

Lecture 29

Understanding

Gradient Descent.

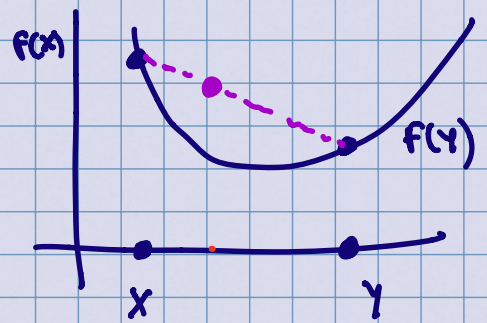
* MIDTERM ALREADY GRADED

* DROP DATE IS 11/18 (WEDNESDAY)

CONVEX FUNCTIONS

A function f is convex if

$$f(\lambda x + (1-\lambda)y) \leq \lambda f(x) + (1-\lambda)f(y)$$

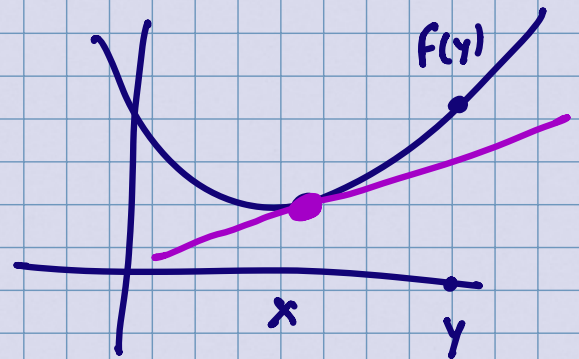


"Function below the chord"

Characterizations (all equivalent, under suitable conditions)

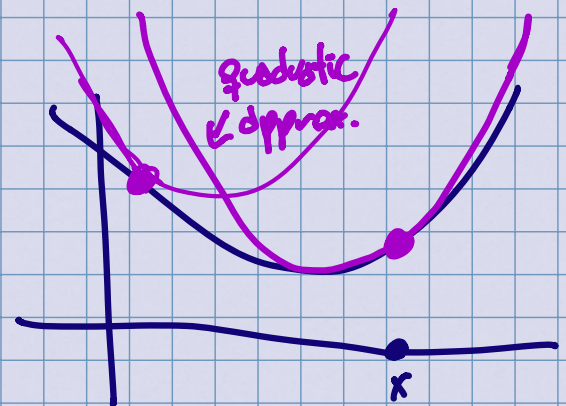
- "first-order": for all x, y

$$f(y) \geq f(x) + \underbrace{\nabla f(x)^T (y-x)}_{\substack{\text{linear approximation} \\ \text{at } x}}$$



- "second order": for all x

$$\underbrace{\nabla^2 f(x)}_{\substack{\text{Hessian} \\ \text{at } x}} \text{ is a } \underline{\text{PSD matrix}}$$



"quadratic approximation has positive curvature".

WHY IS CONVEXITY SO AWESOME?

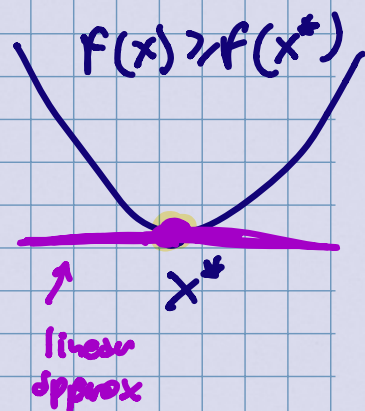


- LOCAL MINIMA ARE ALSO GLOBAL
- SIMPLE CHARACTERIZATION OF OPTIMALITY

→ $\nabla F(x^*) = 0$

"gradient must vanish"

"zero gradient at the optimal point"



- A SIMPLE ALGORITHM ("gradient method")

$$x_{k+1} = x_k - \gamma \nabla F(x_k)$$

next point current point Stepsize gradient at current point

EXAMPLE: SQUARE ROOTS.

HOW TO COMPUTE THE SQUARE ROOT OF A NUMBER?

E.G., $\sqrt{\beta}$? $(\beta > 0)$

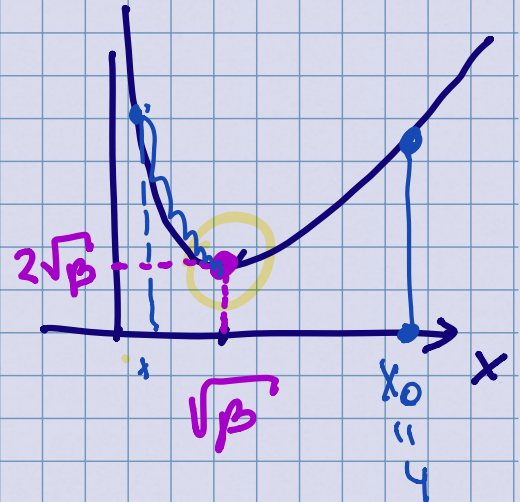
→ HERE'S AN APPROACH^{*}, USING GRADIENT DESCENT

(* DISCLAIMER:
BETTER METHODS EXIST!)

- STEP 1: GIVEN β , FIND A CONVEX FUNCTION WITH $\sqrt{\beta}$ AS THE MINIMUM (OR MINIMIZER)

E.G.: $(x > 0)$

$$F(x) = x + \frac{\beta}{x}$$



why?

$$\nabla F(x) = 1 - \frac{\beta}{x^2} = 0 \Rightarrow x = \pm \sqrt{\beta}$$

• STEP 2: LET'S WRITE GRADIENT DESCENT!

$$\begin{aligned} X_{k+1} &= X_k - \overset{\text{stepsize}}{\gamma} \nabla f(X_k) \\ &= X_k - \gamma \left(1 - \frac{\beta}{X_k^2} \right) \end{aligned}$$

Take, e.g.

$$\beta = 23$$

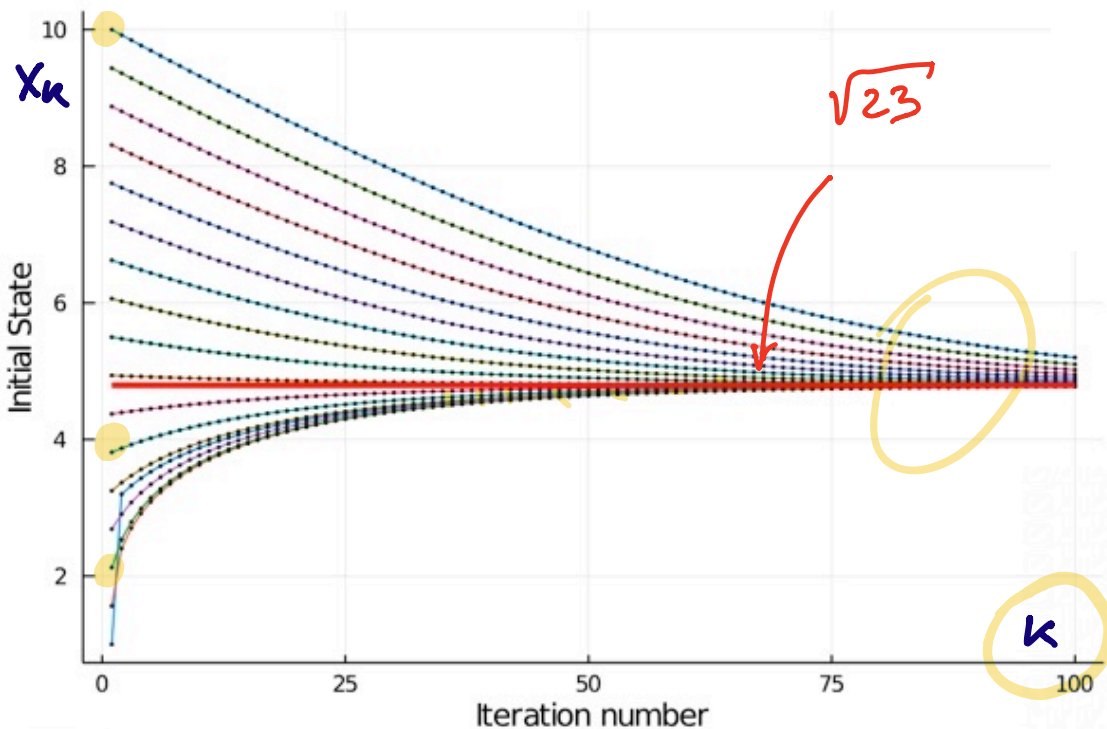
stepsize

$$\gamma = 0.1$$

(from Julia)

```
In [21]: trajectories[:,4]
```

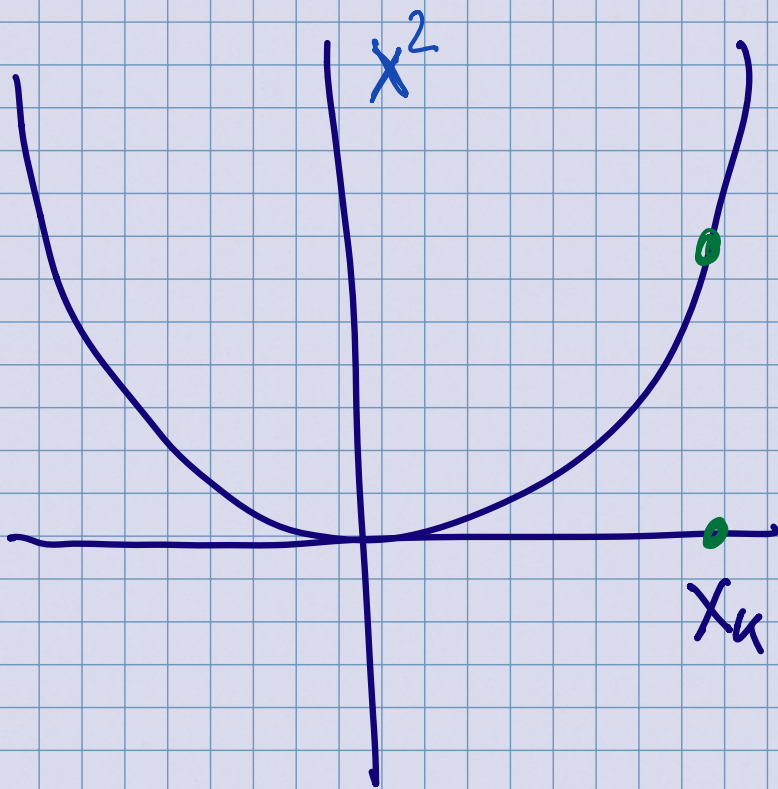
```
Out[21]: 100-element Array{Float64,1}:
 4.0
 4.04375
 4.084406316514934
 4.122276374329674
 4.157624929911346
 4.190681774532284
 4.221647746194094
 4.250699481890863
 4.277993217563307
 4.303667858367747
 4.327847482987608
 4.350643404077565
 4.372155877038263
 ⋮
 4.781262564974172
 4.781872912330582
 4.78245757797536
 4.783017651743386
 4.783554176460985
 4.784068150036249
 4.784560527451245
 4.785032222661161
 4.785484110405099
 4.785917027932993
 4.786331776652848
 4.786729123702271
```



↑
converging
to
 $\sqrt{23}$
 ~ 4.7958

DOES EVERY STEPSIZE WORK?

E.g., in the UNIVARIATE CASE



$$f(x) = \frac{1}{2} x^2$$

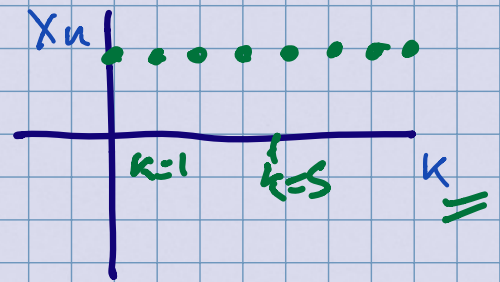
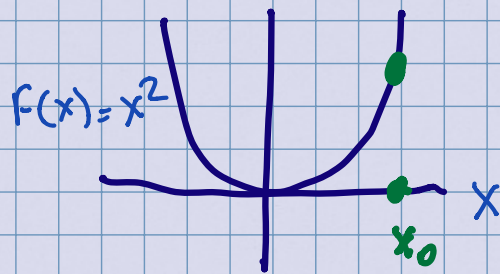
$$\nabla f(x) = x$$

$$x_{k+1} = x_k - \gamma \nabla f(x_k)$$

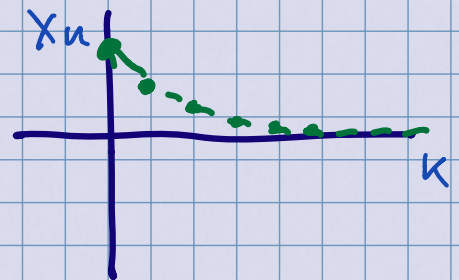
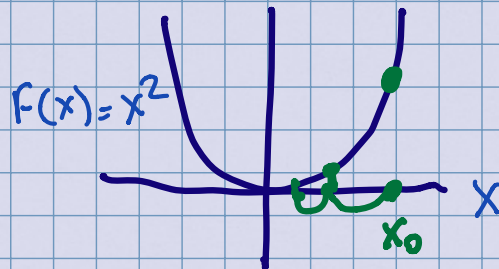
$$\underline{x_{k+1}} = x_k - \gamma x_k = \underbrace{(1 - \gamma)}_{< 1} x_k.$$

$$|1 - \gamma| < 1$$

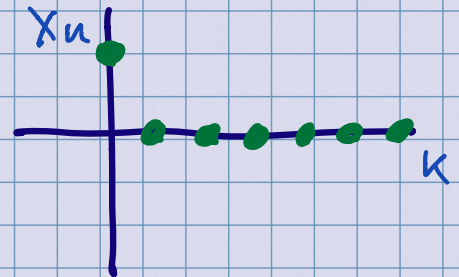
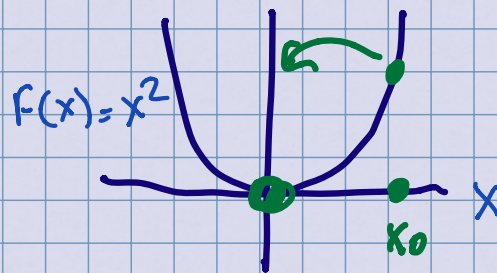
$$\gamma = 0$$



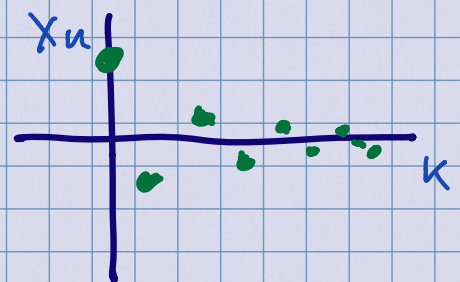
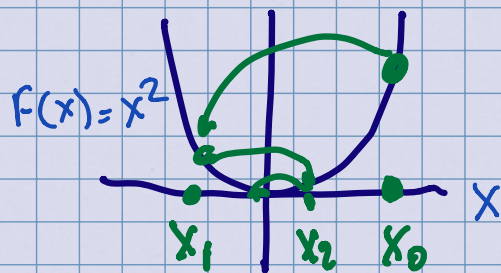
$$\gamma = 0.5$$



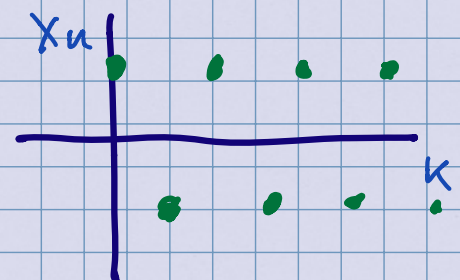
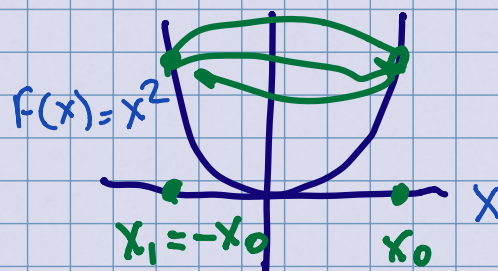
$$\gamma = 1$$



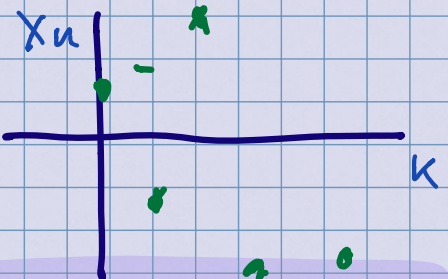
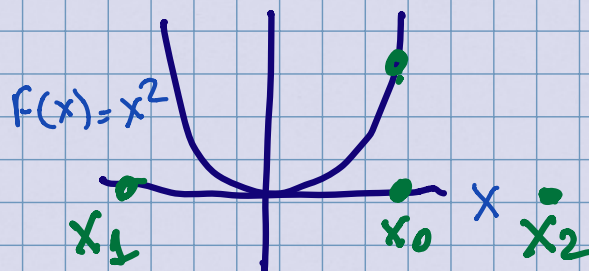
$$\gamma = 1.5$$



$$\gamma = 2$$



$$\gamma > 2$$



What happens for quadratic functions?

$$F(x) = \frac{1}{2} x^T A x$$

(without loss of generality)

A is pd
(so f is strictly convex).

$$\nabla F(x) = Ax$$

eigenvalues

$$\underline{\lambda_1} \geq \lambda_2 \geq \dots \lambda_n > \underline{0}$$

$$\Rightarrow x_{k+1} = x_k - \gamma \nabla F(x_k) = x_k - \gamma A x_k$$

$$x_{k+1} = (I - \gamma A) x_k$$

Q: When does this converge?

$$x_{k+1} = M x_k$$

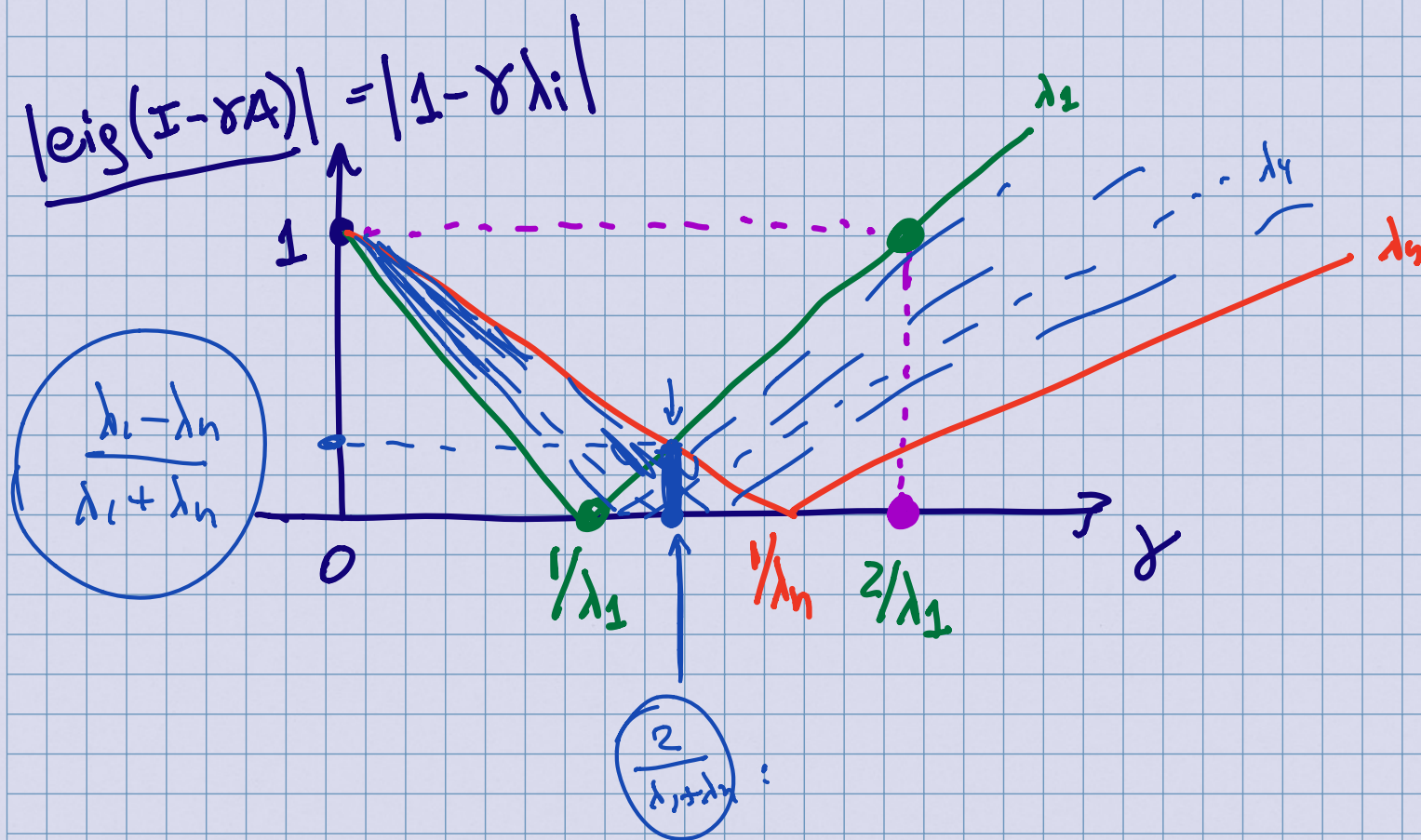
$$M = I - \gamma A$$

convergence

$$|\lambda_i(M)| < 1$$

Q: What is the "best" stepsize?

$$A = \begin{pmatrix} \lambda_1 & & 0 \\ & \ddots & \\ 0 & & \lambda_n \end{pmatrix} \quad I - \gamma A = \begin{pmatrix} \underline{1 - \gamma \lambda_1} & & 0 \\ & \ddots & \\ 0 & & \underline{1 - \gamma \lambda_n} \end{pmatrix}$$



$\Rightarrow$ if $\boxed{0 < \gamma < 2/\lambda_1}$ $\leftarrow$

The gradient method converges.

If we measure convergence rate by

$$R = \|I - \gamma A\| = \max_i |1 - \gamma \lambda_i|,$$

"best" stepsize is

$$\frac{2}{\lambda_1 + \lambda_n}$$

This choice gives a rate

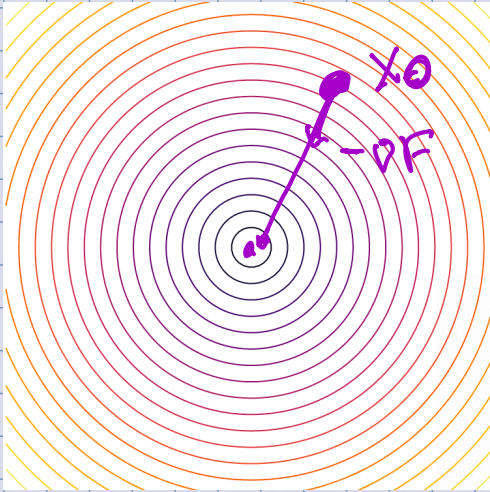
$$R \leq \frac{\lambda_1 - \lambda_n}{\lambda_1 + \lambda_n} = \frac{Q - 1}{Q + 1}$$

$Q = \frac{\lambda_1}{\lambda_n}$ is the condition number

of the matrix A .

Behavior in 2D

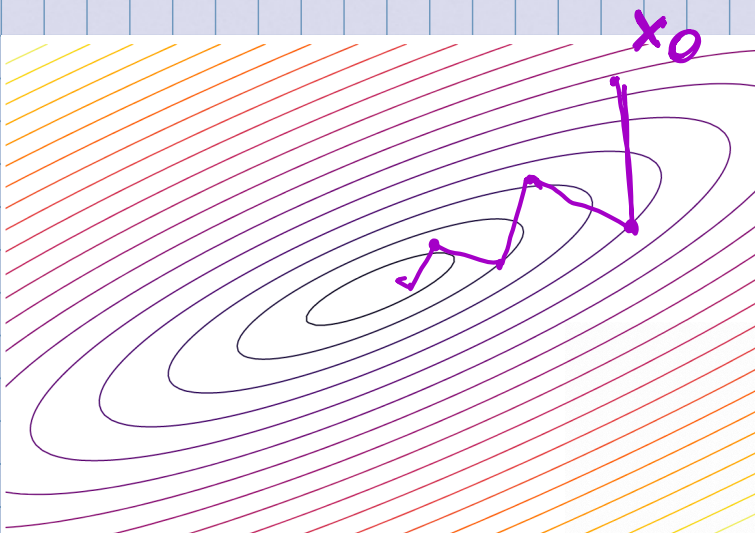
"zigzag behavior"



$$A = \begin{pmatrix} 1 & 0 \\ 0 & 1 \end{pmatrix}$$

$$Q = 1$$

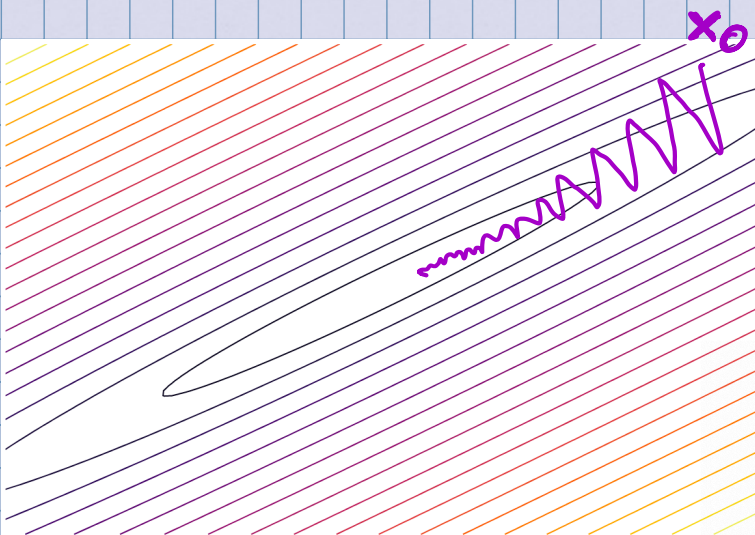
$$f(x, y) = x^2 + y^2$$



$$A = \begin{pmatrix} 1 & 0 \\ 0 & 10 \end{pmatrix}$$

$$Q = 10$$

$$f(x, y) = x^2 + 10y^2$$



$$Q = 100$$

What happens in the general convex case?

- Does it converge?
- For what values of stepsize?

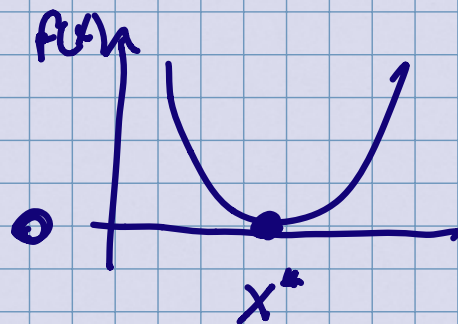
$$x_{k+1} = x_k - \gamma \nabla f(x_k)$$

For simplicity

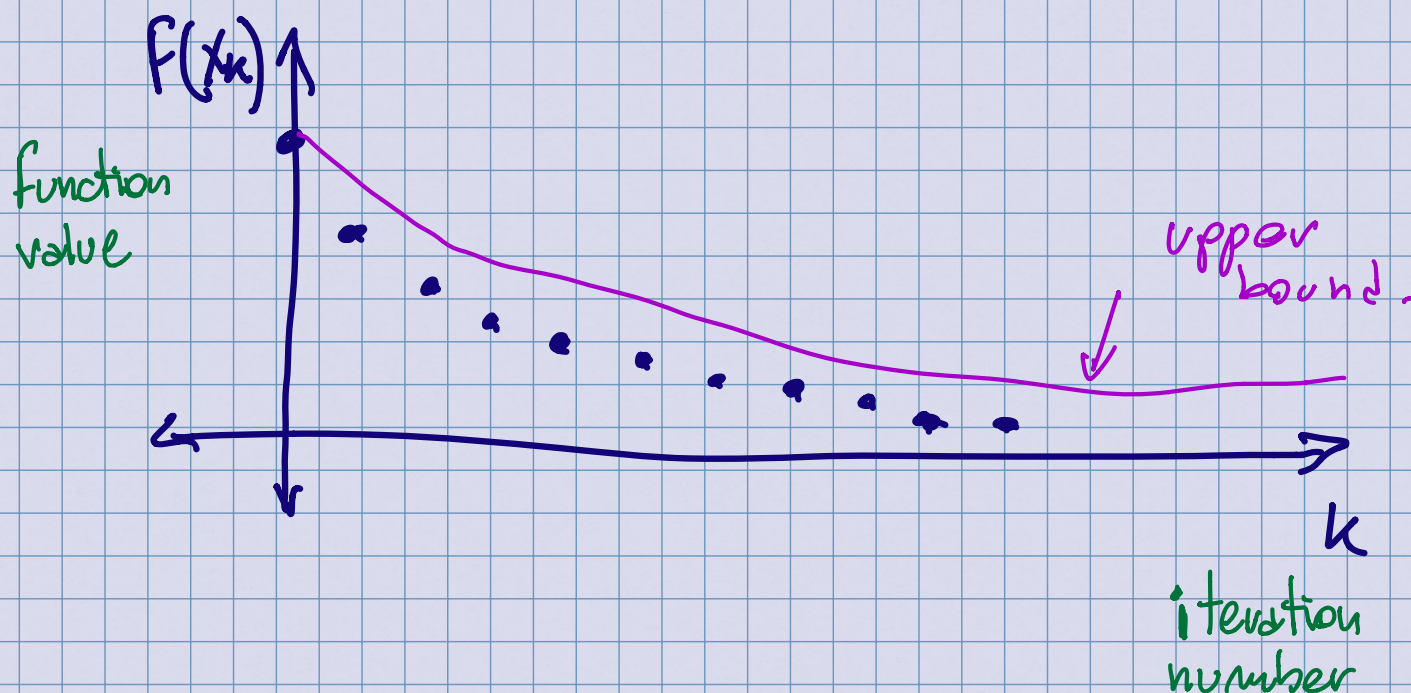
- Assume minimum value is zero:

$$f(x^*) = 0$$

(easy, by shifting the function)
 $f(x) \rightarrow f(x) - f(x^*)$



Q: HOW TO QUANTIFY CONVERGENCE?



Q: WHAT PROPERTIES OF $F(x)$
SHOULD IT DEPEND ON?

Recall:

$F(x)$ CONVEX $\Leftrightarrow$ HESSIAN PSD
(NON NEGATIVE EIGENVALUES)

λ_i eigenvalues of Hessian $\frac{\partial^2 F}{\partial x^2}$
(may depend on the point x)

CASE I: ("Smooth" or "Lipschitz gradient")

λ_i eigenvalues of Hessian

$$0 \leq \lambda_i \leq L$$

(at all points)

UPPER BOUND
ON
CURVATURE

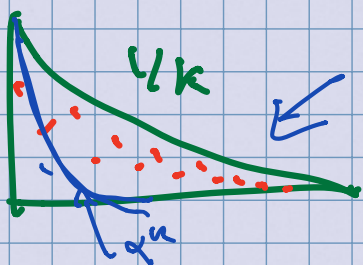
Choose Stepsize

$$\gamma = 1/L$$

Then:

$$f(x_k) \leq \frac{L}{2k} \|x_0 - x^*\|^2$$

goes to zero
as $k \rightarrow \infty$



$\Rightarrow$ GRADIENT DESCENT CONVERGES.

CASE II: ("smooth and strongly convex")

- STRONGER ASSUMPTION: $\left(\begin{array}{l} \lambda_i \text{ are} \\ \text{eigenvalues of} \\ \text{Hessian} \end{array} \right)$

$$m \leq \lambda_i \leq L \quad (\text{at all points})$$

UPPER AND LOWER BOUNDS ON CURVATURE

- Choose Stepsize

$$\gamma = \frac{2}{L+m}$$

Then:

$$f(x_k) \leq \frac{L}{2} \left(\frac{Q-1}{Q+1} \right)^{2k} \|x_0 - x^*\|^2$$

where $Q = L/m$ "condition number". ($Q < 1$)

$\Rightarrow$ CONVERGES MUCH FASTER!

INTUITION / SUMMARY

- Stepsize $\left\{ \begin{array}{l} \text{Too small: slow (but converges)} \\ \text{Too large: faster, but may oscillate, or diverge} \end{array} \right.$

FUNCTION PROPERTIES

CONVERGENCE
(for best stepsize)

(I) SMOOTH $\lambda_1 \leq L$ "slow" $O(1/k)$

(II) SMOOTH + STRONGLY CONVEX $m \leq \lambda_1 \leq L$ "Fast" or "linear" $O(\alpha^k)$
(for some $\alpha < 1$)

- no explicit dependence on dimension

- Condition number $Q = \frac{\lambda_1}{\lambda_n}$ (or $\frac{L}{m}$)

Plays a big role (want Q small)