about
traveling flat areas!

All things they have *not* been designed for!

Part 2

The naive mathematical approach

IMPLICIT REGULARIZATION

- Location of Convergence
 - What gradient methods were
 - Some NN examples
- **(Stable) Implicit Regularization**
- Edge of Stability

SPEAKING MATH

WHERE DO ALGORITHMS GO?

This is not an *optimization* question:
It is a dynamical systems question!

Questions are:

- What are the basins of attraction? Do trajectory recur? (**Long-term behavior**)
- Is the invariant set stable (**Stability in the sense of Lyapunov, ...**)
- How does the behavior changes with the (hyper)parameters (**Bifurcation Analysis**)

THE QUESTION

WHERE DO ALGORITHMS GO?

This is not an *optimization* question:
It is a dynamical systems question!

Previous approaches:
a la Numerical Analysis

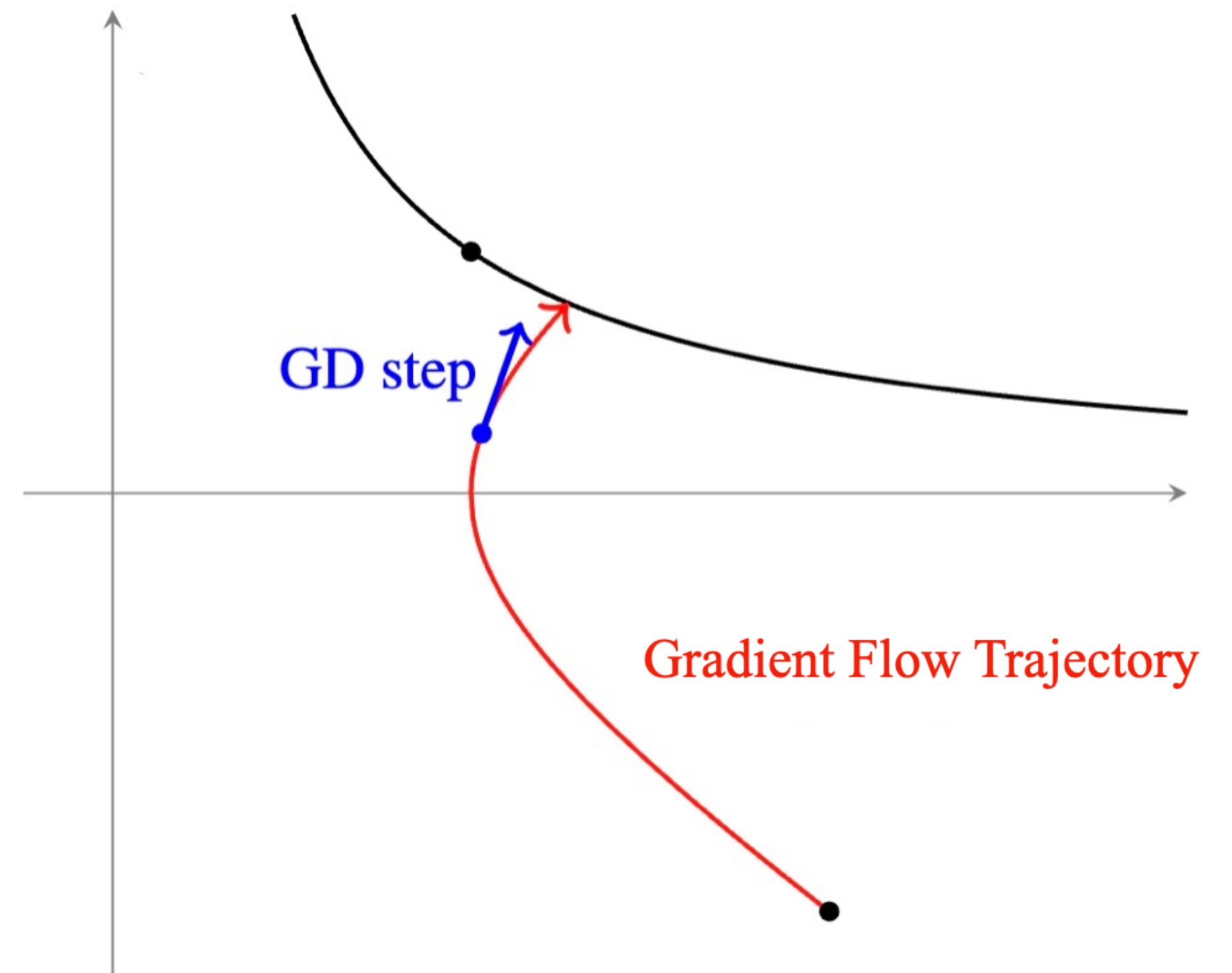
GOAL

What do we want to characterize mathematically

$$\{\theta \text{ s.t. } \nabla L(\theta) = 0\}$$

Not enough!

We want to follow the *trajectory*
and check *where* we go



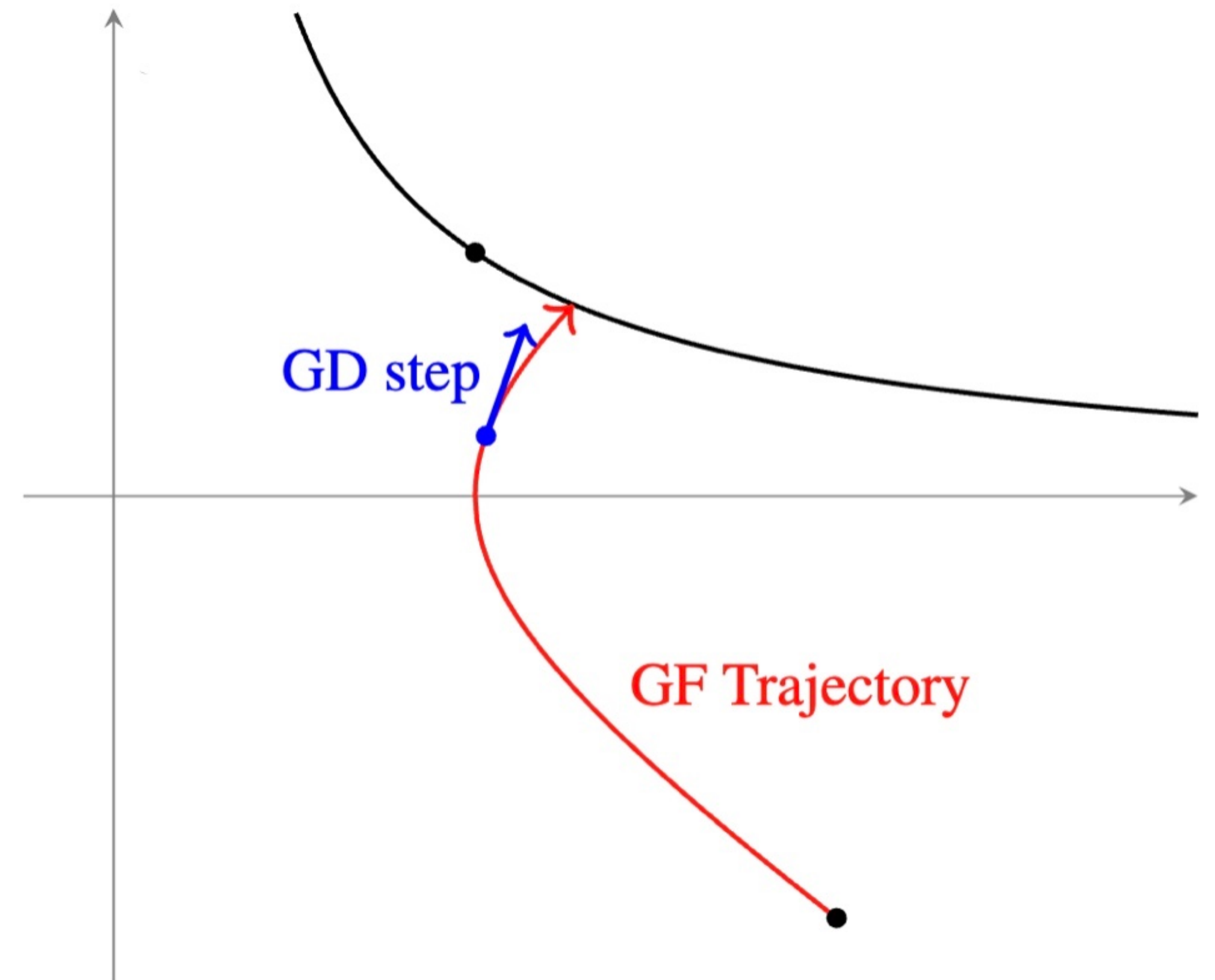
GOAL

What do we want to characterize mathematically

We want to follow the *trajectory* and check *where* we go

How?

- Take two trajectories
- Compute their difference
 - Expanding in Taylor?



RESULTS

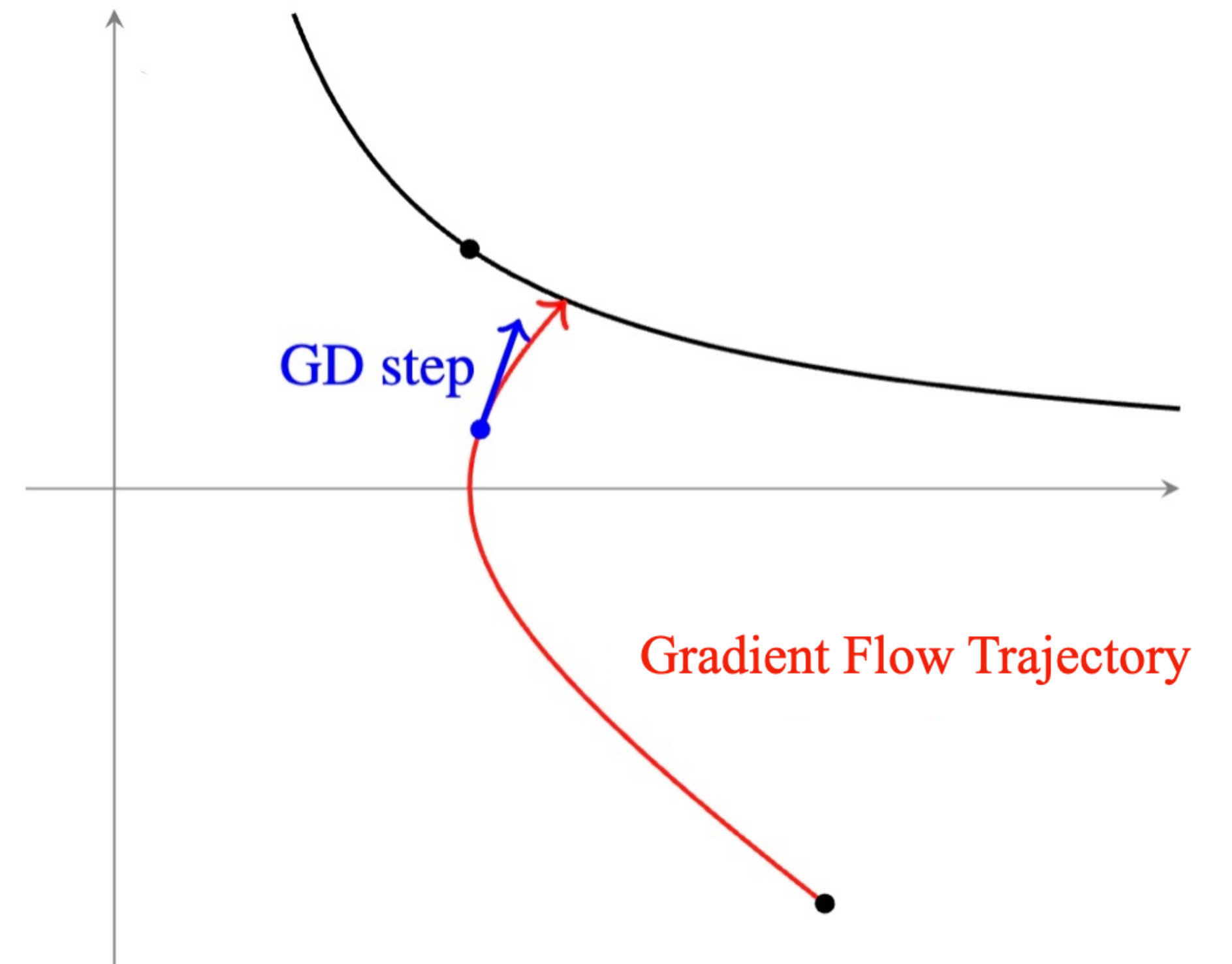
Keywords:

Backward Error Analysis
Implicit Regularization

Discretizing



Implicitly adding a penalty!



SGD vs GD

Let us expand k steps of GD one epoch of SGD trajectories in Taylor and check the difference:

$$\theta_2^{GD} = \theta_0 - \eta \nabla L(\theta_0) - \eta \nabla L(\theta_1)$$

$$\theta_2^{GD} = \theta_0 - 2\eta \nabla L + \eta^2 \nabla^2 L \nabla L$$

Handwritten orange annotations:

- An arrow points from the 2 in θ_2^{GD} to the text $k = \text{EPOCH}$.
- An arrow points from the 2η term to the letter k .
- An arrow points from the η^2 term to the expression $\binom{k}{2} \eta^2$.

Same for SGD with replacement or online.

Comparing SGDs

Let us expand k steps of SGD in Taylor at initialization:

$$\theta^{SGD}(\eta) = \theta - \underbrace{\eta \sum_{B_i \text{ batches}} \nabla L(B_i)}_{\eta k \nabla L} + \underbrace{\eta^2 \sum_{i < j} \nabla^2 L(B_i) \nabla L(B_j)}_{\text{additional}} + \dots$$

- Noisy centered for SGD with
- Deterministic for SGD without

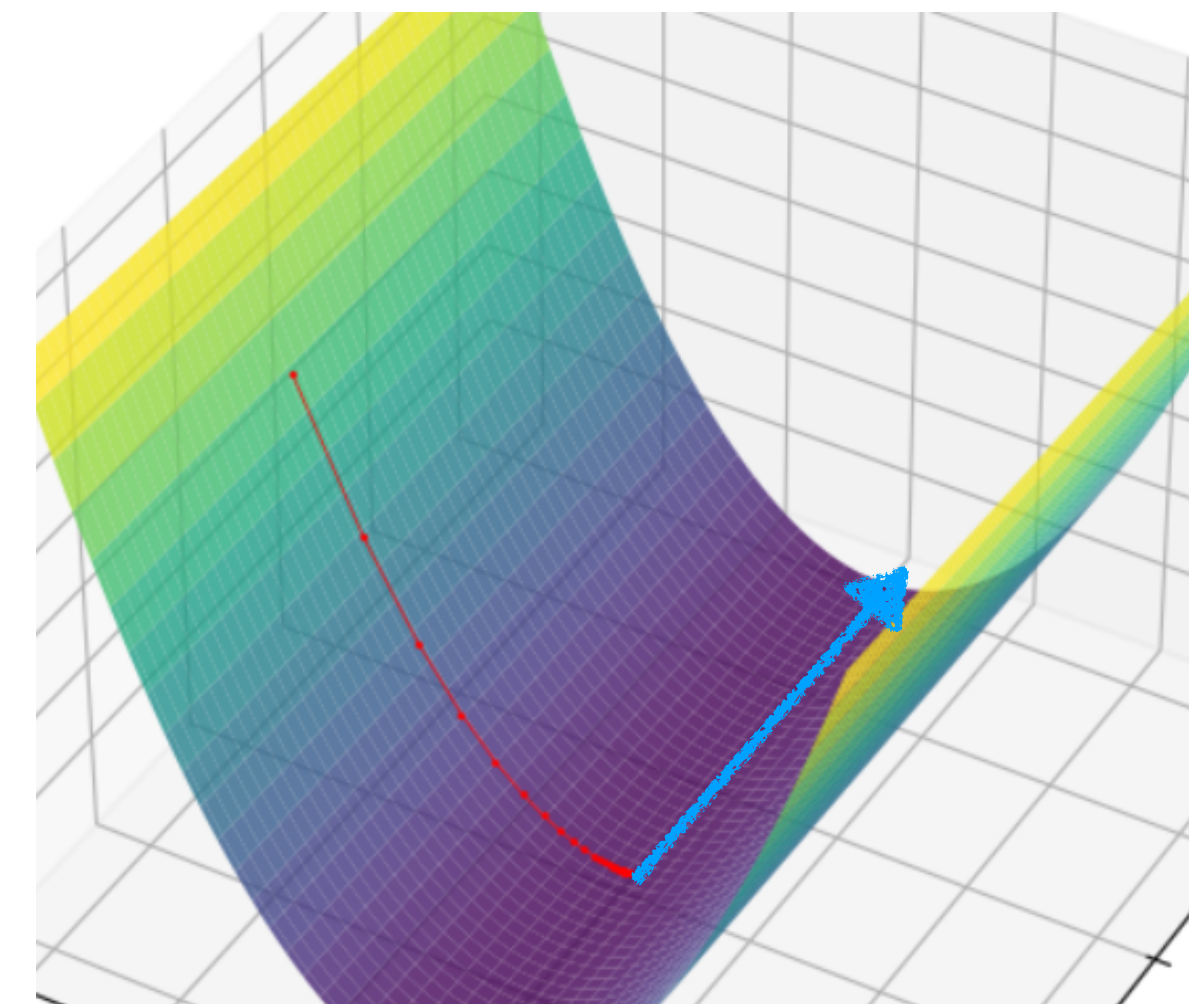
- Noisy centered for SGD with
- Biased for SGD without: batches are dependent

SGD Without Replacement

What the behavior is

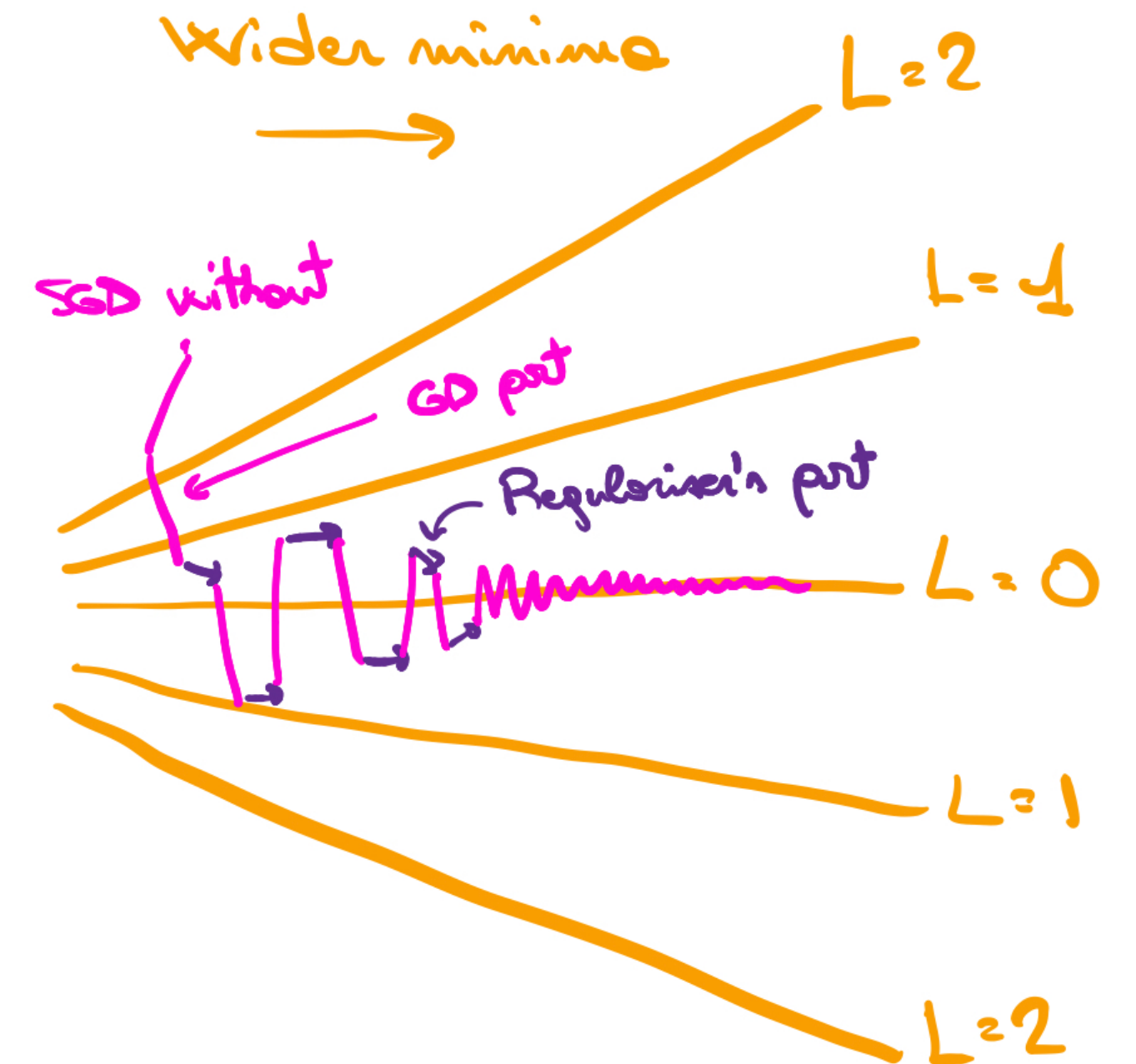
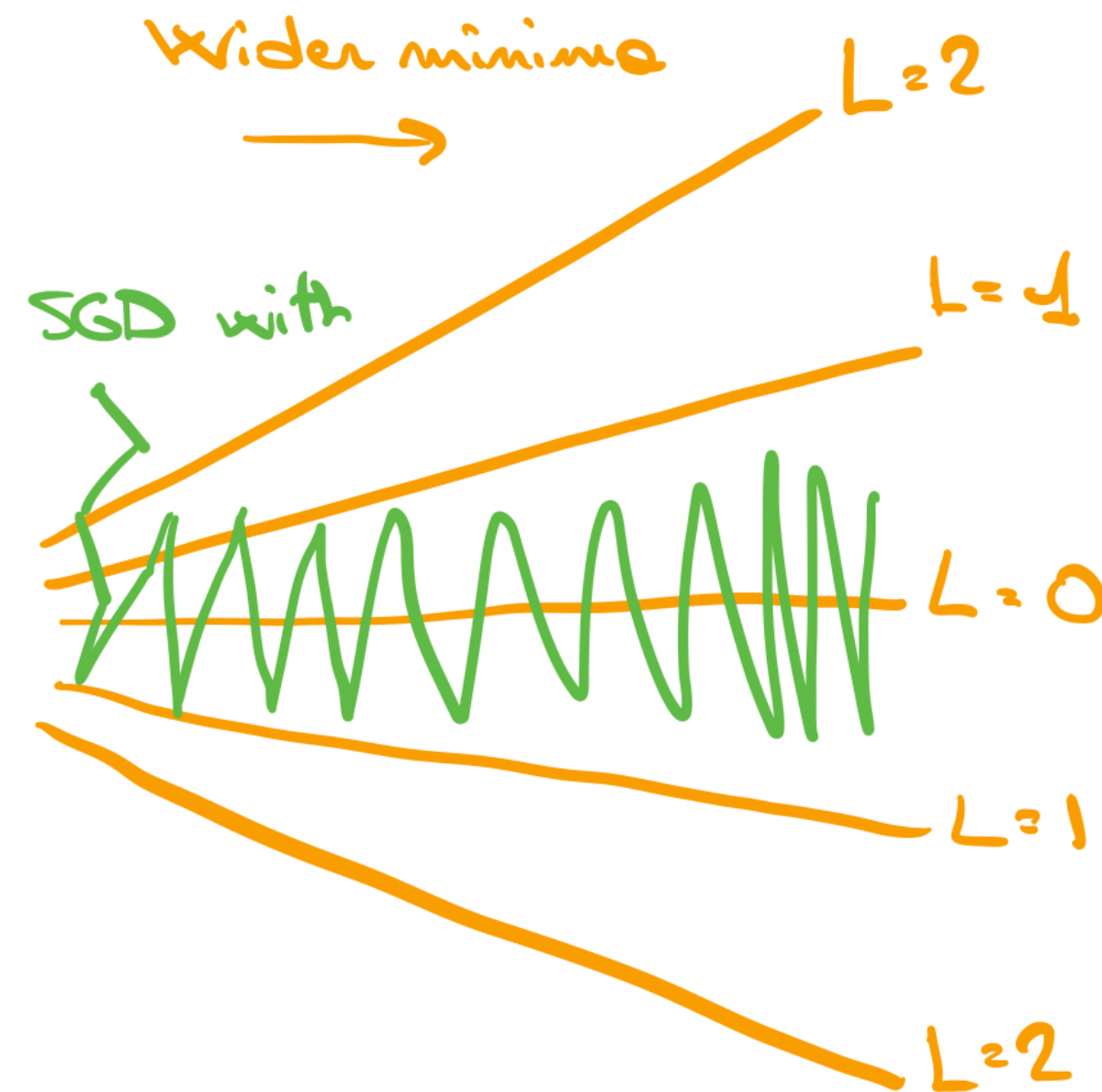
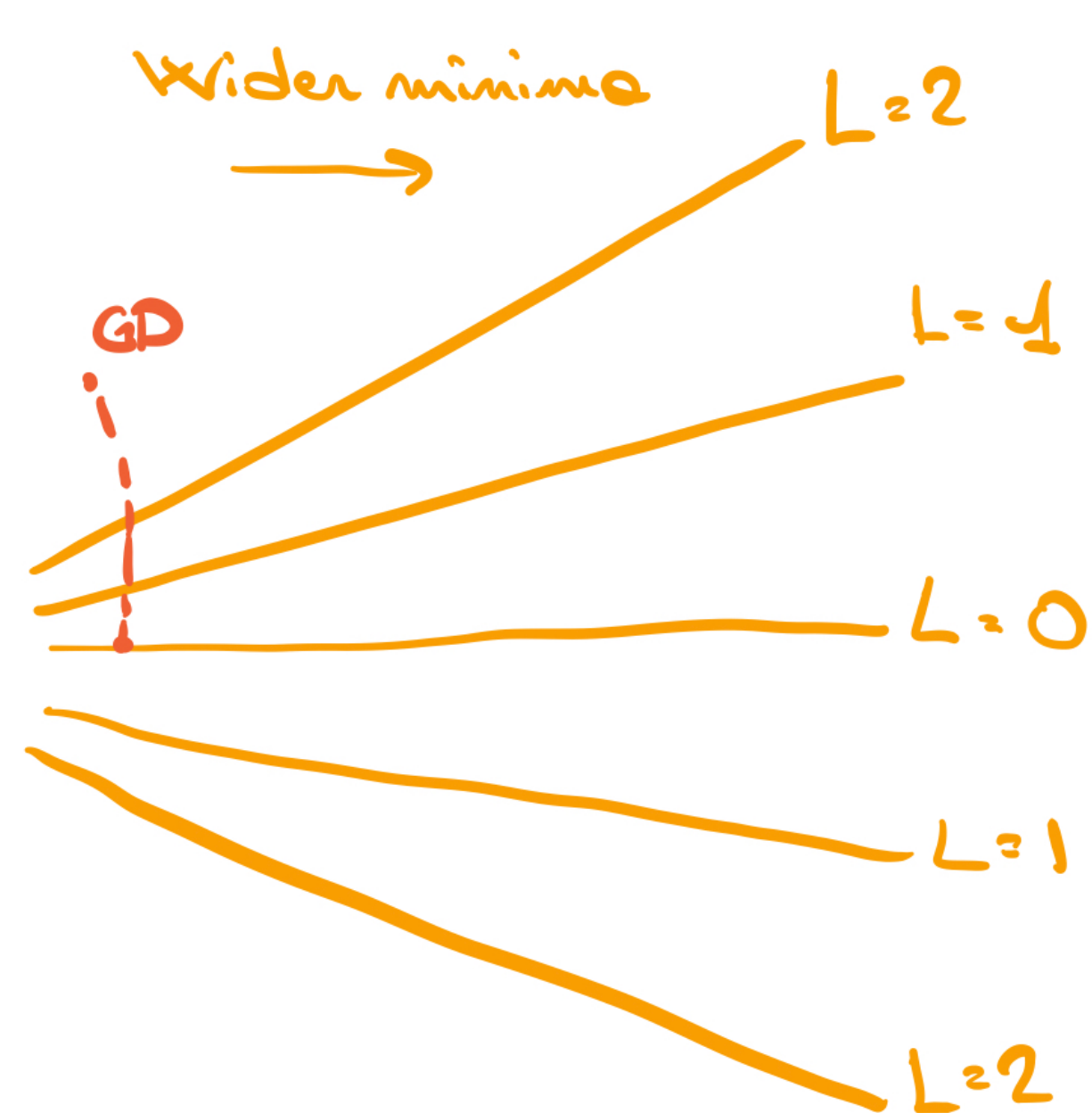
$$SGD = GD + \text{step on Regularizer}$$

- Flat directions: ***Regularizer***
- Sharp directions: ***Follows SGD with***
- More sharp directions: ***Shoots away***



With vs Without

The Bias



RESULTS

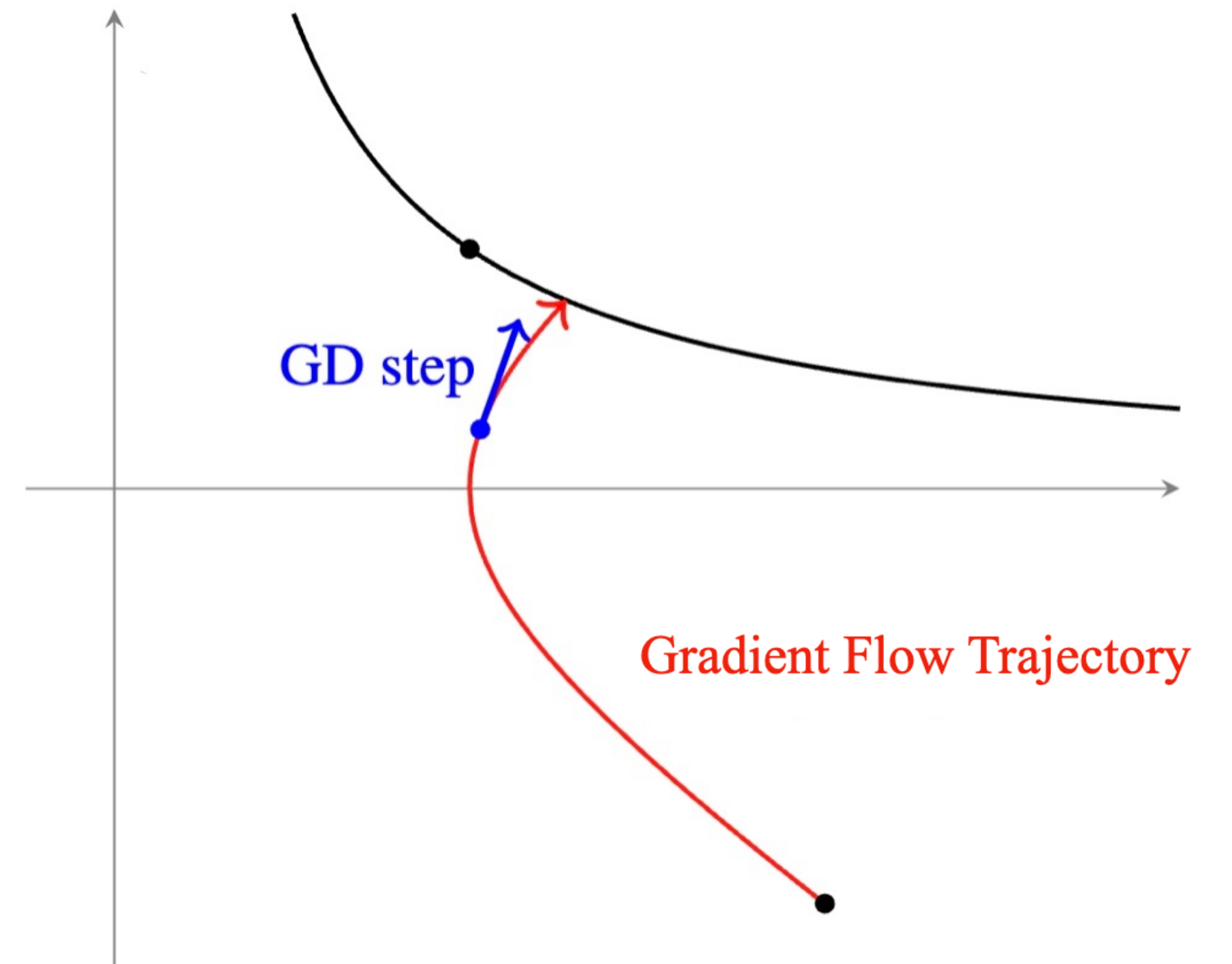
Keywords:

Backward Error Analysis
Implicit Regularization

Discretizing



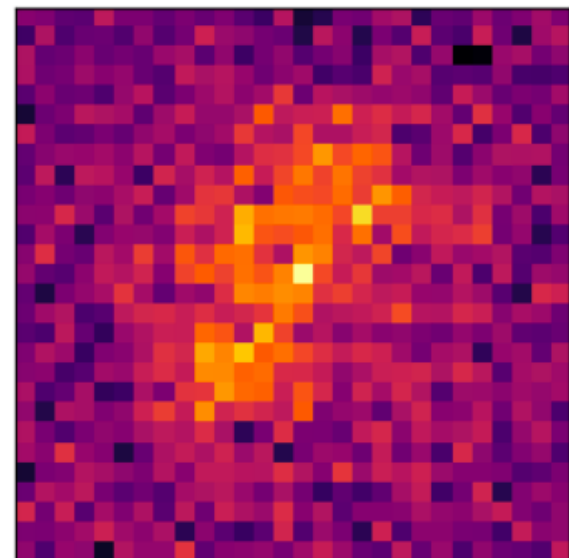
Implicitly adding a penalty!



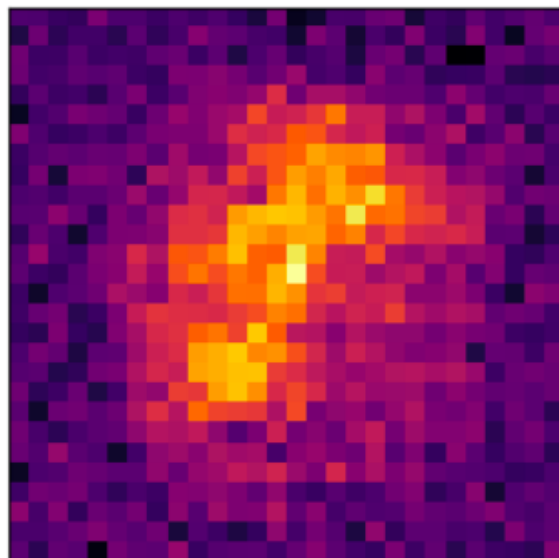
THIS APPROACH

WORKS IN THE CASES ABOVE

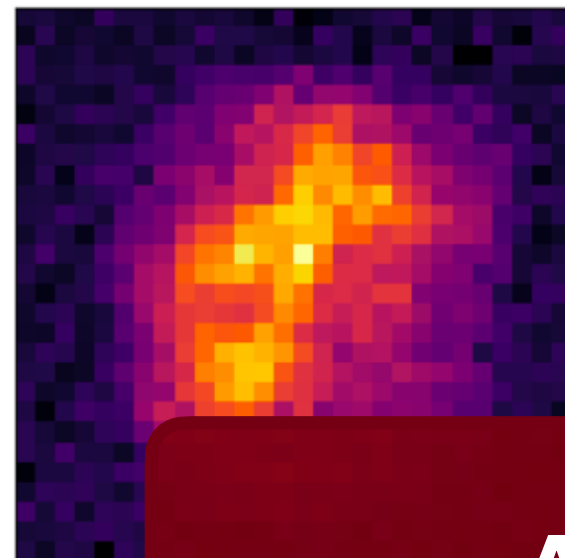
SGD-2048



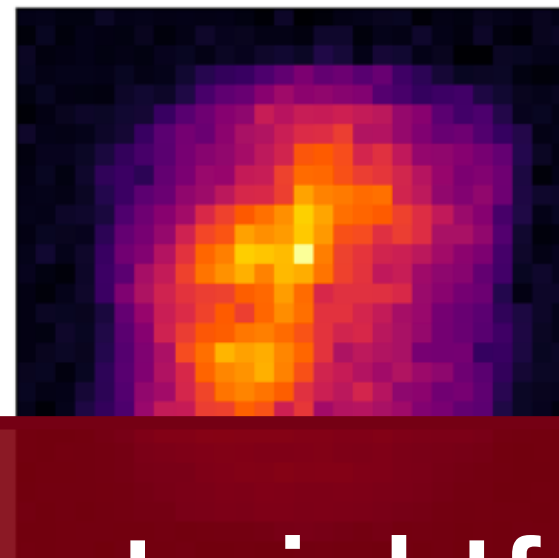
SGD-512



SGD-128



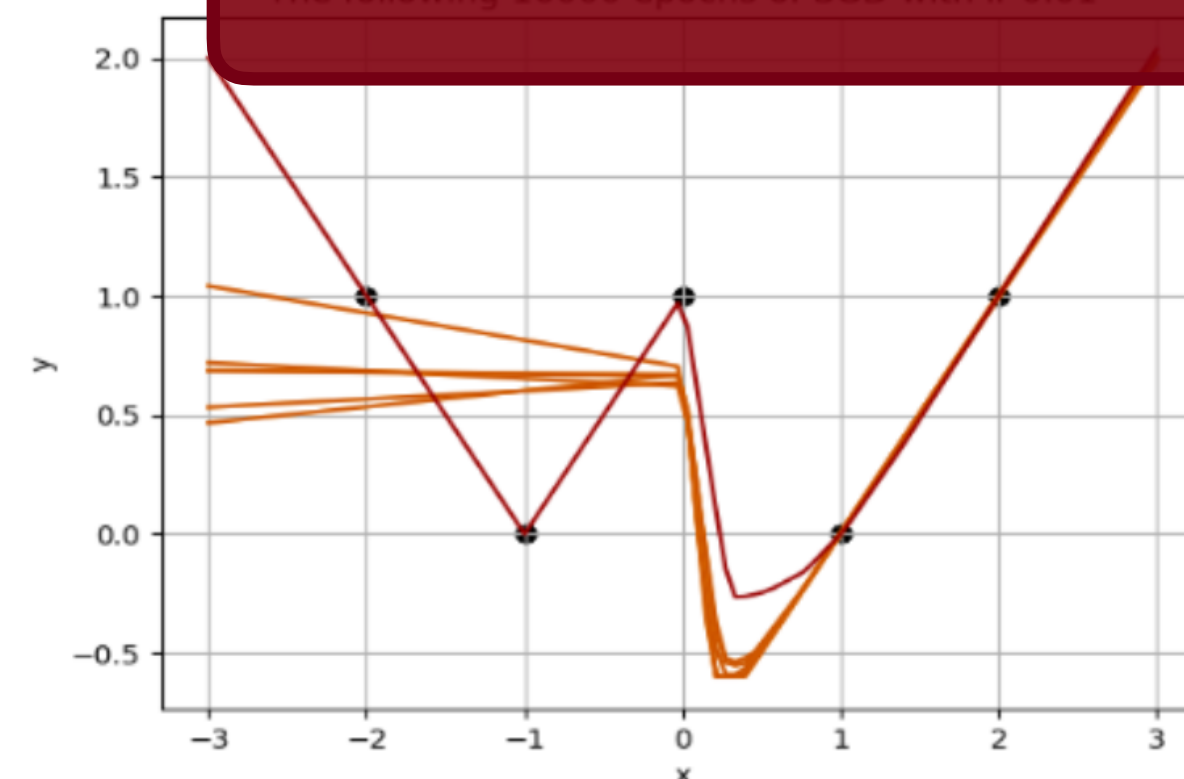
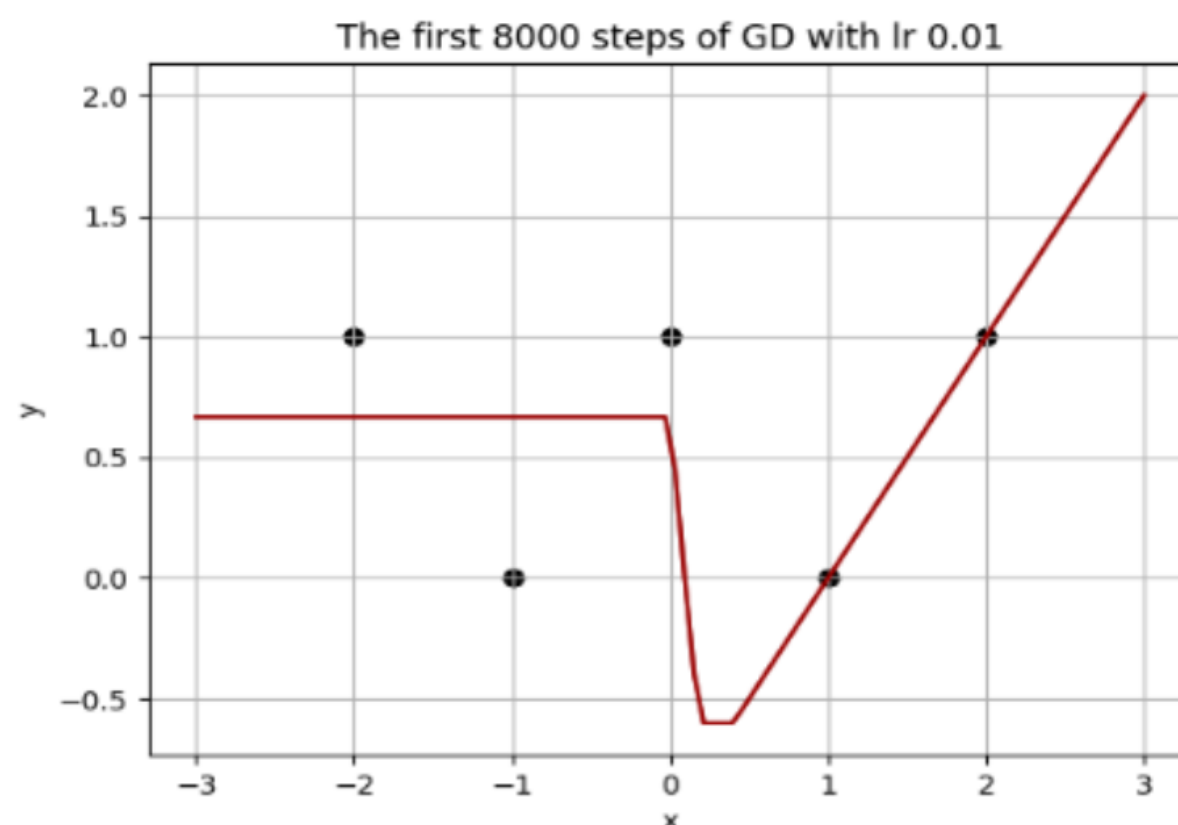
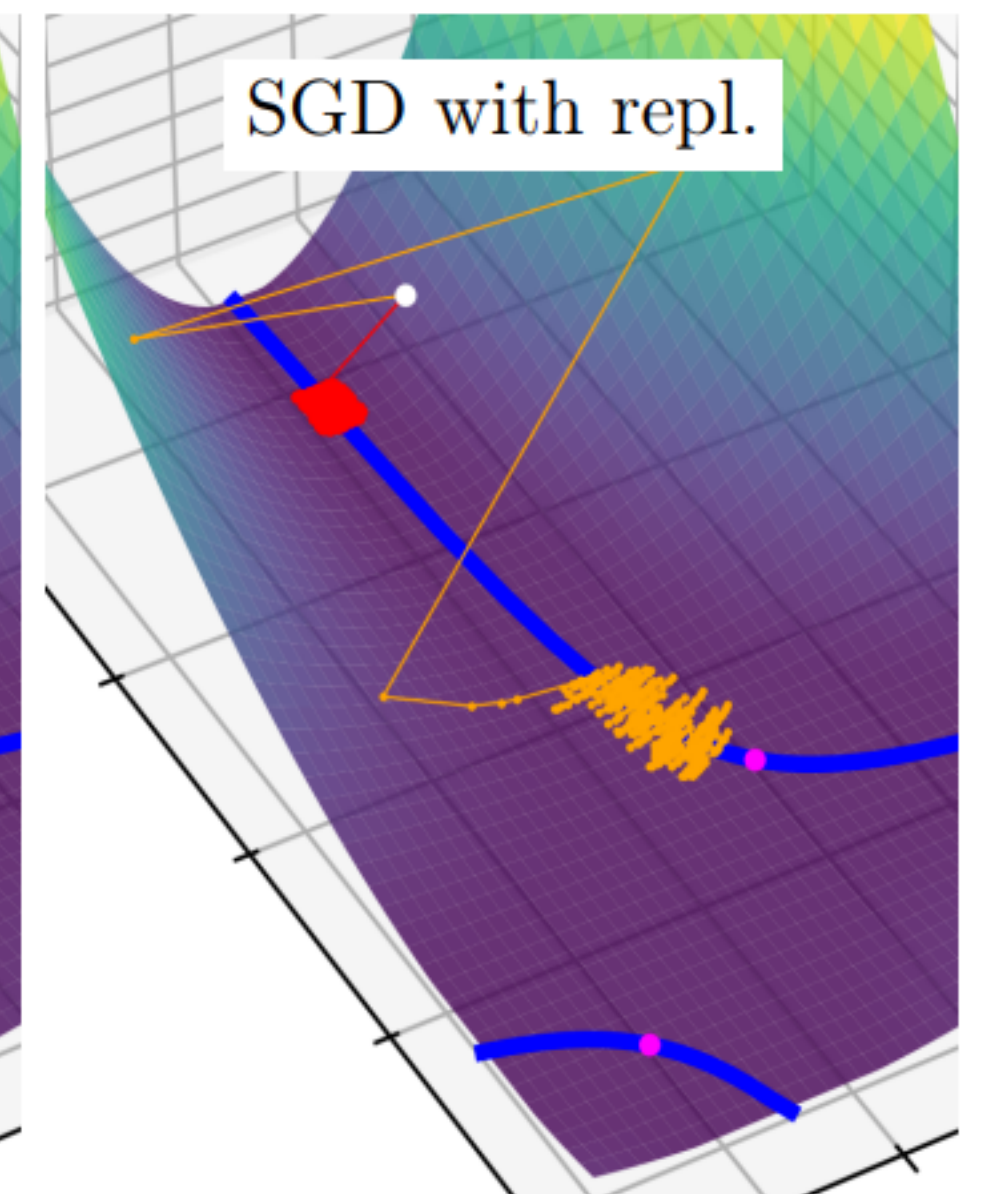
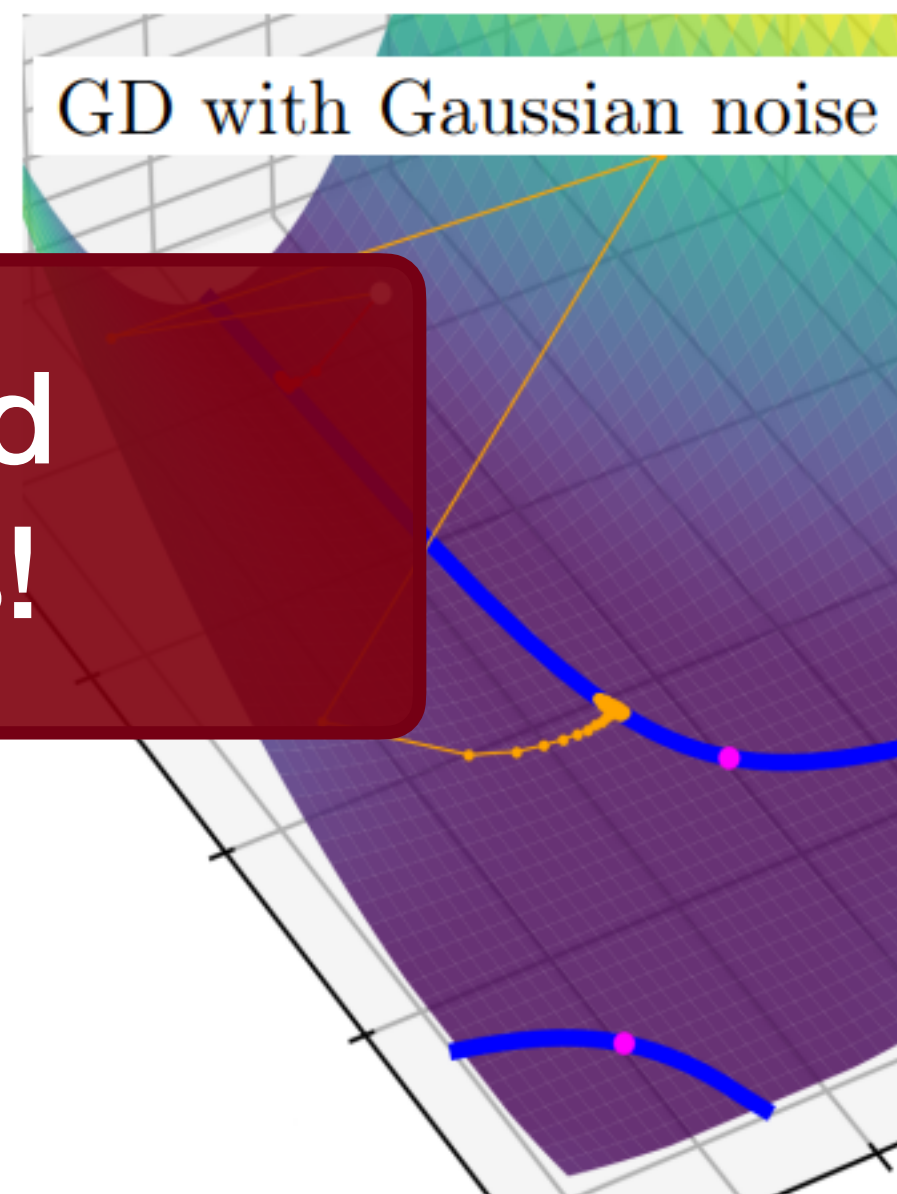
SGD-32



A straightforward approach works!

SGD smaller batch size

The following 16000 epochs of SGD with lr 0.01



Part 3

When does it break?

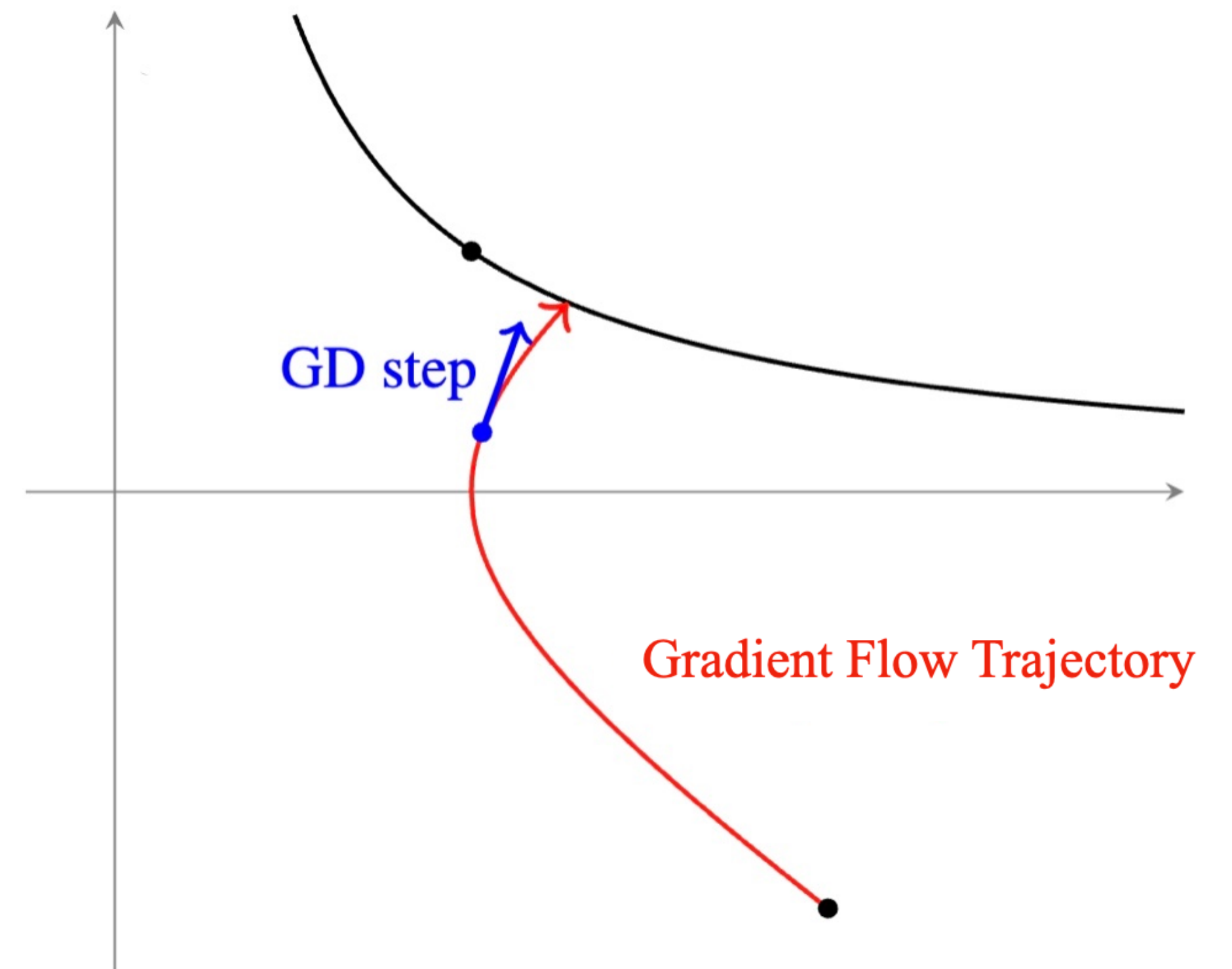
To what extent does it explain?

To what extent can we use it to do the next step: optimality?

LIMITATIONS

Main problems

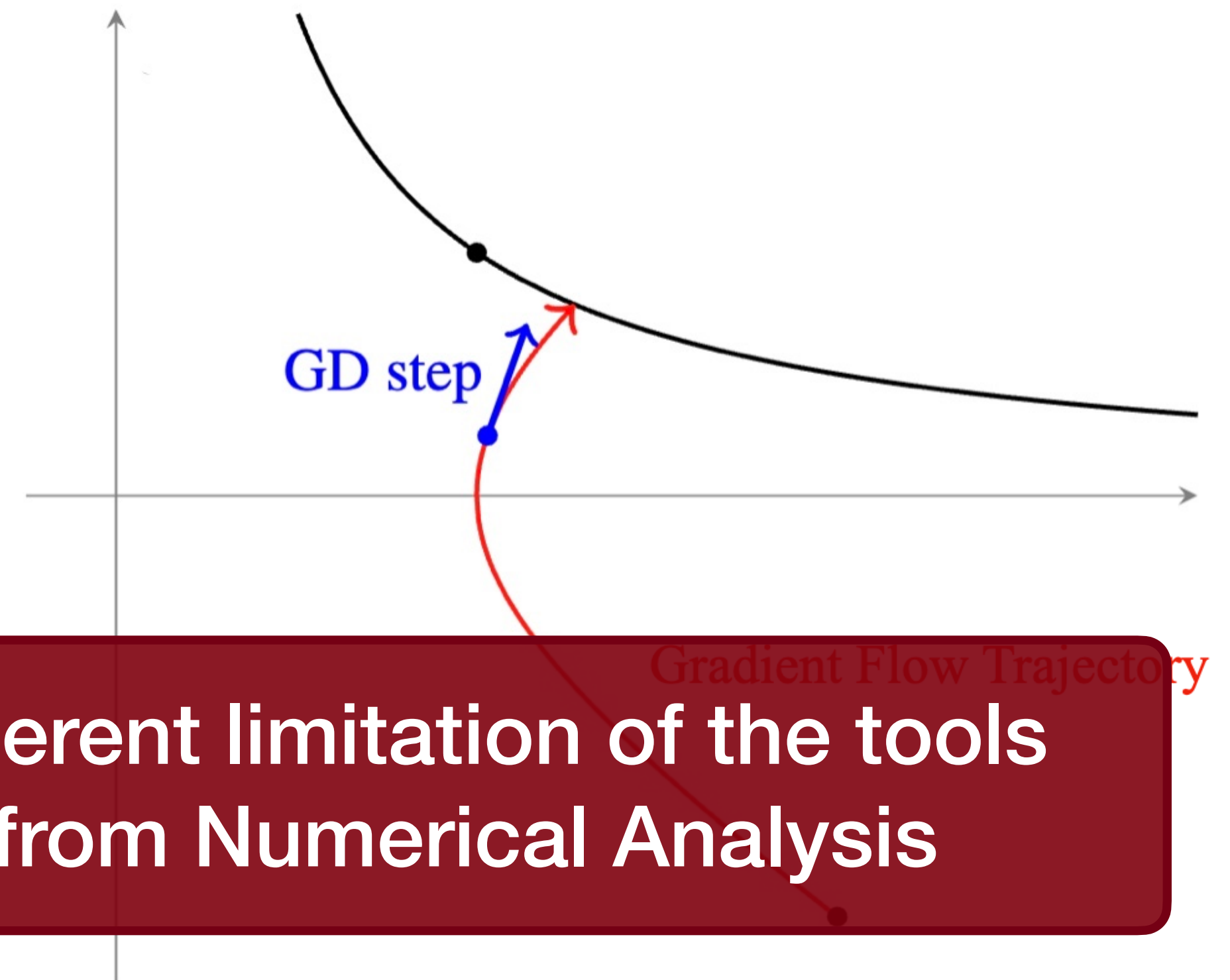
- It works for very small step sizes;
- It may characterize just a few steps;
- It is too complicated to plug in the actual losses and say something sensible;



LIMITATIONS

Main problems

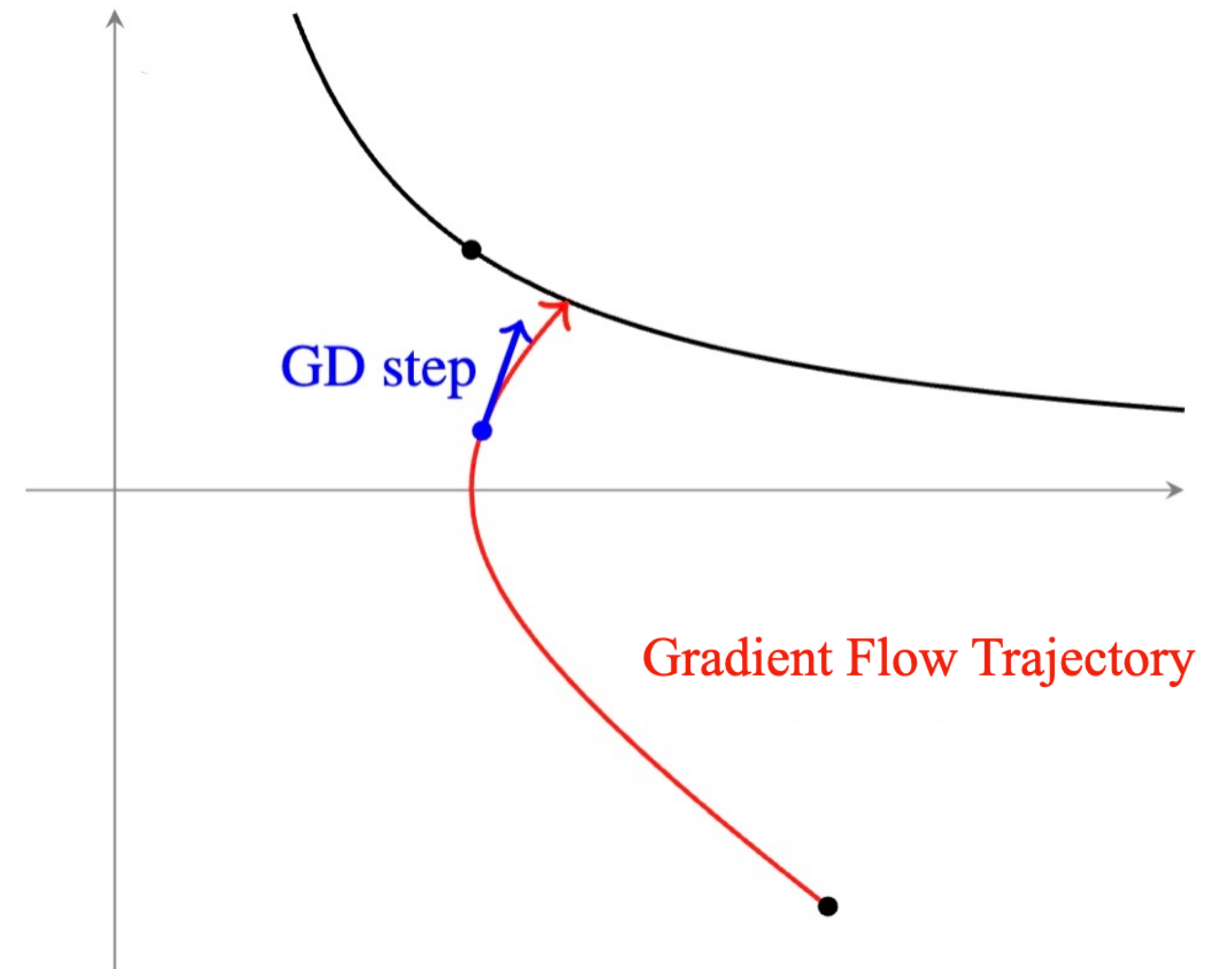
- It works for very small step sizes;
- It may characterize just a few steps;
- It is too complicated to plug in the actual losses and say something sensible;



NEW SET OF QUESTIONS

- It works for very small step sizes.

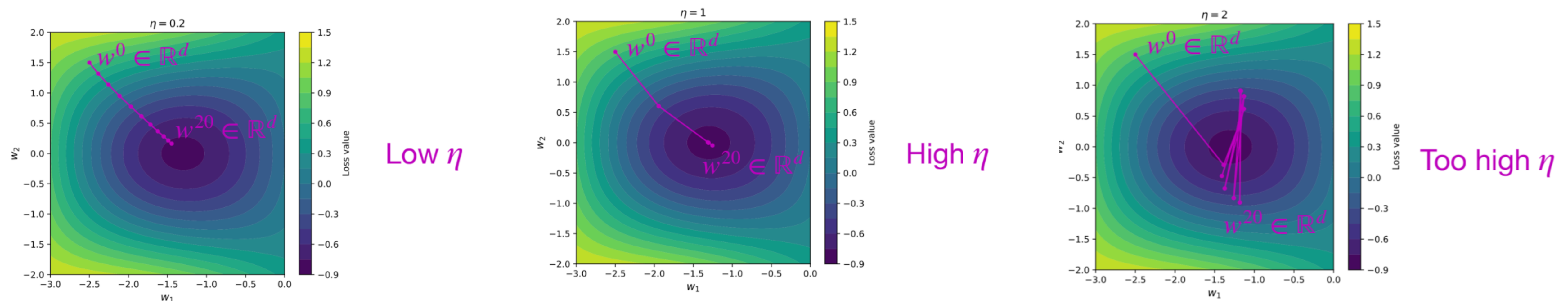
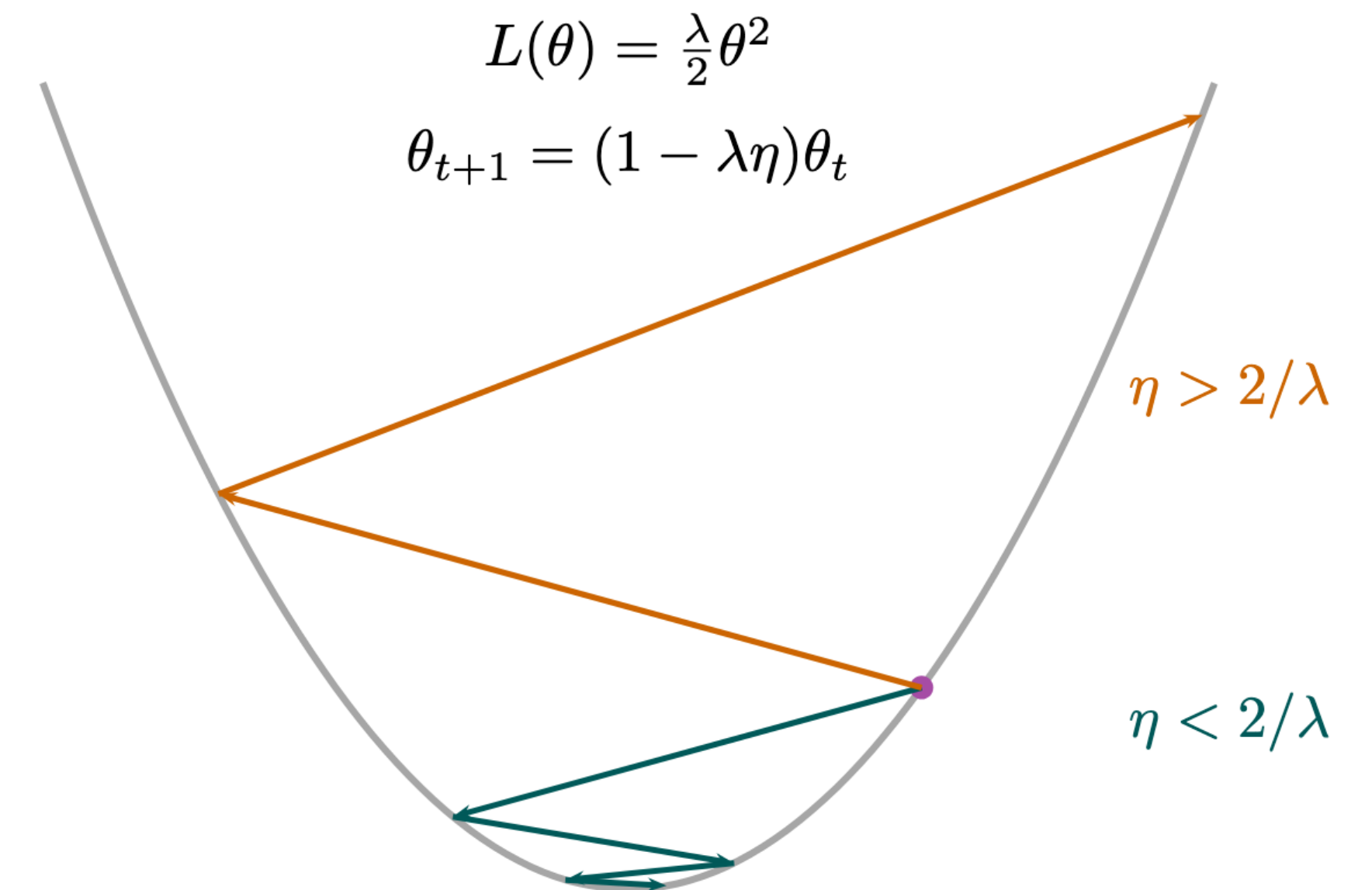
- How large are the steps in practice?
- What does it even mean for the step to be large?



INFINITESIMAL *VS* LOCAL

- How much the *direction* of the gradient *changes*.
- How big it needs to be before it *overshoots*.

2^{nd} Derivative (in the neighborhood)



EDGE OF STABILITY

- Location of Convergence
 - What gradient methods were
 - Some NN examples
- (Stable) Implicit Regularization
- **Edge of Stability**

SURPRISING ANSWER FOR GD

FOR FULL-BATCH GD THESE QUESTIONS HAVE BEEN ANSWERED:

Location of convergence becomes:

How does location of convergence depend on the hyper parameters?

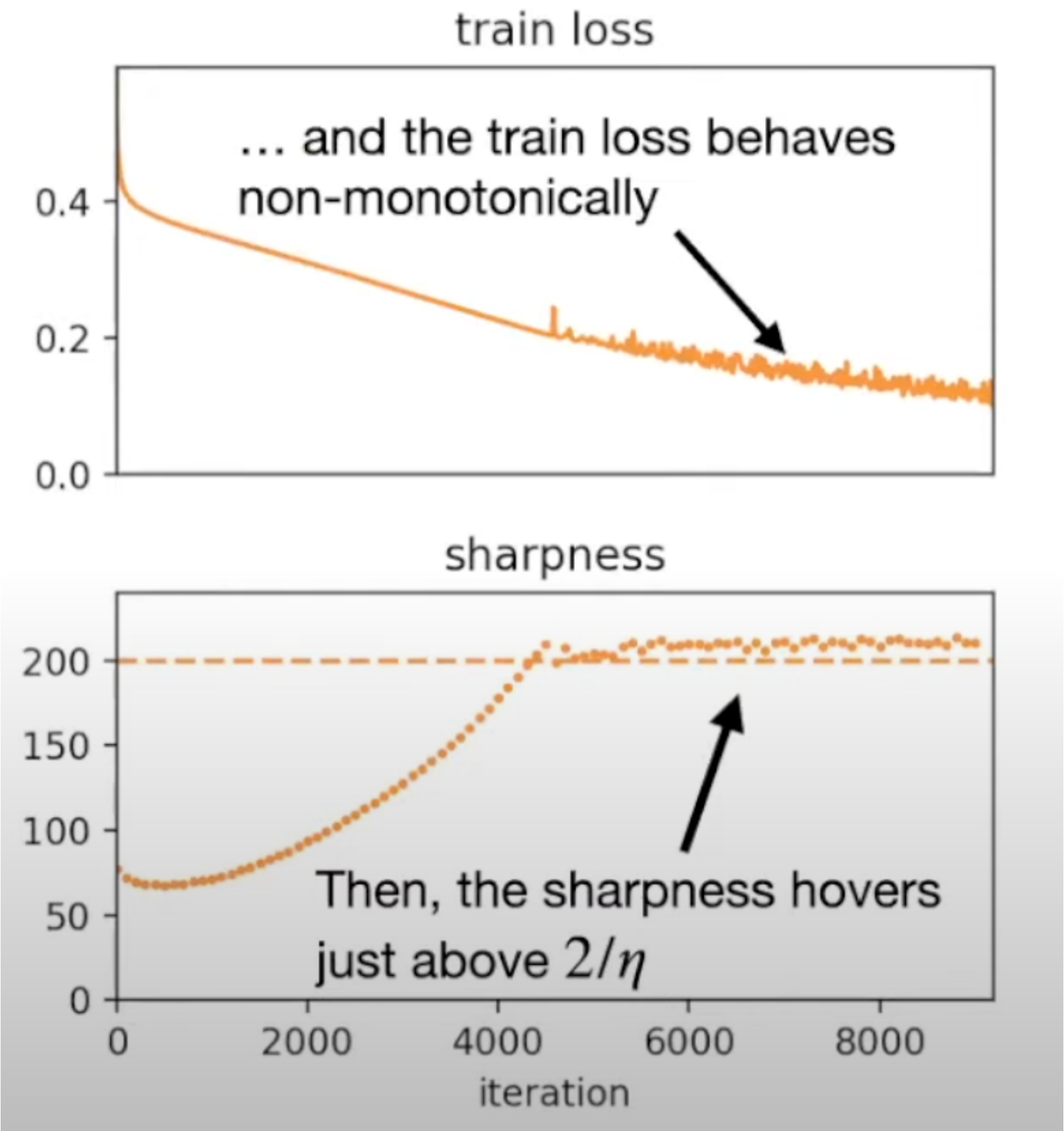
We can answer given this:

Is the step size generally big or small?

THE EDGE OF STABILITY

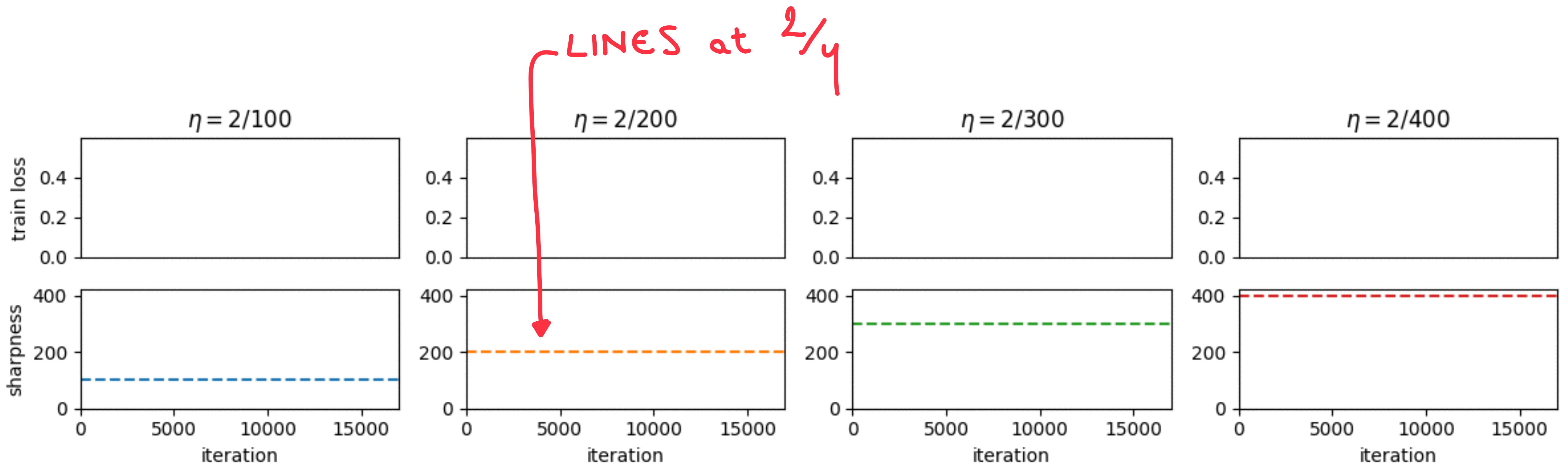
WHAT:

- A *Classification* task.
- Across many different architectures the behavior obeys *a simple rule*.
- The rule is not what you expect.



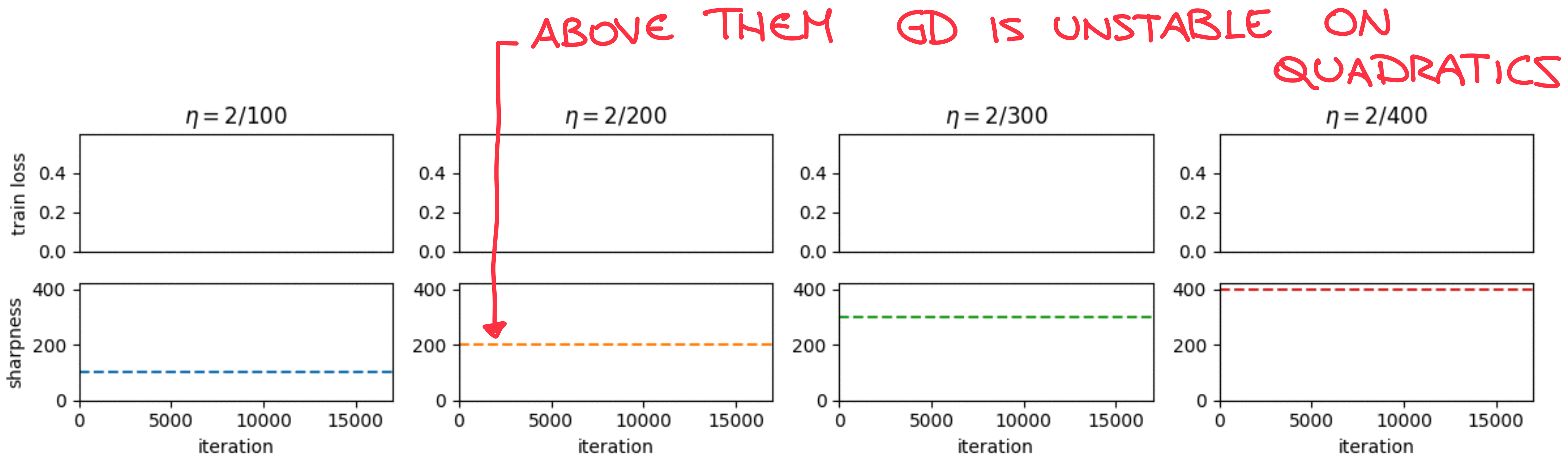
THE EDGE OF STABILITY

Take the *largest eigenvalue of the Hessian*, call it sharpness, plot it.



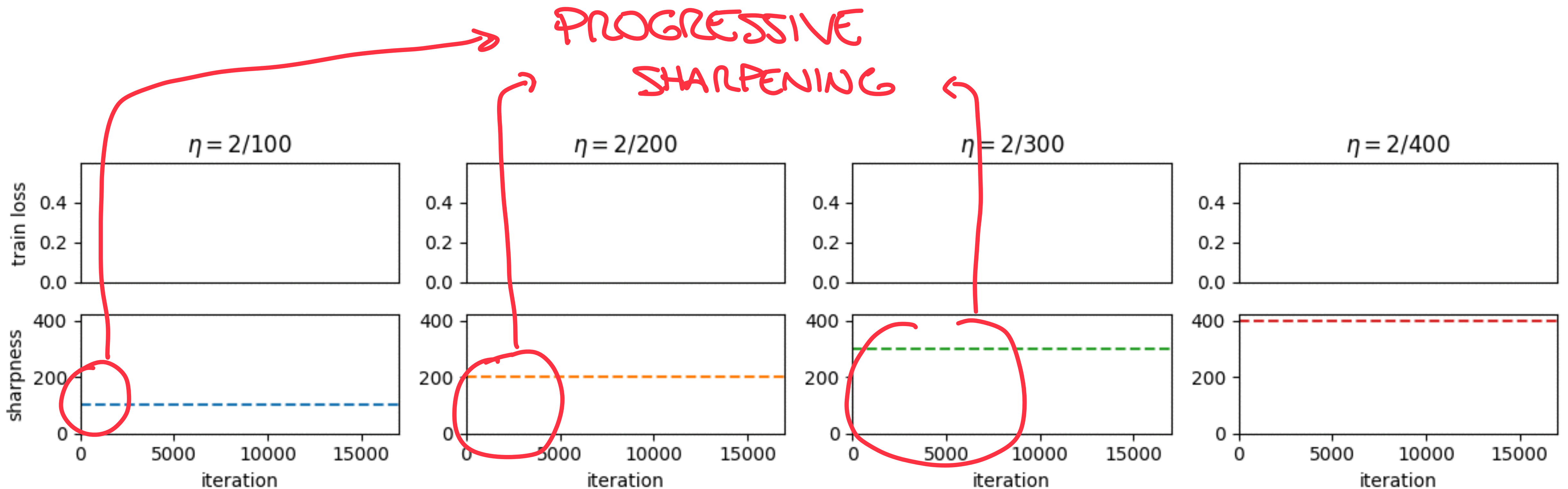
THE EDGE OF STABILITY

Take the *largest eigenvalue of the Hessian*, call it sharpness, plot it.



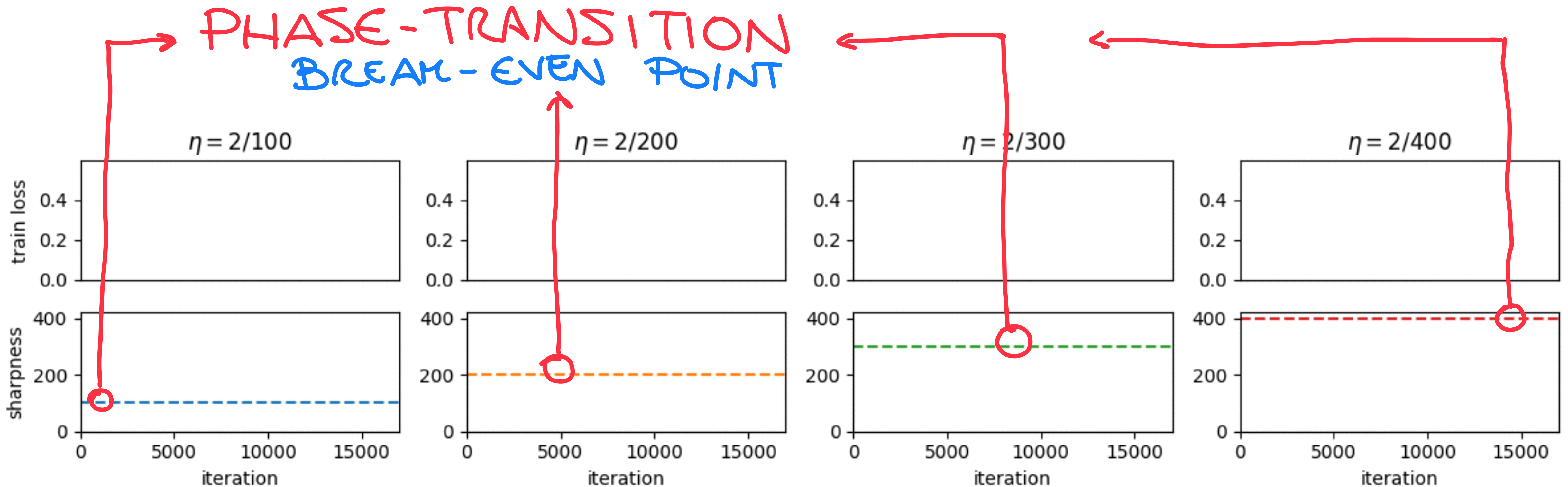
THE EDGE OF STABILITY

Take the *largest eigenvalue of the Hessian*, call it sharpness, plot it.



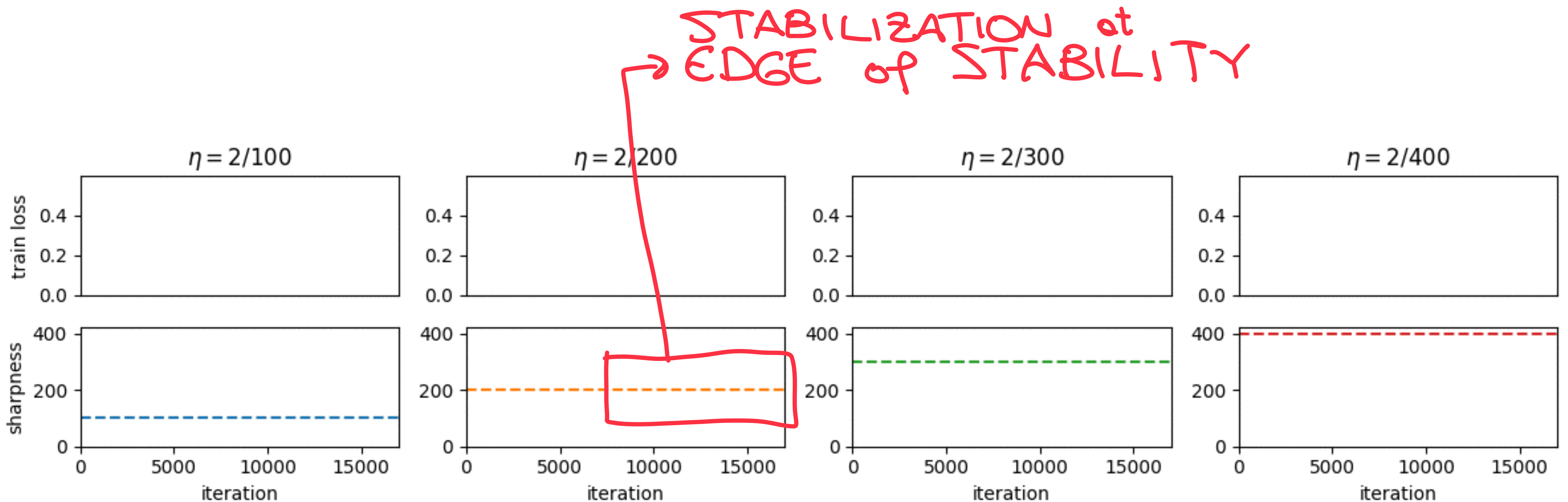
THE EDGE OF STABILITY

Take the *largest eigenvalue of the Hessian*, call it sharpness, plot it.



THE EDGE OF STABILITY

Take the *largest eigenvalue of the Hessian*, call it sharpness, plot it.



NEW SET OF QUESTIONS

FOR FULL-BATCH GD ALL THESE QUESTIONS HAVE BEEN ANSWERED:

Location of convergence becomes:

How does location of convergence depend on the hyper parameters?



We can answer given this:

Is the step size generally big or small?



Answered!

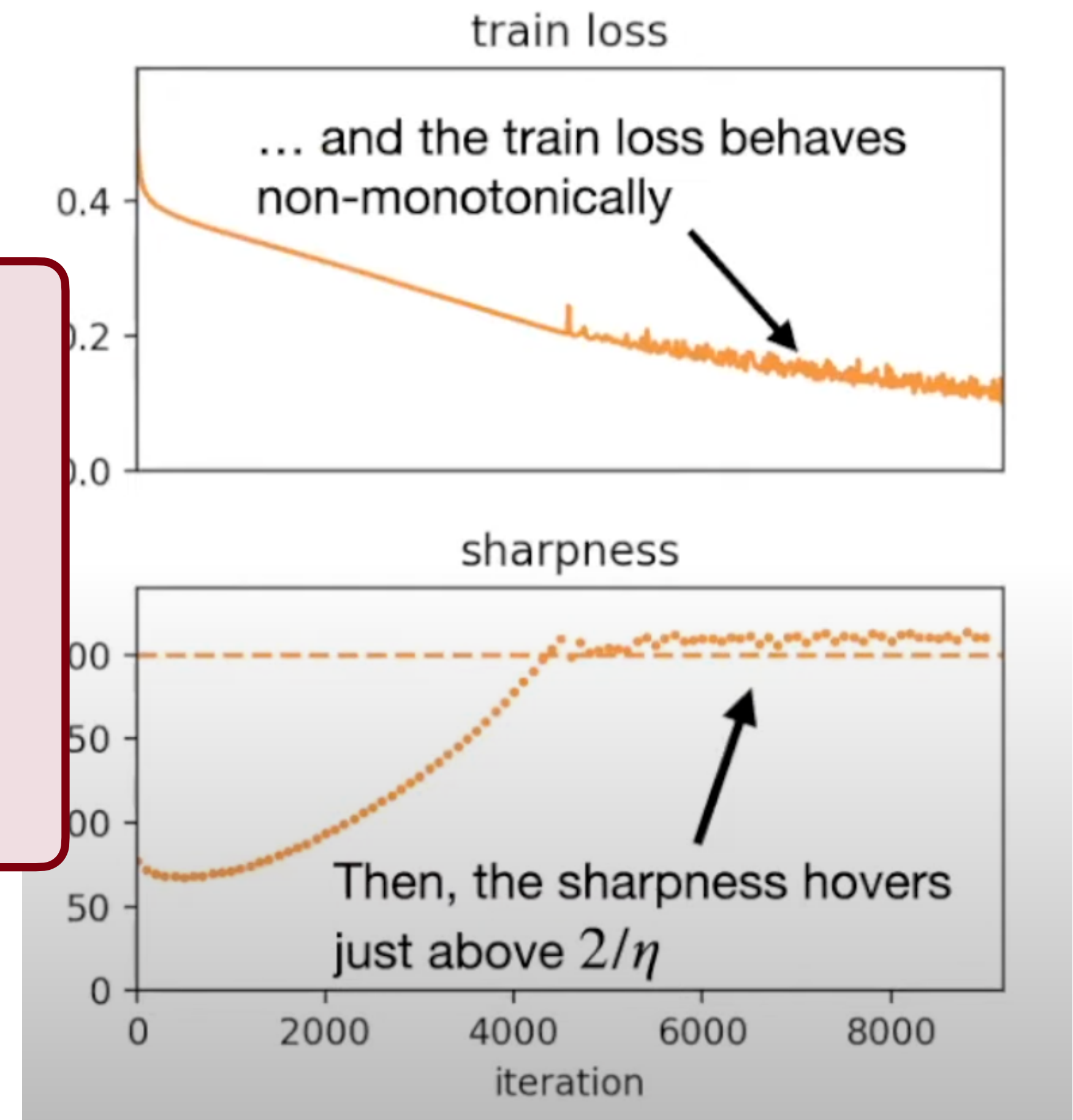
It's just a little heartbreaking that what we did seems useless

THE EDGE OF STABILITY

IMPLICATIONS:

- CONVENTIONAL WISDOM ABOUT STEP SIZE IS WRONG — *SET $\eta \leq 1/L$ IS AN IMPOSSIBLE DREAM*
- YOU HAVE TO GO DEEPER THAN TAYLOR APPROXIMATION WHEN DEALING WITH THE THEORY — *IT WOULD DIVERGE...*
- GD DOES NOT DECREASE MONOTONICALLY, *GRADIENT SIZE INCREASES* — *STILL, GD IS SUCCESSFUL.*

*Not only,
same for most
the convex
optimization*



THE EDGE OF STABILITY

Is this the case of SGD?

THE EDGE OF STABILITY

HOW ABOUT MINI-BATCH SGD?

1000s of works about:
“SGD finding flatter minima”

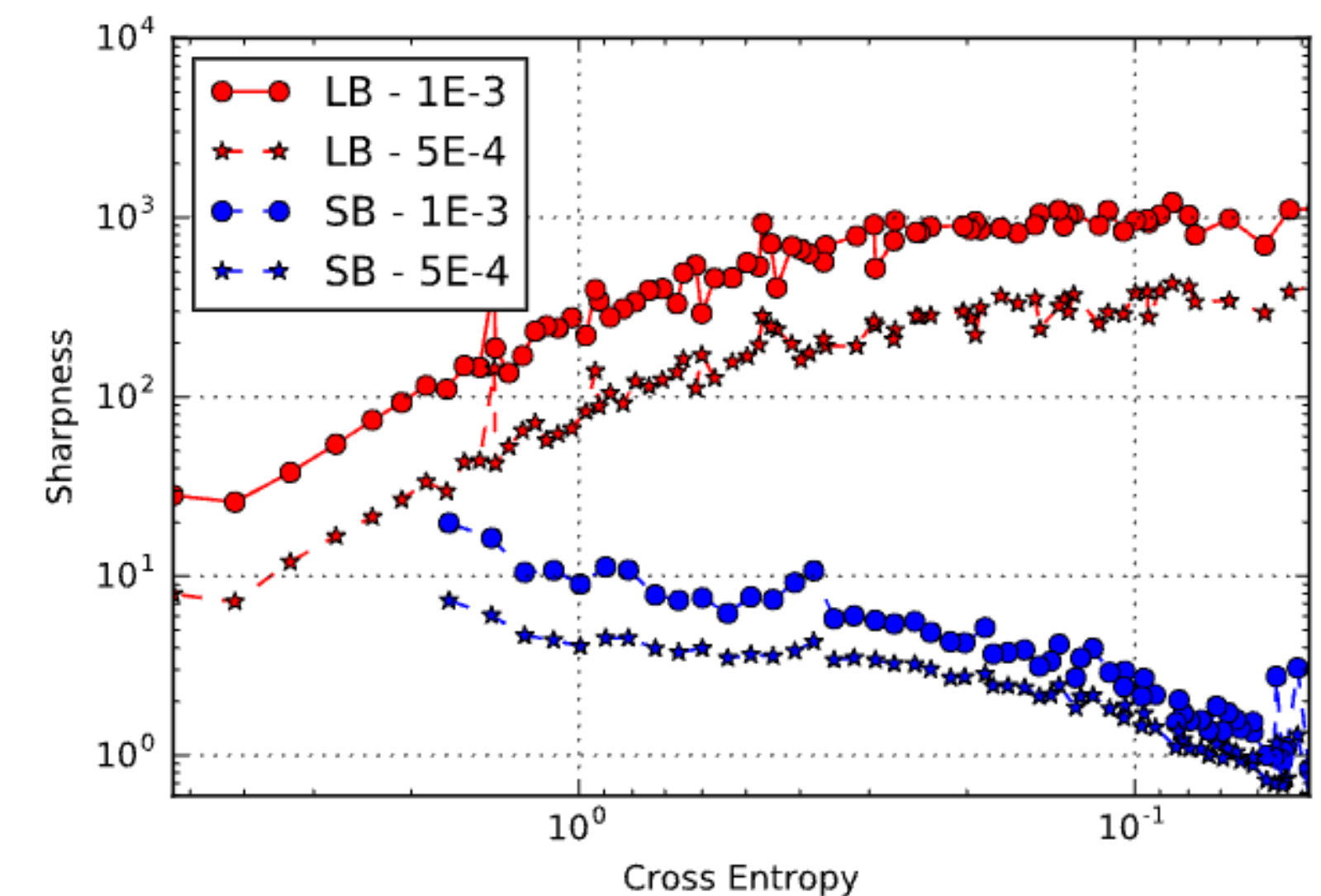
$\lambda_{\max}(H) \leq 2/\eta$ for SGD, usually!

SGD $\implies$ SMALLER $\lambda_{\max}$

KESKAR ET AL. (2016) STARTED THE FIELD

- Smaller batch $\implies$ Smaller Hessian.
- Smaller batch $\implies$ Better Generalization.

- Location Acclimates to Hyperparameters.
- Observed Progressive Sharpening.

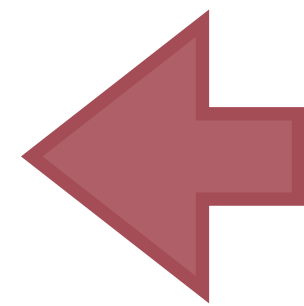


SB	LB
1.23 ± 0.83	205.14 ± 69.52
1.39 ± 0.02	310.64 ± 38.46
28.58 ± 3.13	707.23 ± 43.04

A NEW QUANTITY

$$\textit{Batch Sharpness} := \mathbb{E}_{\mathbf{B}} \left[\frac{\nabla L_{\mathbf{B}}^{\top} \mathbf{H}_{\mathbf{B}} \nabla L_{\mathbf{B}}}{\|\nabla L_{\mathbf{B}}\|_2^2} \right]$$

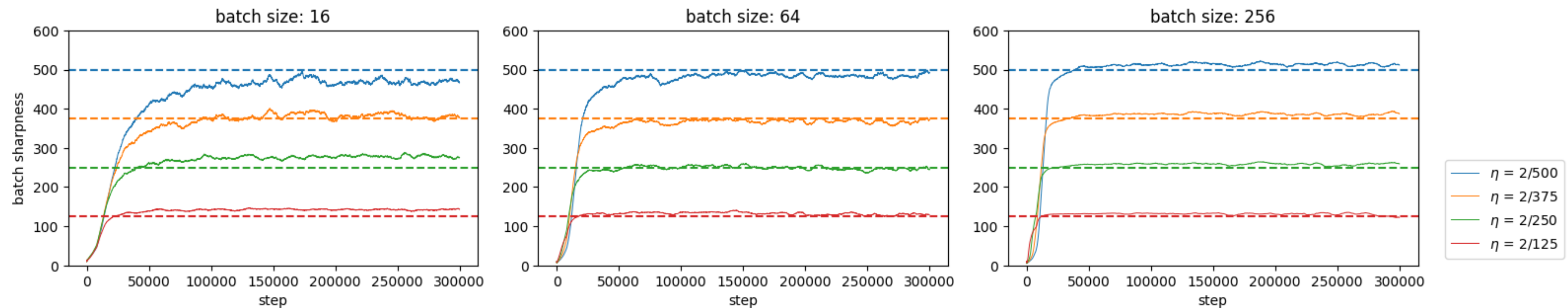
$$\frac{\nabla L_{\mathbf{B}}^{\top} \mathbf{H}_{\mathbf{B}} \nabla L_{\mathbf{B}}}{\|\nabla L_{\mathbf{B}}\|_2^2} \leq \frac{2}{\eta}$$



**This is the EoS threshold
if I were to do full-batch GD on the
dataset B!**

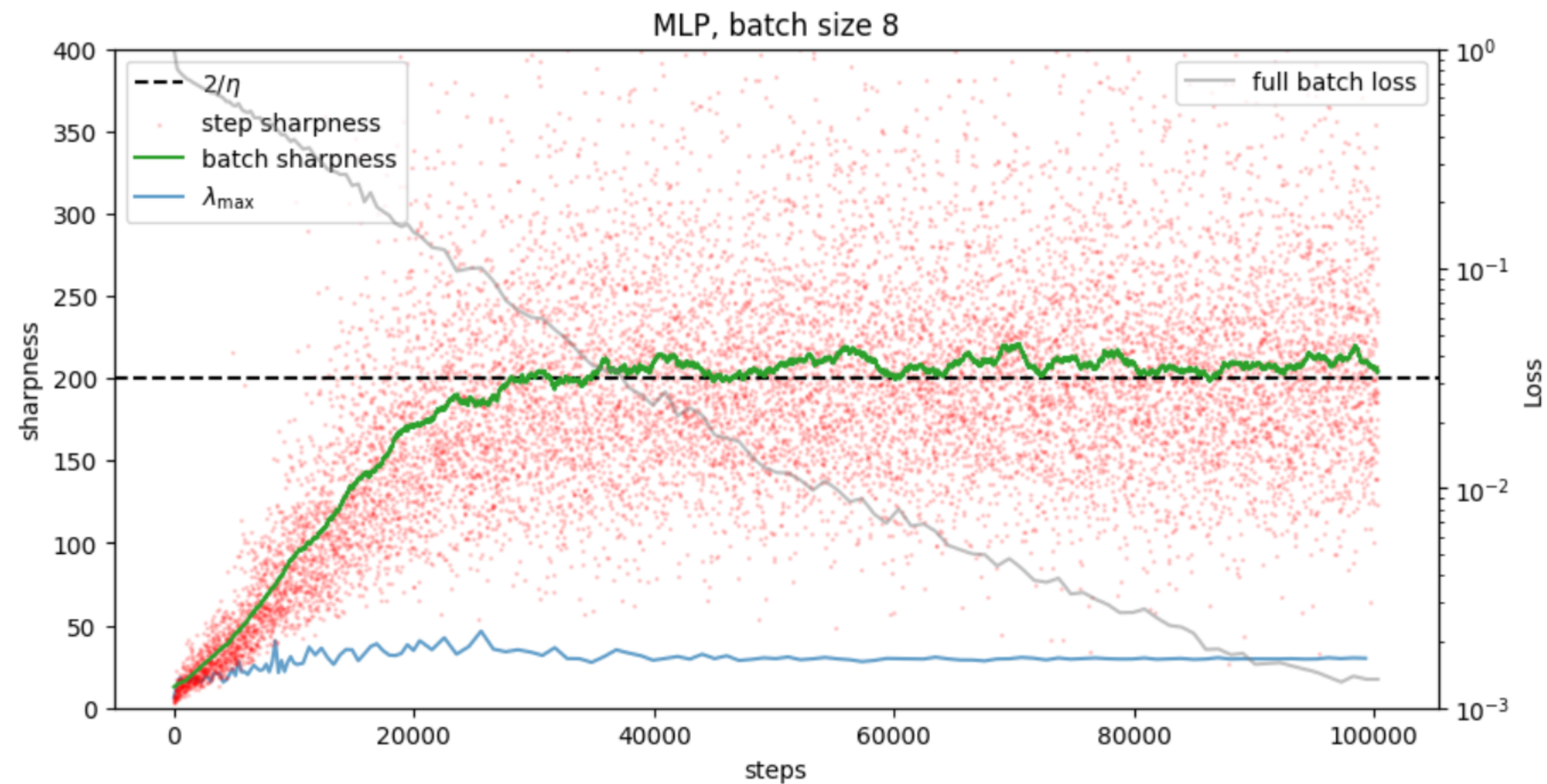
EDGE OF STOCHASTIC STABILITY

$$\textit{Batch Sharpness} := \mathbb{E}_{\mathbf{B}} \left[\frac{\nabla L_{\mathbf{B}}^{\top} \mathbf{H}_{\mathbf{B}} \nabla L_{\mathbf{B}}}{\|\nabla L_{\mathbf{B}}\|_2^2} \right]$$



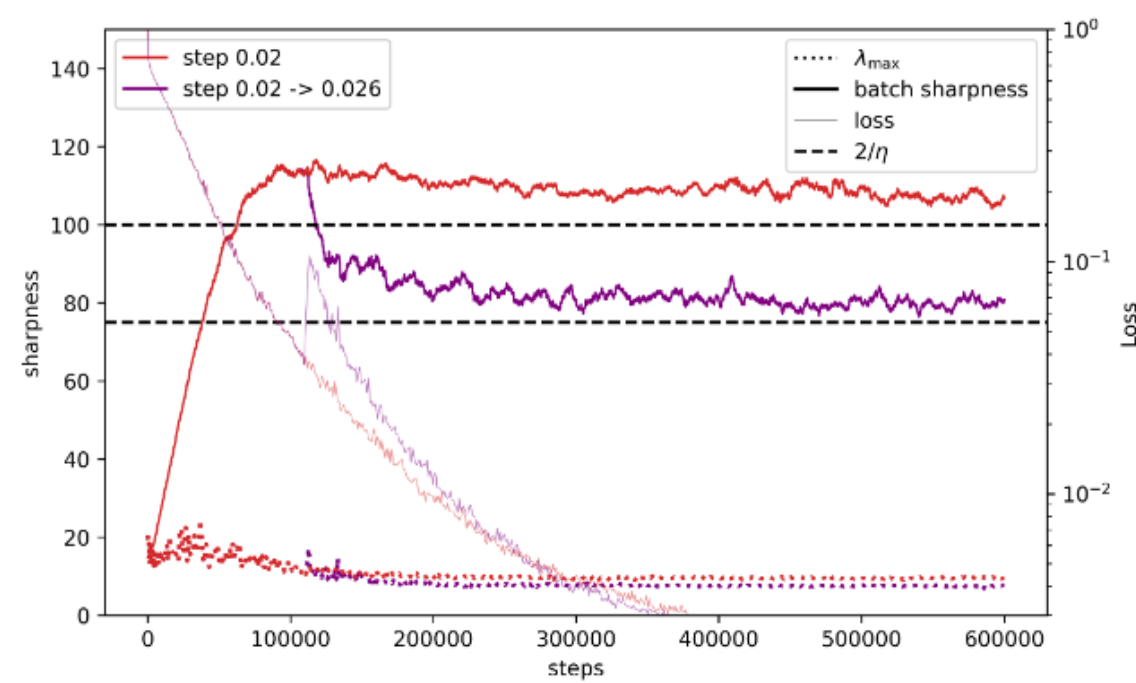
EDGE OF STOCHASTIC STABILITY

- *Batch Sharpness* stabilizes at $2/\eta$ after a long progressive sharpening.

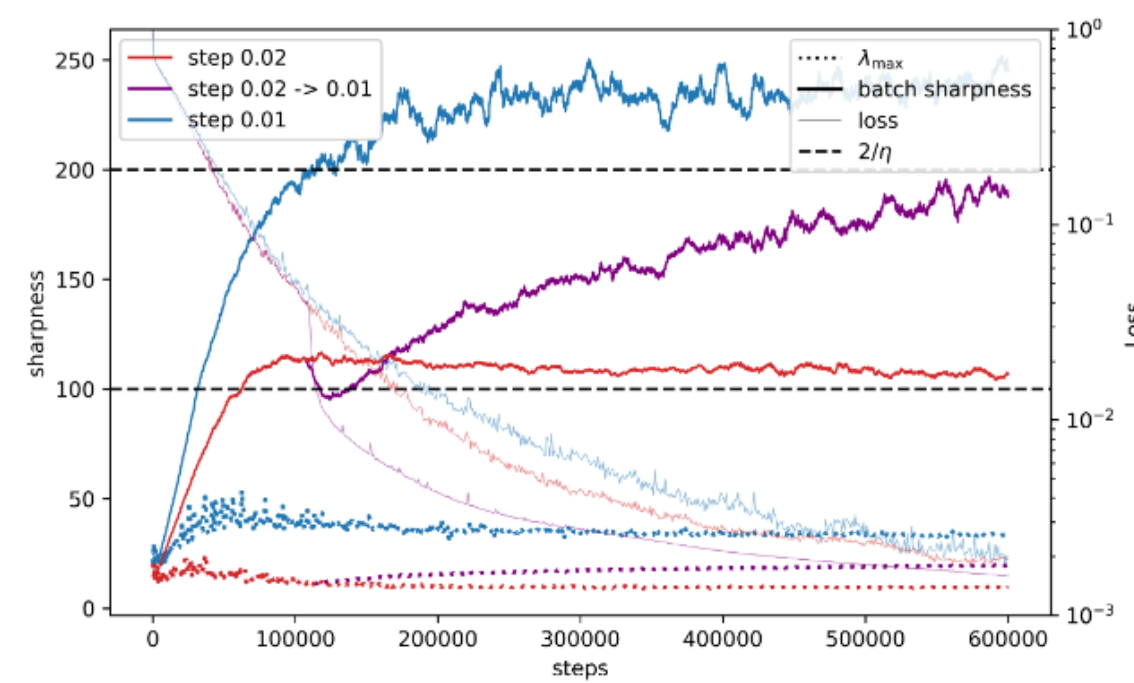


EDGE OF STOCHASTIC STABILITY

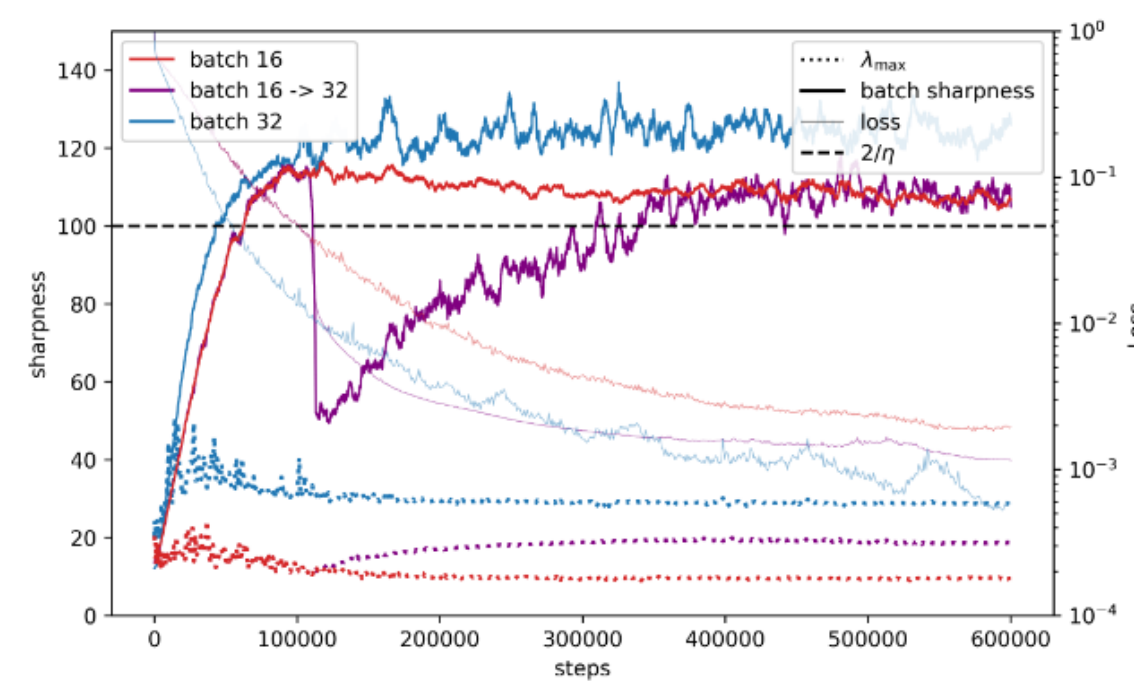
$$\textit{Batch Sharpness} := \mathbb{E}_{\mathbf{B}} \left[\frac{\nabla L_{\mathbf{B}}^{\top} \mathbf{H}_{\mathbf{B}} \nabla L_{\mathbf{B}}}{\|\nabla L_{\mathbf{B}}\|_2^2} \right]$$



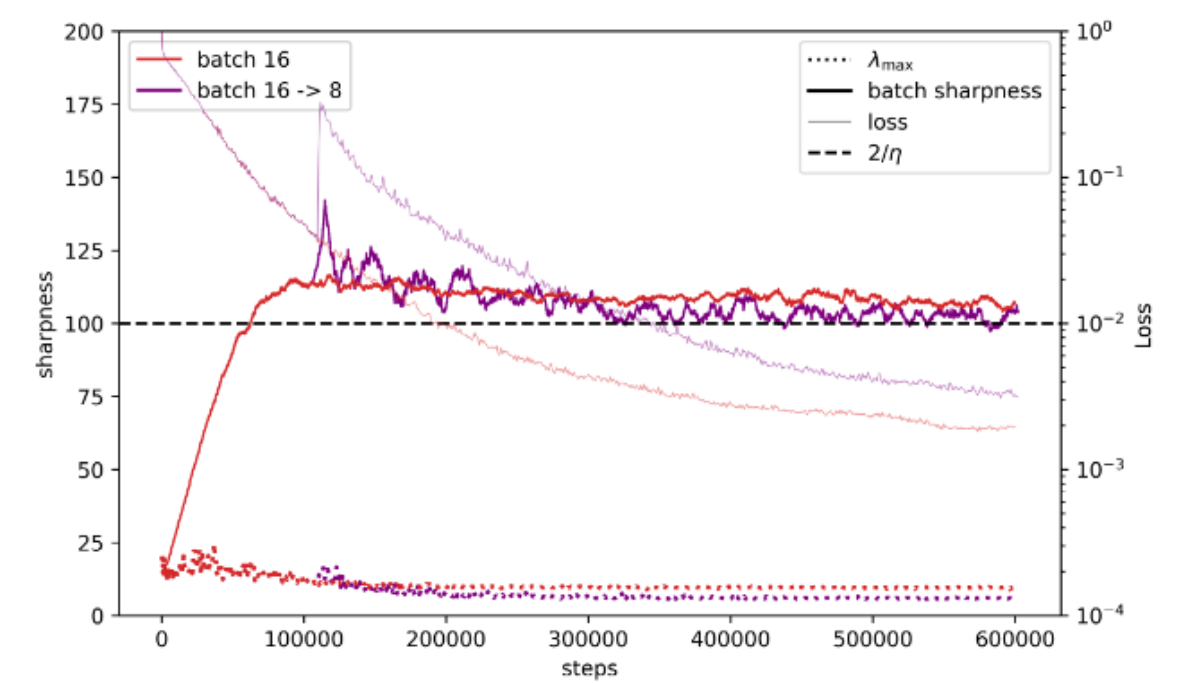
(a) Increasing step size



(b) Decreasing step size



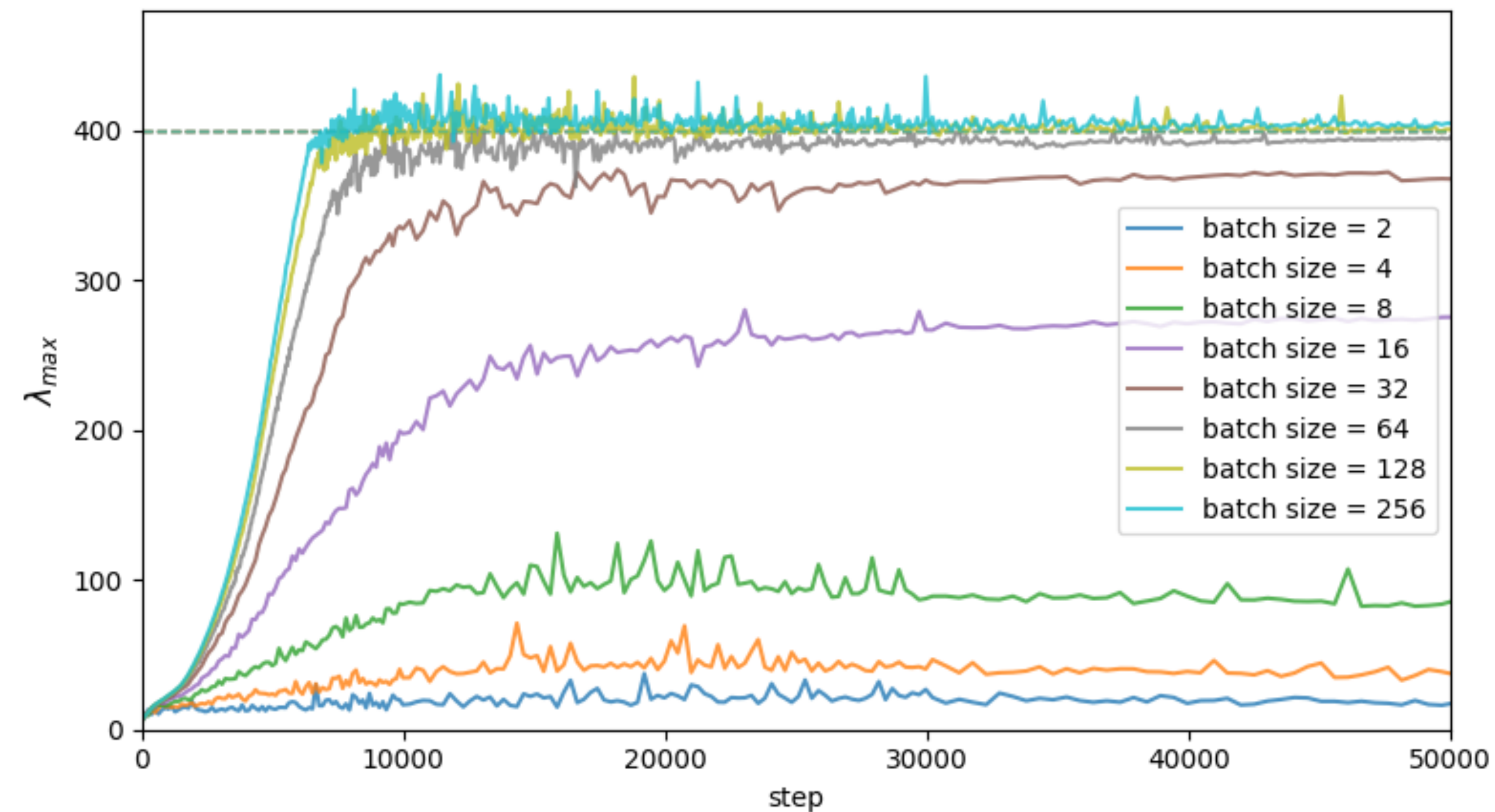
(c) Increasing batch size



(d) Decreasing batch size

EDGE OF STOCHASTIC STABILITY

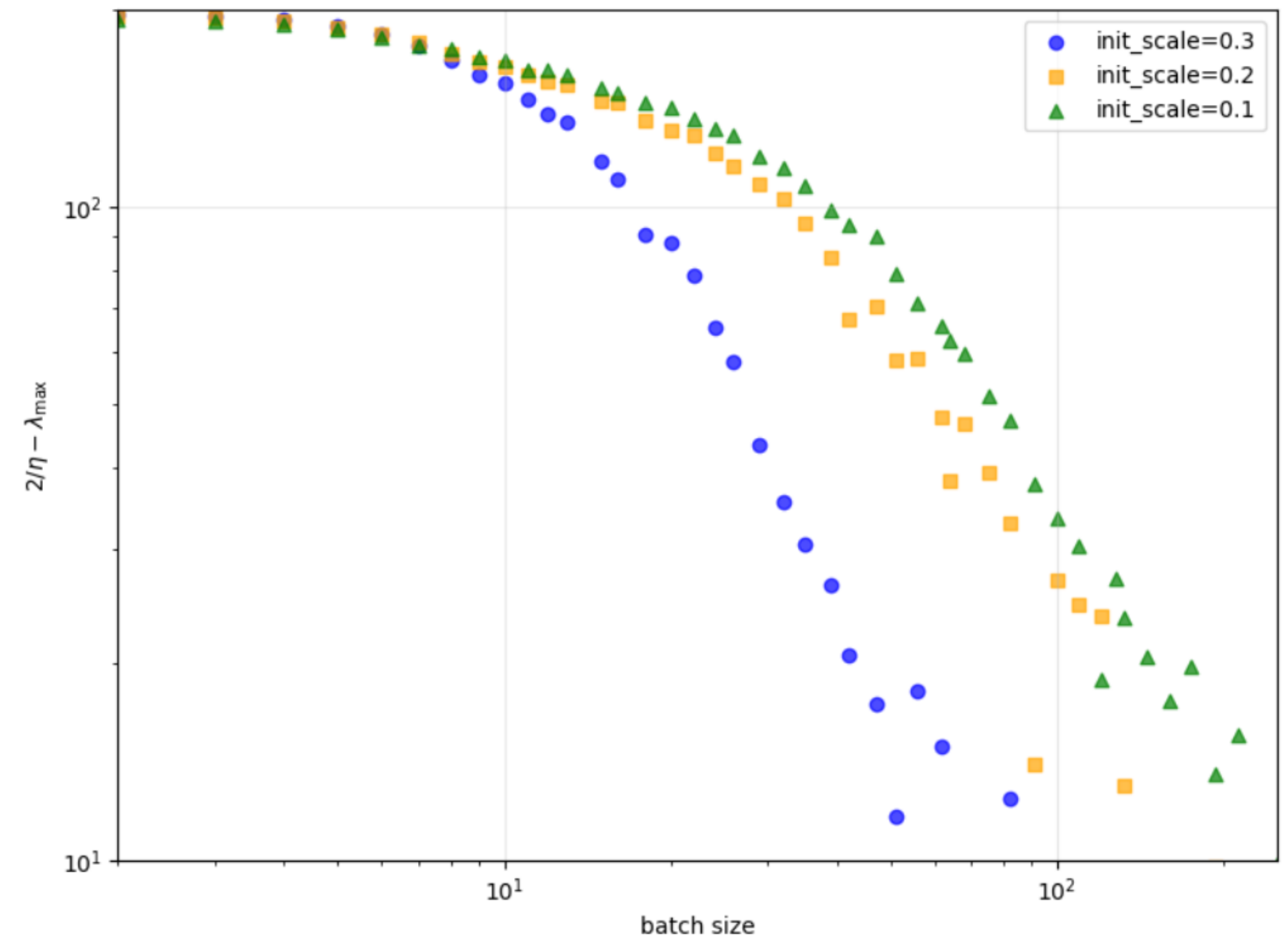
- *Batch Sharpness* stabilizes at $2/\eta$ after a long progressive sharpening.
- Also sharpness $\lambda_{\max} H$ **stabilizes** (usually below —lower for smaller BS).



EDGE OF STOCHASTIC STABILITY

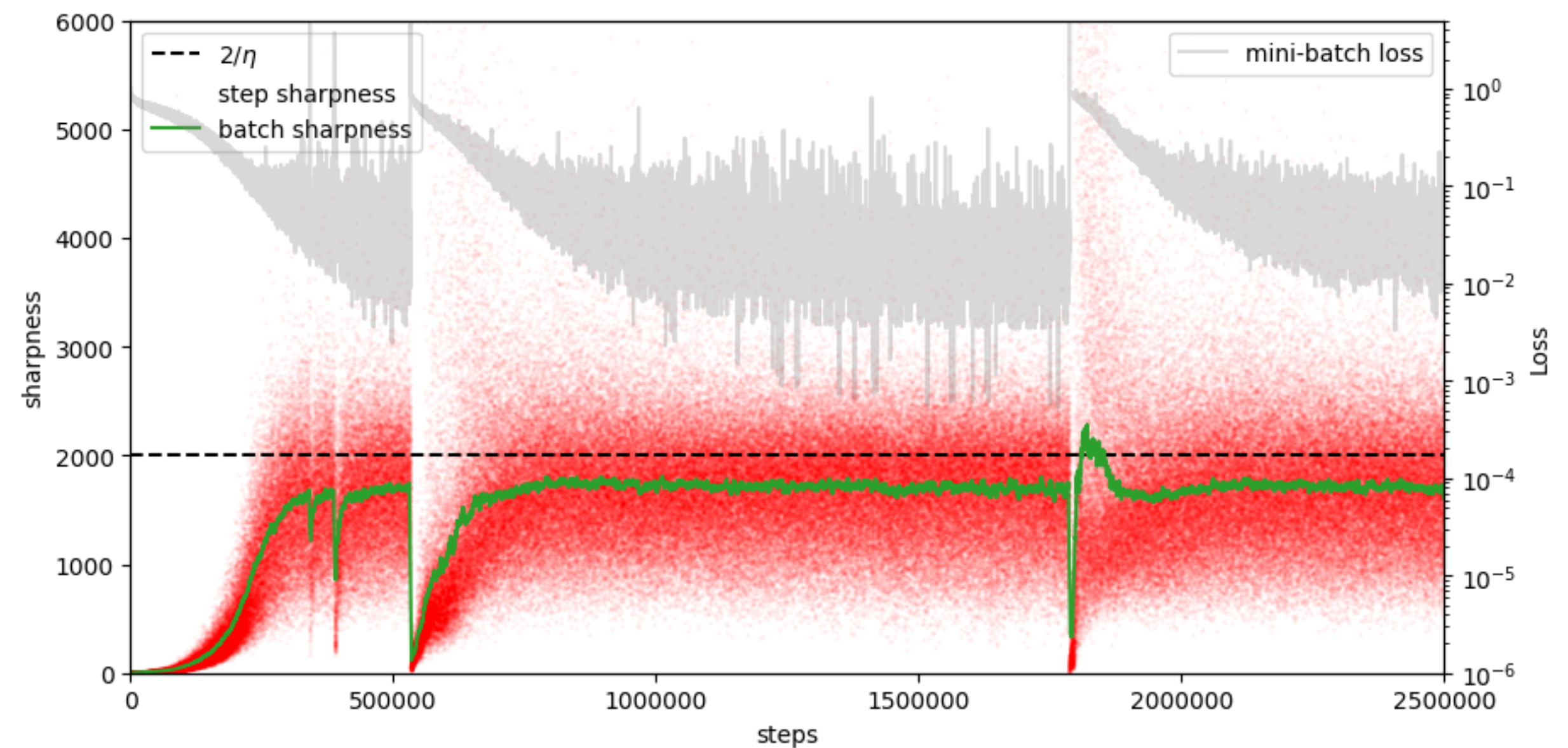
BATCH SHARPNESS, NOT $\lambda_{\max}$!

- *Batch Sharpness* going to $2/\eta$ drives the dynamics
- $\lambda_{\max}$ follows its movements
- Acts inconsistently



EDGE OF STOCHASTIC STABILITY

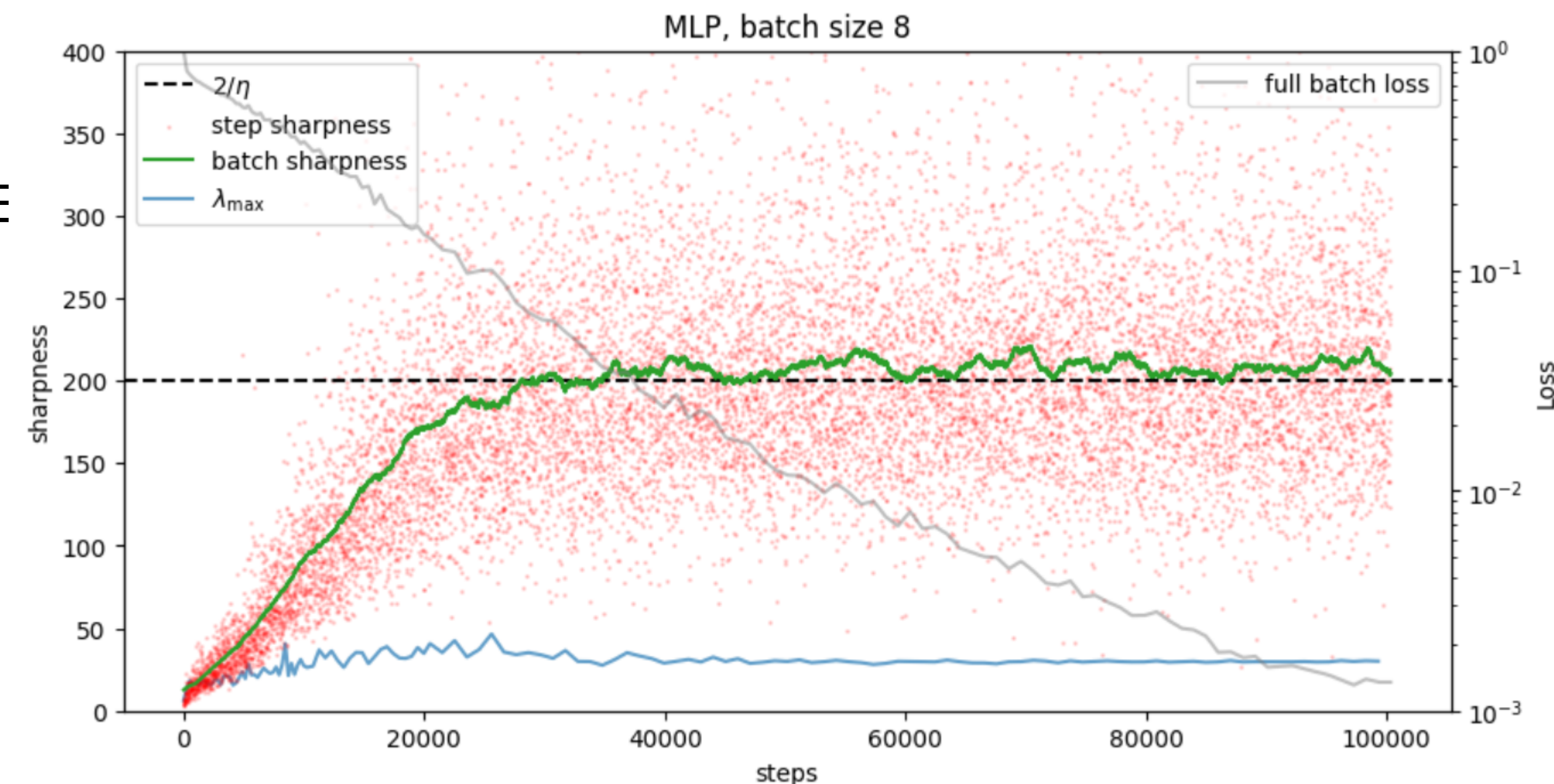
- *Batch Sharpness* stabilizes at $2/\eta$ after a long progressive sharpening.
- Also sharpness $\lambda_{\max} H$ stabilizes.
- The trajectories are subject to catapults when outliers hit!



IMPLICATION 1

THE ONES OF EOS TRANSLATE HERE:

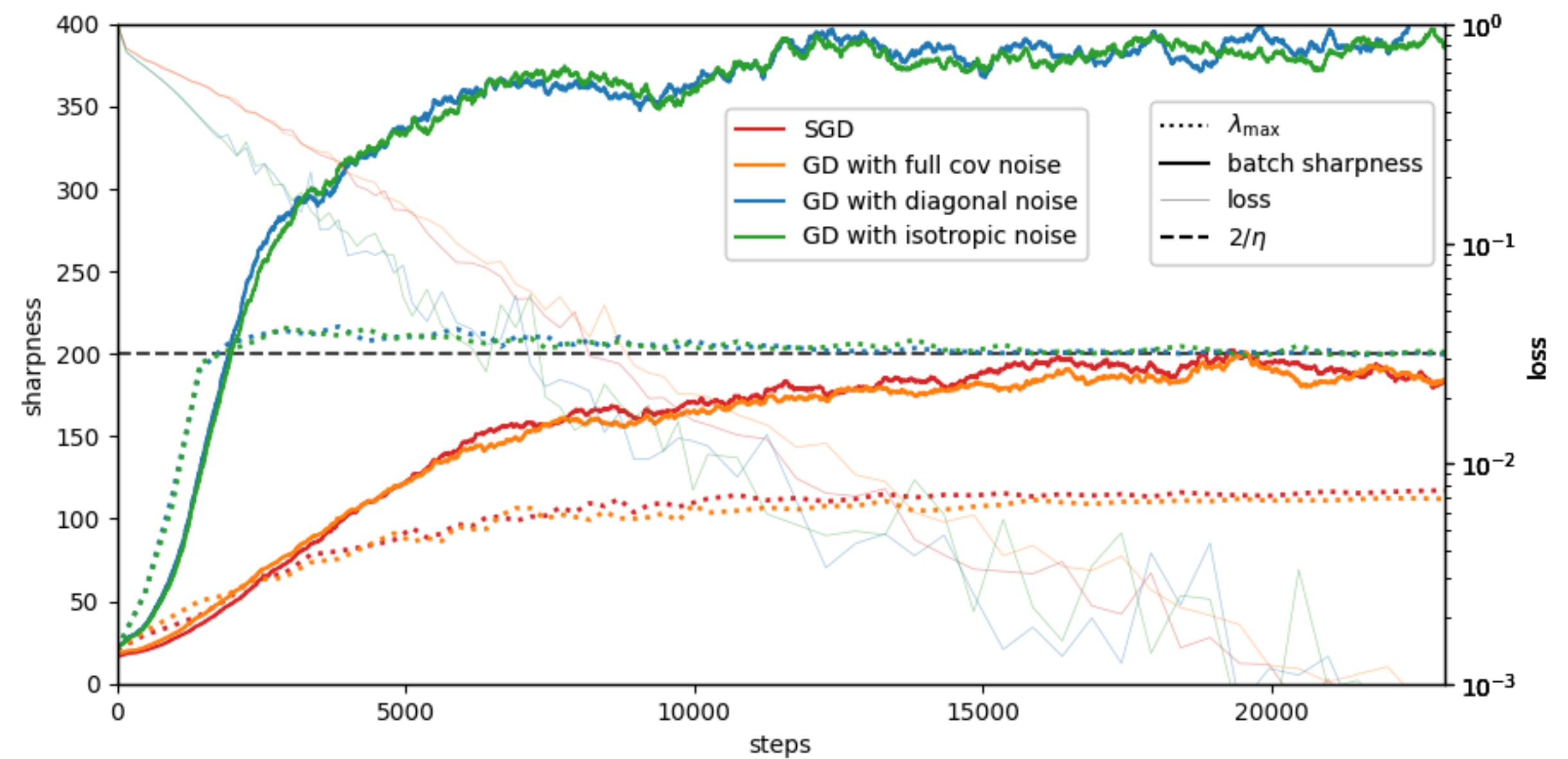
- ✓ CONVENTIONAL WISDOM ABOUT STEP SIZE IS WRONG — *SET $\eta \leq 2/L + \dots$ IS AN IMPOSSIBLE TASK!*
- ✓ YOU HAVE TO GO DEEPER THAN THE ORDER-2 TAYLOR APPROXIMATION WHEN DOING NN THEORY — *IT WOULD DIVERGE...*
- ✓ GD DOES NOT DECREASE MONOTONICALLY THE LOSS, *GRADIENT SIZE INCREASES.* STILL, GD IS SUCCESSFUL.



IMPLICATION

RANDOM HESSIANS

- **Naive SDE modeling of SGD**
full-batch Hessian
- **Noise-injected SGD**
full-batch Hessian
- **Mini-Batch SGD**
distribution of the Hessians

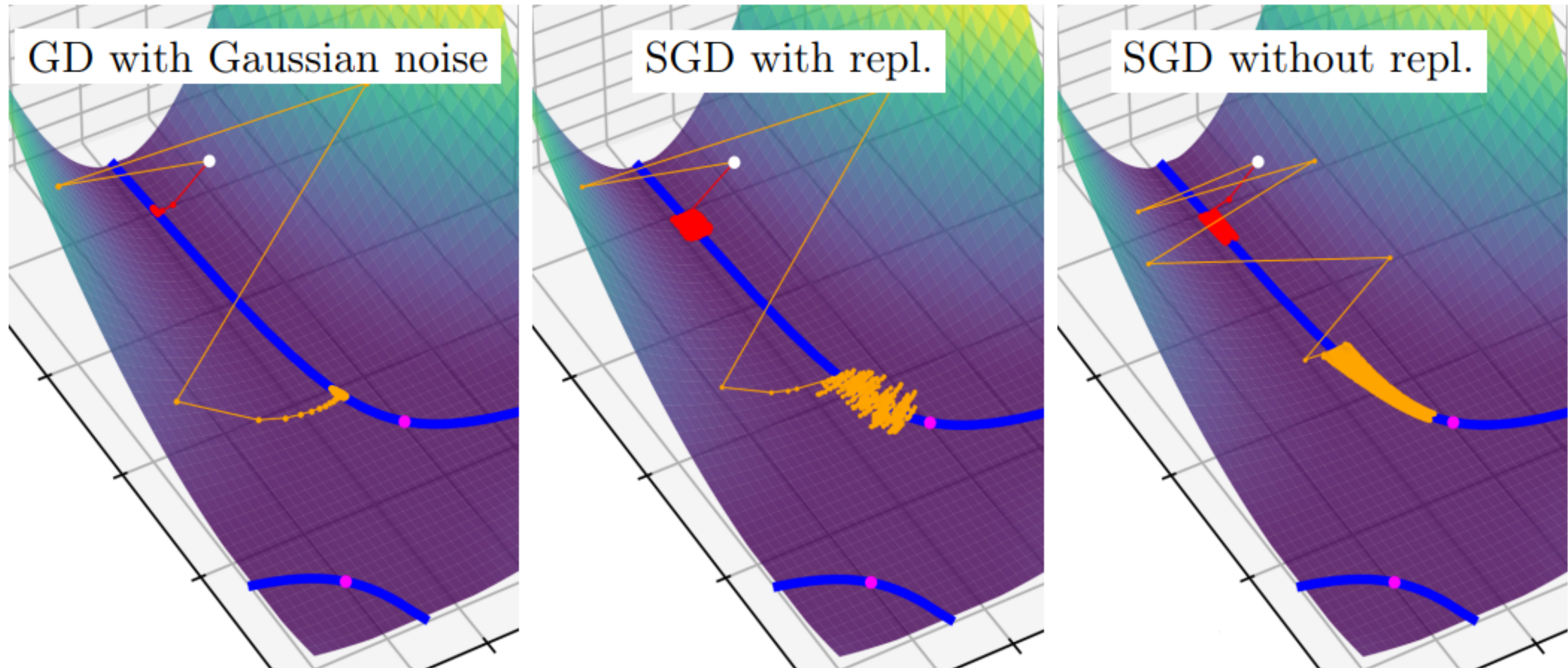


Training Networks

What **qualitatively** matters empirically:

- Deterministic **vs** Stochastic
- What **kind of** stochasticity

Comparing SGDs



$$\frac{1}{2n} \sum_{i=1}^n (\theta_1 \cdot \theta_2 \cdot x_i - y_i)^2 \quad \text{with dataset } \{(x_i, y_i)\}_{i=1}^n \subseteq \mathbb{R} \times \mathbb{R} \text{ and parameters } (\theta_1, \theta_2) \in \mathbb{R}^2$$

CONCLUDING

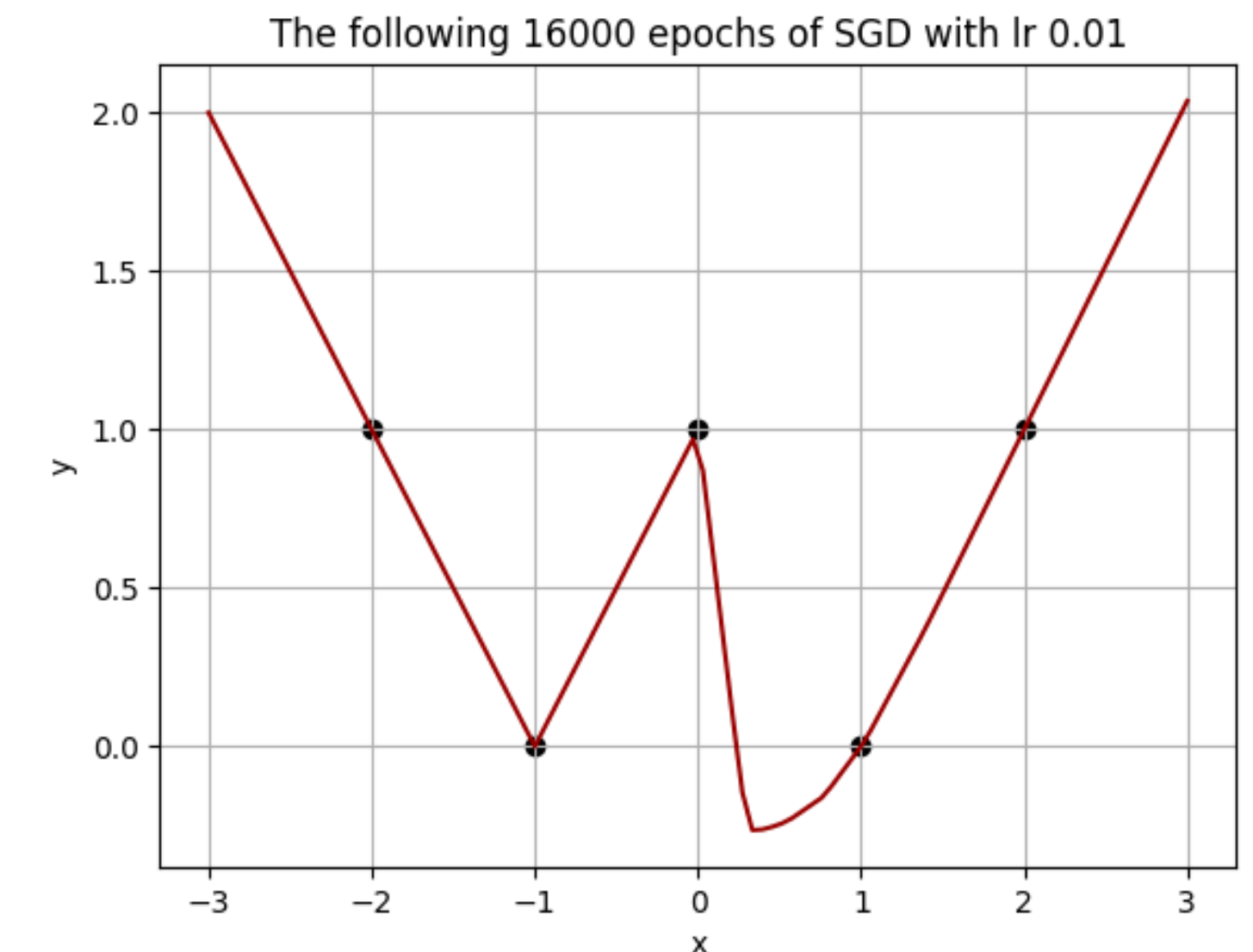
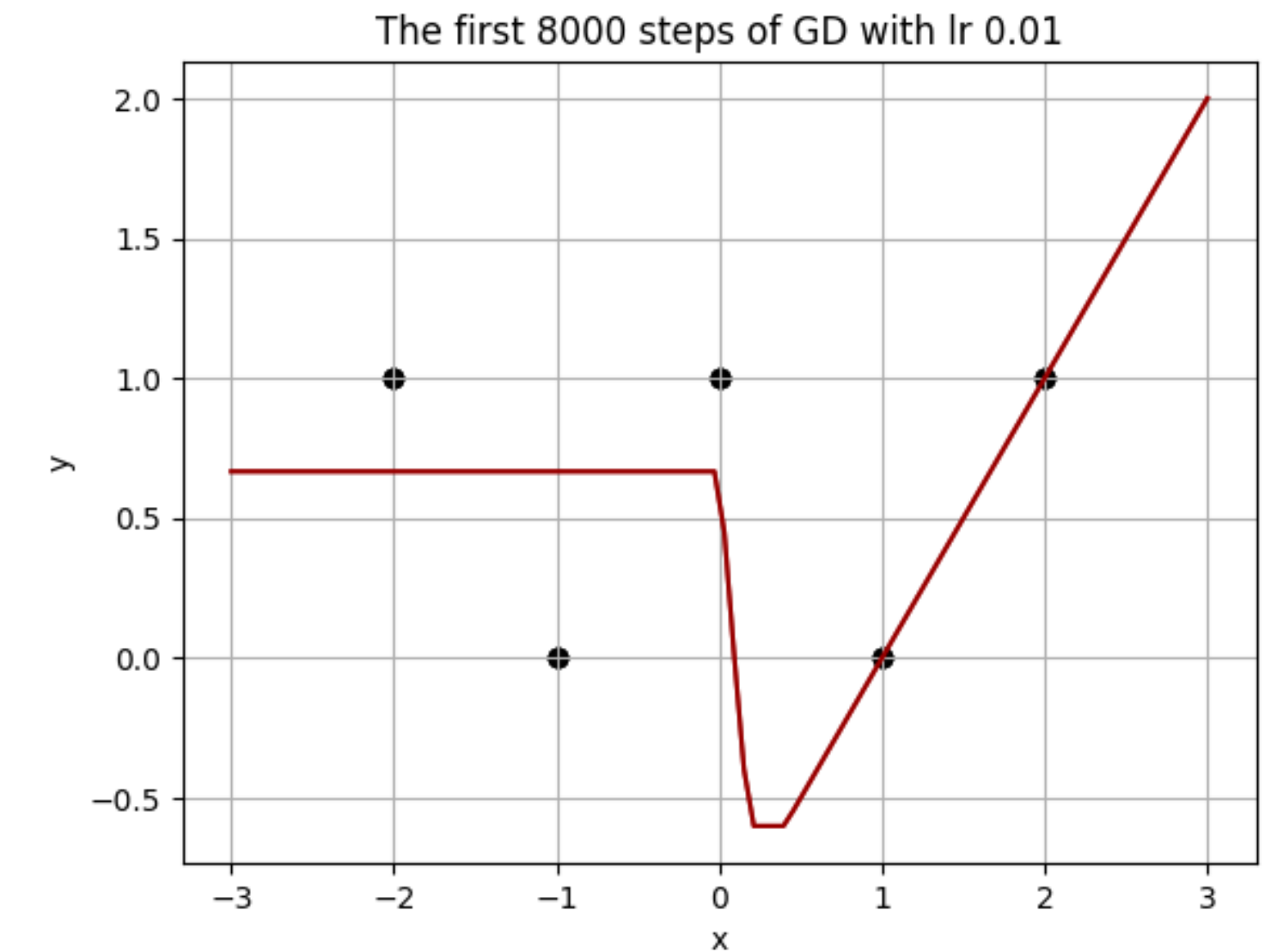
OPTIMIZING *vs* LOCATION

Goal of training neural networks:

Find the one which
Performs best on your task

What we can work with:

- The data we got
- The measure of error $L(\theta)$
- (Within) a compute budget



THE QUESTION

How to go optimally
from local to global

Where does our training algorithm
go in the case of neural networks?

A preliminary

Questions

- We would like to design a **globally optimal** gradient based optimizer.
- We will see that they have been designed for something else:

Convex functions

- We do not have the tools to discuss which one is local to global optimal, yet.
- So we will just explore:

Where they go in NN

What tools can we figure out for the problem

SPEAKING MATH

WHERE DO ALGORITHMS GO?

This is not an *optimization* question:
It is a dynamical systems question!

Questions are:

- What are the basins of attraction? Do trajectory recur? (**Long-term behavior**)
- Is the invariant set stable (**Stability in the sense of Lyapunov, ...**)
- How does the behavior changes with the (hyper)parameters (**Bifurcation Analysis**)

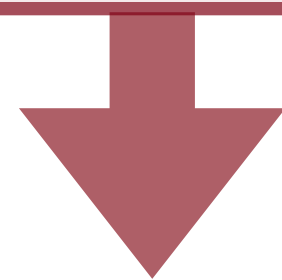
Message 3, Class 17

- We want to pick the **architecture** and the **initialization** to enhance **trainability**.
- Not really for purposes of **approximation** or location of **convergence**!

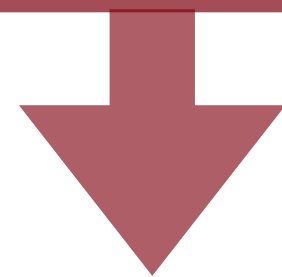
So let's study the *optimizer*
instead of the initialization

THE MOTIVATION

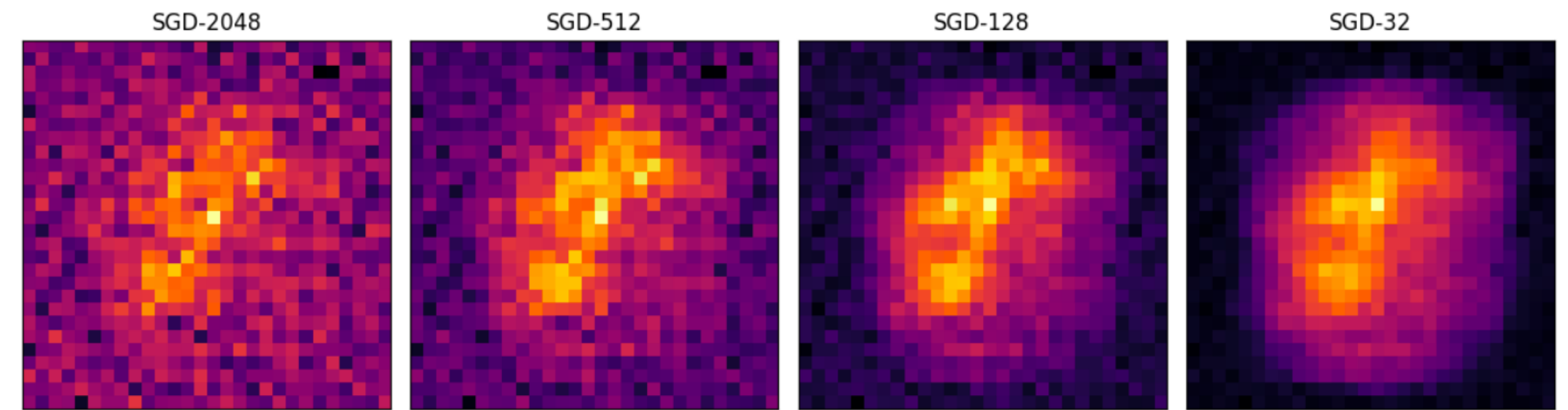
Goal is having best
“performance”



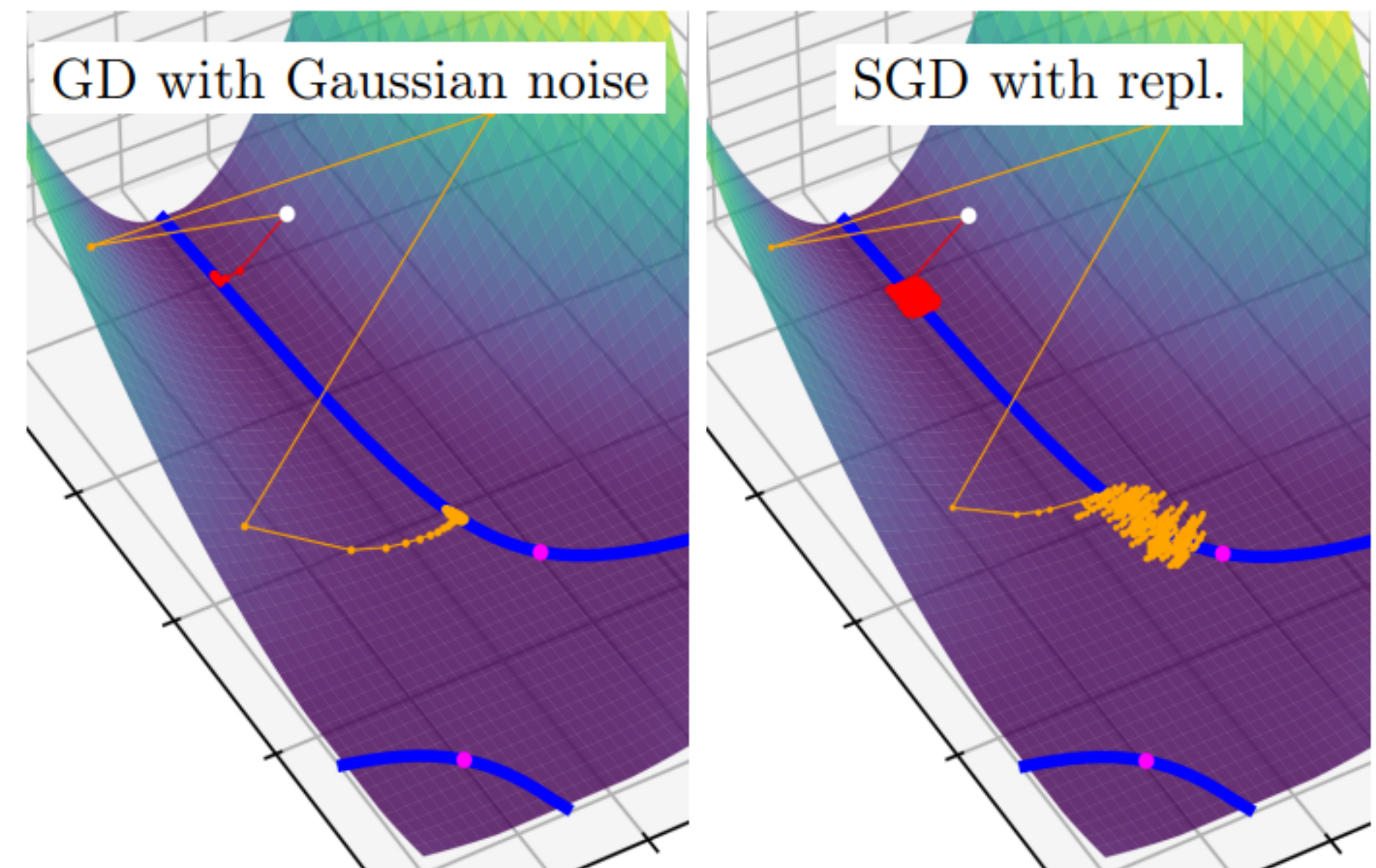
There are many different minima
with different performance



Different algorithms find
different minima, *consistently*



SGD smaller batch size



Training Networks

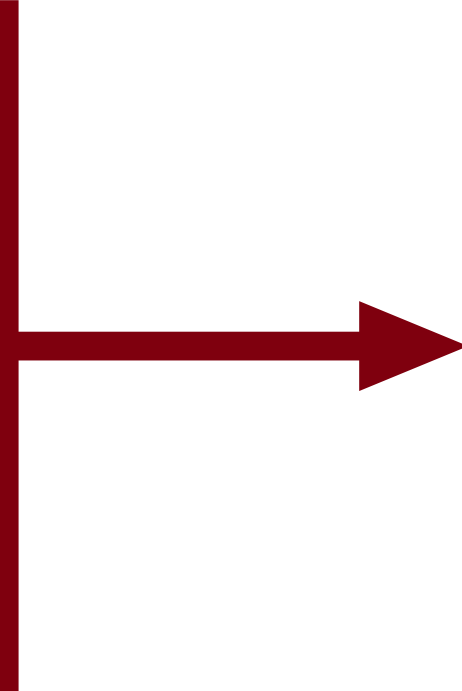
What **qualitatively** matters empirically *(for performance in Regime 2)*:

- Deterministic **vs** Stochastic
- What **kind of** stochasticity
- The choice of **hyperparameters**

Consistently

Questions

- How do they **escape local minima**?
- How do they **escape saddles**?
- How do they **pick** the minimum in the high dimensional manifold?



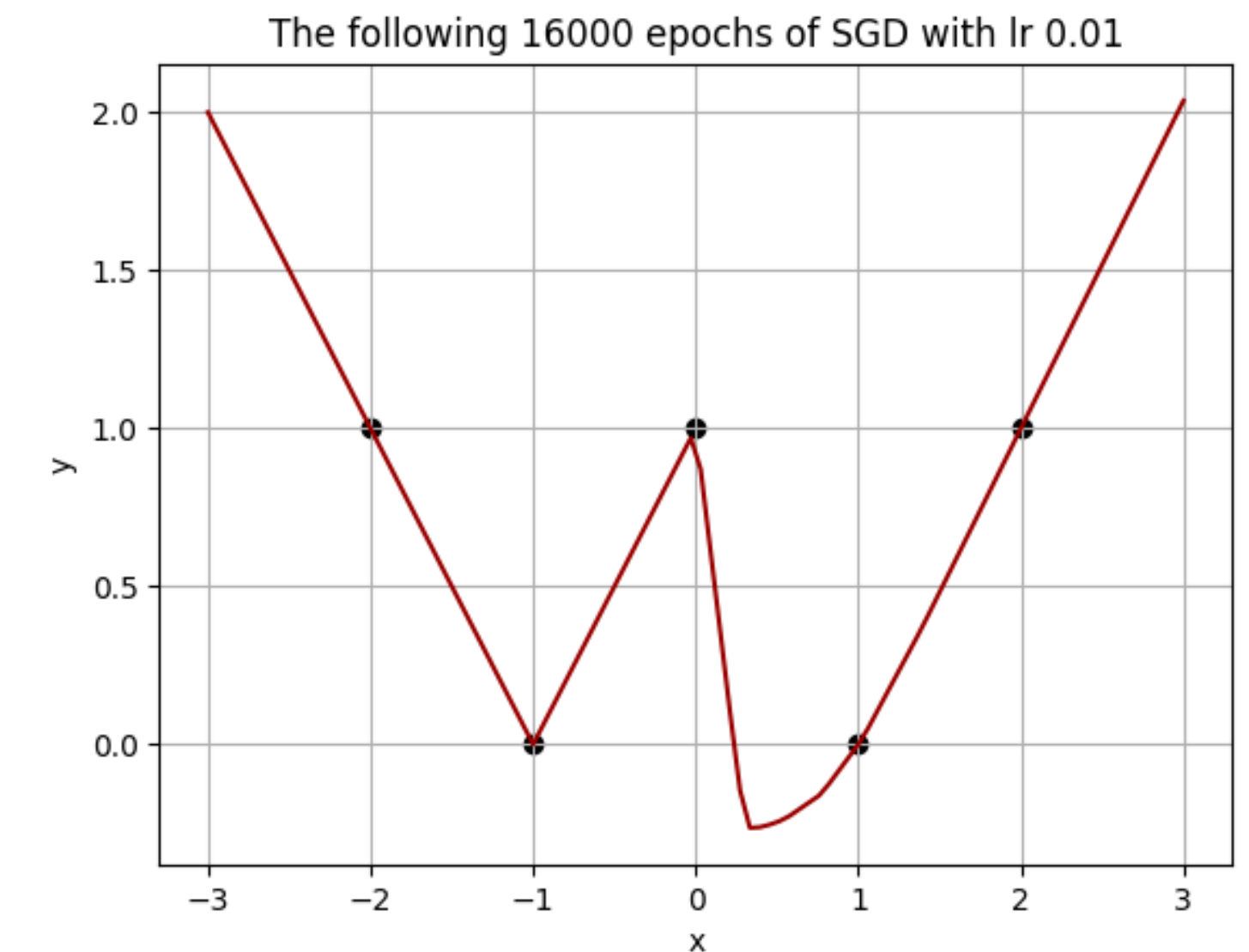
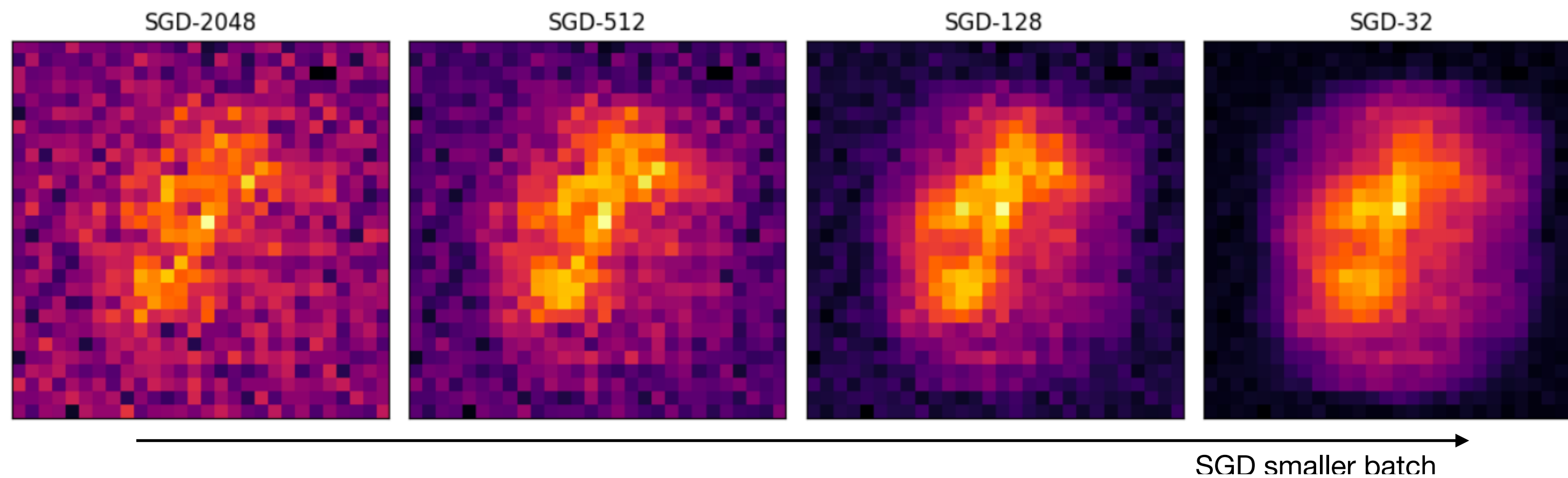
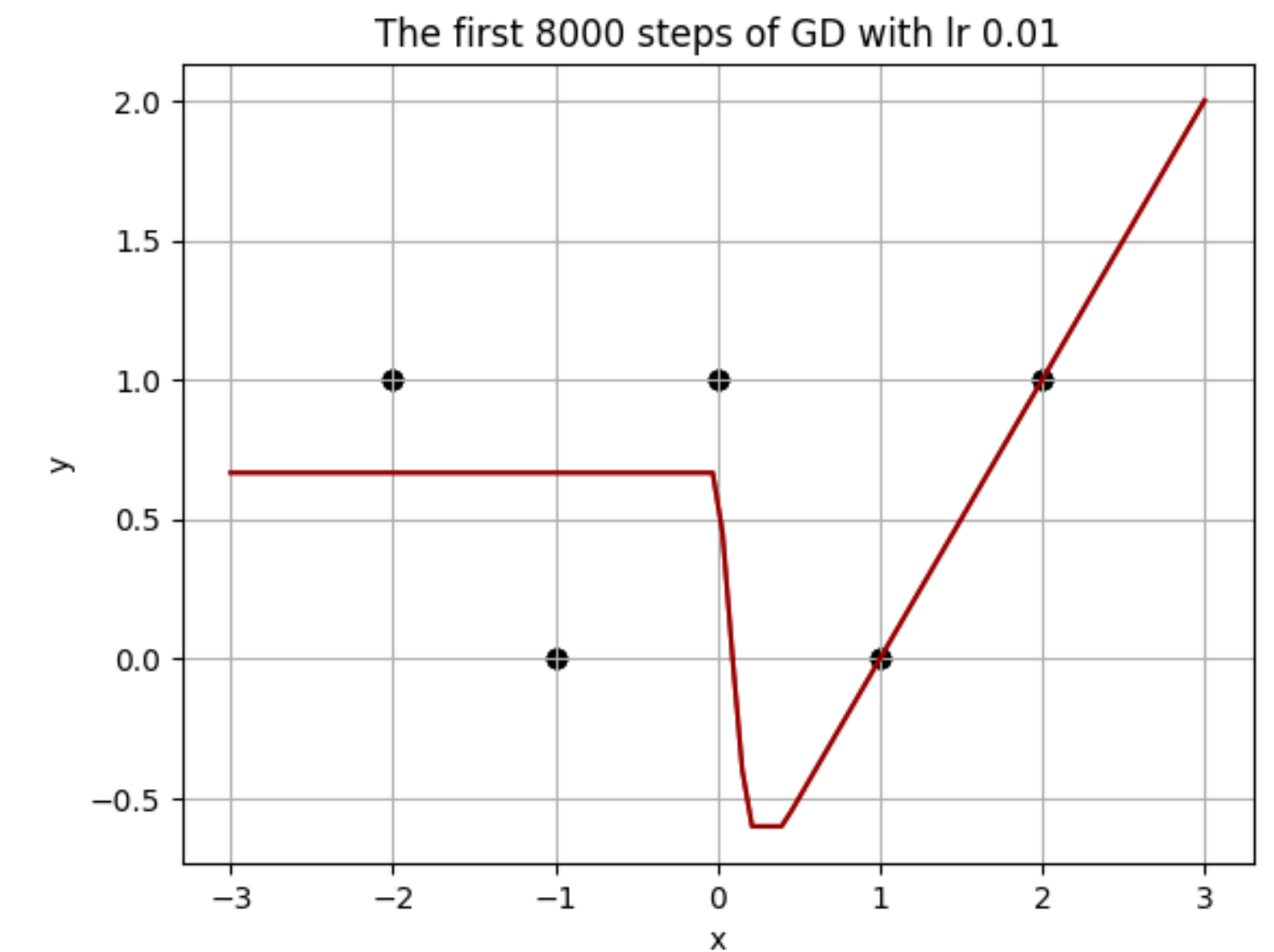
These two are about
traveling flat areas!

All things they have *not* been designed for!

GOAL FOR MATHEMATICIANS

Understanding the location of convergence

- Of different algorithms / hyperparameters



ANSWER

Check the step size / hyper parameters
given some sort of curvature

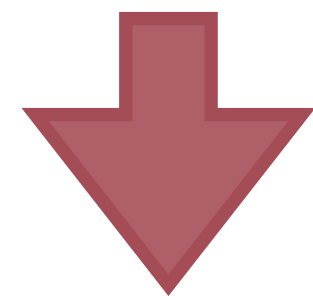
ANSWER

Check the step size / hyper parameters
given some sort of curvature

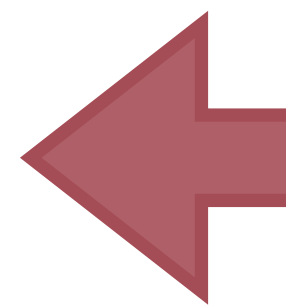
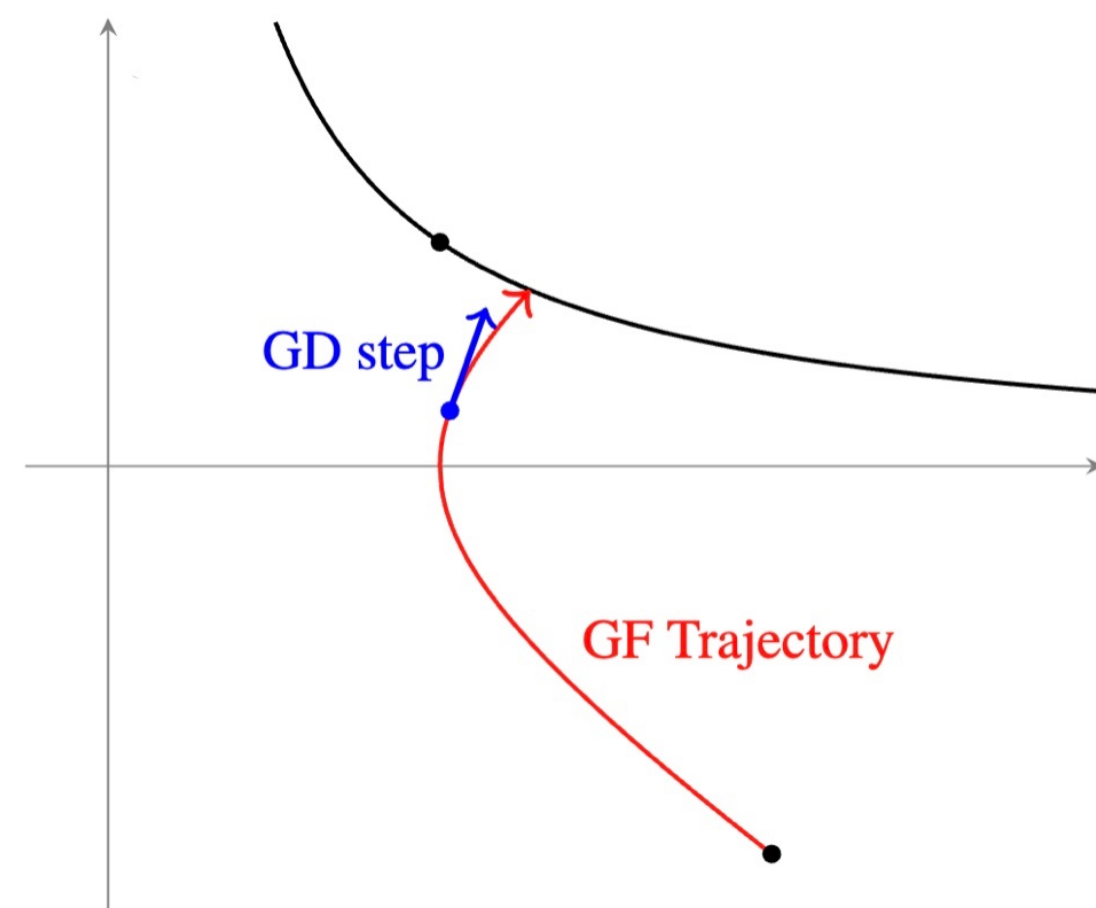
- ***Notions of curvature:***
are the discriminant for large/small step.
- *You start stable:*
Large == Progressive Sharpening is there!
 - ▶ *Why? To what extent? This is open!*

ANSWER

Check the step size / hyper parameters
given some sort of curvature



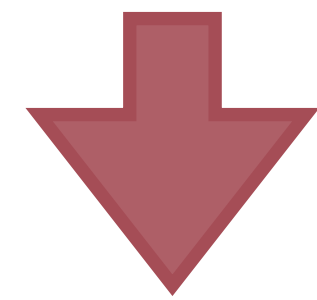
Small



The tools are Numerical
Analysis-ish

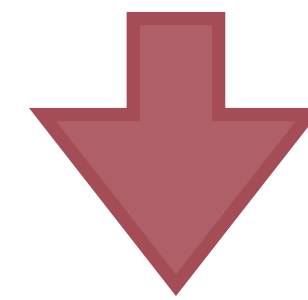
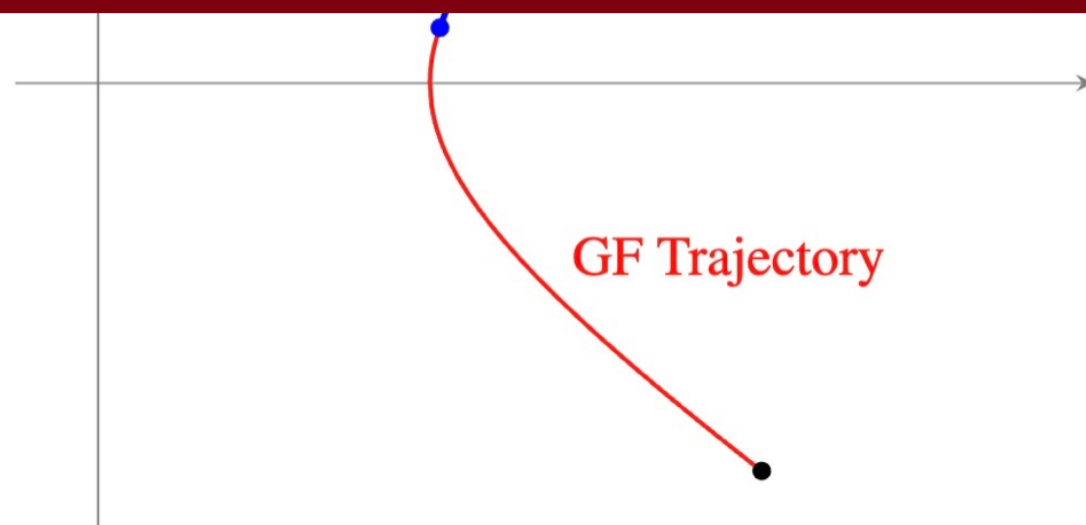
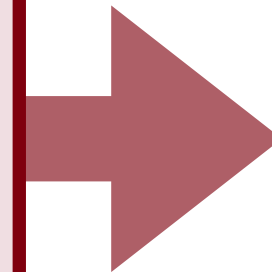
ANSWER

Check the step size / hyper parameters
given some sort of curvature

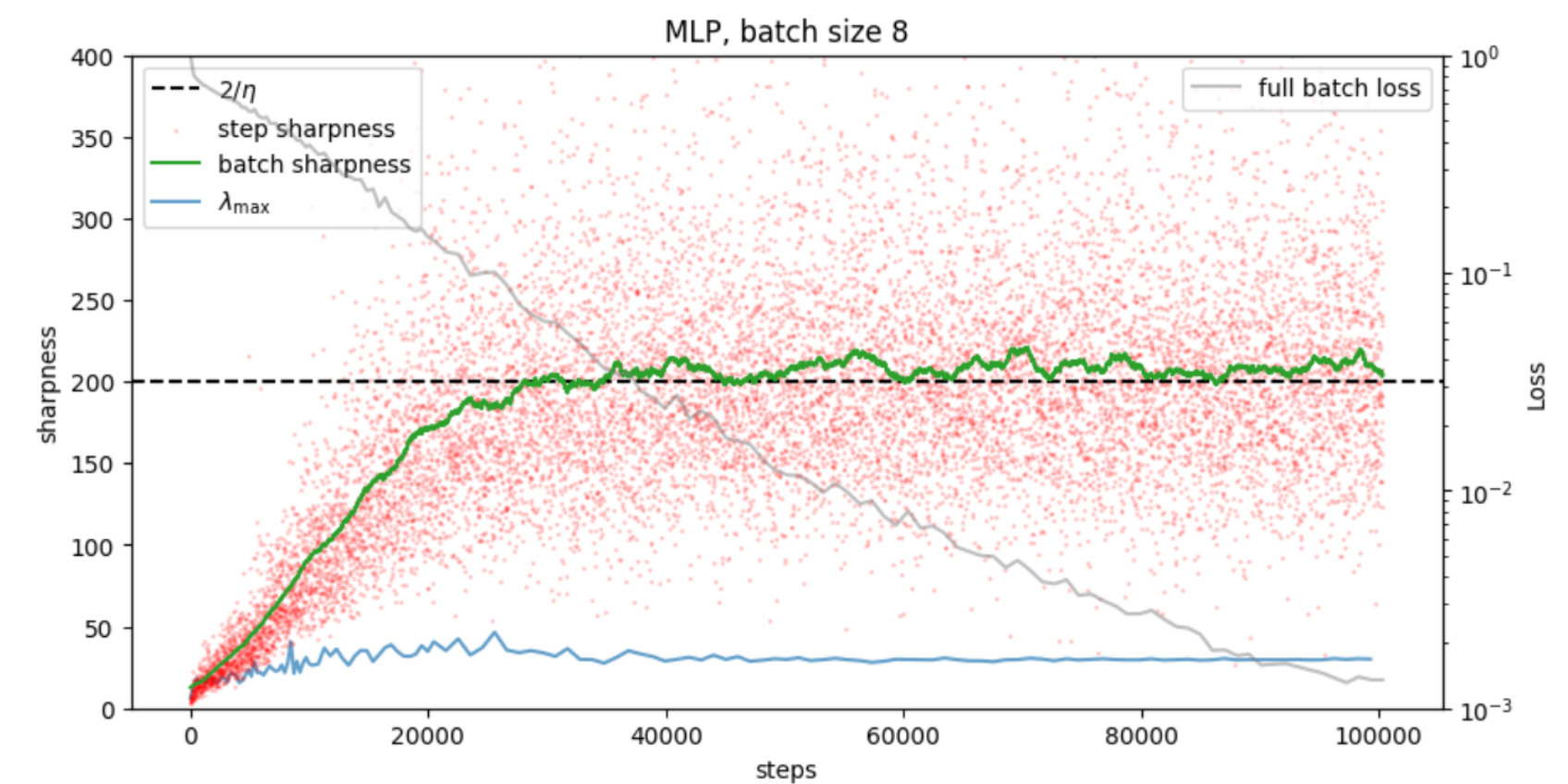


Small

We understood
empirically!

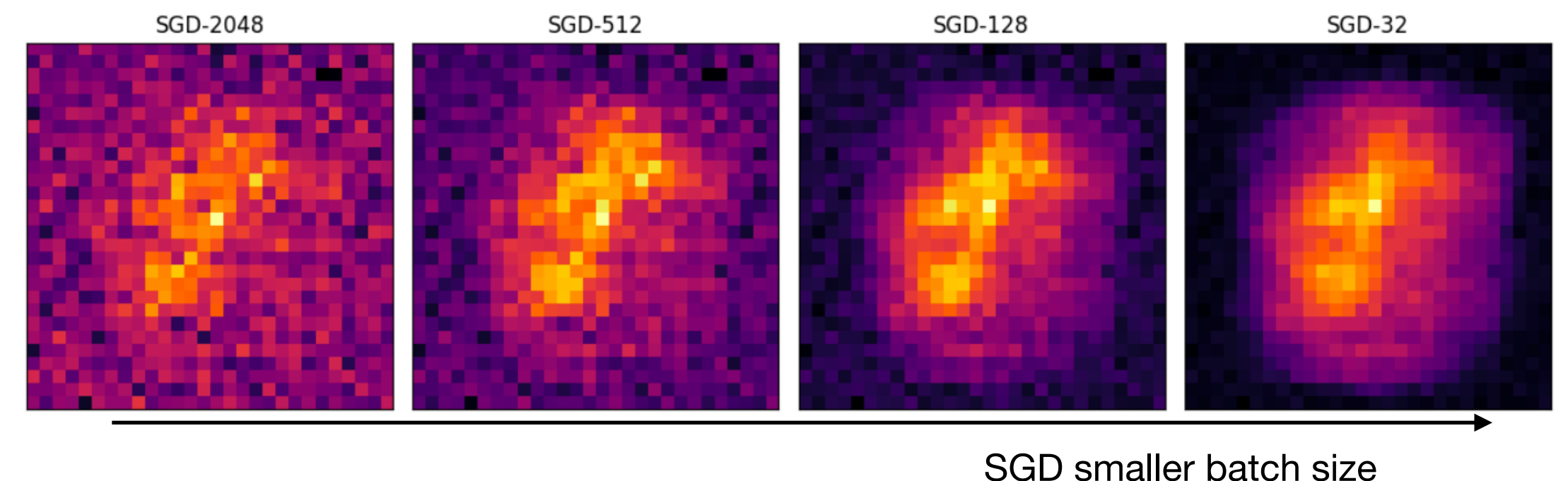


Large



What is Missing

- Many things we didn't discuss as how you escape saddles and local minima.
- How does this translate into what is learned? What do they learn? It changes with the optimizer but how?
- Is there a learning order?



The Broader Question

We have a model (i), some data (ii), a loss (iii), an optimizer (iv) with some hyper parameters (v), a computing budget (vi) and a precise initialization θ_0 (vii):



*What is the the function represented by the neural network with parameters θ^**

And it's about location of stabilization.