

Machine learning

Vladimír Boža

boza@fmph.uniba.sk

M-25 (ask by email about office hours)

compbio.fmph.uniba.sk/vyuka/ml

watch out for announcements on FB

Schedule

Tuesday: 11:30 A

Wednesday: 13:10 B (occasional tutorials in H6)

Teoretical homework (2)	14%
Practical homework (8, points from best 4)	16%
Project (3 days before exam, latest 31.1.2019)	30%
Exam	40%

Need 40% from exam to pass.

Projekt

Pick some application of ML (web page will have some ideas later).

Implement, evaluate

Several options for your output:

- 1 Summary in form of conference paper (5 - 15 pages).
- 2 Small web application demonstrating your project, which includes short summary (methods, results)
- 3 Get at least 10th place in serious (money involved) Kaggle competition + 1 page summary of your methods.

In all cases, you should publish your source codes (Github preferred).

Part of exam would be about methods in your project.

Prereqs

Linear algebra (basics).

Derivatives

Programming (Python)

Being a honeybadger

Applications of ML

Supervised learning

- Speech recognition, image recognition
- Machine translation
- Recommendations of movies, books, ...
- House price prediction
- Marketing predictions (conversion rates, ...)

Unsupervised learning

- Signal decomposition
- Clustering
- Visualization of data

Reinforcement learning

- Games (go, chess, ...)
- Robotics

```
def swap(a, b):  
    return b, a  
  
c, d = swap(a, b)  
  
for i in range(0, 10):  
    if i % 2 == 0:  
        print "parne"
```

Numpy

```
>>> import numpy as np
>>> x = np.array([[1,2],[3,4]], float)
>>> y = np.array([[0,1],[1,0]], float)
>>> x
array([[ 1.,  2.],
       [ 3.,  4.]])
>>> x.shape
(2, 2)
>>> x + y
array([[ 1.,  3.],
       [ 4.,  4.]])
>>> x * y
array([[ 0.,  2.],
       [ 3.,  0.]])
```

Numpy

```
>>> import numpy as np
>>> x = np.array([[1,2],[3,4]], float)
>>> y = np.array([[0,1],[1,0]], float)
>>> x.dot(y)
array([[ 2.,  1.],
       [ 4.,  3.]])
>>> np.exp(x)
array([[ 2.71828183,  7.3890561 ],
       [20.08553692, 54.59815003]])
>>> x * 2
array([[ 2.,  4.],
       [ 6.,  8.]])
```

Supervised learning

Data

Set of n pairs x - input, y - expected output. This is called training set.

Goal

Predict output for new x .

Note

In most cases, the $\vec{x}$ is a vector with m values (**attributes**) and y is scalar value.

Example: house prices

$\vec{x}$			y
Size	# of rooms	Distance from city centre	Price
122	3	0.5	400000
39	1	6	76000
67	3	2	175000
88	2	4	???

Nearest neighbour

- Got a new input $\vec{x}_t$.
- From training examples, pick one $(\vec{x}, y)$ where $\vec{x}$ is the most similar to $\vec{x}_t$. Predict y .
- (Modification: pick k most similar, predict average.)

Good

Good accuracy, when we have a lots of data.

Bad

Slow, bulky (we need to store whole training set in fast memory). Need to define similarity. Sensitive to scaling and irrelevant attributes.

Picking from set of hypothesis

Input

Set of examples $(\vec{x}_1, y_1), \dots, (\vec{x}_n, y_n)$.

Set of hypotheses

$$H \subset \mathbb{R}^m \rightarrow \mathbb{R}$$

Error function

Pick hypothesis $h \in H$, which gives the lowest error.: $\sum_{i=1}^t \text{err}(h(\vec{x}_i), y_i)$, where err is an **error function**.

Simple linear regression

One attribute (flat size).

Hypothesis set: $H = \{h_{\Theta}(x) = \Theta_0 + \Theta_1 x\}$

Error function: $\text{err}(y_p, y) = (y_p - y)^2$

