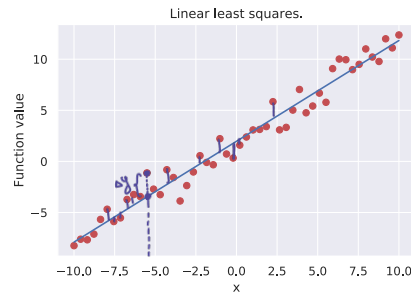
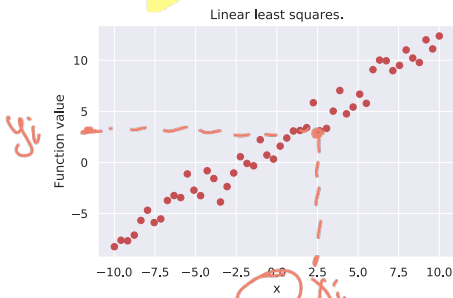


Linear least squares

Problem



$$y_i = f(x_i) \quad f(x) = kx + b$$

$$f(x) = \theta^T x$$

$$x \in \mathbb{R}^n$$

$$\theta \in \mathbb{R}^n$$

In a least-squares, or linear regression, problem, we have measurements $X \in \mathbb{R}^{m \times n}$ and $y \in \mathbb{R}^m$ and seek a vector $\theta \in \mathbb{R}^n$ such that $X\theta$ is close to y . Closeness is defined as the sum of the squared differences:

$$\sum_{i=1}^m (x_i^T \theta - y_i)^2 \rightarrow \min_{\theta \in \mathbb{R}^n} \|X\theta - y\|_2^2 = \langle X\theta - y, X\theta - y \rangle$$

also known as the l_2 -norm squared, $\|X\theta - y\|_2^2$

For example, we might have a dataset of m users, each represented by n features. Each row x_i^T of X is the features for user i , while the corresponding entry y_i of y is the measurement we want to predict from x_i^T , such as ad spending. The prediction is given by $x_i^T \theta$.

We find the optimal θ by solving the optimization problem

$$\|X\theta - y\|_2^2 \rightarrow \min_{\theta \in \mathbb{R}^n}$$

Let θ^* denote the optimal θ . The quantity $r = X\theta^* - y$ is known as the residual. If $\|r\|_2 = 0$, we have a perfect fit.

$$X^T(X\theta - y) = 0$$

$$X^T X \theta - X^T y = 0$$

$$X^T X \theta = X^T y \Rightarrow \underbrace{X^T X}_{n \times n} \cdot \theta = \underbrace{X^T y}_{n \times 1} \quad \tilde{y} = X^T y$$

$$\theta = (X^T X)^{-1} X^T y$$

$$X^+ = (X^T X)^{-1} X^T$$

$$\theta = X^+ y$$

$$\theta = \tilde{X}^{-1} \cdot y$$

не квадратично

$$\text{rank } X = n$$

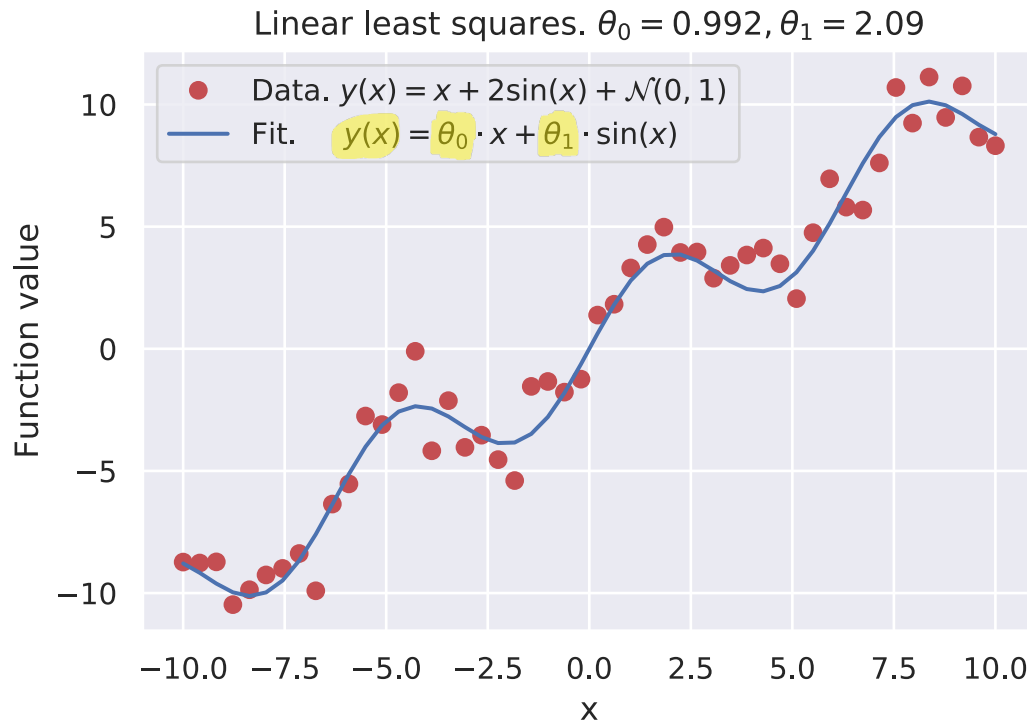
$$m < n \ominus$$

$$n \leq m \oplus$$

$$m \geq n$$



Note, that the function needn't be linear in the argument x but only in the parameters θ that are to be determined in the best fit.



Approaches

Moore-Penrose inverse

If the matrix X is relatively small, we can write down and calculate exact solution:

$$\theta^* = (X^\top X)^{-1} X^\top y = X^\dagger y,$$

where $X^\dagger$ is called [pseudo-inverse](#) matrix. However, this approach squares the condition number of the problem, which could be an obstacle in case of ill-conditioned huge scale problem.

QR decomposition

For any matrix $X \in \mathbb{R}^{m \times n}$ there exists QR decomposition:

$$X = Q \cdot R,$$

where Q is an orthogonal matrix (its columns are orthogonal unit vectors meaning $Q^\top Q = QQ^\top = I$ and R is an upper triangular matrix. It is important to notice, that since $Q^{-1} = Q^\top$, we have:

$$QR\theta = y \longrightarrow R\theta = Q^\top y$$

Now, process of finding theta consists of two steps:

1. Find the QR decomposition of X .
2. Solve triangular system $R\theta = Q^\top y$, which is triangular and, therefore, easy to solve.

Handwritten notes:

$$L = \sum_{i=1}^m (kx_i - y_i)^2$$

$$\frac{\partial L}{\partial k} = \sum_{i=1}^m \frac{\partial}{\partial k} (kx_i - y_i)^2 = \sum_{i=1}^m 2(kx_i - y_i) \cdot x_i = 0$$

$$\rightarrow 2k \sum_{i=1}^m x_i \cdot x_i - 2 \sum_{i=1}^m y_i x_i = 0$$

① $\frac{\partial L}{\partial k} = ?$

② $\frac{\partial L}{\partial k} = 0 \Rightarrow k = ?$

$$K \cdot \sum_{i=1}^m x_i^2 = \sum_{i=1}^m x_i \cdot y_i$$

$$K = \dots$$

Cholesky decomposition

For any positive definite matrix $A \in \mathbb{R}^{n \times n}$ there exists Cholesky decomposition:

$$X^T X = A = L^T \cdot L,$$

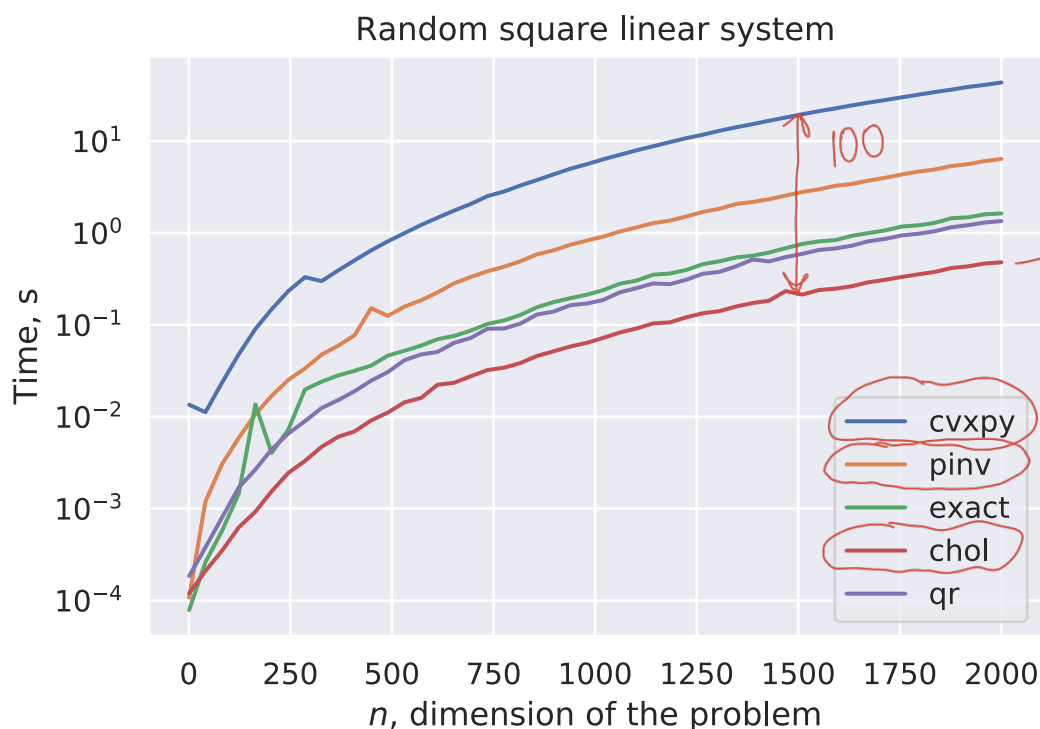
where L is a lower triangular matrix. We have:

$$L^T L \theta = y \quad \longrightarrow \quad L^T z_\theta = y$$

Now, process of finding theta consists of two steps:

1. Find the Cholesky decomposition of $X^T X$.
2. Find the $z_\theta = L \theta$ by solving triangular system $L^T z_\theta = y$
3. Find the θ by solving triangular system $L \theta = z_\theta$

Note, that in this case the error stil proportional to the squared condition number.



np.linalg.solve
np.linalg.lstsq

Code

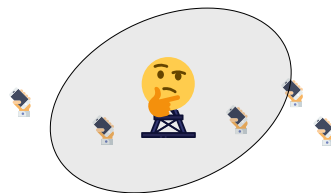
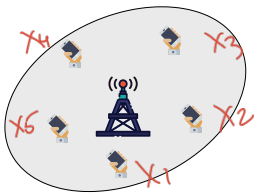
[Open in Colab](#)

References

- [CVXPY documentation](#)
- [Interactive example](#)
- [Jupyter notebook by A. Katrutsa](#)

Minimum volume ellipsoid

Problem

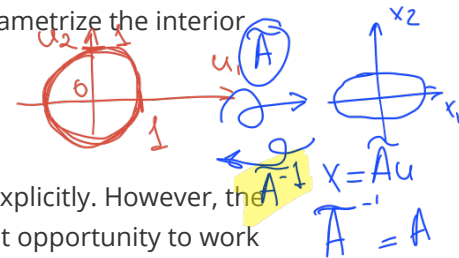


Let $x_1, \dots, x_n$ be the points in $\mathbb{R}^2$. Given these points we need to find an ellipsoid, that contains all points with the minimum volume (in 2d case volume of an ellipsoid is just the square).

$$\|u\|_2^2 \leq 1$$

An invertible linear transformation applied to a unit sphere produces an ellipsoid with the square, that is $\det A^{-1}$ times bigger, than the unit sphere square, that's why we parametrize the interior of ellipsoid in the following way:

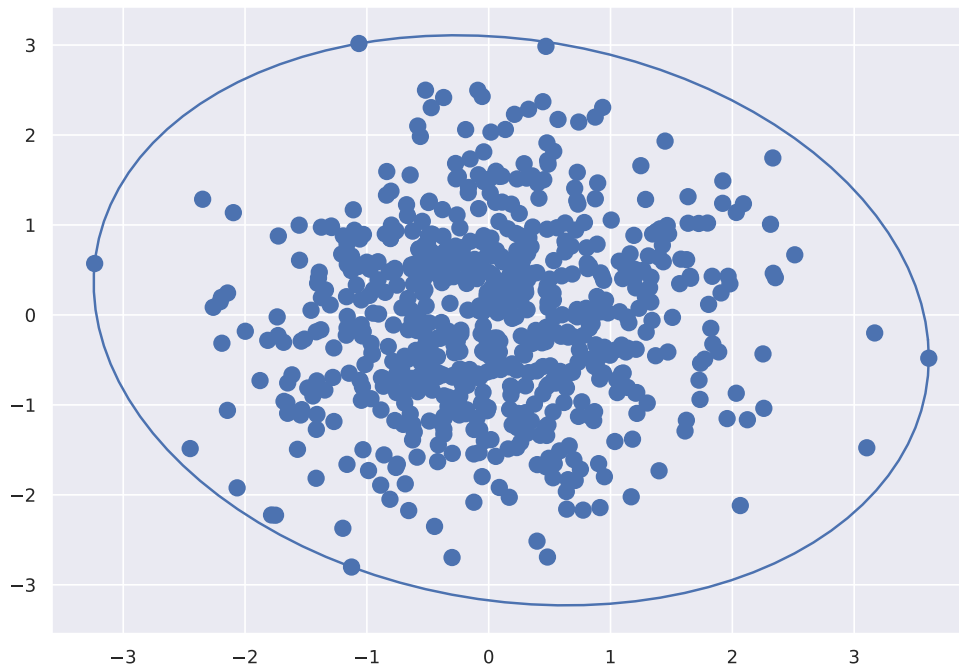
$$S = \{x \in \mathbb{R}^2 \mid u = Ax + b, \|u\|_2^2 \leq 1\}$$



Sadly, the determinant is the function, which is relatively hard to minimize explicitly. However, the function $\log \det A^{-1} = -\log \det A$ is actually convex, which provides a great opportunity to work with it. As soon as we need to cover all the points with ellipsoid of minimum volume, we pose an optimization problem on the convex function with convex restrictions:

$$\begin{aligned} & \min_{A \in \mathbb{R}^{2 \times 2}, b \in \mathbb{R}^2} -\log \det(A) \\ & \text{s.t. } \|Ax_i + b\| \leq 1, i = 1, \dots, n \\ & \quad A \succ 0 \end{aligned}$$

δεν αντιστοιχεί
απλ. βήματα
 $\ln \det(A^{-1}) = -\ln \det(A)$



Code

[Open in Colab](#)

References

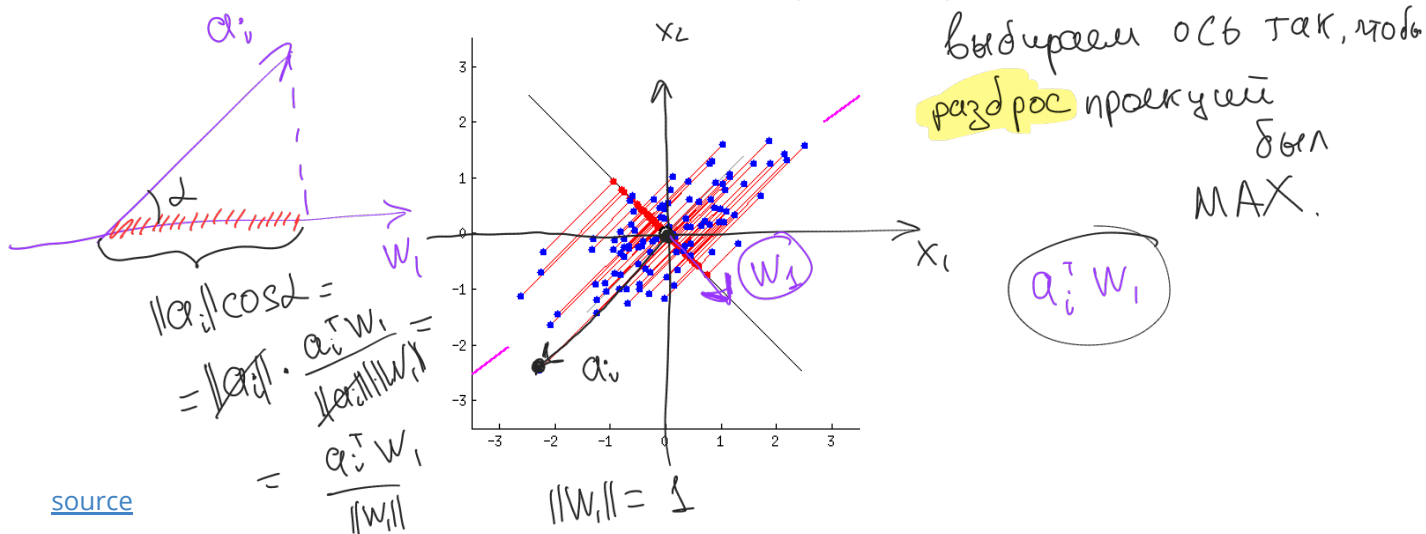
- [Jupyter notebook](#) by A. Katrutsa
- <https://cvxopt.org/examples/book/ellipsoids.html>

Principal component analysis

Intuition

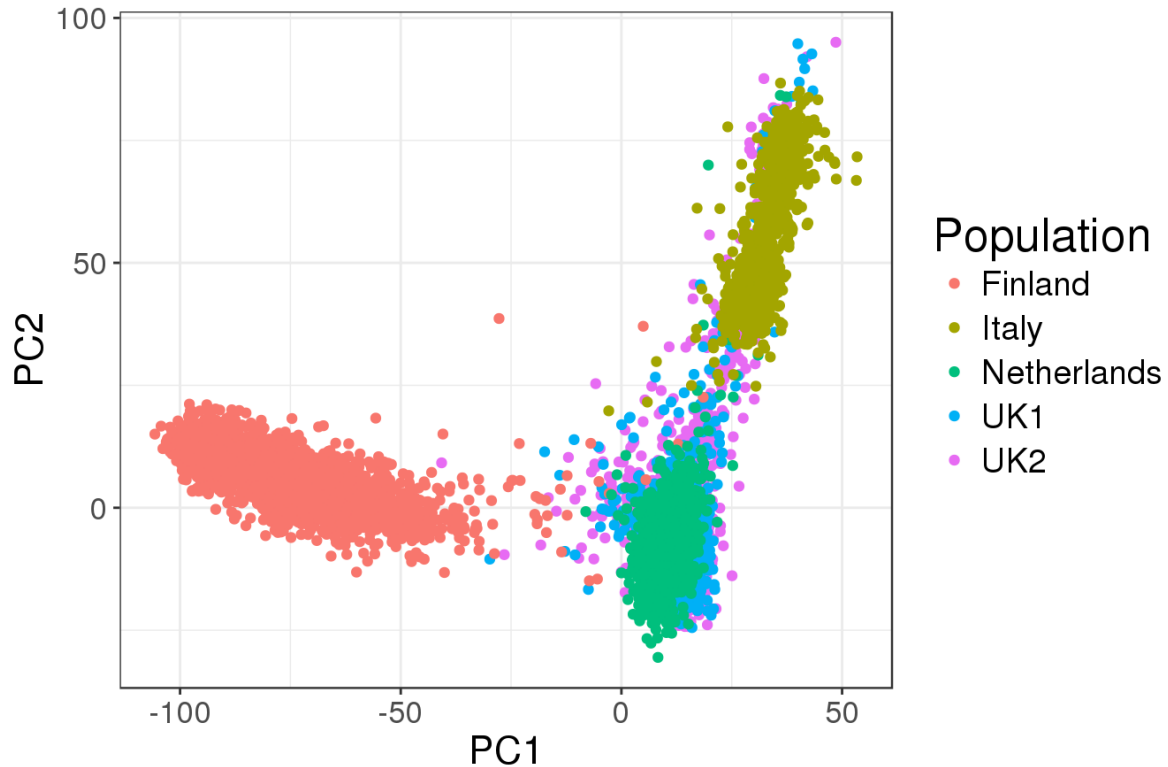
Imagine, that you have a dataset of points. Your goal is to choose orthogonal axes, that describe your data the most informative way. To be precise, we choose first axis in such a way, that maximize the variance (expressiveness) of the projected data. All the following axes have to be orthogonal to the previously chosen ones, while satisfy largest possible variance of the projections.

Let's take a look at the simple 2d data. We have a set of blue points on the plane. We can easily see that the projections on the first axis (red dots) have maximum variance at the final position of the animation. The second (and the last) axis should be orthogonal to the previous one.



This idea could be used in a variety of ways. For example, it might happen, that projection of complex data on the principal plane (only 2 components) bring you enough intuition for clustering. The picture below plots projection of the labeled dataset onto the first two principal components (PC's), we can clearly see, that only two vectors (these PC's) would be enough to differ Finnish people from Italian in particular dataset (celiac disease (Dubois et al. 2010))

Scores of PCA



[source](#)

Problem

The first component should be defined in order to maximize variance. Suppose, we've already normalized the data, i.e. $\sum_i a_i = 0$, then sample variance will become the sum of all squared projections of data points to our vector $\mathbf{w}_{(1)}$, which implies the following optimization problem:

$$\mathbf{w}_{(1)} = \arg \max_{\|\mathbf{w}\|=1} \left\{ \sum_i (\mathbf{a}_{(i)} \cdot \mathbf{w})^2 \right\}$$

or

$$\mathbf{w}_{(1)} = \arg \max_{\|\mathbf{w}\|=1} \{ \|\mathbf{A}\mathbf{w}\|^2 \} = \arg \max_{\|\mathbf{w}\|=1} \{ \mathbf{w}^T \mathbf{A}^T \mathbf{A} \mathbf{w} \}$$

since we are looking for the unit vector, we can reformulate the problem:

$$\mathbf{w}_{(1)} = \arg \max \left\{ \frac{\mathbf{w}^T \mathbf{A}^T \mathbf{A} \mathbf{w}}{\mathbf{w}^T \mathbf{w}} \right\}$$

It is [known](#), that for positive semidefinite matrix $\mathbf{A}^T \mathbf{A}$ such vector is nothing else, but eigenvector of $\mathbf{A}^T \mathbf{A}$, which corresponds to the largest eigenvalue. The following components will give you the same results (eigenvectors).

So, we can conclude, that the following mapping:

$$\Pi = \mathbf{A} \cdot \mathbf{W}$$

where Π is $n \times k$ (data matrix), $\mathbf{A}$ is $n \times d$ (data matrix), $\mathbf{W}$ is $d \times k$ (principal components), $k < d$ (number of principal components).

describes the projection of data onto the k principal components, where $\mathbf{W}$ contains first (by the size of eigenvalues) k eigenvectors of $\mathbf{A}^T \mathbf{A}$.

Now we'll briefly derive how SVD decomposition could lead us to the PCA.

Firstly, we write down SVD decomposition of our matrix:

a_i - значения $a_i \in \mathbb{R}^n$

① нормализуем

$$\sum a_i = 0$$

$$\text{Var } X = \mathbb{E}(X^2) - (\mathbb{E}X)^2$$

$$\sum a_i^T \mathbf{w} = 0$$

$$\|\mathbf{A}\mathbf{w}\|^2 = \langle \mathbf{A}\mathbf{w}, \mathbf{A}\mathbf{w} \rangle = (\mathbf{A}\mathbf{w})^T \mathbf{A}\mathbf{w} = \mathbf{w}^T \mathbf{A}^T \mathbf{A} \mathbf{w}$$

$$\|\mathbf{w}\|^2 = 1 = \langle \mathbf{w}, \mathbf{w} \rangle = \mathbf{w}^T \mathbf{w}$$

$$\mathbf{w}_{(1)} = \text{собств. вект. } (\mathbf{A}^T \mathbf{A})$$

$$\mathbf{Y} \mathbf{e} = \mathbf{W} \mathbf{e}$$

собств. вектор матрицы

сингулярное разложение

$$A = U\Sigma W^T$$

and to its transpose:

$$\begin{aligned} A^T &= (U\Sigma W^T)^T \\ &= (W^T)^T \Sigma^T U^T \\ &= W\Sigma^T U^T \\ &= W\Sigma U^T \end{aligned}$$

Then, consider matrix AA^T :

$$\begin{aligned} A^T A &= (W\Sigma U^T)(U\Sigma V^T) \\ &= W\Sigma I \Sigma W^T \\ &= W\Sigma \Sigma W^T \\ &= W\Sigma^2 W^T \end{aligned}$$

Which corresponds to the eigendecomposition of matrix $A^T A$, where W stands for the matrix of eigenvectors of $A^T A$, while Σ^2 contains eigenvalues of $A^T A$.

At the end:

$$\begin{aligned} \Pi &= A \cdot W = \\ &= U\Sigma W^T W = U\Sigma \end{aligned}$$

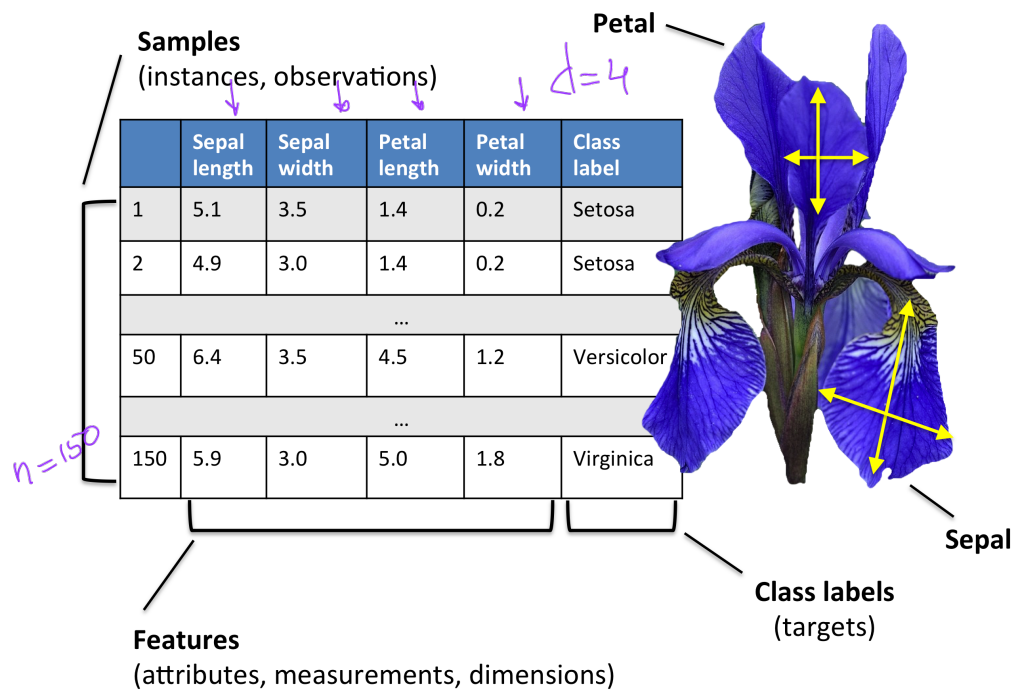
The latter formula provide us with easy way to compute PCA via SVD with any number of principal components:

$$\Pi_r = U_r \Sigma_r$$

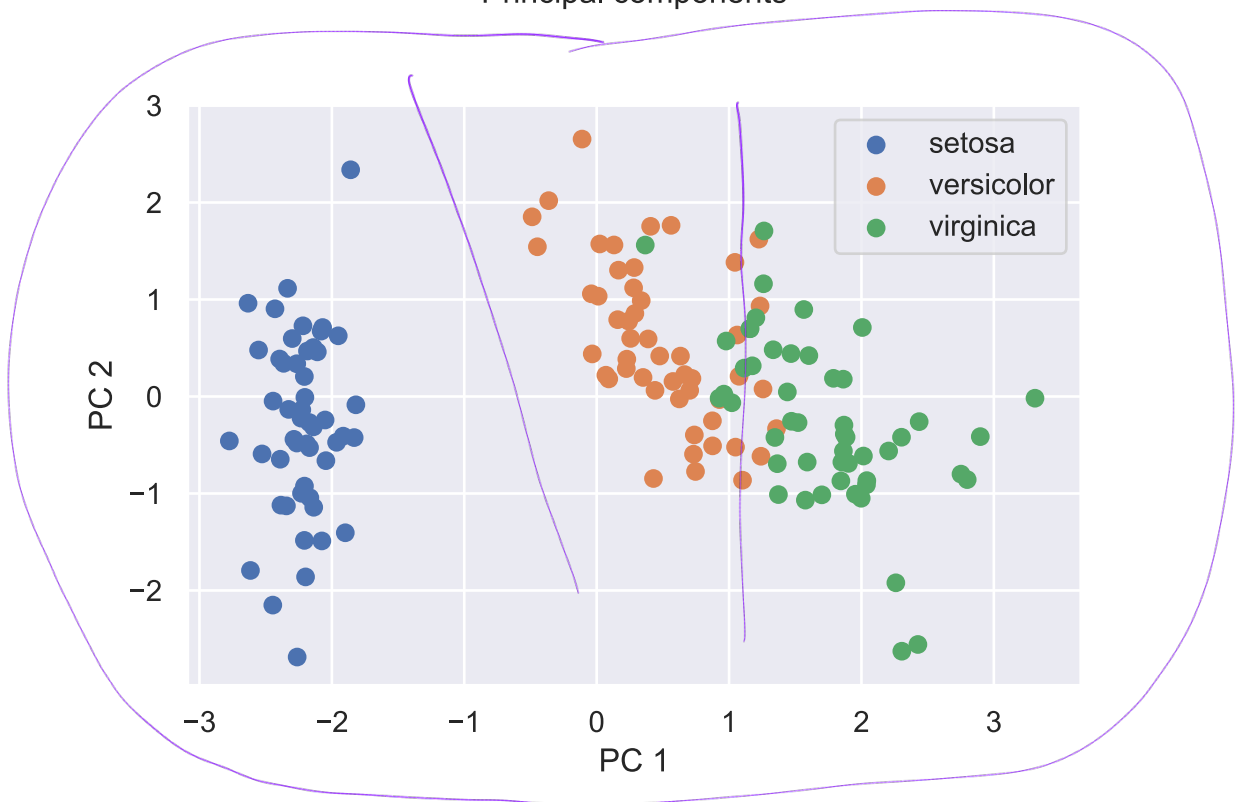
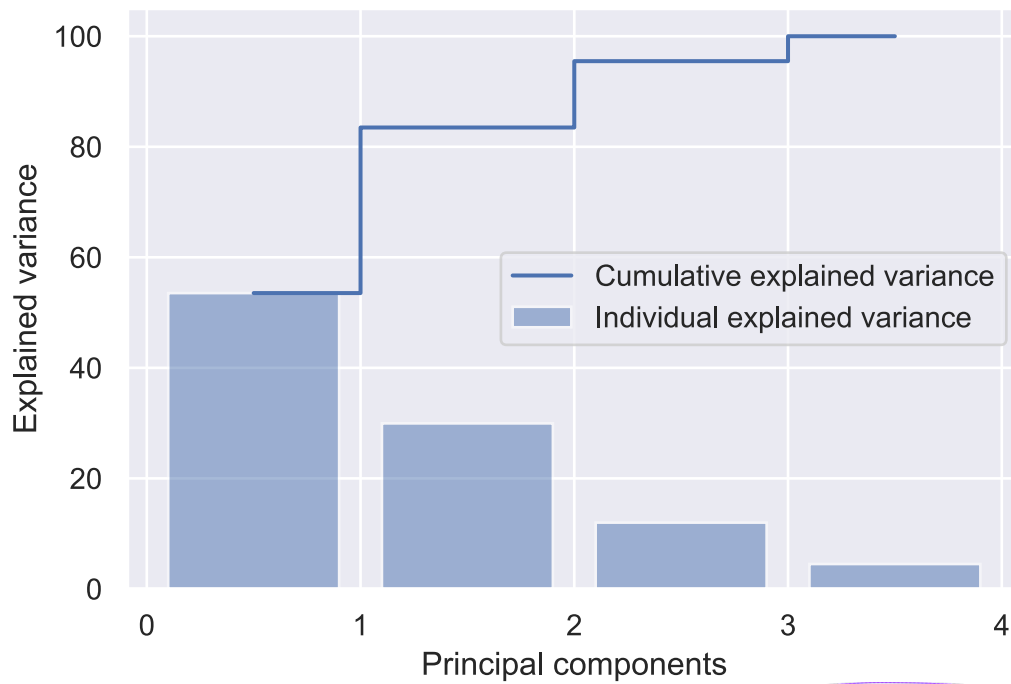
Examples

Iris dataset

Consider the classical Iris dataset



[source](#) We have the dataset matrix $A \in \mathbb{R}^{150 \times 4}$



Code

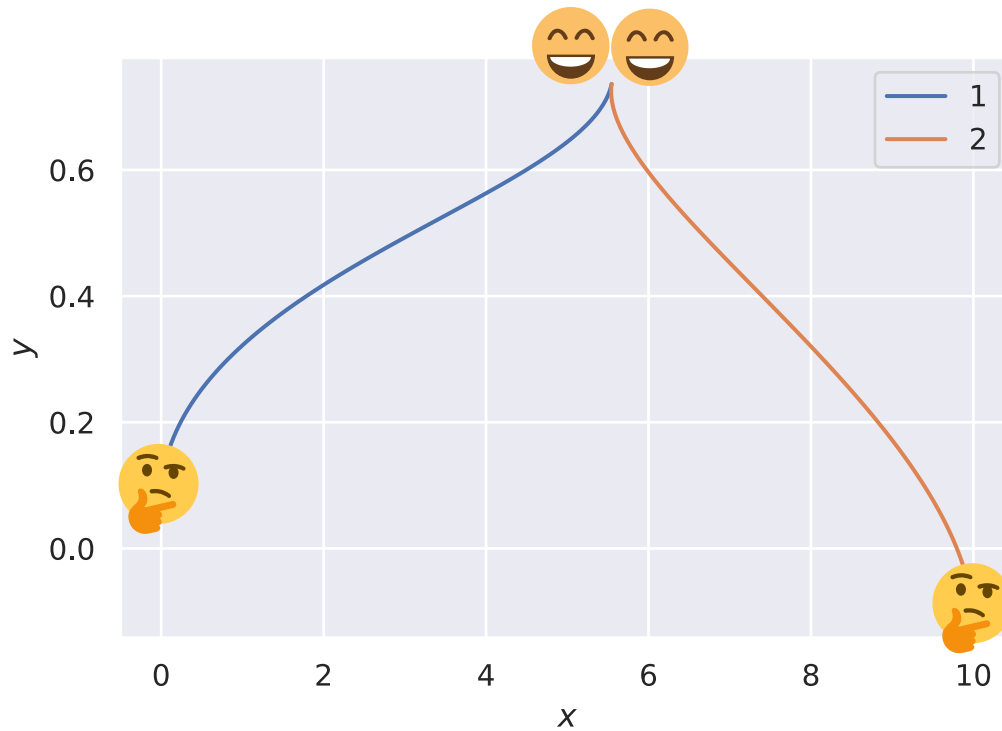
[Open in Colab](#)

Related materials

- [Wikipedia](#)
- [Blog.post](#)
- [Blog.post](#)

Rendezvous problem

Problem



We have two bodies in discrete time: the first is described by its coordinate x_i and its speed v_i , the second has coordinate z_i and speed u_i . Each body has its own dynamics, which we denote as linear systems with matrices A, B, C, D :

$$\begin{aligned}x_{i+1} &= Ax_i + Bu_i \\z_{i+1} &= Cz_i + Dv_i\end{aligned}$$

We want these bodies to meet in future at some point T in such a way, that preserve minimum energy through the path. We will consider only kinetic energy, which is proportional to the squared speed at each point of time, that's why optimization problem takes the following form:

$$\begin{aligned}\min \sum_{i=1}^T \|u_i\|_2^2 + \|v_i\|_2^2 \\ \text{s.t. } x_{t+1} &= Ax_t + Bu_t, \quad t = 1, \dots, T-1 \\ z_{t+1} &= Cz_t + Dv_t, \quad t = 1, \dots, T-1 \\ x_T &= z_T\end{aligned}$$

Problem of this type arise in space engineering - just imagine, that the first body is the spaceship, while the second, say, Mars.

Code

[Open in Colab](#)

References

- [Jupyter notebook](#) by A. Katrutsa

Travelling salesman problem

Problem

Suppose, we have N points in $\mathbb{R}^d$ Euclidian space (for simplicity we'll consider and plot case with $d = 2$). Let's imagine, that these points are nothing else but houses in some 2d village. Salesman should find the shortest way to go through the all houses only once.



That is, very simple formulation, however, implies NP - hard problem with the factorial growth of possible combinations. The goal is to minimize the following cumulative distance:

$$d = \sum_{i=1}^{N-1} \|x_{y(i+1)} - x_{y(i)}\|_2 \rightarrow \min_y,$$

where x_k is the k -th point from N and y stands for the N - dimensional vector of indices, which describes the order of path. Actually, the problem could be [formulated](#) as an LP problem, which is easier to solve.

Genetic (evolution) algorithm

Our approach is based on the famous global optimization algorithm, known as evolution algorithm.

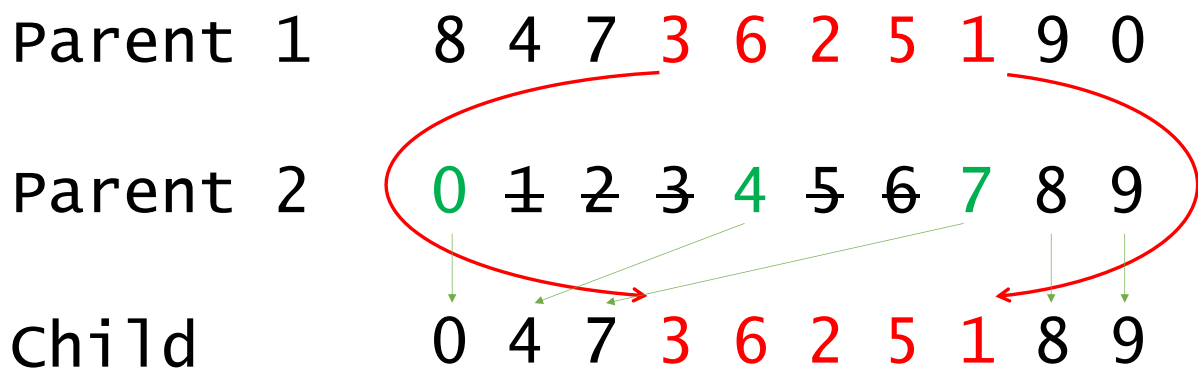
Population and individuals

Firstly we need to generate the set of random solutions as an initialization. We will call a set of solutions $\{y_k\}_{k=1}^n$ as *population*, while each solution is called *individual* (or creature).

Each creature contains integer numbers $1, \dots, N$, which indicates the order of bypassing all the houses. The creature, that reflects the shortest path length among the others will be used as an output of an algorithm at the current iteration (generation).

Crossing procedure

Each iteration of the algorithm starts with the crossing (breed) procedure. Formally speaking, we should formulate the mapping, that takes two creature vectors as an input and returns its offspring, which inherits parents properties, while remaining consistent. We will use [ordered crossover](#) as such procedure.

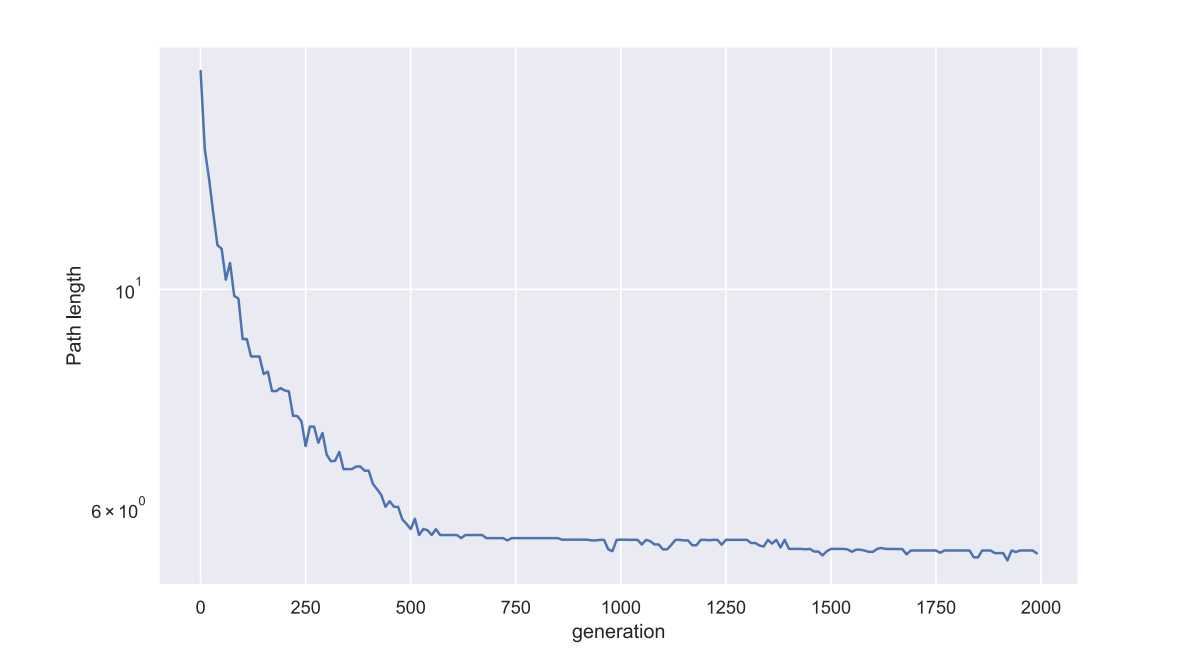
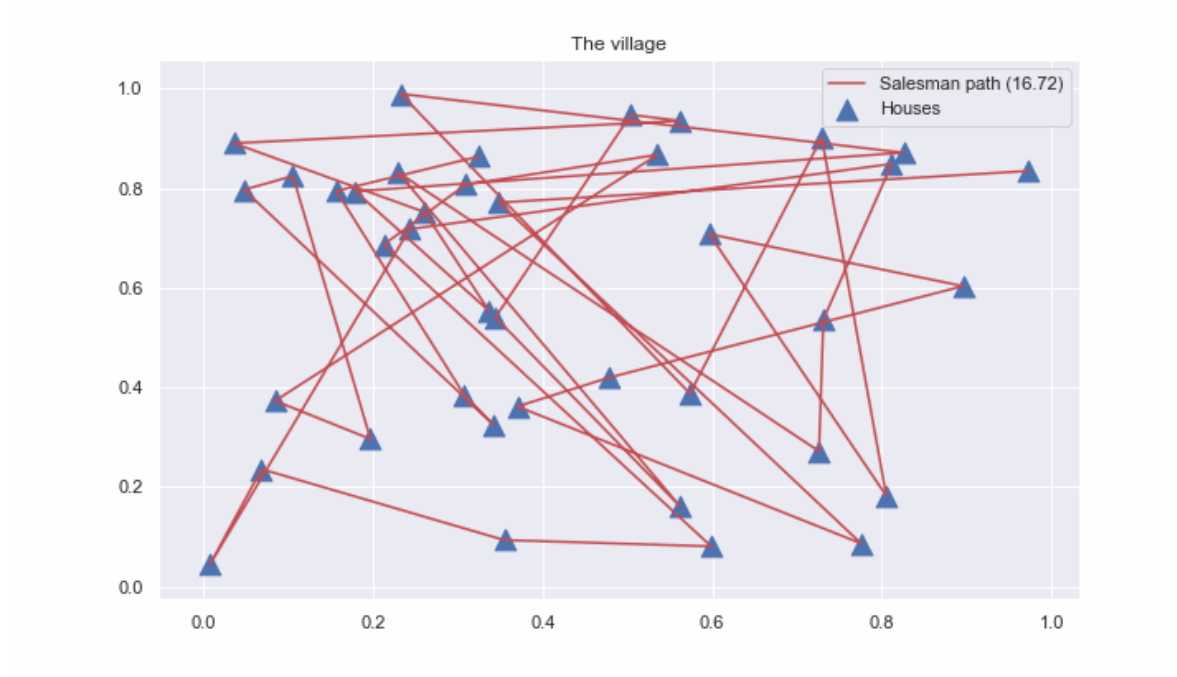


Mutation

In order to give our algorithm some ability to escape local minima we provide it with mutation procedure. We simply swap some houses in an individual vector. To be more accurate, we define mutation rate (say, 0.05). On the one hand, the higher the rate, the less stable the population is, on the other, the smaller the rate, the more often algorithm gets stuck in the local minima. We choose $\text{mutation_rate} \cdot n$ individuals and in each case swap random $\text{mutation_rate} \cdot N$ digits.

Selection

At the end of the iteration we have increased population (due to crossing results), than we just calculate total path distance to each individual and select top n of them.



In general, for any $c > 0$, where d is the number of dimensions in the Euclidean space, there is a polynomial-time algorithm that finds a tour of length at most $(1 + \frac{1}{c})$ times the optimal for geometric instances of TSP in

$$\mathcal{O}\left(N(\log N)^{(\mathcal{O}(c\sqrt{d}))^{d-1}}\right)$$

Code

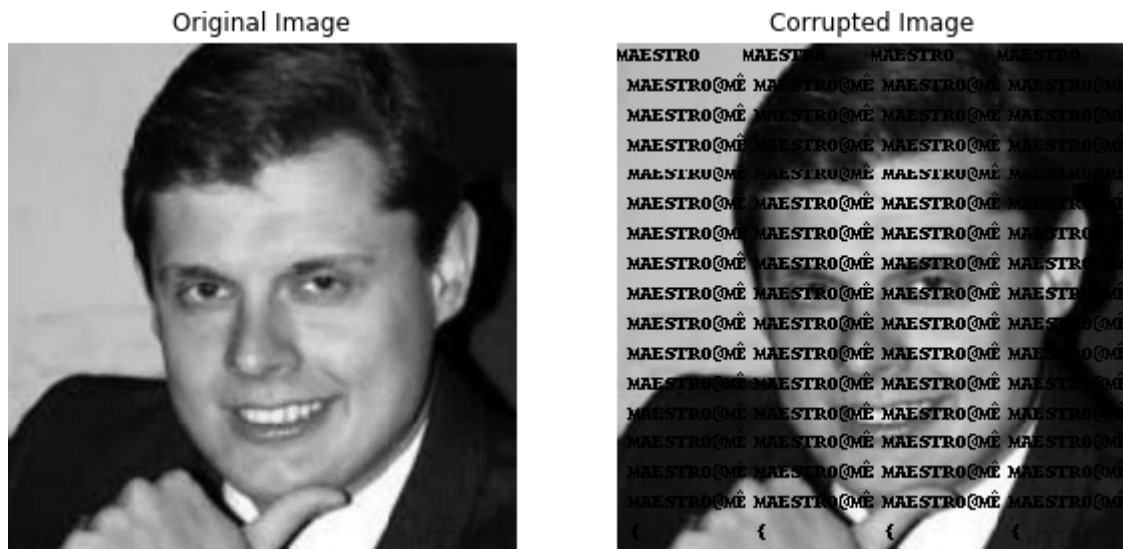
[Open in Colab](#)

References

- [General information about genetic algorithms](#)
- [Wiki](#)

Total variation in-painting

Problem



Grayscale image

A grayscale image is represented as an $m \times n$ matrix of intensities U^{orig} (typically between the values 0 and 255). We are given all the values of corrupted picture, but some of them should be preserved as is through the recovering procedure: $U_{ij}^{corr} \forall (i, j) \in K$, where

$K \subset \{1, \dots, m\} \times \{1, \dots, n\}$ is the set of indices corresponding to known pixel values. Our job is to in-paint the image by guessing the missing pixel values, i.e., those with indices not in K . The reconstructed image will be represented by $U \in \mathbb{R}^{m \times n}$, where U matches the known pixels, i.e. $U_{ij} = U_{ij}^{corr}$ for $(i, j) \in K$.

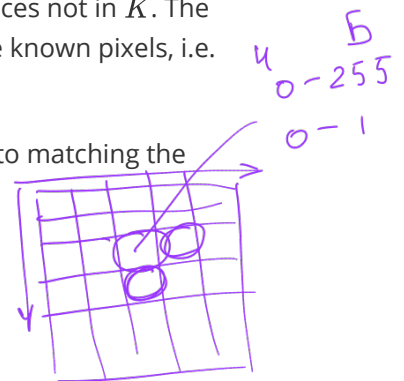
The reconstruction U is found by minimizing the total variation of U , subject to matching the known pixel values. We will use the l_2 total variation, defined as

$$\mathbf{tv}(U) = \sum_{i=1}^{m-1} \sum_{j=1}^{n-1} \left\| \begin{bmatrix} U_{i+1,j} - U_{ij} \\ U_{i,j+1} - U_{ij} \end{bmatrix} \right\|_2.$$

So, the final optimization problem will be written as follows:

$$\begin{aligned} \mathbf{tv}(U) &\rightarrow \min_{U \in \mathbb{R}^{m \times n}} \\ \text{s.t. } &U_{ij} = U_{ij}^{corr}, (i, j) \in K \end{aligned}$$

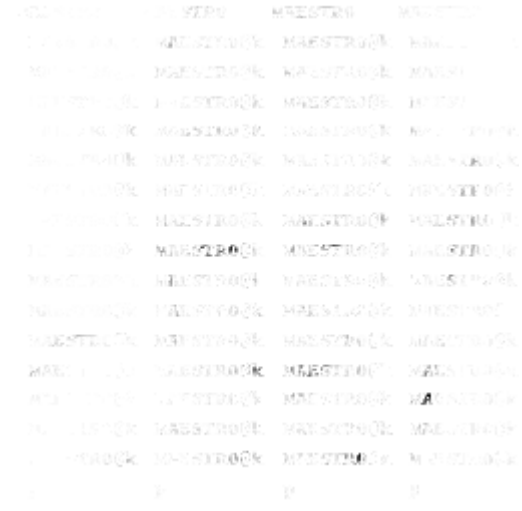
The crucial thing about this problem is defining set of known pixels K . There are some heuristics: for example, we could state, that each pixel with color similar (or exactly equal) to the color of text is unknown. The results for such approach are presented below:



In-Painted Image



Difference Image



Color image

For the color case we consider in-painting problem in a slightly different setting: destroying some random part of all pixels. In this case the image itself is 3d tensor (we convert all others color schemes to the RGB). As it was in the grayscale case, we construct the mask K of known pixels for all color channels uniformly, based on the principle of similarity of particular 3d pixel to the vector $[0, 0, 0]$ (black pixel). The results are quite promising - note, that we have no information about the original picture, but assumption, that corrupted pixels are black. For the color picture we just sum all tv's on the each channel:

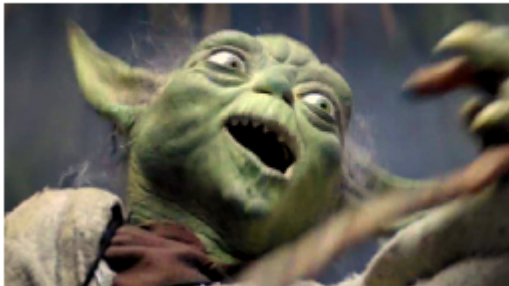
$$\mathbf{tv}(U) = \sum_{k=1}^3 \sum_{i=1}^{m-1} \sum_{j=1}^{n-1} \left\| \begin{bmatrix} U_{i+1,j}^k - U_{ij}^k \\ U_{i,j+1}^k - U_{ij}^k \end{bmatrix} \right\|_2.$$

Then, we need to write down optimization problem to be solved:

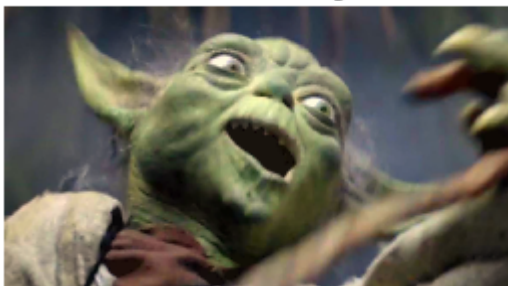
$$\begin{aligned} \mathbf{tv}(U) &\rightarrow \min_{U \in \mathbb{R}^{m \times n \times 3}} \\ \text{s.t. } U_{ij}^k &= U_{ij}^{corr,k}, \quad (i, j) \in K, \quad k = 1, 2, 3 \end{aligned}$$

Results are presented below (these computations are really take time):

Original Image



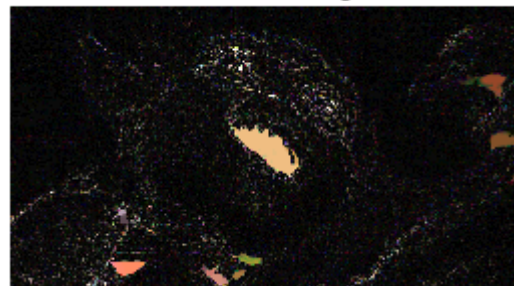
In-Painted Image



Corrupted Image. Corruption level 0.707

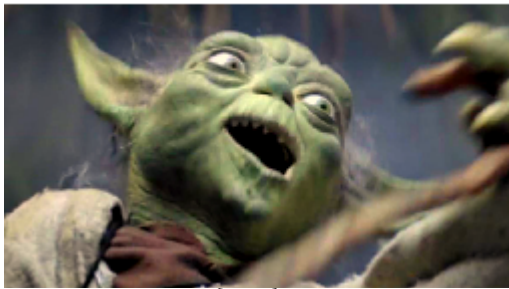


Difference Image



It is not that easy, right?

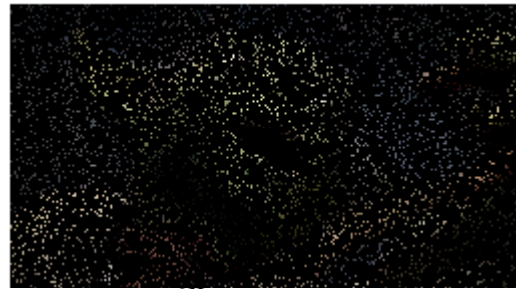
Original Image



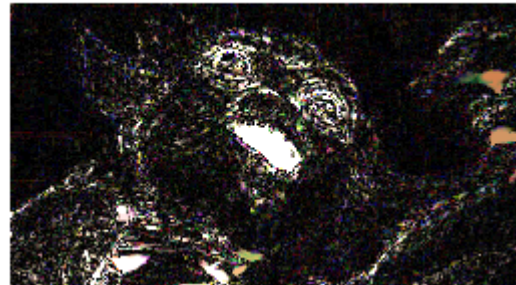
In-Painted Image



Corrupted Image. Corruption level 0.9

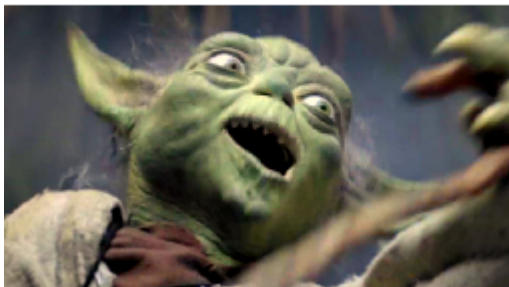


Difference Image



Only 5% of all pixels are left:

Original Image



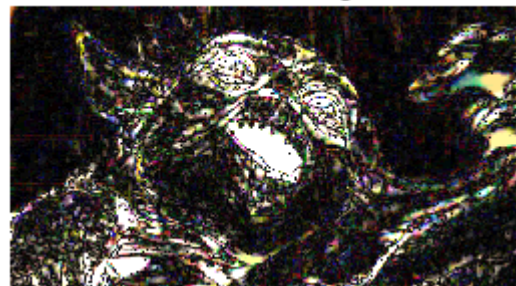
In-Painted Image



Corrupted Image. Corruption level 0.95

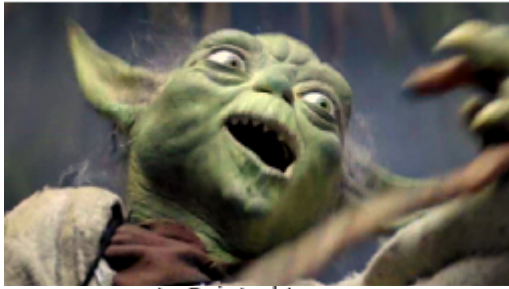


Difference Image



What about 1% of all pixels?

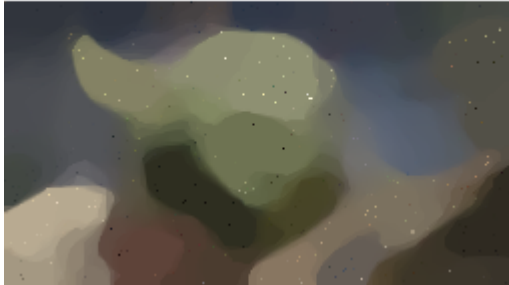
Original Image



Corrupted Image. Corruption level 0.99



In-Painted Image



Difference Image



Code

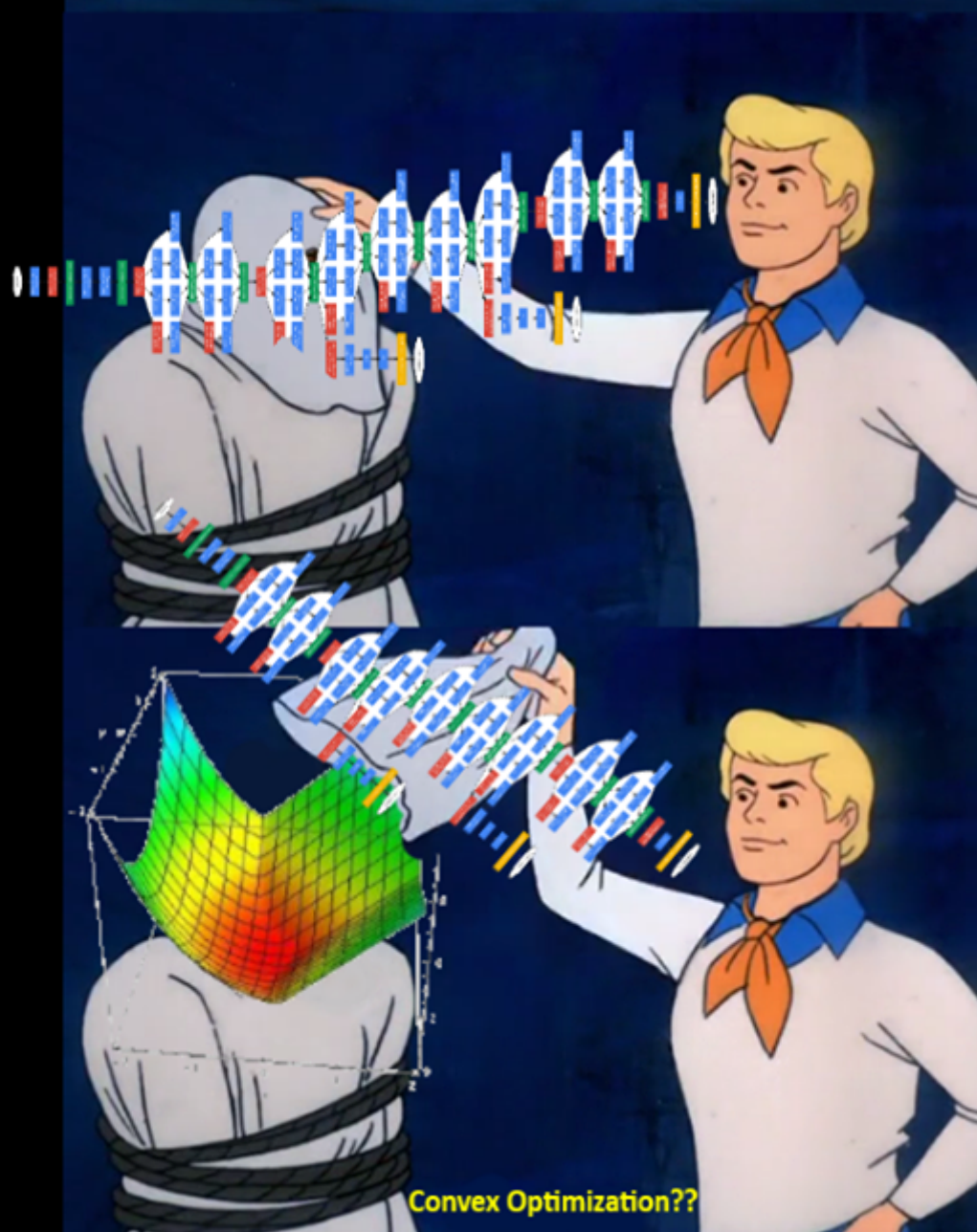
[Open in Colab](#)

References

- [CVXPY documentation](#)
- [Interactive demo](#)

Deep learning

Problem



A lot of practical task nowadays are being solved by the deep learning approach, which is usually implies finding local minimum of a non - convex function, that generalizes well (enough 😊). The goal of this short text is to provide you an importance of the optimization behind neural network training.

Cross entropy

One of the most commonly used loss functions in classification tasks is the normalized categorical cross entropy in K class problem:

$$L(\theta) = -\frac{1}{n} \sum_{i=1}^n (y_i^\top \log(h_\theta(x_i)) + (1 - y_i)^\top \log(1 - h_\theta(x_i))), \quad h_\theta^k(x_i) = \frac{e^{\theta_k^\top x_i}}{\sum_{j=1}^K e^{\theta_j^\top x_i}}$$

Since in Deep Learning tasks the number of points in a dataset could be really huge, we usually use [Stochastic gradient descent based approaches](#) as a workhorse.

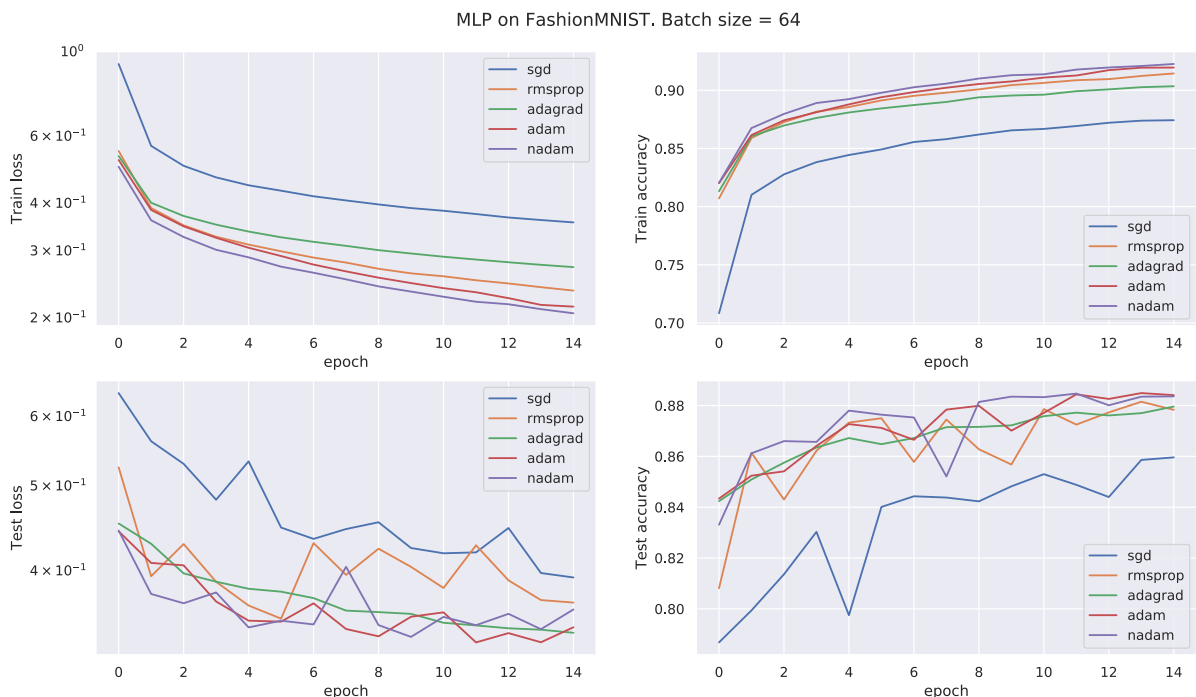
In such algorithms one uses the estimation of a gradient at each step instead of the full gradient vector, for example, in cross entropy we have:

$$\nabla_\theta L(\theta) = \frac{1}{n} \sum_{i=1}^n (h_\theta(x_i) - y_i) x_i^\top$$

The simplest approximation is statistically judged unbiased estimation of a gradient:

$$g(\theta) = \frac{1}{b} \sum_{i=1}^b (h_\theta(x_i) - y_i) x_i^\top \approx \nabla_\theta L(\theta)$$

where we initially sample randomly only $b \ll n$ points and calculate sample average. It can be also considered as a noisy version of the full gradient approach.



Code

[Open in Colab](#)

References

- [Optimization for Deep Learning Highlights in 2017](#)

- [An overview of gradient descent optimization algorithms](#)

Methods

General formulation

$$\begin{aligned} & \min_{x \in \mathbb{R}^n} f(x) \\ & \text{s.t. } g_i(x) \leq 0, i = 1, \dots, m \\ & \quad h_j(x) = 0, j = 1, \dots, k \end{aligned}$$

such that
subject to

целевая функция

$$\min_{x \in S} f(x)$$

допустимое мн-во

Some necessary or/and sufficient conditions are known (See Optimality conditions. KKT and Convex optimization problem' %})

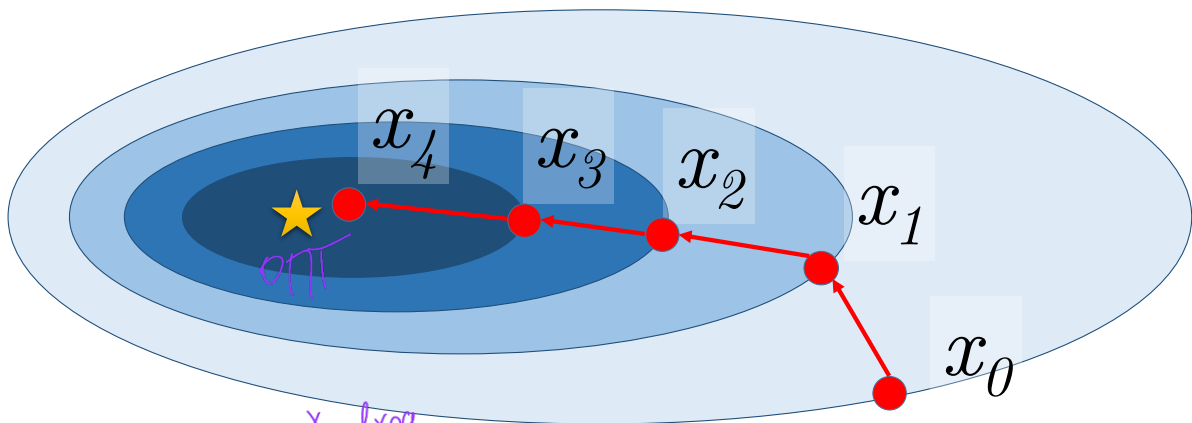
- In fact, there might be very challenging to recognize the convenient form of optimization problem.
- Analytical solution of KKT could be inviable

Iterative methods

Typically, the methods generate an infinite sequence of approximate solutions

$$\{x_t\},$$

which for a finite number of steps (or better - time) converges to an optimal (at least one of the optimal) solution x_* .



```
def GeneralScheme(x, epsilon):
    while not StopCriterion(x, epsilon):
        oracleResponse = RequestOracle(x)
        x = NextPoint(x, oracleResponse)
    return x
```

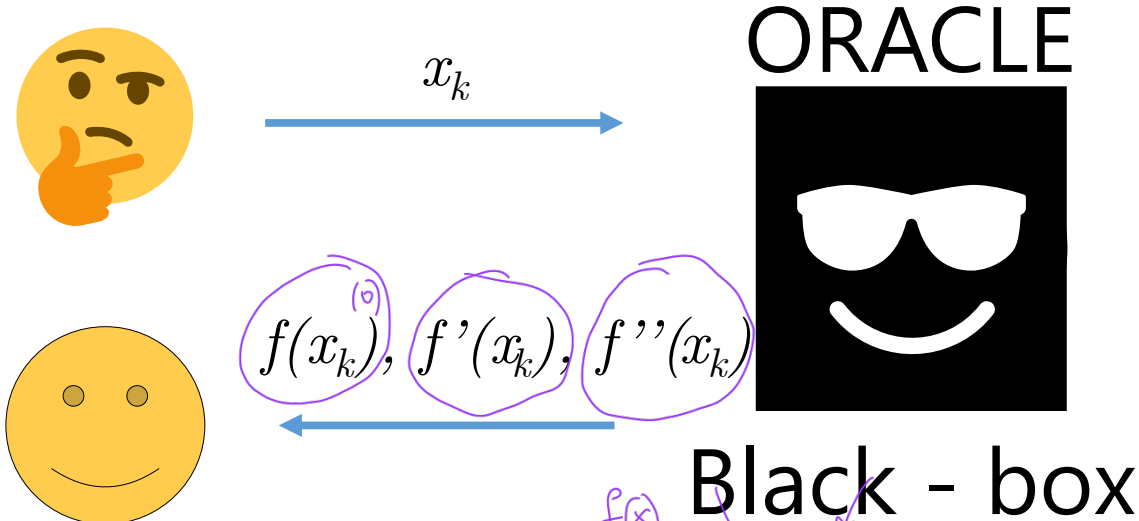
x_0 брoй init

крит. ост.

оракл

итер. метод (GD)

Oracle conception



Complexity Challenges

Unsolvability

In general, **optimization problems are unsolvable.** $\neg(\exists)\forall$

Consider the following simple optimization problem of a function over unit cube:

$$\begin{aligned} \min_{x \in \mathbb{B}^n} f(x) \\ \text{s.t. } x \in \mathbb{B}^n \end{aligned}$$



We assume, that the objective function $f(\cdot) : \mathbb{R}^n \rightarrow \mathbb{R}$ is Lipschitz continuous on $\mathbb{B}^n$.

$$\forall x, y \in \mathbb{B}^n \quad |f(x) - f(y)| \leq L \|x - y\|$$

with some constant L (Lipschitz constant). Here $\mathbb{B}^n$ - the n - dimensional unit cube $\mathbb{B}^n = \{x \in \mathbb{R}^n \mid 0 \leq x_i \leq 1, i = 1, \dots, n\}$.

Our goal is to find such $\tilde{x} : |f(\tilde{x}) - f^*| \leq \varepsilon$ for some positive ε . Here f^* is the global minimizer of the problem. Uniform grid with p points on each dimension guarantees at least this quality

$$\frac{|f(x) - f(y)|}{\|x - y\|} \leq L \quad \|\tilde{x} - x_*\|_\infty \leq \frac{1}{2p}$$

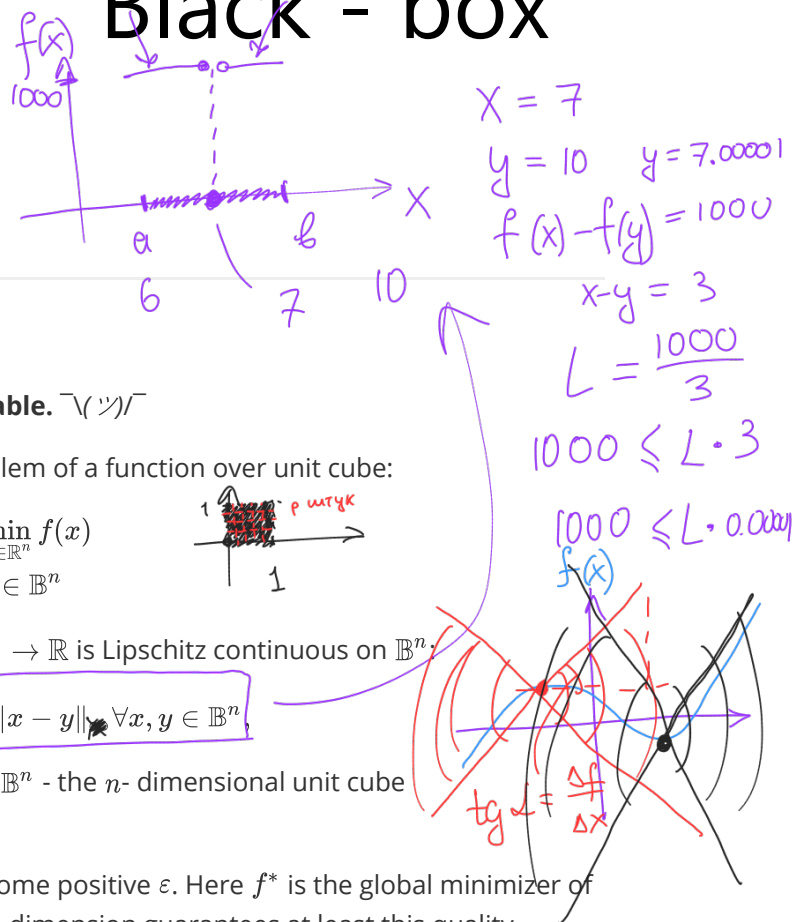
which means, that

$$|f(\tilde{x}) - f(x_*)| \leq \frac{L}{2p} = \varepsilon = 0.01$$

Our goal is to find the p for some ε . So, we need to sample $(\frac{L}{2\varepsilon})^n$ points, since we need to measure function in p^n points. Doesn't look scary, but if we'll take $L = 2, n = 11, \varepsilon = 0.01$, computations on the modern personal computers will take 31,250,000 years.

$$\left(\frac{L}{2\varepsilon}\right)^n$$

31,250,000
годебуны
n = 12



Stopping rules

- Argument closeness:

$$\|x_k - x_*\|_2 < \varepsilon$$

- Function value closeness:

$$\|f_k - f^*\|_2 < \varepsilon$$

- Closeness to a critical point

$$\|f'(x_k)\|_2 < \varepsilon$$

But x_* and $f^* = f(x_*)$ are unknown!

Sometimes, we can use the trick:

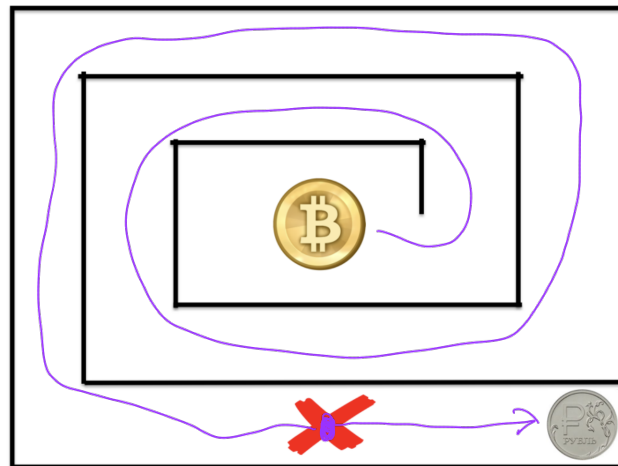
$$\|x_{k+1} - x_k\| = \|x_{k+1} - x_k + x_* - x_*\| \leq \|x_{k+1} - x_*\| + \|x_k - x_*\| \leq 2\varepsilon$$

Note: it's better to use relative changing of these values, i.e.

$$\frac{\|x_{k+1} - x_k\|_2}{\|x_k\|_2}$$

крит. точка
MAX
MIN
сег. точки
 $\leq \varepsilon$

Local nature of the methods



Problem classifications

Methods classifications

Speed of convergence

- Sublinear

$$\|x_{k+1} - x_*\|_2 \leq Ck^\alpha,$$

where $\alpha < 0$ and $0 < C < \infty$

- Linear

$$\|x_{k+1} - x_*\|_2 \leq Cq^k \quad \text{or} \quad \|x_{k+1} - x_*\|_2 \leq q\|x_k - x_*\|_2,$$

where $q \in (0, 1)$ and $0 < C < \infty$

- Superlinear

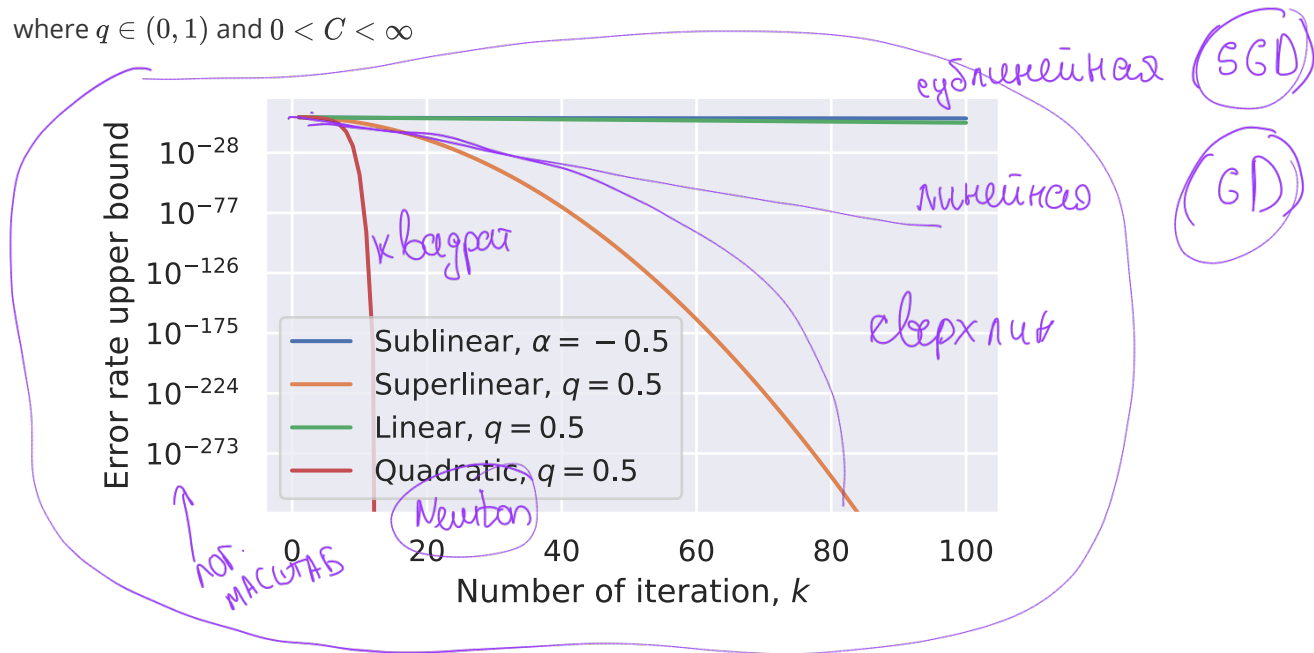
$$\|x_{k+1} - x_*\|_2 \leq Cq^{k^2} \quad \text{or} \quad \|x_{k+1} - x_*\|_2 \leq C_k\|x_k - x_*\|_2,$$

where $q \in (0, 1)$ or $0 < C_k < \infty, C_k \rightarrow 0$

- Quadratic

$$\|x_{k+1} - x_*\|_2 \leq Cq^{2^k} \quad \text{or} \quad \|x_{k+1} - x_*\|_2 \leq C\|x_k - x_*\|_2^2,$$

where $q \in (0, 1)$ and $0 < C < \infty$



Root test

Пусть $\{r_k\}_{k=m}^{\infty}$ — последовательность неотрицательных чисел, сходящаяся к нулю, и пусть

$$\alpha = \limsup_{k \rightarrow \infty} r_k^{1/k}$$

- Если $0 \leq \alpha < 1$, то $\{r_k\}_{k=m}^{\infty}$ имеет линейную сходимость с константой α .
- В частности, если $\alpha = 0$, то $\{r_k\}_{k=m}^{\infty}$ имеет сверхлинейную сходимость.
- Если $\alpha = 1$, то $\{r_k\}_{k=m}^{\infty}$ имеет сублинейную сходимость.
- Случай $\alpha > 1$ невозможен.

Ratio test

Пусть $\{r_k\}_{k=m}^{\infty}$ — последовательность строго положительных чисел, сходящаяся к нулю. Пусть

$$\alpha = \lim_{k \rightarrow \infty} \frac{r_{k+1}}{r_k}$$

- Если существует α и при этом $0 \leq \alpha < 1$, то $\{r_k\}_{k=m}^{\infty}$ имеет линейную сходимость с константой α
- В частности, если $\alpha = 0$, то $\{r_k\}_{k=m}^{\infty}$ имеет сверхлинейную сходимость
- Если α не существует, но при этом $q = \lim_{k \rightarrow \infty} \sup \frac{r_{k+1}}{r_k} < 1$, то $\{r_k\}_{k=m}^{\infty}$ имеет линейную сходимость с константой, не превосходящей q .
- Если $\lim_{k \rightarrow \infty} \inf \frac{r_{k+1}}{r_k} = 1$, то $\{r_k\}_{k=m}^{\infty}$ имеет сублинейную сходимость.
- Ситуация $\lim_{k \rightarrow \infty} \inf \frac{r_{k+1}}{r_k} > 1$ невозможна.
- Во всех остальных случаях (т. е. когда $\lim_{k \rightarrow \infty} \inf \frac{r_{k+1}}{r_k} < 1 \leq \lim_{k \rightarrow \infty} \sup \frac{r_{k+1}}{r_k}$) нельзя утверждать что-либо конкретное о скорости сходимости $\{r_k\}_{k=m}^{\infty}$

References

- [CMC seminars \(ru\)](#)
- [CVXPY documentation](#)

- [Colab notebook](#) with basic examples