

§ Multidimensional optimisation (without constraints)

I. Reminders ... Math

1D

$$f: \mathbb{R} \rightarrow \mathbb{R}$$

f' intuition $\rightarrow$ slope of the tangent
 f''
 $\vdots$

Taylor expansion formula (order m)

$$f(x_0+h) = \sum_{k=0}^m \frac{f^{(k)}(x_0)}{k!} \times h^k + o(h^k)$$

polynomial (h)

$\downarrow$
order 1

$$f(x_0+h) = f(x_0) + \underbrace{f'(x_0) \cdot h}_{(*)} + o(h)$$

order 1

linear fct

tangent

tangent

1D

~

mD

mD — optimization

$$f: \mathbb{R}^n \rightarrow \mathbb{R}$$

/ general

$$f: \mathbb{R}^n \rightarrow \mathbb{R}^m$$

Order 1 derivatives

$$f' = \lim_{h \rightarrow 0} \frac{f(\vec{x}_0 + \vec{h}) - f(\vec{x}_0)}{\vec{h}}$$

Partial derivatives

$$\frac{\partial f}{\partial x_i}(\vec{x}_0)$$

$\downarrow$

$$f(x_1, \dots, x_n)$$

all variables are fixed except ith one

$$x_i \mapsto f(x_1^0, \dots, x_{i-1}^0, (x_i), x_{i+1}^0, \dots, x_n^0)$$

$$\mathbb{R} \rightarrow \mathbb{R}$$

Differential

(*) f differentiable:

$$f(\vec{x}_0 + \vec{h}) = f(\vec{x}_0) +$$

$$Df(\vec{x}_0) \cdot \vec{h}$$

linear fct

differential fct

linear fct ...

$Df(\vec{x}_0)$ linear ...

matrix

$$Df(\vec{x}_0): \mathbb{R}^n \rightarrow \mathbb{R}$$

(basis $\rightarrow$ n vectors $\{e_1, \dots, e_n\}$)
 size of its matrix?

$$n \times n \begin{pmatrix} \cdot & \cdot & \cdot & \cdot \end{pmatrix} \begin{matrix} \uparrow \\ \text{1 line} \end{matrix}$$

$\downarrow \quad \downarrow \quad \downarrow \quad \downarrow$
 $Df(e_1) \quad \dots \quad Df(e_n) \in \mathbb{R}$

$Dg(\vec{x}_0) \longleftrightarrow 1 \times m$ matrix

canonical basis:

$$\left(\frac{\partial g}{\partial x_1} \quad \dots \quad \frac{\partial g}{\partial x_m} \right)$$

$$\begin{matrix} 1D \\ g' \longleftrightarrow \nabla g \end{matrix}$$

$$\nabla g(\vec{x}_0) = \begin{pmatrix} \frac{\partial g}{\partial x_1} \\ \vdots \\ \frac{\partial g}{\partial x_m} \end{pmatrix}$$

$$\begin{matrix} \text{matrix of} \\ Dg \\ \downarrow \\ Dg^t \end{matrix}$$

Remainders

$\vec{u}, \vec{v}$ vectors

$\langle \vec{u}, \vec{v} \rangle$ $\xrightarrow{2 \text{ notations}}$ $\vec{u} \cdot \vec{v}$

$\parallel$

$$\vec{u}^t \times \vec{v}$$

}

$$(A \times B)$$

u^t

$$(\dots\dots\dots)$$

v

$$\begin{pmatrix} \cdot \\ \cdot \\ \cdot \\ \cdot \\ \cdot \end{pmatrix}$$

$$1 \times 1 \leftrightarrow \mathbb{R}$$

$$= (\cdot)$$

Taylor formula (order 1) $\xrightarrow{\text{linear}}$

$$g(\vec{x}_0 + \vec{h}) = g(\vec{x}_0) + \overbrace{Dg(\vec{x}_0)}^{\parallel} \cdot \vec{h} + o(\|\vec{h}\|)$$

$$\parallel \nabla g(\vec{x}_0)^t \times \vec{h}$$

$$\parallel \langle \nabla g(\vec{x}_0), \vec{h} \rangle$$

Intuition behind the gradient

Taylor formula

$$g(\vec{x}_0 + \vec{h}) \sim g(\vec{x}_0) + \langle \nabla g(\vec{x}_0), \vec{h} \rangle$$

$\Leftrightarrow$

$$g(\vec{x}_0 + \vec{h}) - g(\vec{x}_0) \sim$$

if $\nabla g > 0$
then $g \nearrow$

if $\nabla g < 0$
then $g \searrow$

if $\nabla g = 0$
then g is constant

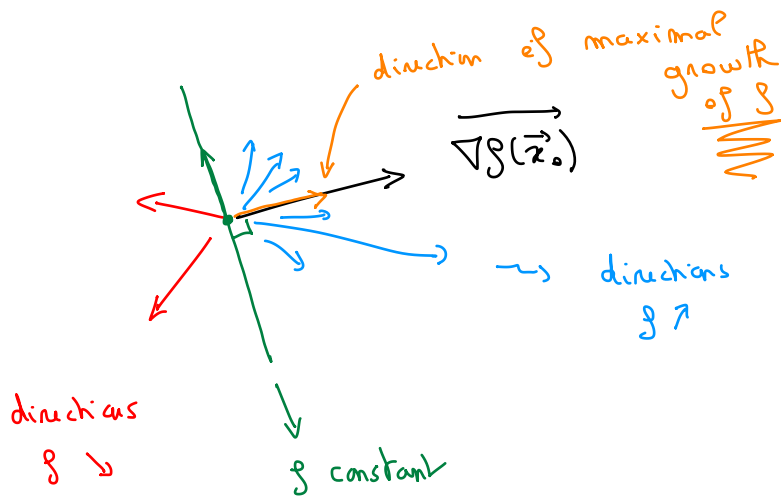
$$\langle \overbrace{\nabla g(\vec{x}_0)}^{\text{vector}}, \vec{h} \rangle$$

$$\|\vec{h}\| = 1$$

maximum
for

$\vec{h}$ and
 $\nabla g(\vec{x}_0)$ aligned

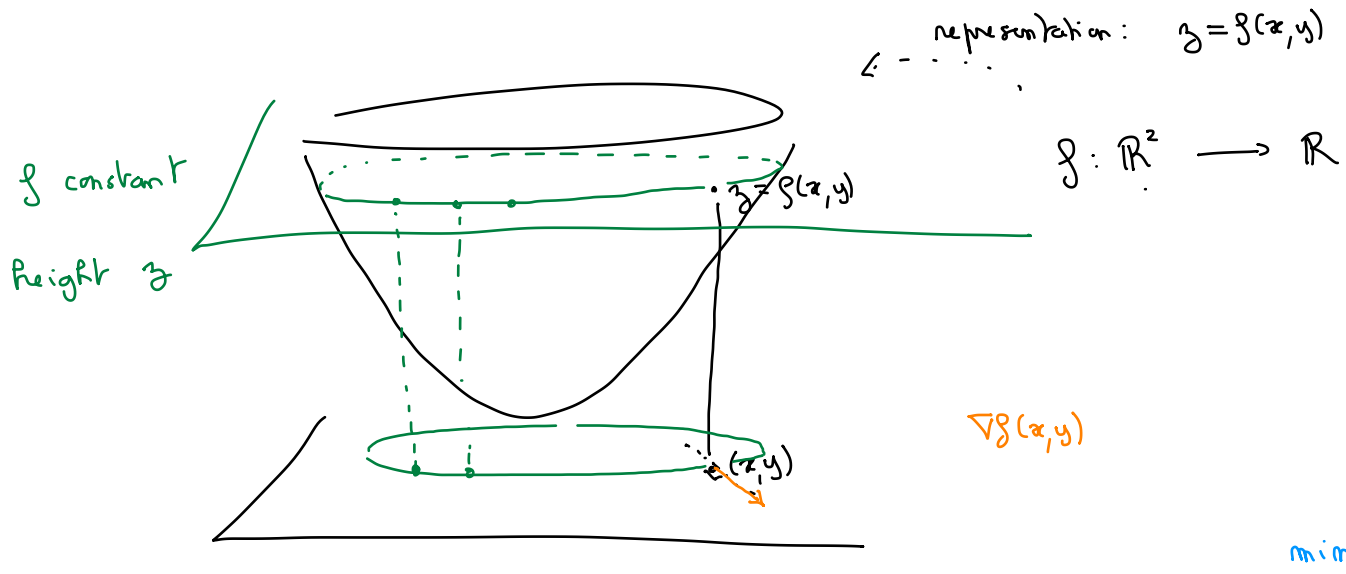
sign?



$$\vec{h} \text{ st } \langle \nabla f(\vec{z}_0), \vec{h} \rangle > 0$$

$$\vec{h} \text{ st } \dots < 0$$

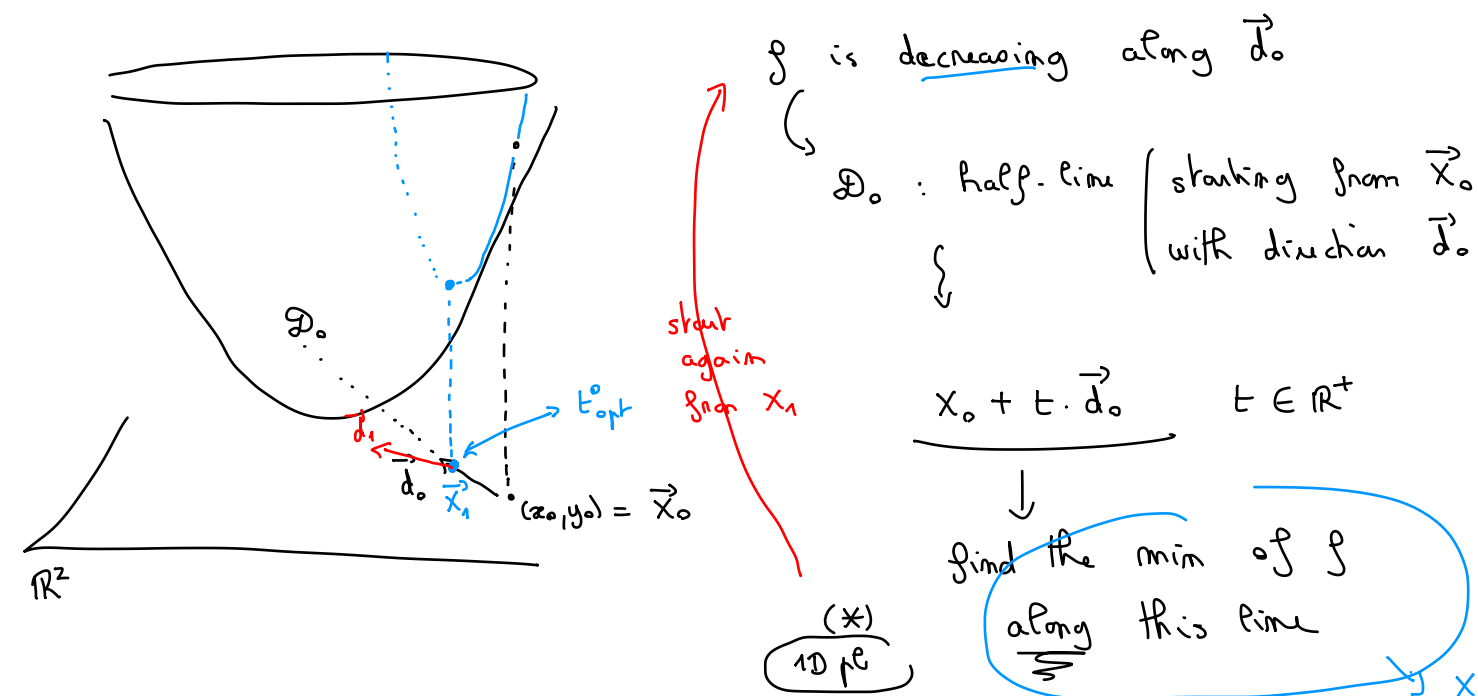
$$\vec{h} \text{ st } \dots = 0$$



min f ?

II. Descent algorithms

Example $f: \mathbb{R}^2 \rightarrow \mathbb{R}$ represented by a surface

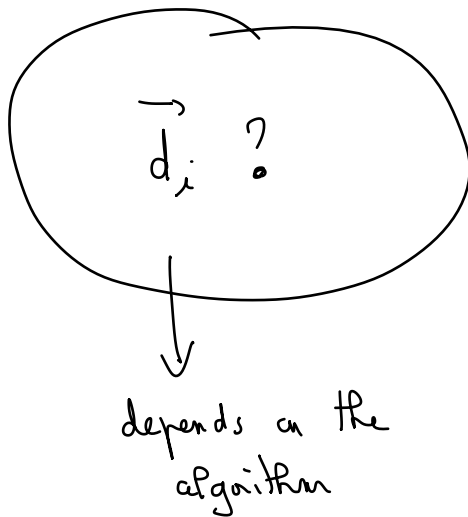
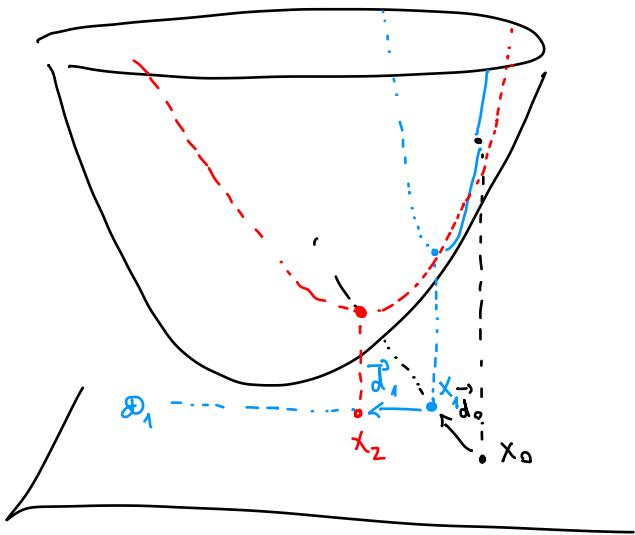


(*) find t s.t. $\underbrace{g(\vec{x}_0 + t \cdot \vec{d}_0)}_{\text{minimum}}$

$$\underbrace{g(\vec{x}_0 + t \cdot \vec{d}_0)}_{\text{minimum}}$$

$$\begin{array}{l} t \mapsto g(\vec{x}_0 + t \cdot \vec{d}_0) \quad \text{--- min?} \\ \mathbb{R} \longrightarrow \mathbb{R} \end{array}$$

1D optimization



Most basic

gradient descent algorithm

$$\vec{d}_i = - \nabla g(\vec{x}_i)$$

→ greedy algorithm

- not the fastest — many iterations

④ numerically very stable

flat functions - challenging for opt. algorithms...

conjugate gradient
descent algorithm

④ faster

⊖ may fail with non convex functions ...

order 1 methods

Order 2
methods

— based on Newton-Raphson

Newton-Raphson method

Quasi-Newton methods

approximations of Hessian

$$x_{i+1} = x_i - H(g)(x_i)^{-1} \times \nabla g(x_i)$$

Newton

$g: \mathbb{R} \rightarrow \mathbb{R}$
↓
golden + parabolic
Newton on g'
 $g'(x) = 0$
↓
 $x_{i+1} = x_i - \frac{g'(x_i)}{g''(x_i)}$

1D

intuition

$\mathbb{R}^m$

$$x_{i+1} = x_i - \frac{\nabla g(x_i)}{\text{second order derivative of } g}$$

matrix - Hessian matrix
 $H(g)(x_i)$

$$\left(\frac{\partial^2 g}{\partial x_i \partial x_j} \right)$$

① Gradient descent alg.

gradient-descent ($g, \nabla g, x_0, \epsilon, \delta, N$)

$x \leftarrow x_0$ // x : current point

iter $\leftarrow 0$

$x_{prev} \leftarrow x + ones(size(x))$ // just to enter the while loop

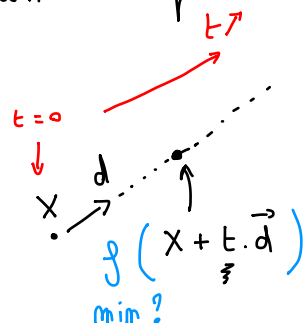
while ($(\|x - x_{prev}\| > \epsilon) \&\& (iter < N)$)

{
 $d \leftarrow -\nabla g(x)$
 // Optimize along the Ray-line (x, d)

$\varphi \leftarrow @(\epsilon) g(x + t \cdot \vec{d})$

// 1D optim ($\rightarrow$ Brent (parabolic + golden) $\rightarrow$ $gminbnd$
here: 1Doptim($\varphi, \underline{\hspace{2cm}}$ interval s.t. the fun is unimodal

// Sweeping method to init the interval



$x_i \rightsquigarrow$ direction: $d_i = -\nabla g(x_i)$

$$\varphi(t) = g(x_i + t \cdot d_i)$$

$$x_{i+1} \dots \min \rightarrow t_{opt}$$

$$x_{i+1} = x_i + t_{opt} \cdot d_i \leftarrow$$

$$\min \text{ of } \varphi$$

$$\varphi'(t_{opt}) = 0$$

$$\varphi'(t) = \overbrace{g'(x_i + t \cdot d_i)}^{\text{vec}} \otimes \overbrace{d_i}^{\text{vec}}$$

$\downarrow$
•

$g: \mathbb{R}^n \rightarrow \mathbb{R}$

intuitively $\sim 1D$

$$(g \circ g)' = g' \circ g \times g'$$

$$\varphi'(t) = \nabla g(x_i + t \cdot d_i) \cdot d_i$$

$$\text{At } t_{opt} \Rightarrow \varphi'(t_{opt}) = \nabla g(\underbrace{x_i + t_{opt} \cdot d_i}_{x_{i+1}}) \cdot \vec{d}_i$$

$\downarrow$
 $\nabla g(x_{i+1}) = -\vec{d}_{i+1}$

$$\Leftrightarrow \boxed{0 = \underbrace{d_{i+1}}_{\downarrow} \cdot \underbrace{d_i}_{\downarrow}}$$

Convergence

$\int \mathbb{I} g$ g is $\mathcal{C}^2$ and elliptic
then the algorithm converges

$$\left(\begin{array}{l} \exists \alpha \text{ st. } \forall x, y \\ \|\nabla g(x) - \nabla g(y)\| \geq \alpha \|x - y\| \end{array} \right)$$

$\nwarrow$
"curvature"



Reminder - quadratic forms / polynomials of d^2 exactly

ex: $\mathbb{R}^3$ $g(x,y,z) = x^2 - y^2 + 2z^2 + \boxed{xz} + \boxed{2}yz$

$\updownarrow$
 $X^t A X$

can be represented by a matrix

$A = \begin{pmatrix} x & y & z \\ y & & \\ z & & \end{pmatrix}$

symmetric

squared terms

$\rightarrow xz \rightarrow$ split symmetrically

$\begin{matrix} x \\ y \\ z \end{matrix} \begin{pmatrix} 1 & 0 & \frac{1}{2} \\ 0 & -1 & 1 \\ \frac{1}{2} & 1 & 2 \end{pmatrix}$

$g(x,y,z) = g(X) = X^t \cdot A \cdot X$

$g(X) = A \cdot X$

Gradient of such a function:

$\nabla g = 2AX$

! is A symmetric

$\nabla g(X) = A^t$

If A not symmetric

$\nabla g = (A + A^t) \cdot X$

Problem: practical

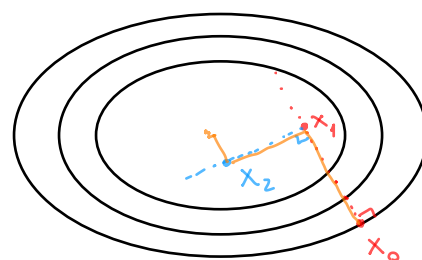
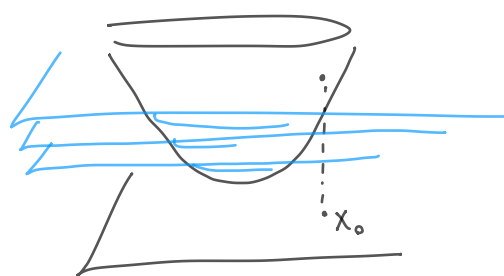
$g_1(x,y) = \alpha x^2 + \beta y^2$

$\alpha, \beta > 0$

steepest gradient descent converges to min

for this BASIC function

Expected: STOP at the min



steps

The problem comes from:

$\rightarrow \vec{d}_i \perp \vec{d}_{i+1} \leadsto$ the choice is "limited"

$\leadsto$ always the same directions

$\leadsto$ if $\vec{d}_0$ does not cross the min
no $\vec{d}_i$ will cross it ...

$\rightarrow$ the alg. is greedy ...

Conjugate gradient descent improves the convergence

$\hookrightarrow$ designed for quadratic functions

Formulae come from such functions — polynomial $d^{\circ} 2$

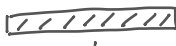
$\hookrightarrow$ can be extended to $\mathbb{C}^2$ elliptic functions

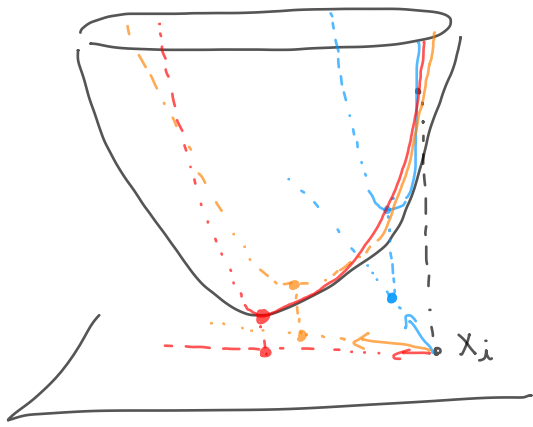
② Conjugate gradient descent

Idea : f is a quadratic function

$$f(x) = \underbrace{x^t \cdot A \cdot x}_{\text{order 2}} + \underbrace{t B \cdot x}_{\substack{\text{order 1} \\ \text{linear}}} + \underbrace{\widehat{c}}_{\substack{\text{order 0} \\ \text{constant}}}$$

$\{e_1, \dots, e_m\}$ basis
 $f: \mathbb{R}^m \rightarrow \mathbb{R}$
 (matrix $\begin{pmatrix} \cdot & \cdot \\ \vdots & \vdots \end{pmatrix}$ $\downarrow \dots \downarrow$ m cols
 $f(e_i) \in \mathbb{R}$)

linear
 $\downarrow$
 matrix $\times x$
 size?
 $1 \times m$

 $\downarrow t$
 B



each $\vec{d}$ $\rightarrow$ a min
(the min of g
along line $(x_i, \vec{d})$)

$\downarrow$

which $\vec{d}$ gives the best minimum?

$\nabla g(x_i)$

For quadratic function we can compute an exact formula for this "optimal" $\vec{d}$

This "optimal" $\vec{d}$ is a linear combination

- $\nabla g(x_i)$
- previous direction $\vec{d}_{i-1}$

Optimal direction:

$$\vec{d}_{i+1} = \nabla g(x_{i+1}) + \frac{\|\nabla g(x_{i+1})\|^2}{\|\nabla g(x_i)\|^2} \vec{d}_i$$

g quadratic $g(x) = \frac{1}{2} x^T A x + B^T x + c$
 x_0 : close to the minimum ...

APg:

$$x \leftarrow x_0$$

$$x_{\text{prev}} \leftarrow x + \begin{pmatrix} 1 \\ \vdots \\ 1 \end{pmatrix}$$

$$\text{iter} \leftarrow 0$$

$$d \leftarrow \nabla g(x) \quad \triangle! \quad d : - \text{descent direction}$$

$$\text{while } (\|x - x_{\text{prev}}\| > \epsilon) \quad \&\& \quad (\text{iter} < N)$$

$$\left\{ \begin{array}{l} \varphi(t) = g(x - t \cdot d) \\ // \text{min } \varphi? \end{array} \right.$$

// if g is quadratic

$$t_{\text{opt}} = \frac{\|\nabla g(x_i)\|^2}{x_i^T A x_i}$$

// if g $\mathbb{R}^2$ elliptic
 // 1D optim $\leadsto$ find min φ

$$t_0 = 0$$

$$t_1 = \delta$$

$$y_0 = \varphi(t_0)$$

$$y_1 = \varphi(t_1)$$

Gradient Steepest descent
 $\vec{d}_i = - \nabla g(x_i)$
 $\varphi(t) = g(x_i - t \cdot \vec{d}_i)$

while ($y_0 > y_1$)

$$t_0 = t_1$$

$$y_0 = y_1$$

$$t_1 = t_0 + \delta$$

$$y_1 = f(t_1)$$

end

$$t_{\text{opt}} = \text{IDoptim}(\varphi, \phi, t_1)$$

$$X_{\text{prev}} \leftarrow X$$

$$X \leftarrow X - t_{\text{opt}} \cdot d$$

$$d \leftarrow \nabla f(X) + \frac{\|\nabla f(X)\|^2}{\|\nabla f(X_{\text{prev}})\|^2} \cdot d$$

end

Convergence ...
 → case f quadratic
 → case f $\mathbb{R}^2$ elliptic → the alg. converges

When f quadratic and A definite positive:

all eigenvalues are > 0

$\Leftrightarrow$

$$X^T A X > 0 \quad \forall X \neq 0$$

Property

$\forall i \quad \forall j < i \quad \vec{d}_i, \vec{d}_j$ are conjugate with respect to A
 ↓
 $\langle A \vec{d}_j, \vec{d}_j \rangle = 0$

conjugate $\sim$ orthogonal
 for the A -dot product

$\vec{d}_i$ conjugate with $\vec{d}_0, \dots, \vec{d}_{i-1}$

Question

$$\text{is } f: \mathbb{R}^n \rightarrow \mathbb{R}$$

quadratic elliptic / A sym. definite positive

the algorithm ~~converges~~ in at most n steps.
 steps

For C^2 functions $\Rightarrow$ faster than gradient descent

Help... $g(x) = \frac{1}{2} \cdot x^T \cdot A \cdot x + B^T \cdot x + c$

$$\nabla g(x) = A \cdot x + B$$

In the definition of $\vec{d}$:

$$\|\nabla g(x_i)\|^2 = \|Ax_i + B\|^2 = \langle Ax_i + B, Ax_i + B \rangle$$

$$g(x) = x^T \cdot A \cdot x + B^T x + c$$

$\downarrow$

$$\nabla g(x) = 2AX + BX$$

Libraries $\leadsto$ many variants

pre-conditioned | conjugate gradients

bi-conjugate gradients

III . Order 2 methods $\equiv$ Newton

① Newton in $\mathbb{R}^m$ $\rightarrow$ Newton-Raphson

Goal $\rightarrow$ compute the zeros of g

$$g: \mathbb{R}^m \rightarrow \mathbb{R}^m$$

$g: \mathbb{R}^m \rightarrow \mathbb{R}$
Optimize $g \mid x: \min$

$\Downarrow$ N.C.

$$\nabla g(x) = 0$$

$$\nabla g: \mathbb{R}^m \rightarrow \mathbb{R}^m$$

$$x \mapsto \nabla g(x)$$

$$g = \nabla g$$

ex:

$$g \begin{pmatrix} x \\ y \\ z \end{pmatrix} = \begin{pmatrix} x^2 + \cos(yz) \\ 2x + \ln(z) \\ xyz \end{pmatrix}$$

Newton $\longleftrightarrow$ solve non linear systems

$$\begin{aligned} &\Updownarrow \\ g \begin{pmatrix} x \\ y \\ z \end{pmatrix} = \vec{0} &\iff \begin{cases} x^2 + \cos(yz) = 0 \\ 2x + \ln(z) = 0 \\ xyz = 0 \end{cases} \end{aligned}$$

Algorithm

Intuitively 1D $\leadsto$ mD

1D

$$g(x) = 0$$

Newton iteration:

$$x_{i+1} = x_i - \frac{g(x_i)}{g'(x_i)}$$

mD

$$\vec{x}_{i+1} = \vec{x}_i - \frac{\vec{g}(\vec{x}_i)}{g'(\vec{x}_i)} \quad \text{vector}$$

$$g: \mathbb{R}^n \rightarrow \mathbb{R}^m$$

~~g~~

$$\frac{\partial g}{\partial x_i} \in \mathbb{R}^m$$

Jacobian: $J(g)$

$$g: \mathbb{R}^n \rightarrow \mathbb{R} \quad \nabla g = \begin{pmatrix} \frac{\partial g}{\partial x_1} \\ \vdots \\ \frac{\partial g}{\partial x_n} \end{pmatrix}$$

Jacobian — "set" of partial derivatives
↓
matrix

$$g: \mathbb{R}^n \rightarrow \mathbb{R}^m$$

$g(\vec{x})$ vector dim m

$$\begin{pmatrix} g_1(\vec{x}) \\ \vdots \\ g_m(\vec{x}) \end{pmatrix}$$

$$g_i: \mathbb{R}^n \rightarrow \mathbb{R}$$

$$J(g) = \left(\frac{\partial g_i}{\partial x_j} \right)_{i,j \in \{1 \dots m\}} = \begin{pmatrix} \nabla g_1^T & \dots \\ \vdots & \\ \nabla g_m^T & \dots \end{pmatrix}$$

special case
 $g: \mathbb{R}^n \rightarrow \mathbb{R}$
Dg linear $\mathbb{R}^n \rightarrow \mathbb{R}^m$
↕
matrix
J(g) — Jacobian
↓
Gradient^T — ∇g

ex:

$$\mathbb{R}^3 \rightarrow \mathbb{R}^3 \quad g \begin{pmatrix} x \\ y \\ z \end{pmatrix} = \begin{pmatrix} x^2 + \cos(yz) \\ 2x + \ln(z) \\ xyz \end{pmatrix} \begin{matrix} \leftarrow g_1 \\ \leftarrow g_2 \\ \leftarrow g_3 \end{matrix}$$

$$g_i: \mathbb{R}^3 \rightarrow \mathbb{R} \quad \nabla g_i$$

Back to mD-Newton:

mD

$$\vec{x}_{i+1} = \vec{x}_i - \frac{\vec{g}(\vec{x}_i)}{g'(\vec{x}_i)} \quad \text{vector}$$

~~$g'(\vec{x}_i)$~~

$$J(g)(\vec{x}_i)$$

matrix

$$J(g)(\vec{x}_i)^{-1} \times \dots$$

Newton - Raphson — $g: \mathbb{R}^n \rightarrow \mathbb{R}^n$ / $\boxed{g(x) = \vec{0}}$

x_0 : close to the solution

while $(\|x_i - x_{i-1}\| > \epsilon)$

$$x_{i+1} = x_i - \underbrace{J(g)(x_i)^{-1}} \times g(x_i)$$

end

$$\sim \frac{g(x_i)}{g'(x_i)} \text{ 1D ...}$$

Newton for optimization

$$f: \mathbb{R}^n \rightarrow \mathbb{R}$$

$\leadsto \min f$?

$\downarrow$ necessary condition

$$\nabla f(x) = \vec{0}$$

$g = \nabla f$

$J(\nabla f)$?

$$\nabla f(x) = \begin{pmatrix} \frac{\partial f}{\partial x_1} \\ \vdots \\ \frac{\partial f}{\partial x_n} \end{pmatrix} \begin{matrix} \leftarrow g_1 \\ \vdots \\ \leftarrow g_n \end{matrix}$$

$$\leadsto g_i = \frac{\partial f}{\partial x_i}$$

$$g = \nabla f$$

$$J(\nabla f) = \begin{pmatrix} \frac{\partial g_i}{\partial x_j} \end{pmatrix} = \begin{pmatrix} \frac{\partial^2 f}{\partial x_i \partial x_j} \end{pmatrix}_{i,j \in \{1 \dots n\}} = H(f)$$

derivate f'

Hessian matrix
of f

Newton for optimization

$$f: \mathbb{R}^n \rightarrow \mathbb{R}$$

$\leadsto \min f$?

$\downarrow$ necessary condition

$$\nabla f(x) = \vec{0}$$

$g = \nabla f$

x_0 — close from the minimum

while $(\|x_i - x_{i-1}\| > \epsilon)$

$$x_{i+1} = x_i - \underbrace{H(f)(x_i)^{-1}} \times \nabla f(x_i)$$

end

input a $n \times n$ matrix

Quasi-Newton methods ...
→ BFGS
→
heuristics

approximate $H(g)^{-1}$
from ∇g