

<https://introml.mit.edu>

6.390 Intro to Machine Learning

Lecture 3: Gradient Descent Methods

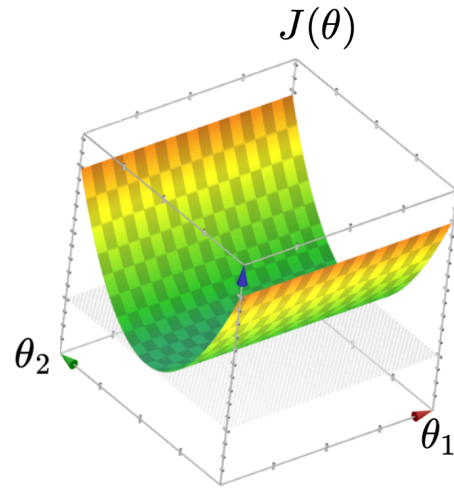
Shen Shen

Sep 21, 2026

2:30pm, Room 10-250

[Slides and Lecture Recording](#)

Recall:

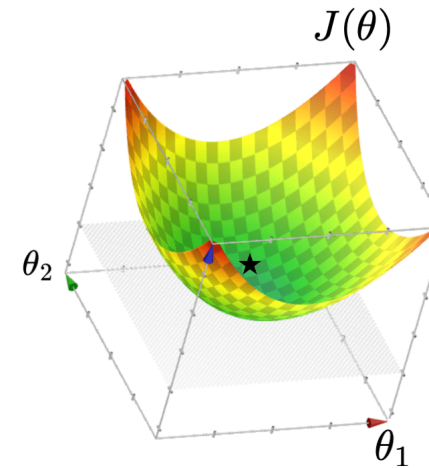


When X is not full column rank

- $J(\theta)$ has a "flat" bottom
- That formula is not well-defined 🙅
- Infinitely many optimal hyperplanes

No way yet to get any θ^*

θ^* numerically
sensitive



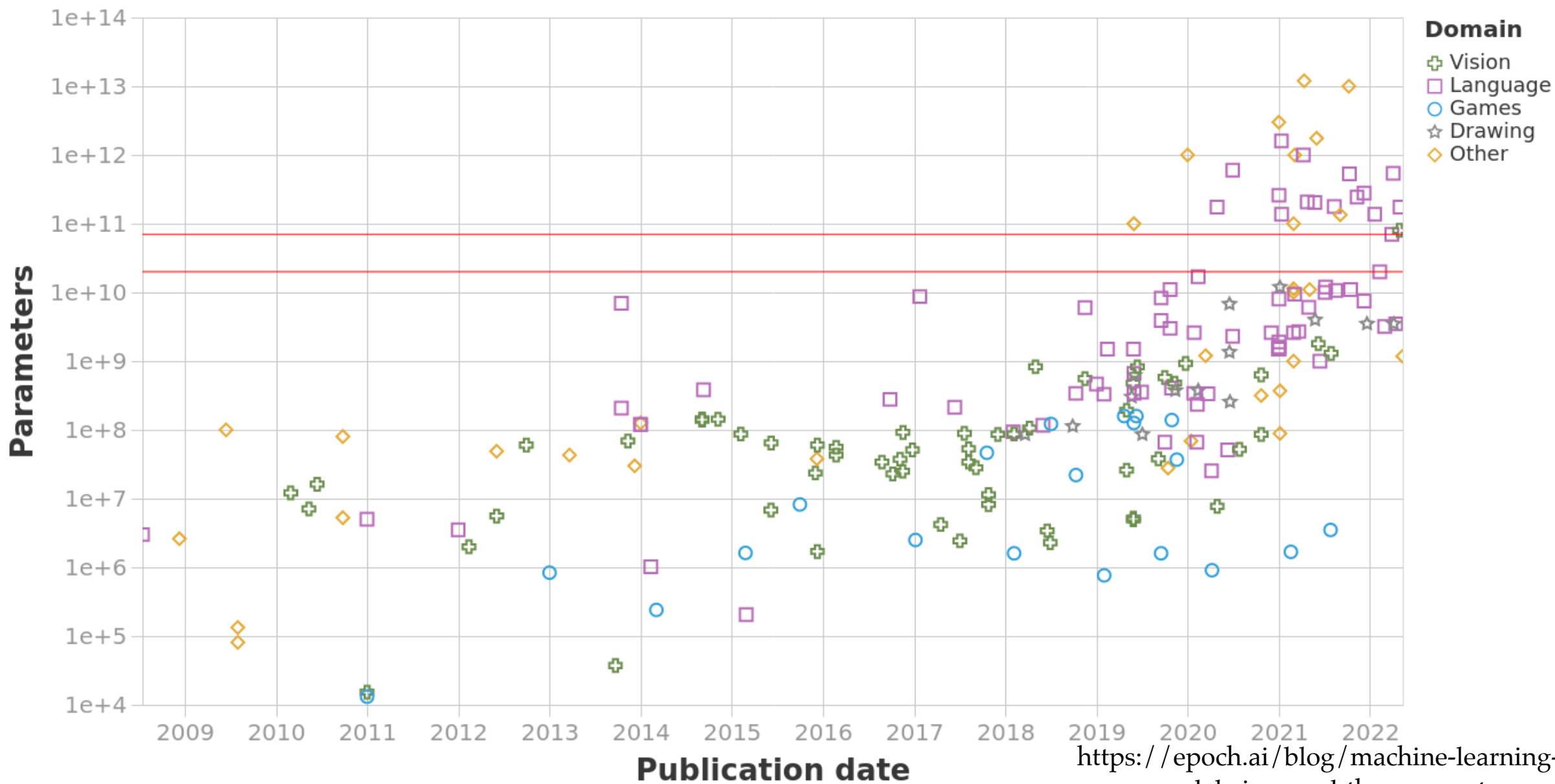
Typically, X is full column rank

- $J(\theta)$ "curves up" everywhere
- $\theta^* = (X^\top X)^{-1} X^\top Y$
- unique optimal hyperplane

θ^* can be costly to compute
(lab2 Q2.8)

Parameters of milestone Machine Learning systems over time

n = 203



<https://epoch.ai/blog/machine-learning-model-sizes-and-the-parameter-gap> ₃

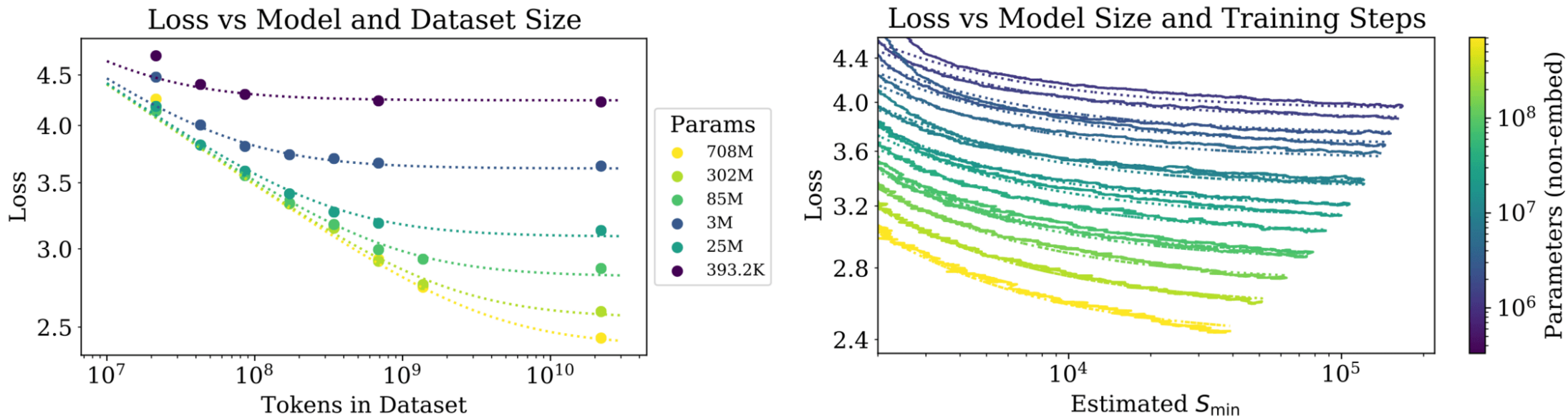


Figure 4 **Left:** The early-stopped test loss $L(N, D)$ varies predictably with the dataset size D and model size N according to Equation (1.5). **Right:** After an initial transient period, learning curves for all model sizes N can be fit with Equation (1.6), which is parameterized in terms of $S_{\min}$, the number of steps when training at large batch size (details in Section 5.1).

<https://losslandscape.com/gallery/>

In the real world,

- the number of parameters is huge
- the number of training data points is huge
- hypothesis class is typically highly nonlinear
- loss function is rarely as simple as squared error

Need a more **efficient** and **general** algorithm to train our ML system

⇒ gradient descent methods

Outline

- Gradient descent (GD)
 - The gradient vector
 - GD algorithm
 - Gradient descent properties
- Stochastic gradient descent (SGD)

For $f : \mathbb{R}^m \rightarrow \mathbb{R}$, its *gradient* $\nabla f : \mathbb{R}^m \rightarrow \mathbb{R}^m$ is defined at the point $p = (x_1, \dots, x_m)$ as:

$$\nabla f(p) = \begin{bmatrix} \frac{\partial f}{\partial x_1}(p) \\ \vdots \\ \frac{\partial f}{\partial x_m}(p) \end{bmatrix}$$

1. The gradient generalizes the concept of a derivative to multiple dimensions.
2. By construction, the gradient has the same dimension as the function's input.

The gradient may not always exist or be well-behaved.
Today, we have nice gradients unless otherwise specified.

3. The gradient can be symbolic or numerical.

e.g., $f(x, y, z) = x^2 + y^3 + z$

$$\nabla f(p) = \begin{bmatrix} \frac{\partial f}{\partial x_1}(p) \\ \vdots \\ \frac{\partial f}{\partial x_m}(p) \end{bmatrix}$$

its *symbolic* gradient:

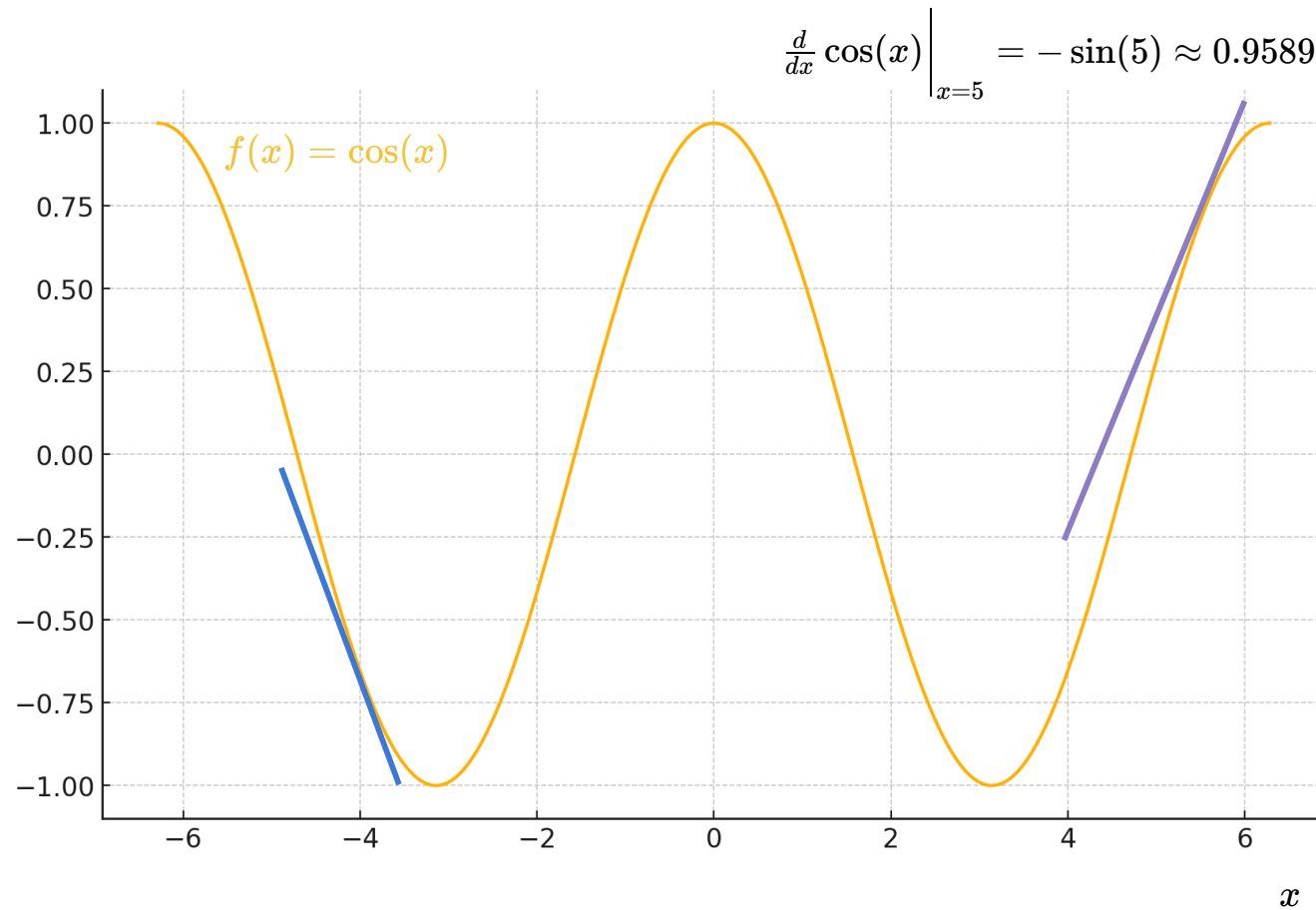
$$\nabla f(x, y, z) = \begin{bmatrix} 2x \\ 3y^2 \\ 1 \end{bmatrix}$$

evaluating the symbolic gradient at a point gives a *numerical* gradient:

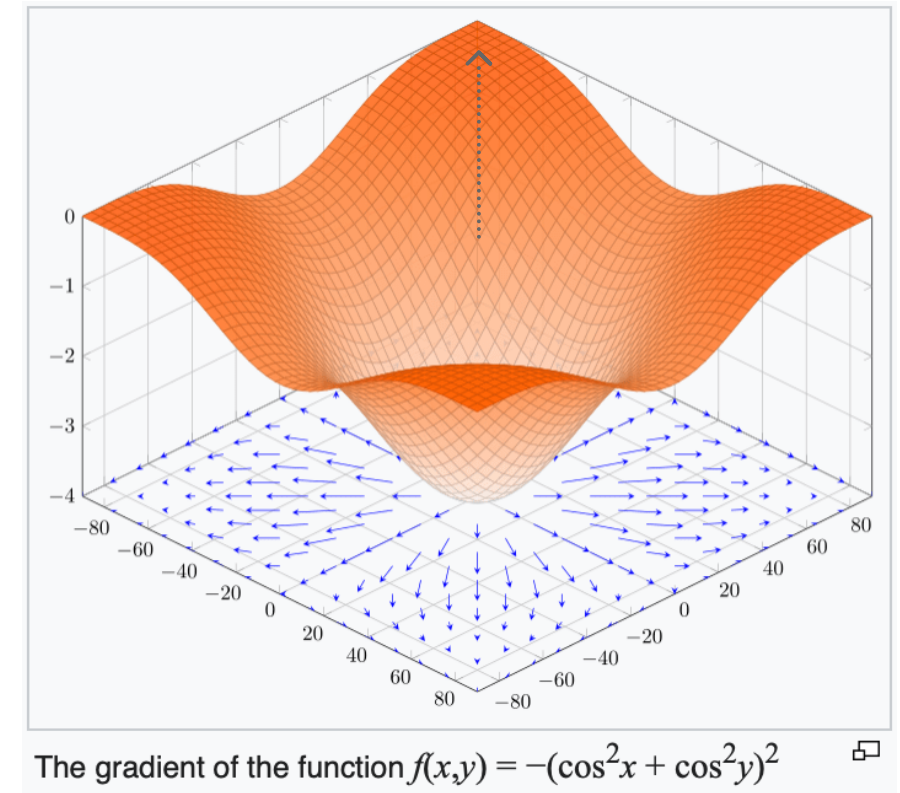
$$\nabla f(3, 2, 1) = \nabla f(x, y, z) \Big|_{(x,y,z)=(3,2,1)} = \begin{bmatrix} 6 \\ 12 \\ 1 \end{bmatrix}$$

just like a derivative can be a function or a number.

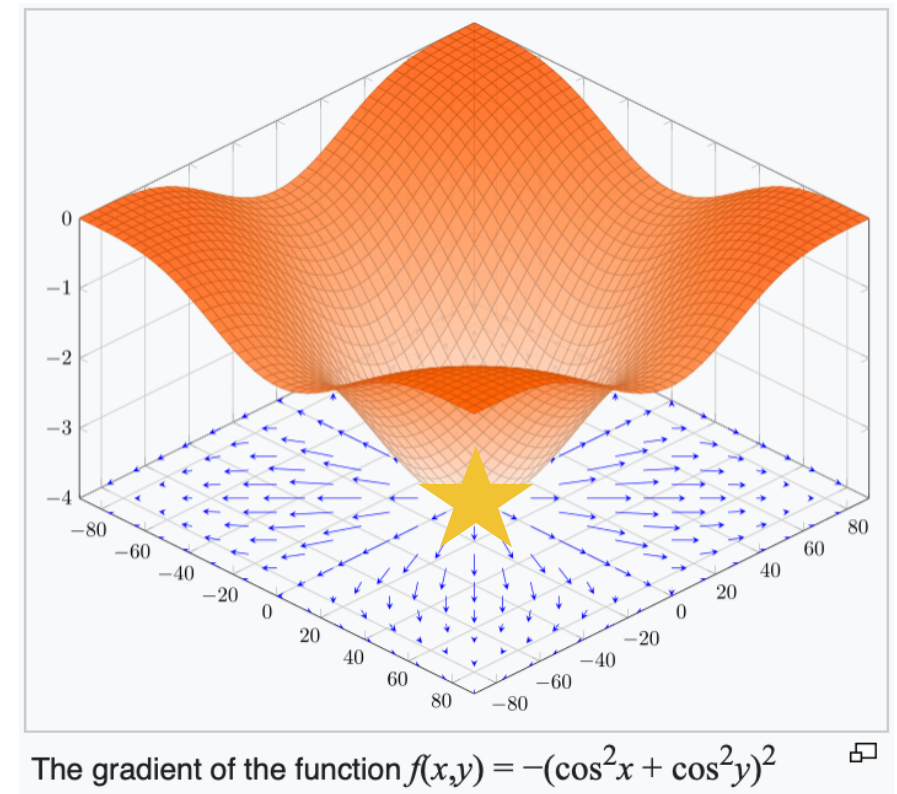
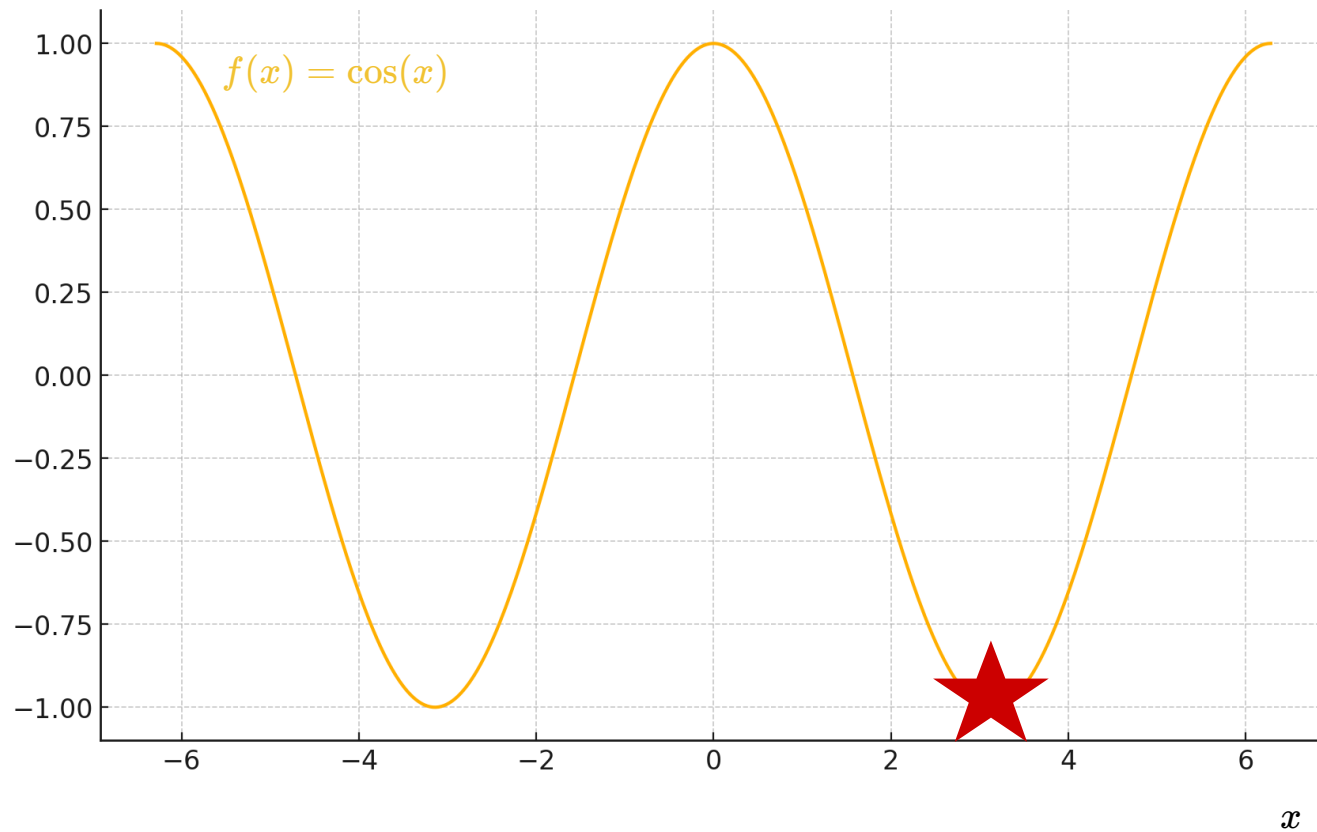
4. The gradient points in the direction of the (steepest) *increase* in the function value.



$$\left. \frac{d}{dx} \cos(x) \right|_{x=-4} = -\sin(-4) \approx -0.7568$$



5. For f defined on all of $\mathbb{R}^m$, the gradient at a minimizer is *necessarily* zero.



For $f : \mathbb{R}^m \rightarrow \mathbb{R}$, its *gradient* $\nabla f : \mathbb{R}^m \rightarrow \mathbb{R}^m$ is defined at the point $p = (x_1, \dots, x_m)$ as:

$$\nabla f(p) = \begin{bmatrix} \frac{\partial f}{\partial x_1}(p) \\ \vdots \\ \frac{\partial f}{\partial x_m}(p) \end{bmatrix}$$

1. The gradient generalizes the concept of a derivative to multiple dimensions.
2. By construction, the gradient has the same dimension as the function's input.
3. The gradient can be symbolic or numerical.
4. The gradient points in the direction of the (steepest) *increase* in the function value.
5. For f defined on all of $\mathbb{R}^m$, the gradient at a minimizer is *necessarily* zero.

The gradient may not always exist or be well-behaved.

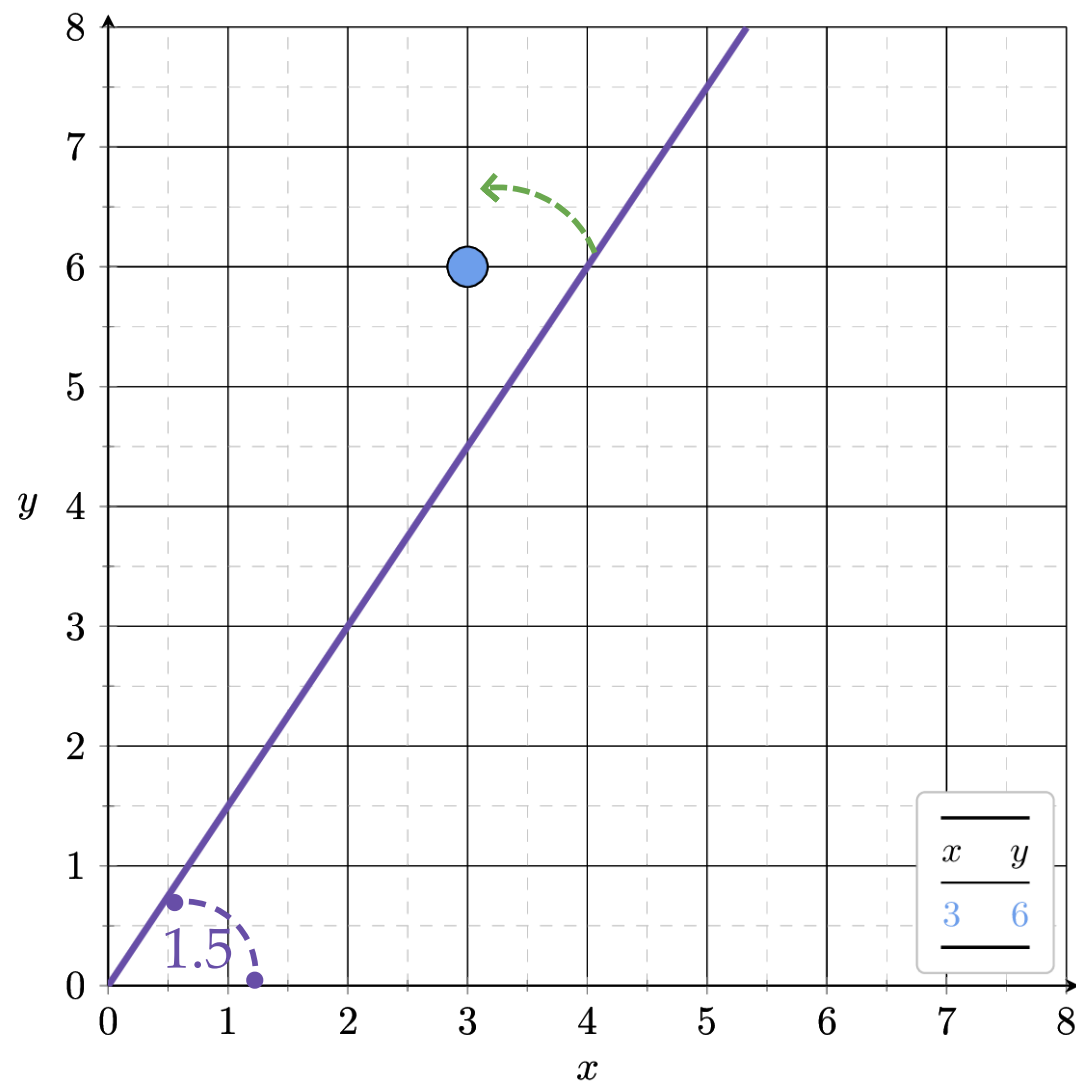
In our context, J plays the role of f , θ plays p .

Outline

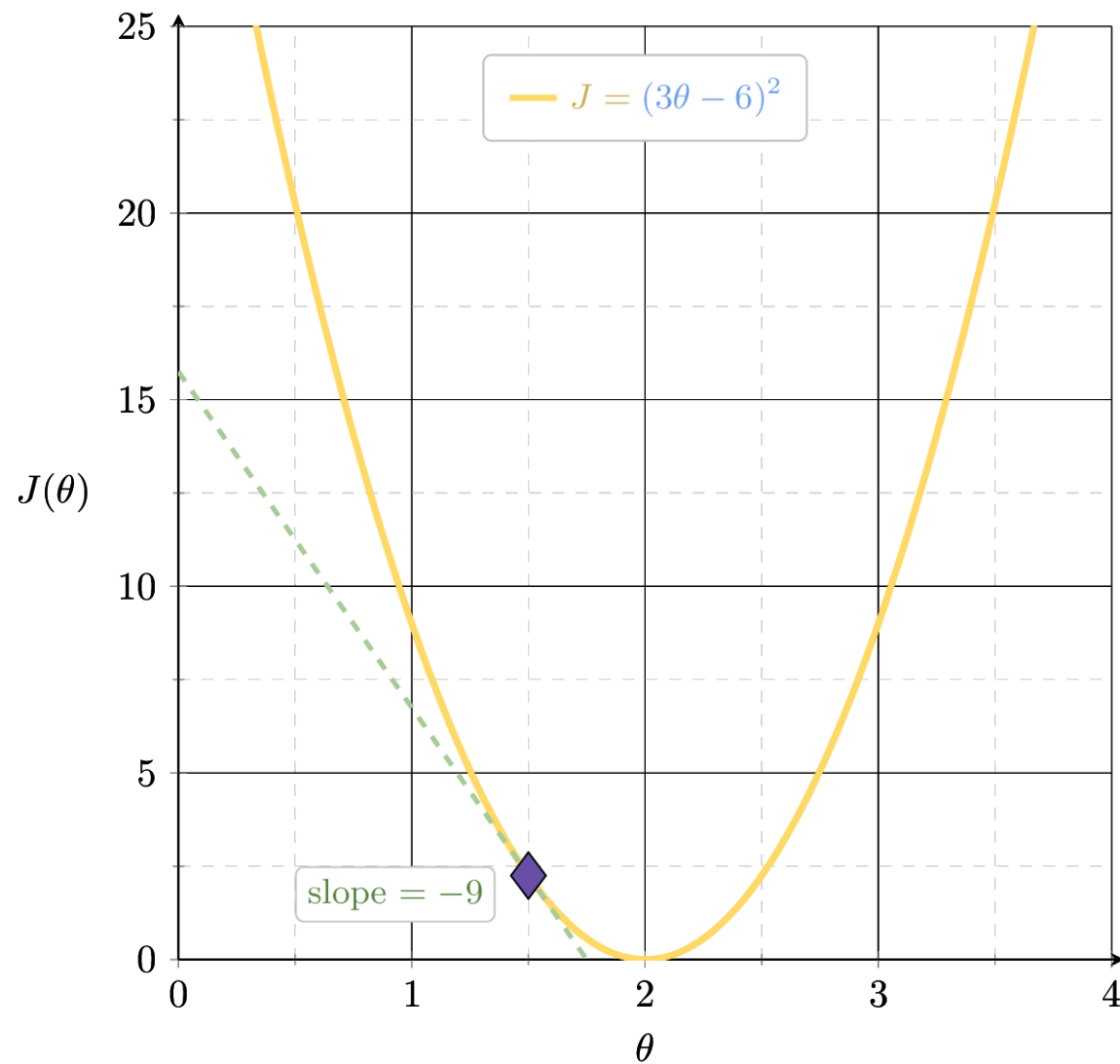
- Gradient descent (GD)
 - The gradient vector
 - GD algorithm
 - Gradient descent properties
- Stochastic gradient descent (SGD)

Example 1: fit a line (no offset) to minimize MSE

Suppose we fit $h = 1.5x$

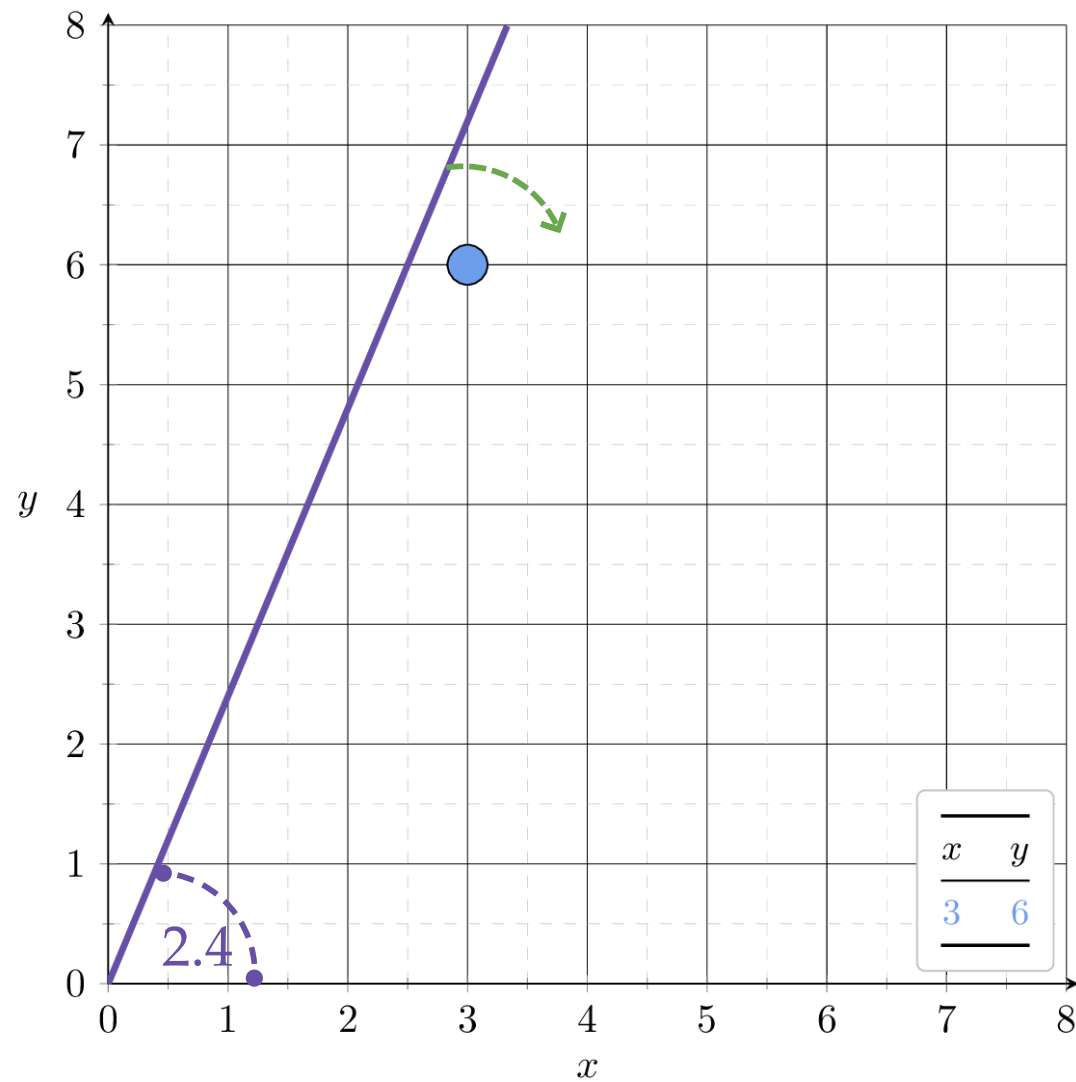


MSE could get better, by leveraging the gradient

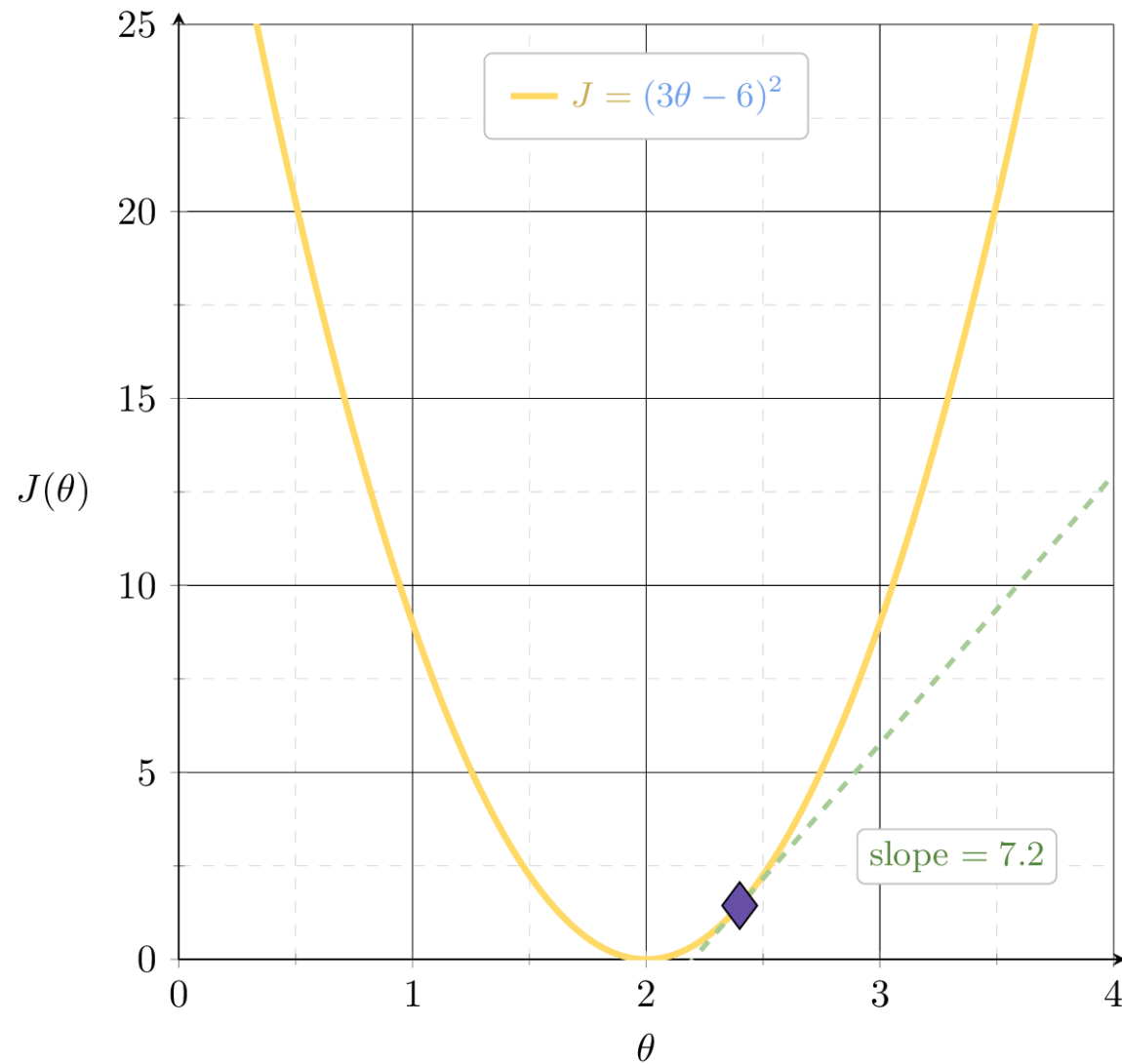


Example 1: fit a line (no offset) to minimize MSE

Suppose we fit $h = 2.4x$



MSE could get better, by leveraging the gradient



hyperparameters

initial guess of parameters

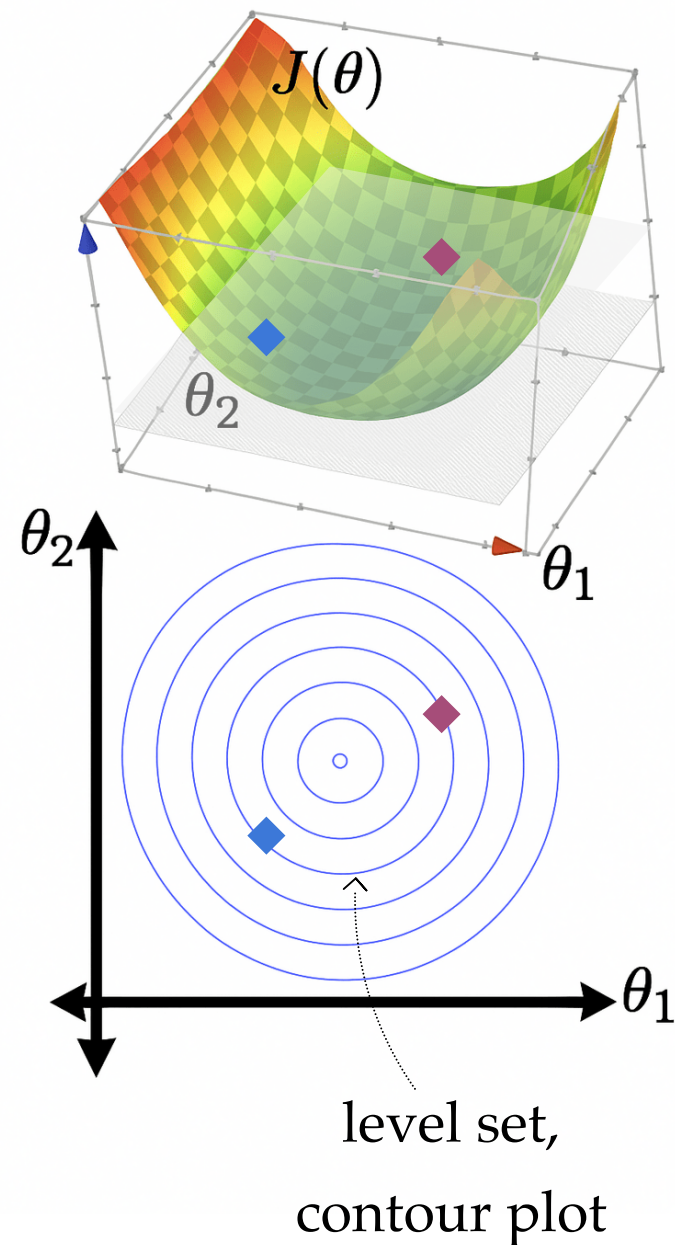
precision

Algorithm 1 Gradient Descent(θ_{init} , η , J , ϵ)

- 1: Initialize $\theta^{(0)} = \theta_{\text{init}}$
 - 2: Initialize $t = 0$
 - 3: **repeat**
 - 4: $t = t + 1$
 - 5: $\theta^{(t)} = \theta^{(t-1)} - \eta \nabla_{\theta} J(\theta^{(t-1)})$
 - 6: **until** $|J(\theta^{(t)}) - J(\theta^{(t-1)})| < \epsilon$
 - 7: **return** $\theta^{(t)}$
-

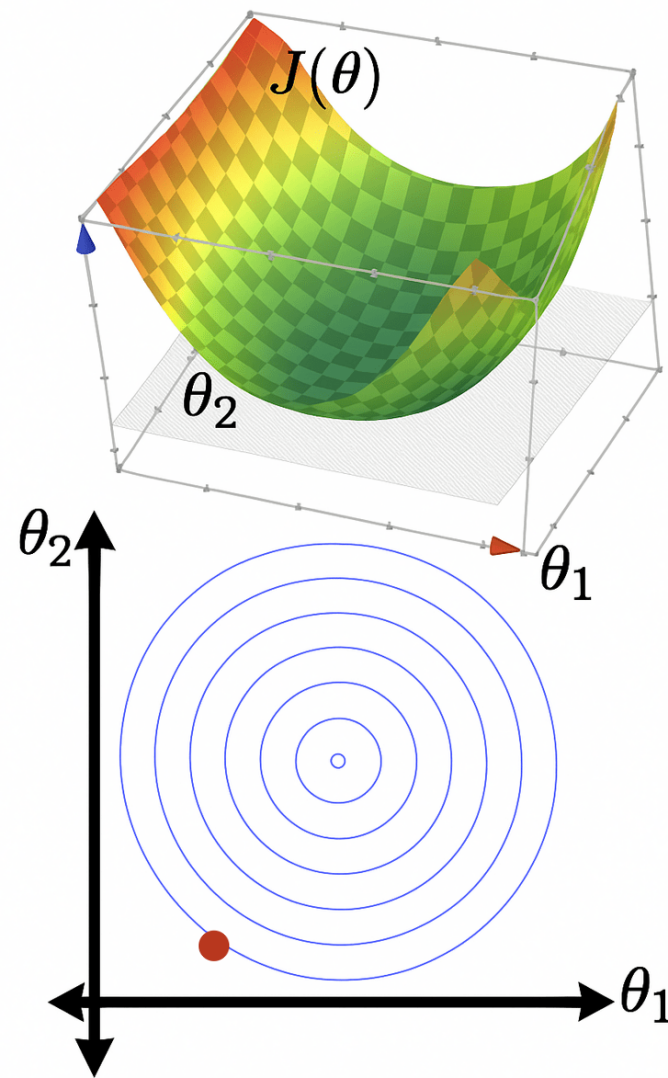
Algorithm 1 Gradient Descent($\theta_{\text{init}}, \eta, J, \epsilon$)

- 1: Initialize $\theta^{(0)} = \theta_{\text{init}}$
 - 2: Initialize $t = 0$
 - 3: **repeat**
 - 4: $t = t + 1$
 - 5: $\theta^{(t)} = \theta^{(t-1)} - \eta \nabla_{\theta} J(\theta^{(t-1)})$
 - 6: **until** $|J(\theta^{(t)}) - J(\theta^{(t-1)})| < \epsilon$
 - 7: **return** $\theta^{(t)}$
-



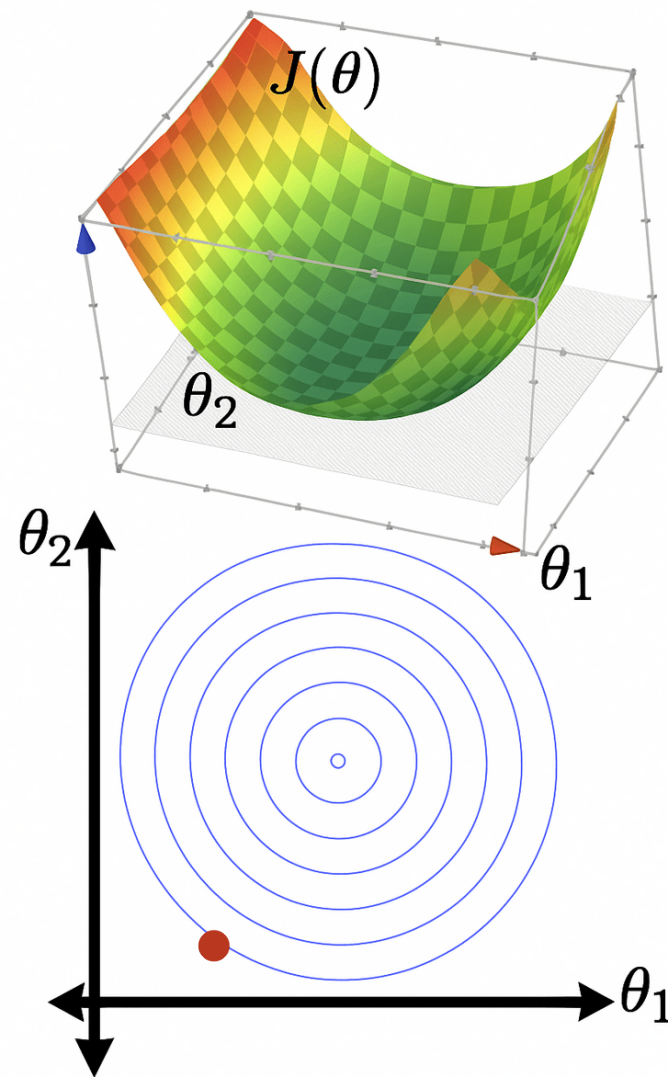
Algorithm 1 Gradient Descent($\theta_{\text{init}}, \eta, J, \epsilon$)

- 1: Initialize $\theta^{(0)} = \theta_{\text{init}}$
 - 2: Initialize $t = 0$
 - 3: **repeat**
 - 4: $t = t + 1$
 - 5: $\theta^{(t)} = \theta^{(t-1)} - \eta \nabla_{\theta} J(\theta^{(t-1)})$
 - 6: **until** $|J(\theta^{(t)}) - J(\theta^{(t-1)})| < \epsilon$
 - 7: **return** $\theta^{(t)}$
-



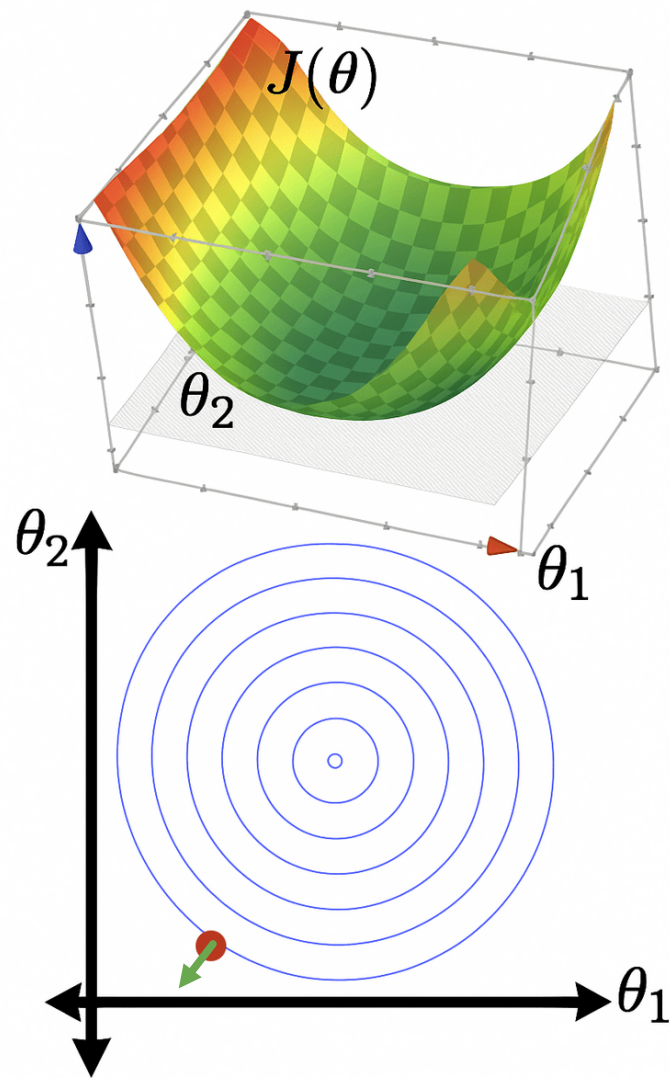
Algorithm 1 Gradient Descent($\theta_{\text{init}}, \eta, J, \epsilon$)

- 1: Initialize $\theta^{(0)} = \theta_{\text{init}}$
 - 2: Initialize $t = 0$ iteration counter
 - 3: **repeat**
 - 4: $t = t + 1$
 - 5: $\theta^{(t)} = \theta^{(t-1)} - \eta \nabla_{\theta} J(\theta^{(t-1)})$
 - 6: **until** $|J(\theta^{(t)}) - J(\theta^{(t-1)})| < \epsilon$
 - 7: **return** $\theta^{(t)}$
-



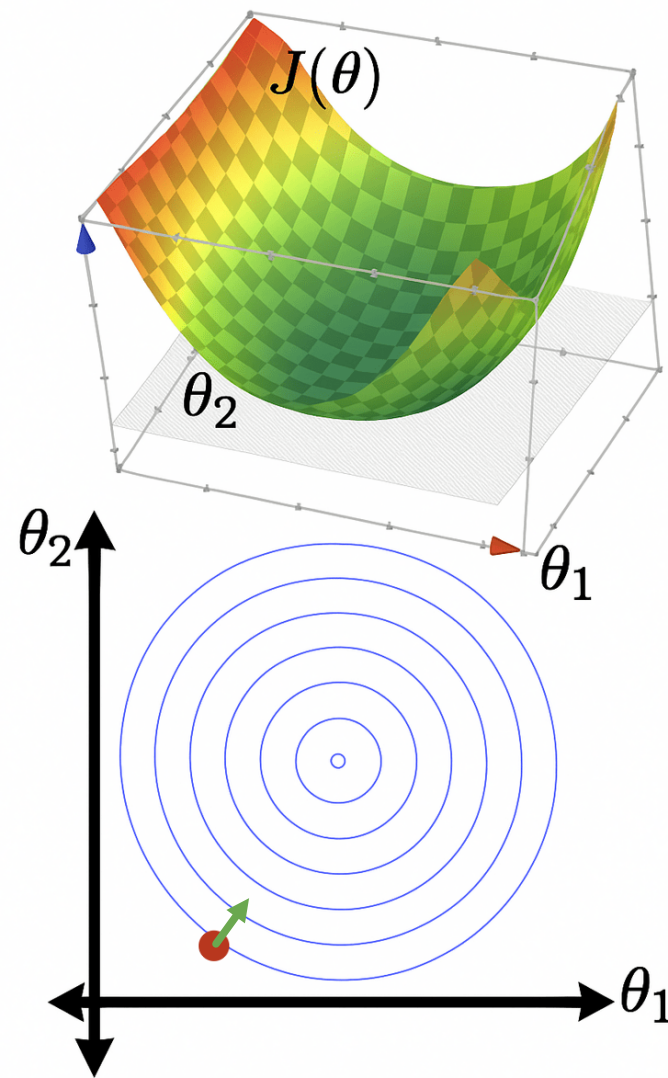
Algorithm 1 Gradient Descent($\theta_{\text{init}}, \eta, J, \epsilon$)

- 1: Initialize $\theta^{(0)} = \theta_{\text{init}}$
 - 2: Initialize $t = 0$
 - 3: **repeat**
 - 4: $t = t + 1$
 - 5: $\theta^{(t)} = \theta^{(t-1)} - \eta \nabla_{\theta} J(\theta^{(t-1)})$
 - 6: **until** $|J(\theta^{(t)}) - J(\theta^{(t-1)})| < \epsilon$
 - 7: **return** $\theta^{(t)}$
-



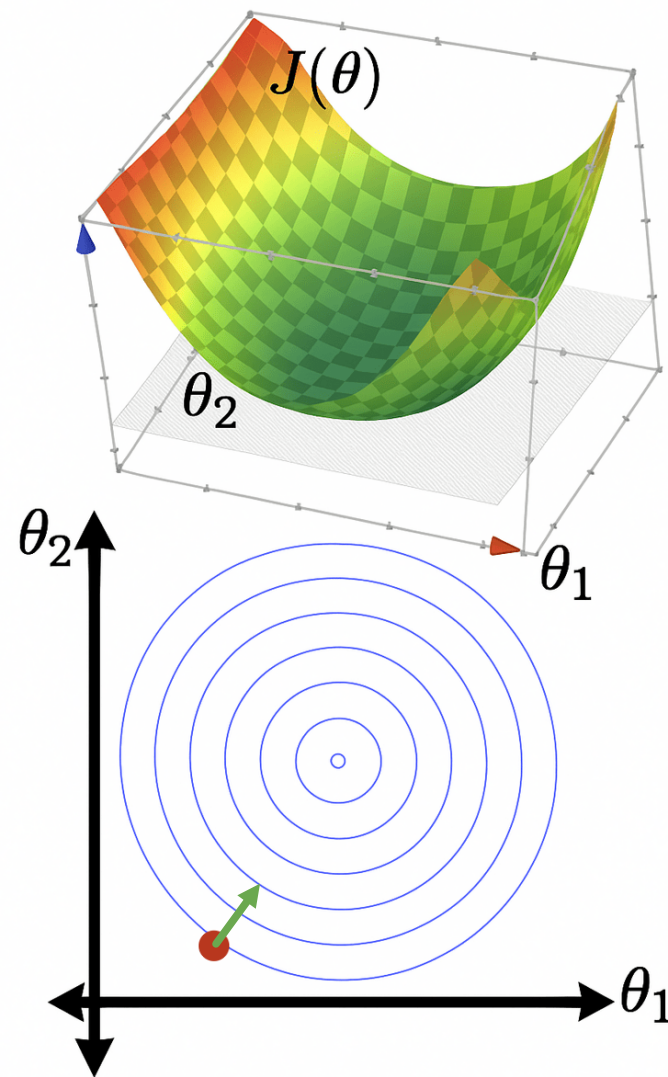
Algorithm 1 Gradient Descent($\theta_{\text{init}}, \eta, J, \epsilon$)

- 1: Initialize $\theta^{(0)} = \theta_{\text{init}}$
 - 2: Initialize $t = 0$
 - 3: **repeat**
 - 4: $t = t + 1$
 - 5: $\theta^{(t)} = \theta^{(t-1)} - \eta \nabla_{\theta} J(\theta^{(t-1)})$
 - 6: **until** $|J(\theta^{(t)}) - J(\theta^{(t-1)})| < \epsilon$
 - 7: **return** $\theta^{(t)}$
-



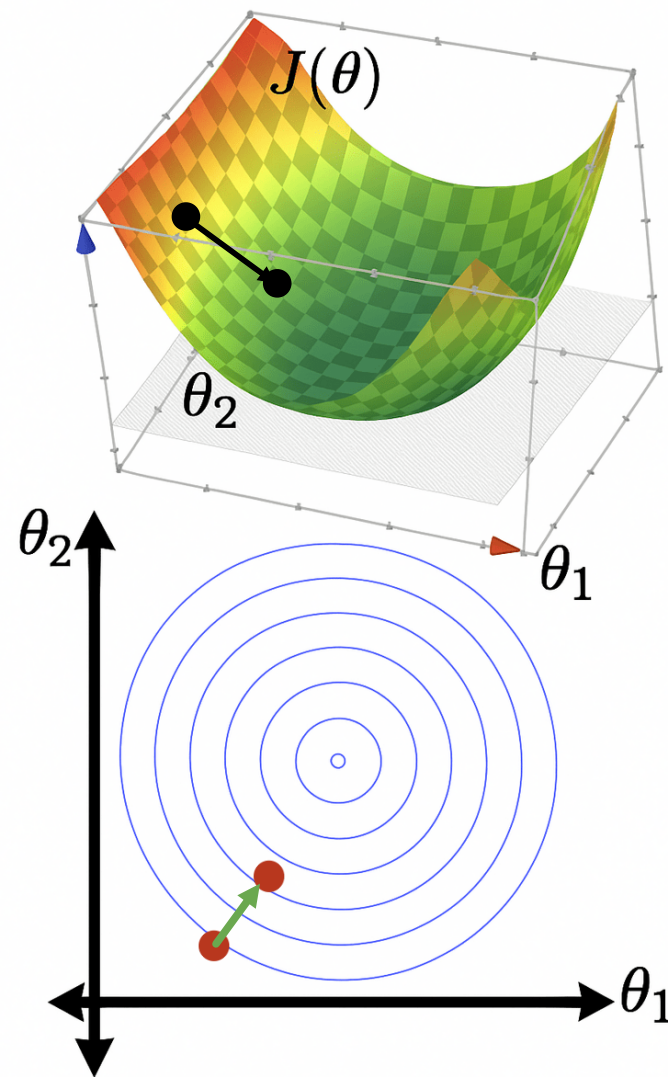
Algorithm 1 Gradient Descent($\theta_{\text{init}}, \eta, J, \epsilon$)

- 1: Initialize $\theta^{(0)} = \theta_{\text{init}}$
 - 2: Initialize $t = 0$
 - 3: **repeat**
 - 4: $t = t + 1$
 - 5: $\theta^{(t)} = \theta^{(t-1)} - \eta \nabla_{\theta} J(\theta^{(t-1)})$
 - 6: **until** $|J(\theta^{(t)}) - J(\theta^{(t-1)})| < \epsilon$
 - 7: **return** $\theta^{(t)}$
-



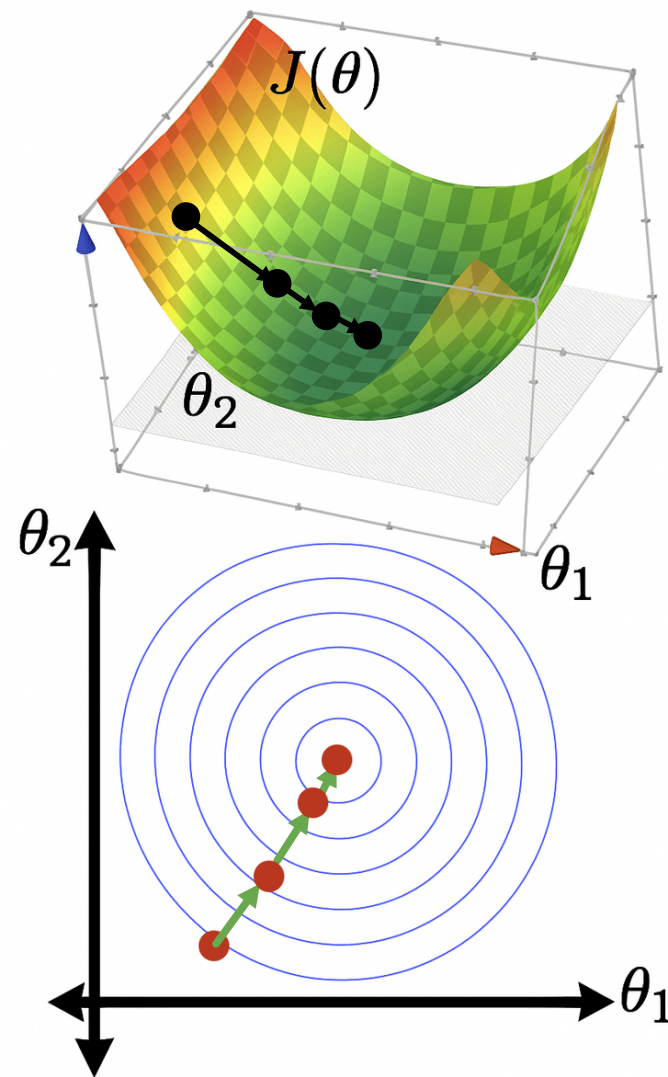
Algorithm 1 Gradient Descent($\theta_{\text{init}}, \eta, J, \epsilon$)

- 1: Initialize $\theta^{(0)} = \theta_{\text{init}}$
 - 2: Initialize $t = 0$
 - 3: **repeat**
 - 4: $t = t + 1$
 - 5: $\theta^{(t)} = \theta^{(t-1)} - \eta \nabla_{\theta} J(\theta^{(t-1)})$
 - 6: **until** $|J(\theta^{(t)}) - J(\theta^{(t-1)})| < \epsilon$
 - 7: **return** $\theta^{(t)}$
-



Algorithm 1 Gradient Descent($\theta_{\text{init}}, \eta, J, \epsilon$)

- 1: Initialize $\theta^{(0)} = \theta_{\text{init}}$
 - 2: Initialize $t = 0$
 - 3: **repeat**
 - 4: $t = t + 1$
 - 5: $\theta^{(t)} = \theta^{(t-1)} - \eta \nabla_{\theta} J(\theta^{(t-1)})$
 - 6: **until** $|J(\theta^{(t)}) - J(\theta^{(t-1)})| < \epsilon$
 - 7: **return** $\theta^{(t)}$
-



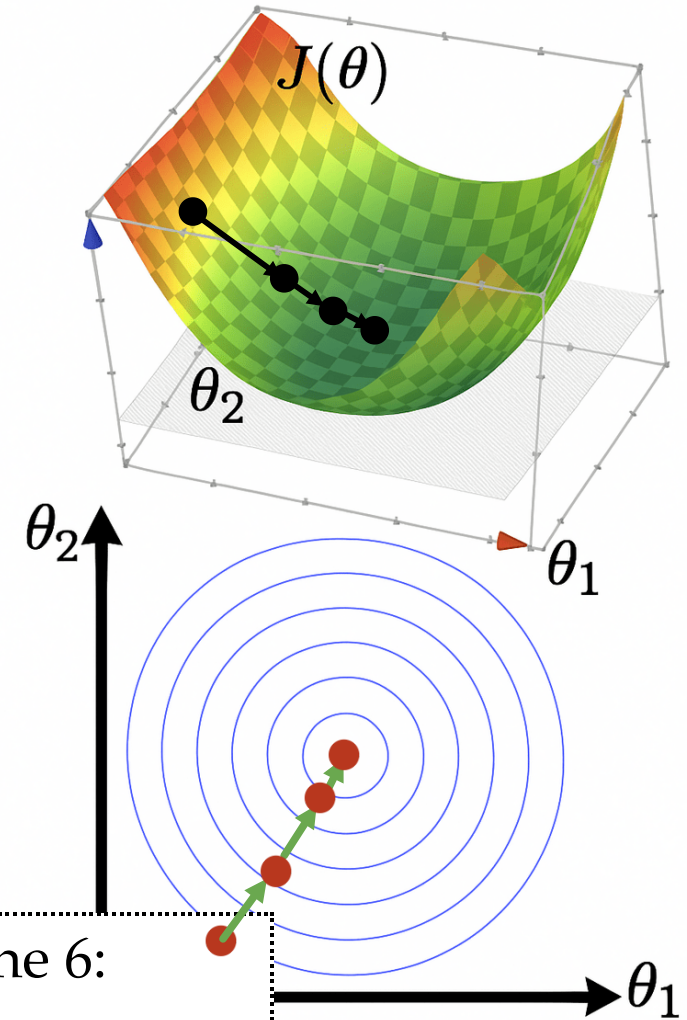
Algorithm 1 Gradient Descent($\theta_{\text{init}}, \eta, J, \epsilon$)

- 1: Initialize $\theta^{(0)} = \theta_{\text{init}}$
- 2: Initialize $t = 0$
- 3: **repeat**
- 4: $t = t + 1$
- 5: $\theta^{(t)} = \theta^{(t-1)} - \eta \nabla_{\theta} J(\theta^{(t-1)})$
- 6: **until** $|J(\theta^{(t)}) - J(\theta^{(t-1)})| < \epsilon$
- 7: **return** $\theta^{(t)}$

objective improvement
is nearly zero.

Other possible stopping criterion for line 6:

- Small parameter change: $\|\theta^{(t)} - \theta^{(t-1)}\| < \epsilon$, or
- Small gradient norm: $\|\nabla_{\theta} J(\theta^{(t-1)})\| < \epsilon$

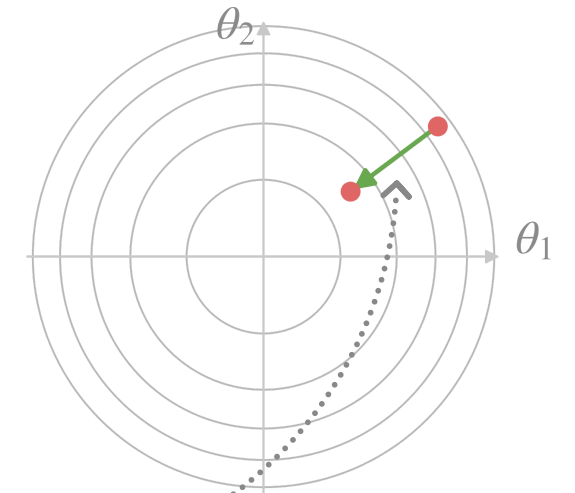
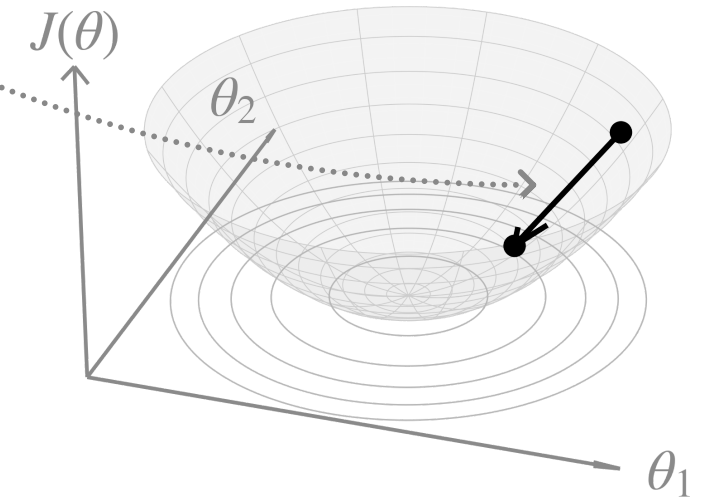


What does this 3d vector represent?
anything in the pseudocode?

Algorithm 1 Gradient Descent($\theta_{\text{init}}, \eta, J, \epsilon$)

- 1: Initialize $\theta^{(0)} = \theta_{\text{init}}$
 - 2: Initialize $t = 0$
 - 3: **repeat**
 - 4: $t = t + 1$
 - 5: $\theta^{(t)} = \theta^{(t-1)} - \eta \nabla_{\theta} J(\theta^{(t-1)})$
 - 6: **until** $|J(\theta^{(t)}) - J(\theta^{(t-1)})| < \epsilon$
 - 7: **return** $\theta^{(t)}$
-

What does this 2d vector represent?
anything in the pseudocode?

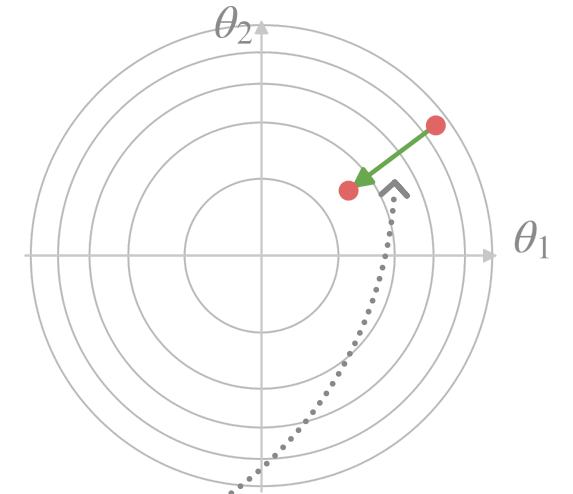
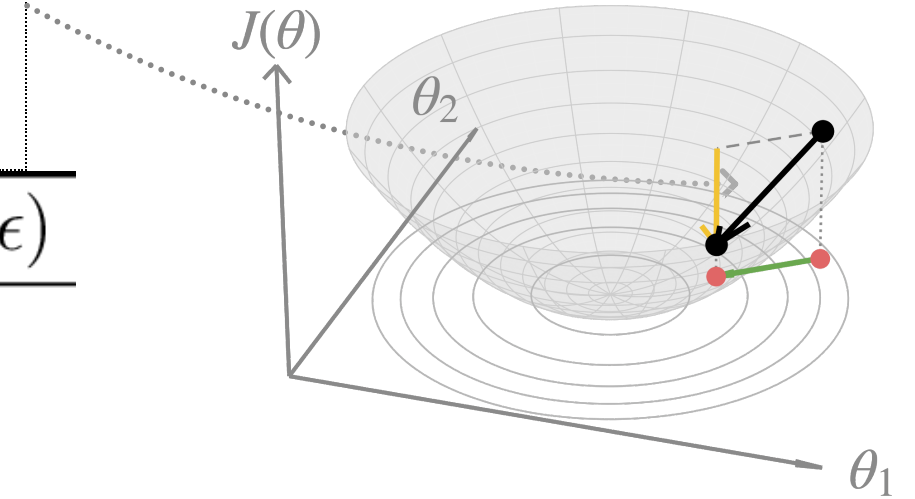


this 3d vector isn't directly in the pseudocode;

- its 2d horizontal projection is the step taken in line 5
- its 1d vertical projection is the J change in line 6

Algorithm 1 Gradient Descent($\theta_{\text{init}}, \eta, J, \epsilon$)

```
1: Initialize  $\theta^{(0)} = \theta_{\text{init}}$ 
2: Initialize  $t = 0$ 
3: repeat
4:    $t = t + 1$ 
5:    $\theta^{(t)} = \theta^{(t-1)} - \eta \nabla_{\theta} J(\theta^{(t-1)})$ 
6: until  $|J(\theta^{(t)}) - J(\theta^{(t-1)})| < \epsilon$ 
7: return  $\theta^{(t)}$ 
```



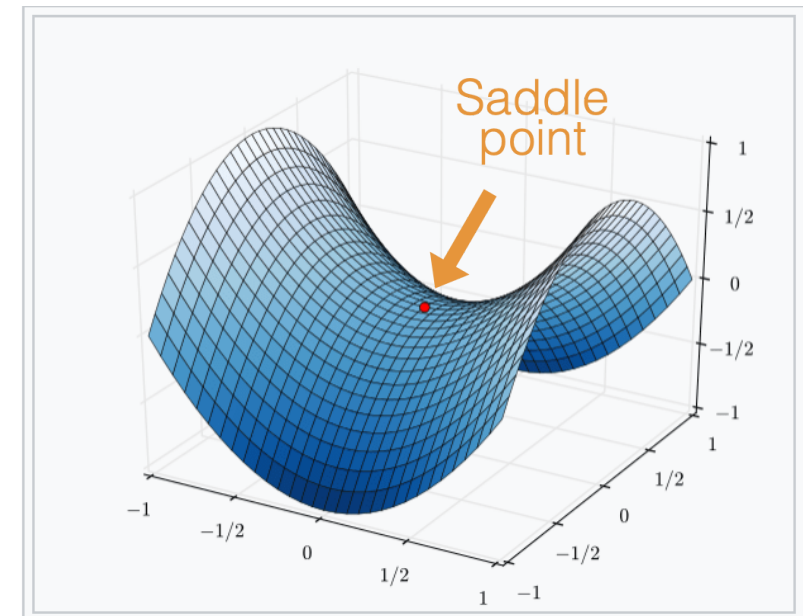
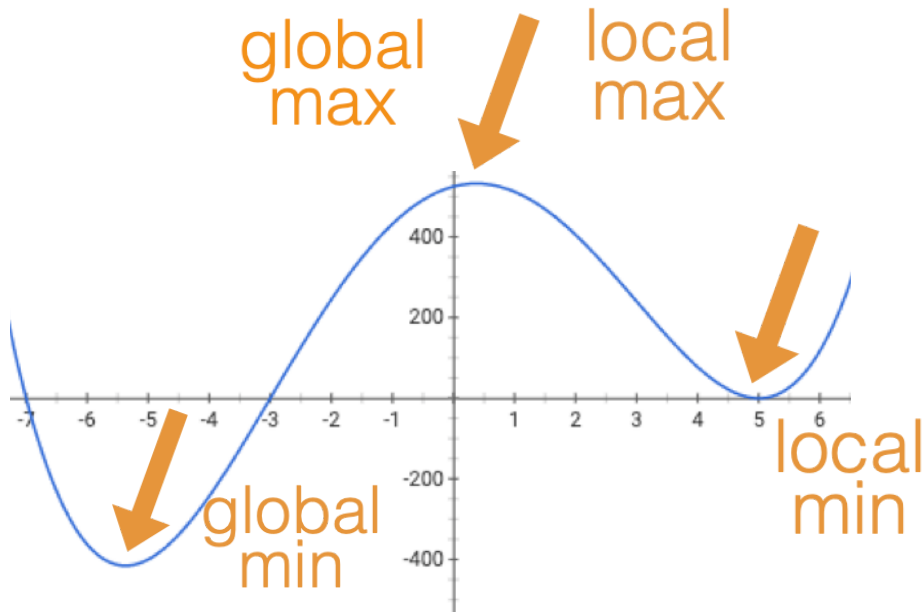
this 2d vector represents the step taken in line 5

Outline

- Gradient descent (GD)
 - The gradient vector
 - GD algorithm
 - Gradient descent properties
- Stochastic gradient descent (SGD)

When minimizing a function, we aim for a global minimizer.

At a global minimizer $\Rightarrow$ the gradient vector is zero $\Leftarrow$ gradient descent can achieve this (to arbitrary precision)



When minimizing a function, we aim for a global minimizer.

At a global minimizer $\Leftarrow \begin{cases} \text{the gradient vector is zero} \\ \text{the objective function is convex} \end{cases}$

A function f is *convex* if:

any line segment connecting two points of the graph of f lies above or on the graph.

- f is concave if $-f$ is convex.
- Convex functions are the largest well-understood class of functions where optimization theory guarantees convergence and efficiency

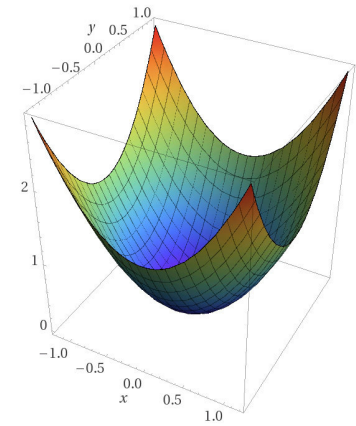
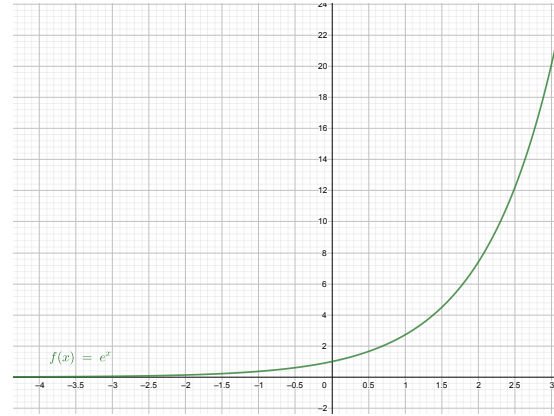
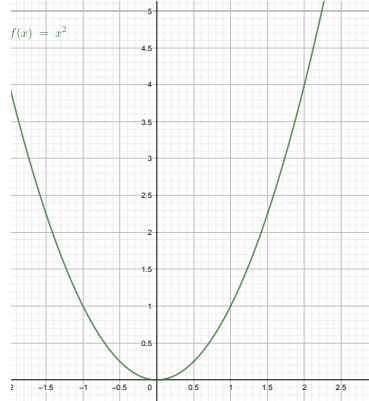
<https://shenshen.mit.edu/demos/convex-obj.html?embed>

- J_{MSE} is always convex
- J_{ridge} with $\lambda > 0$ is always (strongly) convex

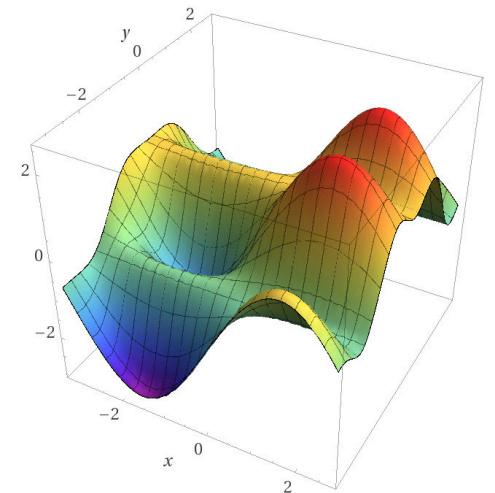
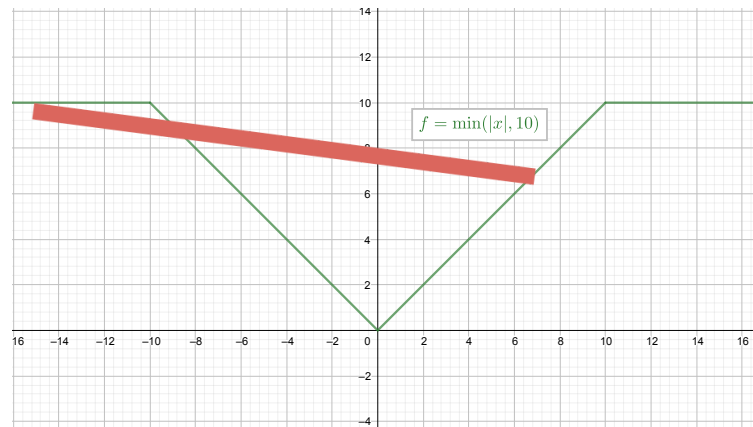
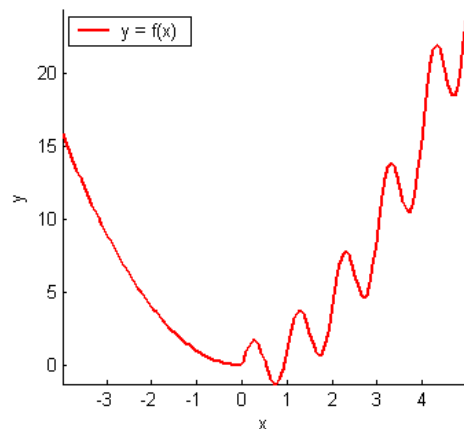
convexity is why we can claim the point
whose gradient is zero is a global minimizer.

More examples

Convex functions



Non-convex functions



Gradient Descent Performance

- Assumptions:
 - f is sufficiently "smooth"
 - f is convex
 - f has at least one global minimum
 - Run gradient descent for sufficient iterations
 - η is sufficiently small
- Conclusion:
 - Gradient descent converges arbitrarily close to a global minimizer of f .

Gradient Descent Performance

- Assumptions:

- f is sufficiently "smooth"
- f is convex
- f has at least one global minimum
- Run gradient descent for sufficient iterations

- η is sufficiently small

if violated, may not have gradient,
can't run gradient descent

- Conclusion:

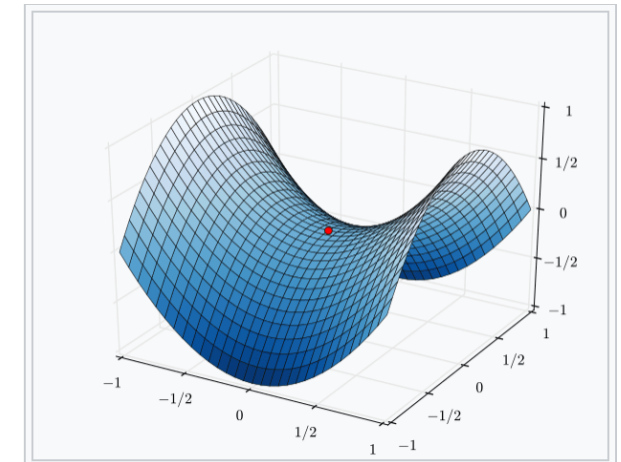
- Gradient descent converges arbitrarily close to a global minimizer of f .

Gradient Descent Performance

- Assumptions:

- f is sufficiently "smooth"
- f is convex
- f has at least one global minimum
- Run gradient descent for sufficient iterations
- η is sufficiently small

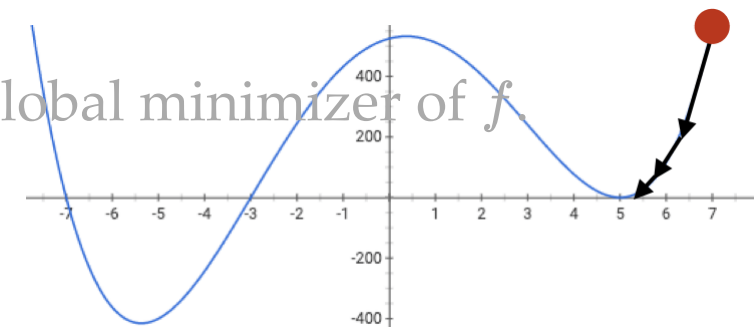
if violated, may get stuck at a saddle point



- Conclusion:

- Gradient descent converges arbitrarily close to a global minimizer of f .

or a local minimum



Gradient Descent Performance

- Assumptions:

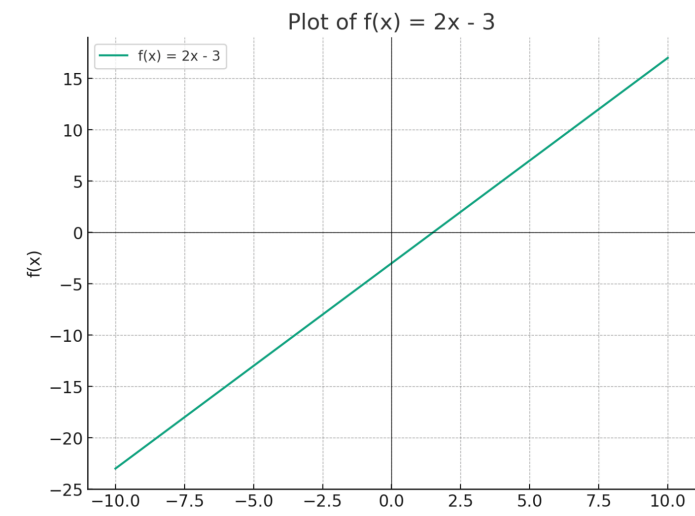
- f is sufficiently "smooth"
- f is convex
- f has at least one global minimum
- Run gradient descent for sufficient iterations
- η is sufficiently small

- Conclusion:

- Gradient descent converges arbitrarily close to a global minimizer of f .

if violated:

may not terminate / no
minimum to converge to



Gradient Descent Performance

- Assumptions:

- f is sufficiently "smooth"
- f is convex
- f has at least one global minimum
- Run gradient descent for sufficient iterations
- η is sufficiently small

if violated:

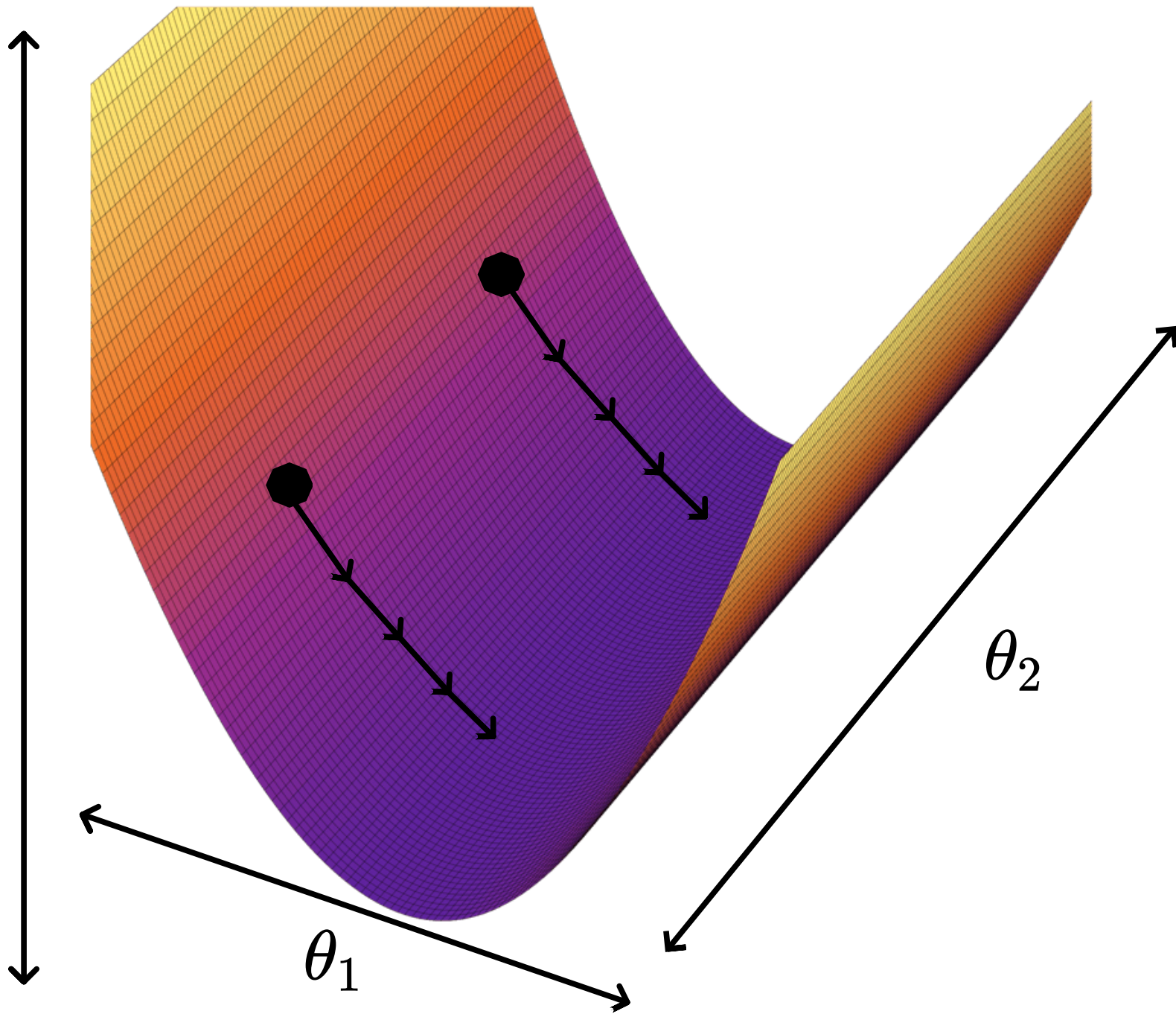
see demo on next slide,
also lab / hw

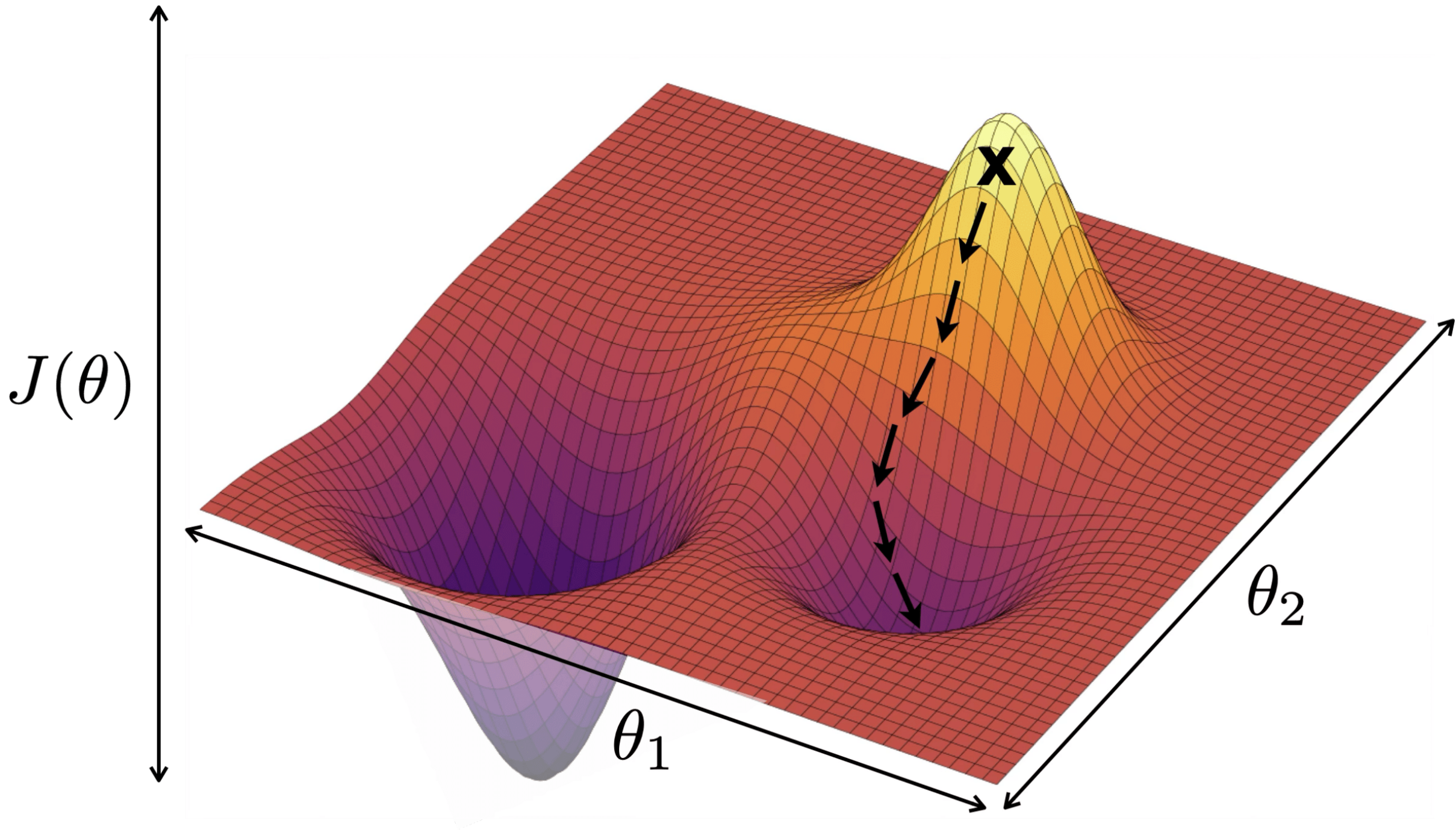
- Conclusion:

- Gradient descent converges arbitrarily close to a global minimizer of f .

<https://shenshen.mit.edu/demos/gd/gd3d.html>

$J(\theta)$

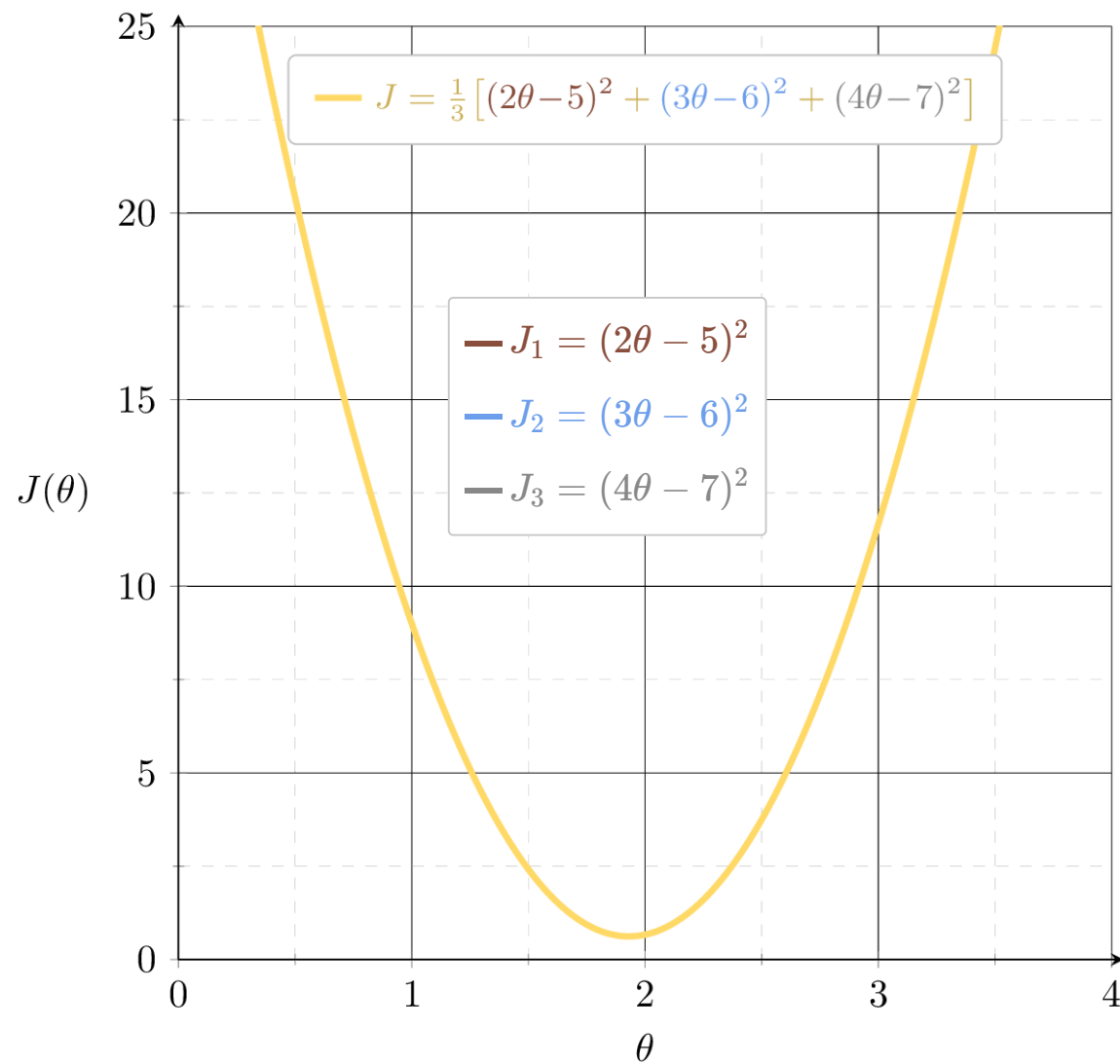
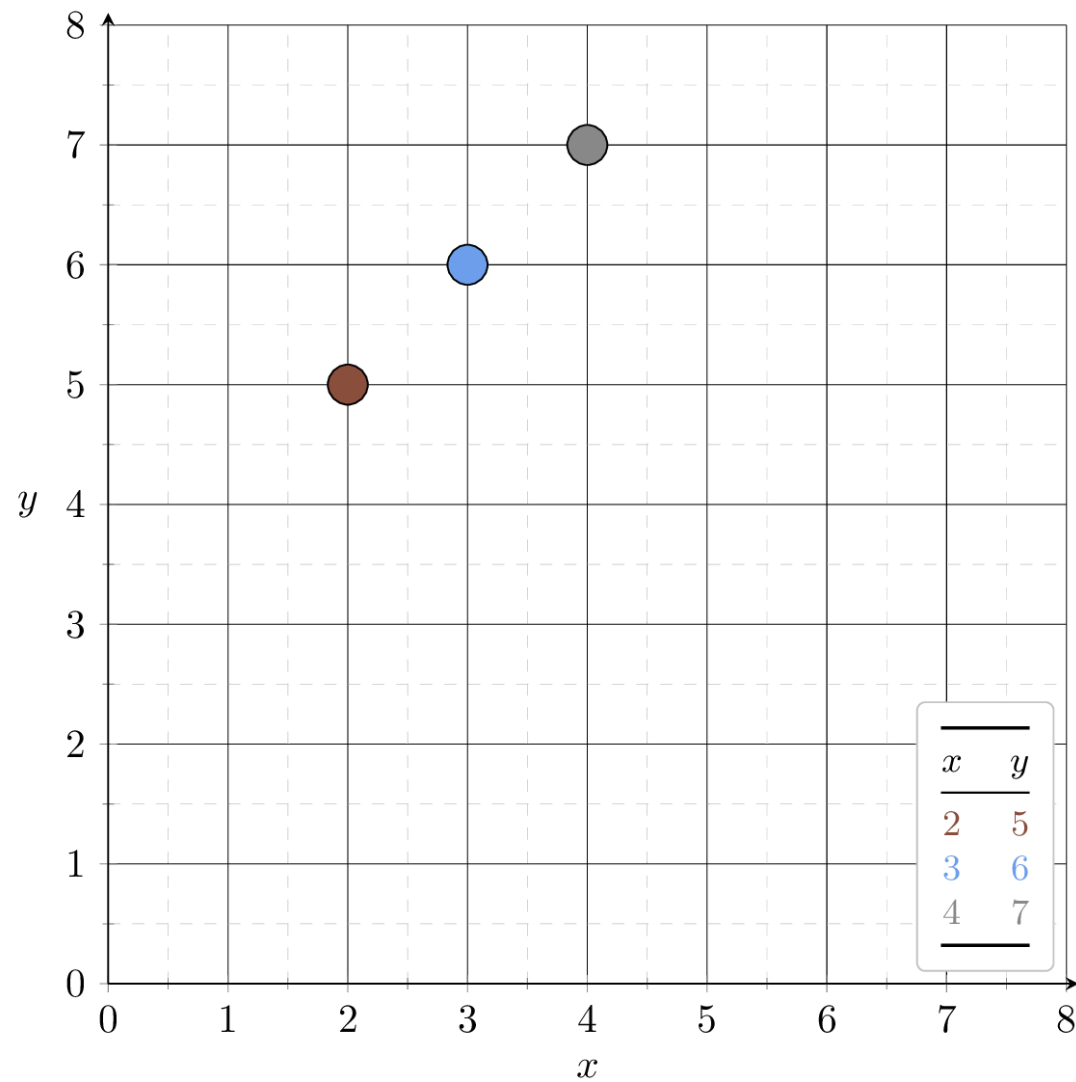




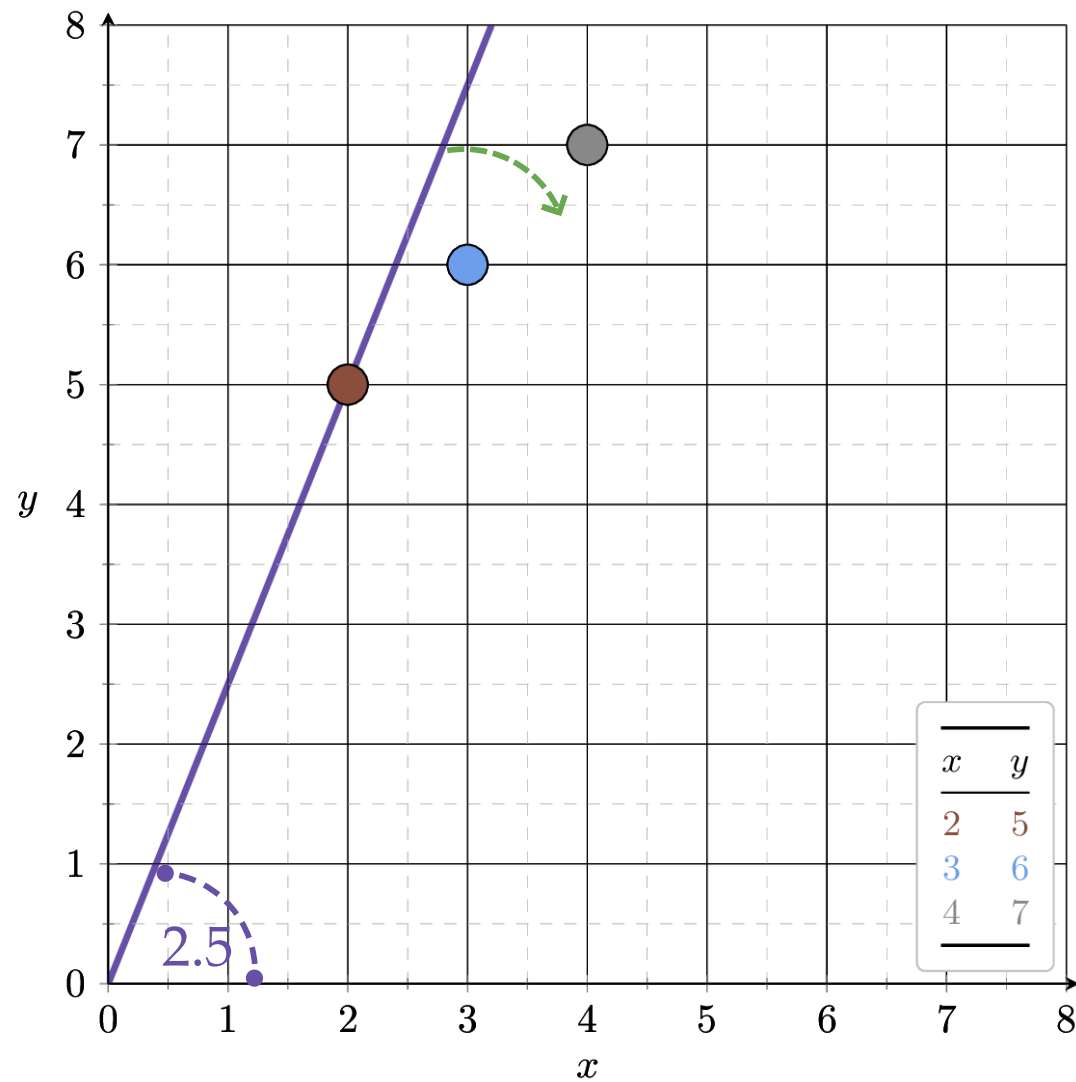
Outline

- Gradient descent (GD)
- Stochastic gradient descent (SGD)
 - SGD algorithm and setup
 - SGD vs. GD

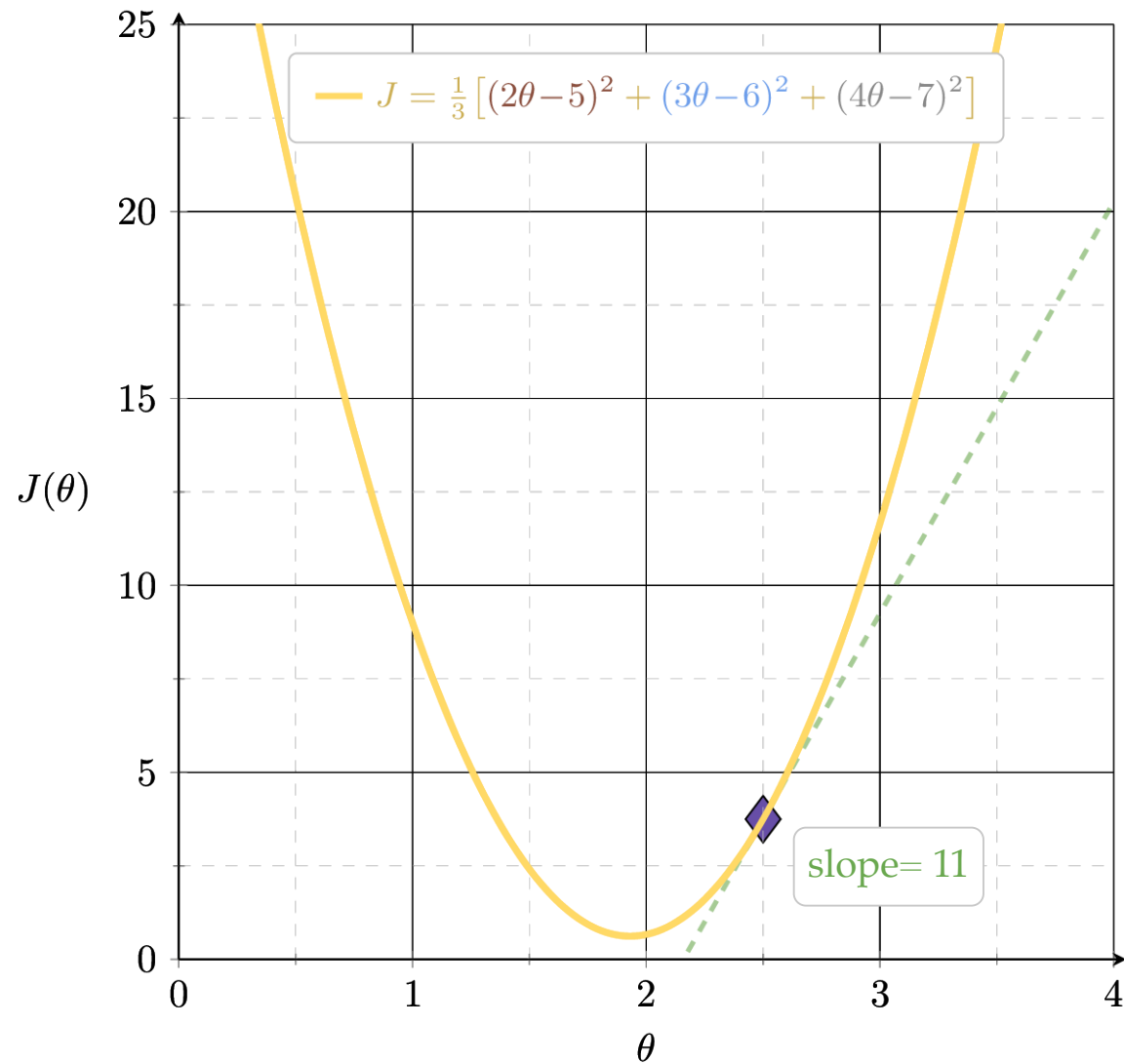
Example 2: fit a line (no offset) to minimize MSE (3 data points)



Suppose we fit $h = 2.5x$



MSE could get better, by leveraging the gradient

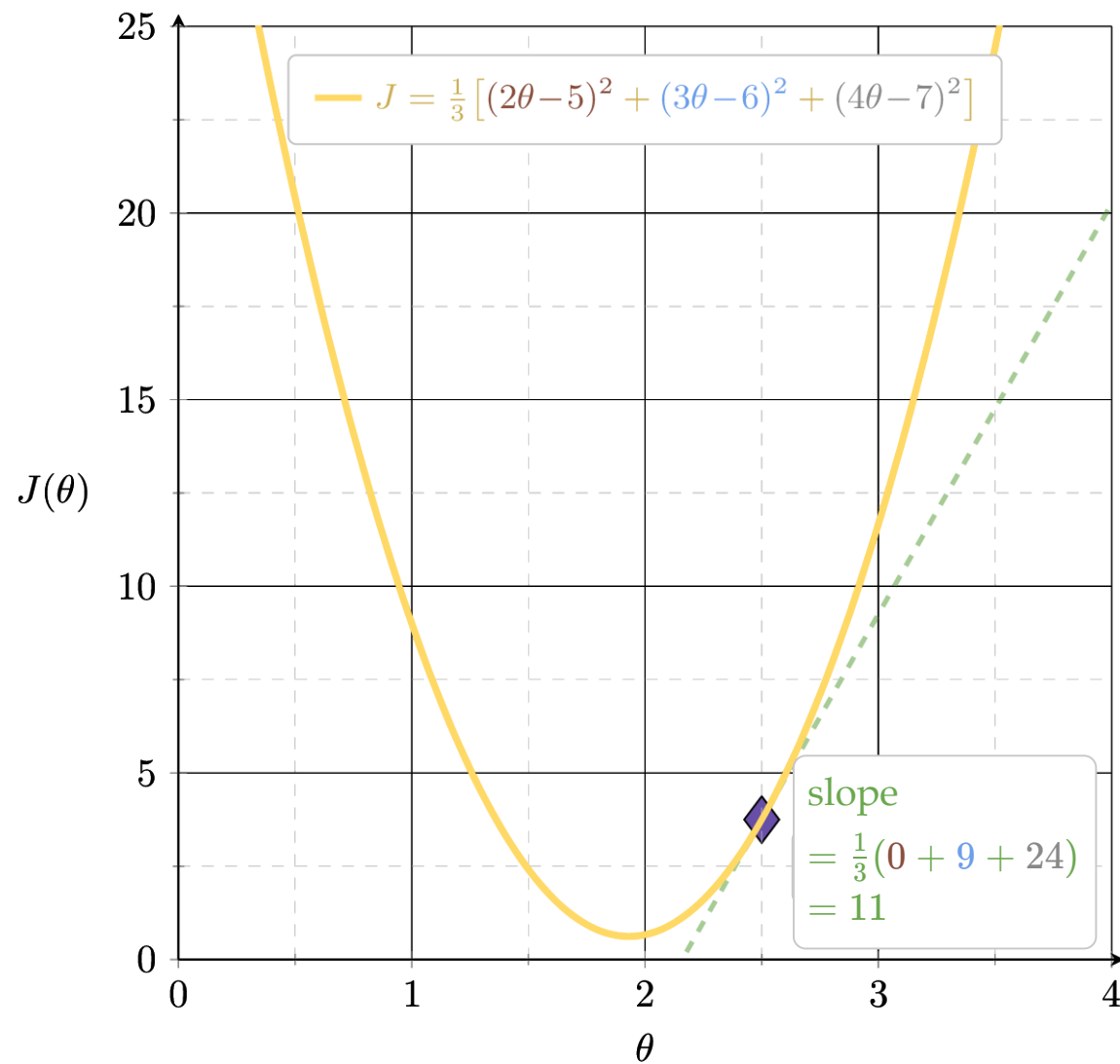
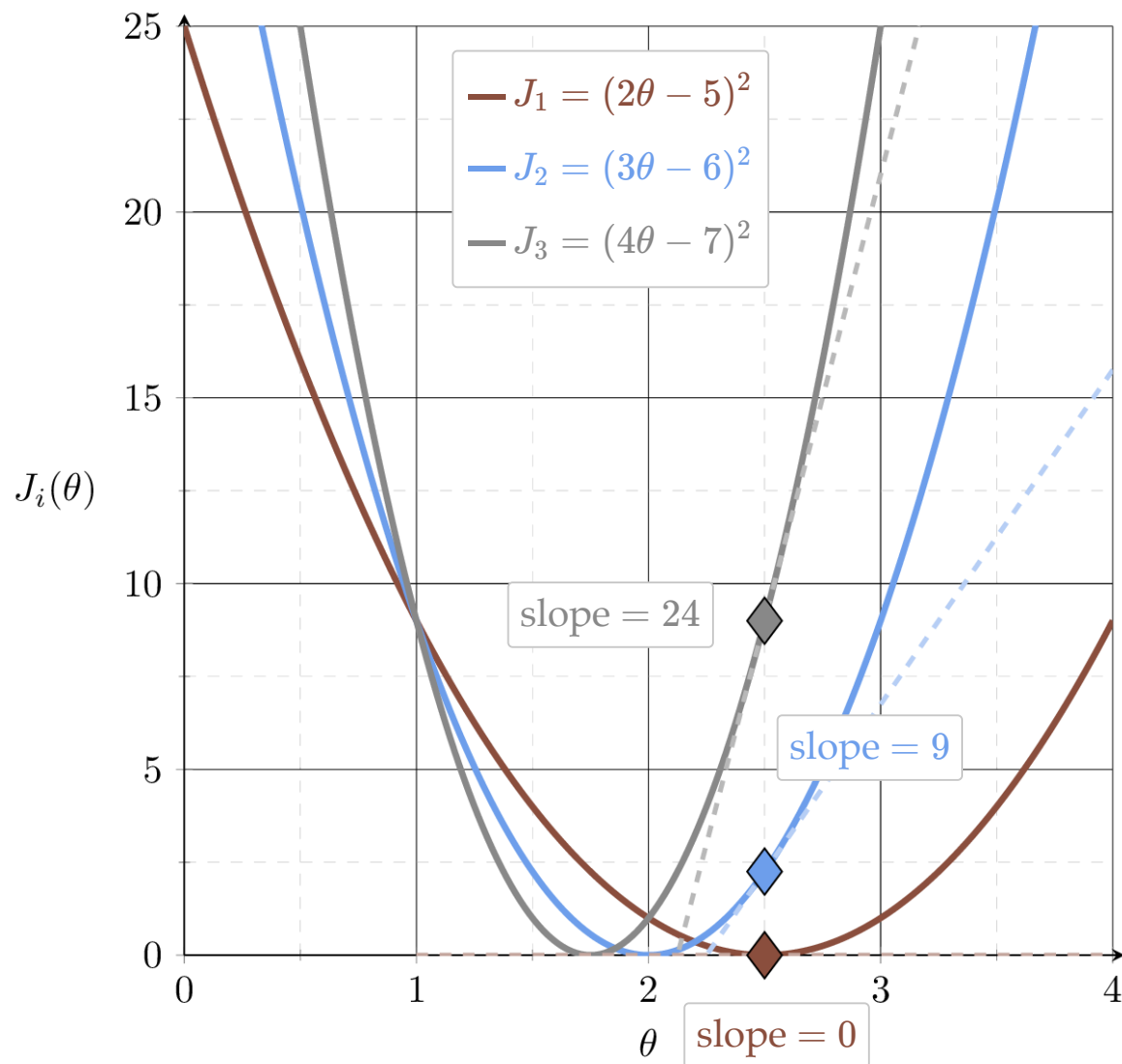


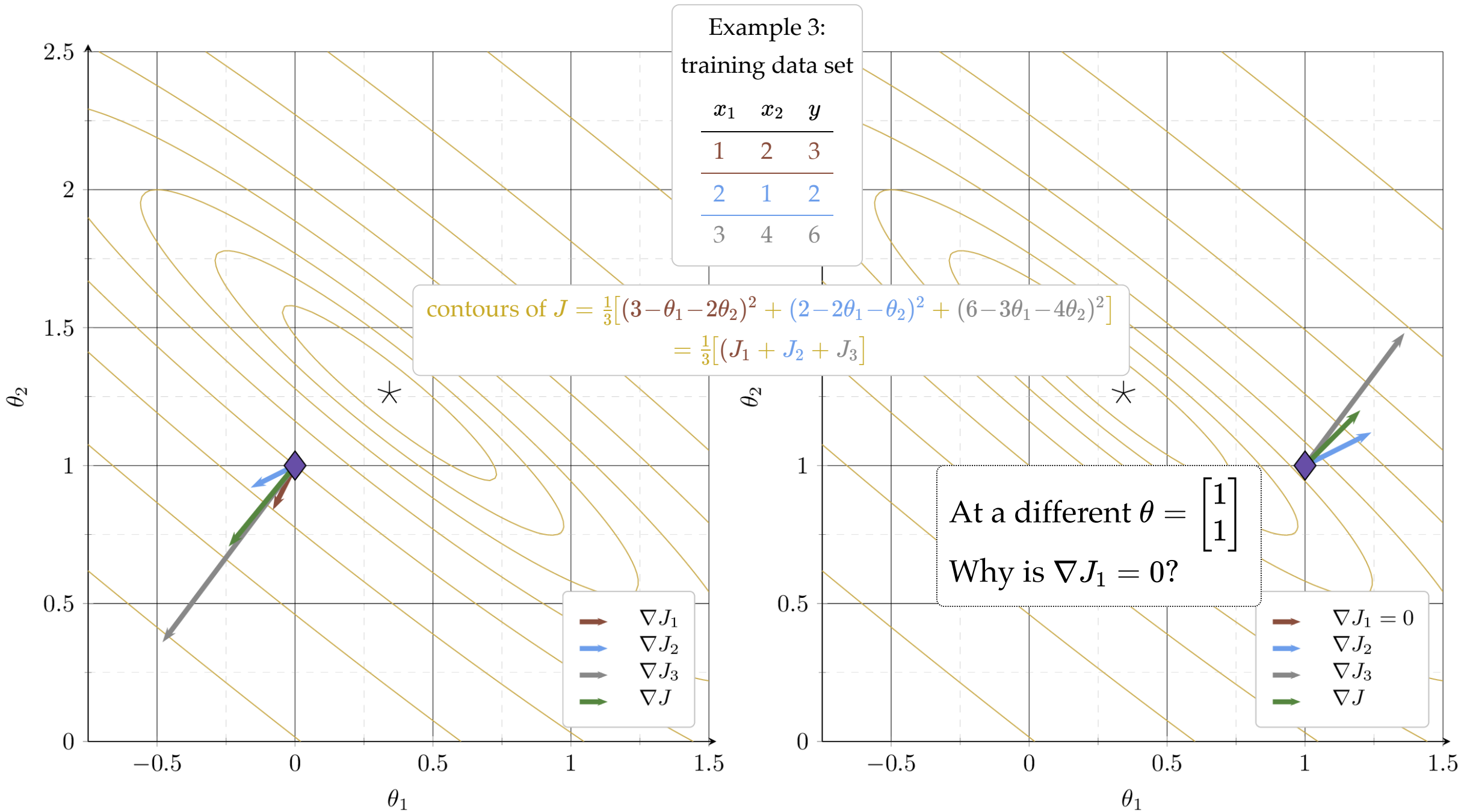
Besides

$$J = \frac{1}{3}[(J_1 + J_2 + J_3)]$$

we also have

$$\nabla_{\theta} J = \frac{1}{3} \left[\nabla_{\theta} J_1 + \nabla_{\theta} J_2 + \nabla_{\theta} J_3 \right]$$





Gradient of an ML objective

- the MSE of a linear hypothesis:

$$J(\theta) = \frac{1}{n} \sum_{i=1}^n \left(\theta^\top x^{(i)} - y^{(i)} \right)^2$$

- and its gradient w.r.t. θ :

$$\nabla_{\theta} J(\theta) = \frac{1}{n} \sum_{i=1}^n 2 \left(\theta^\top x^{(i)} - y^{(i)} \right) x^{(i)}$$

In general,

- An ML objective function is a finite sum

$$J(\theta) = \frac{1}{n} \sum_{i=1}^n J_i(\theta)$$

- and its gradient w.r.t. θ :

$$\nabla_{\theta} J(\theta) = \nabla \left(\frac{1}{n} \sum_{i=1}^n J_i(\theta) \right) = \frac{1}{n} \sum_{i=1}^n \nabla_{\theta} J_i(\theta)$$



👉 (gradient of the sum) = (sum of the gradient)

Gradient of an ML objective

In general,

- An ML objective function is a finite sum

$$J(\theta) = \frac{1}{n} \sum_{i=1}^n J_i(\theta)$$

loss incurred on a single i^{th} data point

- and its gradient w.r.t. θ :

$$\nabla_{\theta} J(\theta) = \nabla \left(\frac{1}{n} \sum_{i=1}^n J_i(\theta) \right) = \frac{1}{n} \sum_{i=1}^n \nabla_{\theta} J_i(\theta)$$

need to add n of these, each $\nabla_{\theta} J_i \in \mathbb{R}^d$

Costly in practice!

gradient info from the i^{th} data point's loss

Algorithm 1 Gradient Descent($\theta_{\text{init}}, \eta, J, \epsilon$)

```
1: Initialize  $\theta^{(0)} = \theta_{\text{init}}$ 
2: Initialize  $t = 0$ 
3: repeat
4:    $t = t + 1$ 
5:    $\theta^{(t)} = \theta^{(t-1)} - \eta \nabla_{\theta} J(\theta^{(t-1)})$ 
6: until  $|J(\theta^{(t)}) - J(\theta^{(t-1)})| < \epsilon$ 
7: return  $\theta^{(t)}$ 
```

Algorithm 2 Stochastic Gradient Descent($\theta_{\text{init}}, \eta, \{J_i\}_{i=1}^n, \epsilon$)

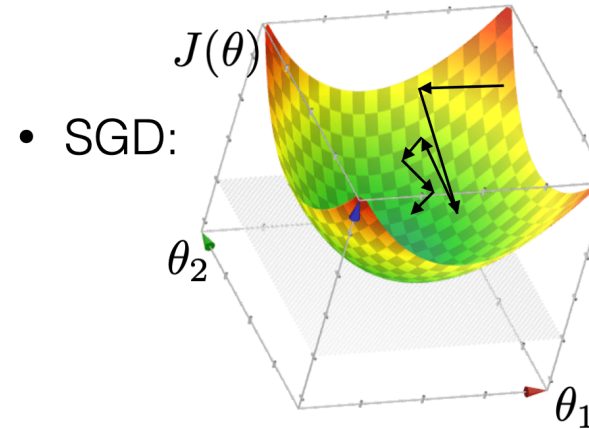
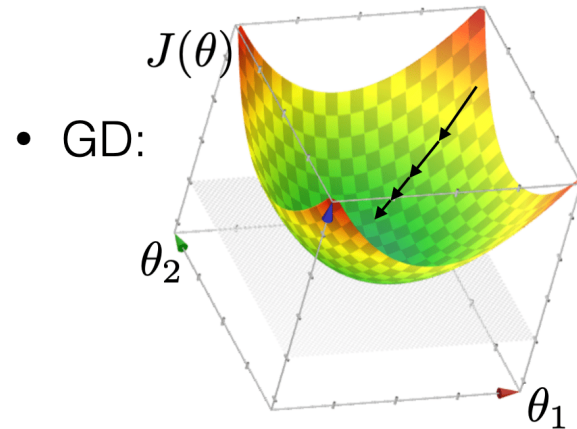
```
1: Initialize  $\theta^{(0)} = \theta_{\text{init}}$ 
2: Initialize  $t = 0$ 
3: repeat
4:    $t = t + 1$ 
5:    $i = \text{randomly sample from } \{1, \dots, n\}$ 
6:    $\theta^{(t)} = \theta^{(t-1)} - \eta \nabla_{\theta} J_i(\theta^{(t-1)})$ 
7: until  $|J(\theta^{(t)}) - J(\theta^{(t-1)})| < \epsilon$ 
8: return  $\theta^{(t)}$ 
```

$$\nabla_{\theta} J(\theta) = \frac{1}{n} \sum_{i=1}^n \nabla_{\theta} J_i(\theta) \approx \nabla_{\theta} J_i(\theta)$$

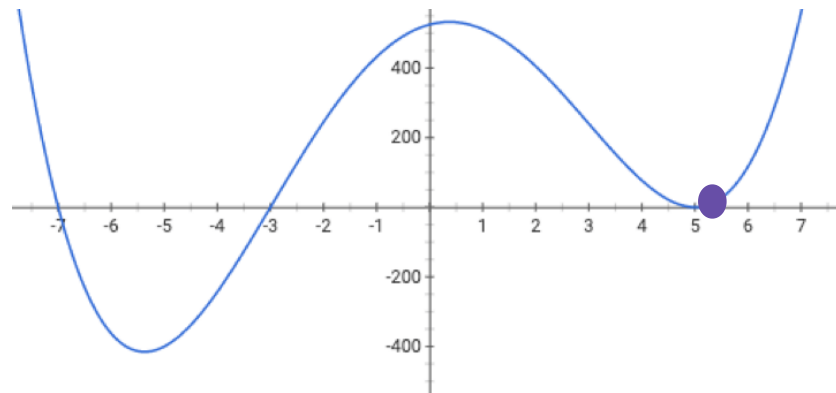
Compared with GD, SGD

$$\nabla_{\theta} J(\theta) = \frac{1}{n} \sum_{i=1}^n \nabla_{\theta} J_i(\theta) \approx \nabla_{\theta} J_i(\theta)$$

is cheaper per step



is much noisier



may get us out of a local minimum

Stochastic gradient descent performance

- Assumptions:
 - f is sufficiently "smooth"
 - f is convex
 - f has at least one global minimum
 - Run gradient descent for sufficient iterations
 - η is sufficiently small and satisfies additional "scheduling" condition¹
- Conclusion:
 - Stochastic gradient descent converges arbitrarily close to a global minimum of f with probability 1.

¹ $\sum_{t=1}^{\infty} \eta(t) = \infty$ and $\sum_{t=1}^{\infty} \eta(t)^2 < \infty$, e.g., $\eta(t) = 1/t$

See Robbins & Monro "A Stochastic Approximation Method" for more details

Algorithm 3 Mini-batch Gradient Descent($\theta_{\text{init}}, \eta, b, \{J_i\}_{i=1}^n, \epsilon$)

- 1: Initialize $\theta^{(0)} = \theta_{\text{init}}$
 - 2: Initialize $t = 0$
 - 3: **repeat**
 - 4: $t = t + 1$
 - 5: $B = \text{random mini-batch of size } b \text{ from } \{1, \dots, n\}$
 - 6: $\theta^{(t)} = \theta^{(t-1)} - \eta \nabla_{\theta} J_B(\theta^{(t-1)})$
 - 7: **until** $|J(\theta^{(t)}) - J(\theta^{(t-1)})| < \epsilon$
 - 8: **return** $\theta^{(t)}$
-

↖ batch size

SGD

$$\nabla_{\theta} J_i(\theta)$$

mini-batch GD

$$\frac{1}{b} \sum_{i=1}^b \nabla_{\theta} J_i(\theta)$$

GD

$$\frac{1}{n} \sum_{i=1}^n \nabla_{\theta} J_i(\theta) = \nabla_{\theta} J(\theta)$$

→
😊 more accurate gradient, stronger theoretical guarantee

😞 more costly per parameter update step

Summary

$$\theta^{(t)} = \theta^{(t-1)} - \eta \nabla_{\theta} J(\theta^{(t-1)})$$

- Most ML problems require optimization; closed-form solutions don't always exist or scale.
- Gradient descent iteratively updates θ in the direction of steepest descent of J .
- With a convex J and small enough η , GD converges to a global minimum.
- SGD follows one data point's gradient, so each step is noisy but cheap.
- Mini-batch averages a small batch of data points' gradients; it achieves the practical middle ground and is the default in modern practice.