

Python recap

①

Objects $\rightarrow$ Mutable, $a = 5$
 $\rightarrow$ Immutable, $a = (1, 3, 5)$, dictionaries

Char. of objects: $\left\{ \begin{array}{l} \text{type : number, string, list, ...} \\ \text{value : data value contained} \\ \text{identity : related with CS} \end{array} \right.$

Attributes: data, functions associated with them. Name of attr follows the name of obj, following a dot.

$\rightarrow$ $\left\{ \begin{array}{l} \text{Data attributes: Data attached to an object.} \\ \text{Methods: Function attached to an object. Typically it performs a function on the obj.} \end{array} \right.$

Example: `import numpy as np`

`x = np.array([1, 3, 5])`

~~`np.array([1, 3, 5])`~~

Data attribute: `x.shape` $\rightarrow$ # elmts in x
no operation

$\rightarrow$ Method: `x.mean()` : mean of values
there is operation

Modules: Libraries of code, eg math, numpy, ...

`dir(str)` gives a list of all available methods for a string

`help(name.upper)` gives info about the ^{method} `upper` applied to the string name.

Numbers: Integer, floating point, complex.

• $2 + 5$, $2 * 5$, $2 ** 5 = 2^5$, $6 / 7$, $6 // 7 = \text{parte entera de } 6/7$

• In the interactive mode, $-$ is the value of the last value returned.

• `math.factorial(6) = 6!`

• `import random`, `random.choice([---])` gives a random entry of `[---]`

• Booleans: `True`, `False`, `and`, `or`, `not`, `<`, `<=`, `==`, `>=`, `!=` ($\neq$)

In lists, `==` compares element by element, whereas `is` compares whether the objects are the same.

Sequences: lists, tuples, range objects, ~~strings~~

`S[0]`, = location 1st, `S[0:2]` = list w/ loc. 1st & 2nd

• Lists: Mutable sequences, strings are not.

`S = [2, 4, 6, 8]`

~~mutable~~ `S.append(10) = [2, 4, 6, 8, 10]` ↖ new S

`T = [10, 14, 16]`

`S + T = [2, 4, ..., 16]`

`S.reverse()`, `S = [10, 8, ..., 2]`

`S.sort()`, `S = [8, 2, 10, 6, 4]`

(There is a function also for this: `sorted()`)

`len(S) = 5` (length) ~~length~~

• Tuples: ~~Base~~ immutable sequences.

$T = (1, 3, 5, 7)$, can ~~not~~⁺ and access $T[0] = 1$.

$x = 12.15$
 $y = 75.2$, ~~the~~ coordinate $= (x, y)$ is a tuple.

I can do $(c1, c2, c3, c4) = T$, and then $c1 = 1$
 $c2 = 3$

Say now that $P = [(0, 0), (5, 9), \dots, (4, 21)]$ list of tuples:

Can do: $\left| \begin{array}{l} \text{for } (a, b) \text{ in } P: \\ \quad \text{print}(a, b) \end{array} \right| \quad \left(\text{for } a, b \text{ in } P \right)$ ~~also~~ works

$c = (2)$ is an integer, $c = (2,)$ is a tuple w/ 1 object

$T.count(7)$ counts the number of 7 in T.

$sum(T) = 1 + 3 + 5 + 7$

• Range: immutable, it is a type!!

$range(5)$ gives $range(0, 5)$. To see the content do

$list(range(5)) = [0, \dots, 4]$

$list(range(1, 6)) = [1, \dots, 5]$

$list(range(1, 13, 2)) = [1, 3, 5, \dots, 11]$

Strings: Immutable sequences of characters

$S = \text{"Python"} , \quad \text{len}(S) = 6 , \quad S[0] = 'P' , \quad S[-1] = 'n'$

$S[0:3] = \text{'Pyt'} \quad S[-3:0] = \text{'hon'}$ (last 3 char)

$\text{"y" in } S : \text{True}$

$\text{"hello"} + \text{"word"} = \text{"helloworld"}$

$\text{name} = \text{"Tina Fey"} , \quad \text{name.replace("T", "t")} = \text{"tine Fey"}$

$\text{names} = \text{name.split(" ")} \quad (\text{splits when there is "space"})$

$\text{names} = [\text{"Tine"}, \text{"Fey"}]$

$\text{"tone".upper()} = \text{'TONE'}$
 lower

$\text{" ".join(list of words)} = \text{word}_1 \text{ word}_2 \text{ word}_3 \dots$

Sets: Unordered collections of hashable objects
means that ~~the~~ you can use them for immutable obj.

↳ Frozen set: immutable

Set: mutable

$\text{ids} = \text{set}([1, \dots, 7])$

turns the list $[1, \dots, 7]$ into a set w/ 7 elements

$\text{ids} : \{1, 7, \dots\}$

$\text{ids.add}(8) = \{1, \dots, 8\}$

③

Say now $ids = \text{set}(\text{range}(10))$, $ids = \{0, \dots, 9\}$

males = set([1, 3, 6, 7])

$$f_{\text{female}} = \{0, 2, 4, \dots\}$$
$$(1 := 0, \quad - := 1, \quad 2 := 0)$$
$$(1 := 0, \quad - := 1, \quad 2 := 0)$$

wand = "antidestabilisierungs" -

key len (letters). gives the # of unique letters.

$x.$ is subset (y) determines whether $x \subseteq y$.

Dictionaries: Maps from key objects to value objects.

↳ curly braces

$$\text{age} = \{ \text{"Tim"}: 29, \text{"Jim"}: 31, \text{"Pam"}: 27, \text{"Sam"}: 35 \}$$

```
age["Pam"] = 27
```

$$\text{age}["Tim"] = \text{age}["Tim"] + 1$$
$$1 + = 1$$

names = age.keys() , it is a "dict-key" type object.

age["Tom"] = 50 adds a new entry in the dict

!! In names "Tom" will be included! "View-object"

ages = age.values()

Tom in age : True.

age.items() # returns a dict-key

object, essentially a list of tuples

[(name, age)]

Caution !! (Dynamic typing):

x = 3

y = x

x = 2

Result:

$\begin{cases} x = 2 \\ y = 3 \end{cases}$

Because integers are immutable objects

Strings also immutable
Tuples also immutable

x = [1, 2, 3]

y = x

x[0] = 7

Result

x = [7, 2, 3]

y = [7, 2, 3]



Lists are mutable objects

Dictionaries also mutable

To keep y as before, two possibilities $\begin{cases} y = \text{list}(x) \\ y = x[:] \end{cases}$

In both cases it creates a new object.

Remarks : $\begin{cases} == & \text{equality element-wise} \\ \text{is} & \text{equality as objects} \end{cases}$

$x = [1, 2, 3]$		$x = [1, 2, 3]$
$y = [1, 2, 3]$		$y = x$
$x == y$: True		$x \text{ is } y$: True
$x \text{ is } y$: False		

Statements : Return, import, pass (do nothing, just in some cases for syntactical reasons)

if ---
elif ---
else ---

For : $\begin{array}{l} \text{for } n \text{ in range}(5): \\ \quad \text{print}(n) \end{array} \quad \left. \begin{array}{l} 0 \\ 1 \\ 2 \\ 3 \\ 4 \end{array} \right\}$

$\text{names} = ["Pepe", "Juan", \dots]$

$\begin{array}{l} \text{for name in names:} \\ \quad \text{print}(*\text{name}) \end{array} \quad \left. \begin{array}{l} \text{'Pepe'} \\ \text{'Juan'} \\ \vdots \end{array} \right\}$

$\begin{array}{l} \text{for } i \text{ in range}(\text{len}(\text{names})): \\ \quad \text{print}(\text{names}[i]) \end{array} \quad \left. \begin{array}{l} \text{'Pepe'} \\ \text{'Juan'} \\ \vdots \end{array} \right\}$

In other languages we do (because iterating over a list is not possible)

Can iterate over a dictionary as follows:

```
age = { "Ivan": 55, "Pepe": 20, ... }
```

```
for name in age:
    print(name, age[name])
```

} Ivan 55
Pepe 20

* Can also do ... in age.keys()

Or if I want the names to be showed alphabetically, ... in sorted(age.keys()).

sorted(____, reverse=True) will order the list in reverse alphabetically and

List comprehension: $l = [a**2 \text{ for } a \text{ in range}(10) \text{ if } a \% 2 == 0]$.

Reading and writing files

```
for line in open("oro.txt"):
    print(line)
```

} soy
un
oro

← w/ extra spaces.
This is because in the
txt file it actually is

soy\n
un\n
oro

This can be avoided using rstrip() method:

```
for line in open("oro.txt"):
    line = line.rstrip()
    print(line)
```

} soy
un
oro

`split("a")` method takes a string and divides it wherever there is a "a".

Create a txt file:

`F = open("output.txt", "w")`

↖ write

w+ : write & read

`F.write("Python\n")`

a+ : append

`F.close()` ↖ to close

for `i` in range(1,11):

`F.write("The square of " + str(i) + " is " + str(i**2) + "\n")`

`F.read()` returns a single string with the whole text (also `\n`).

`F.readlines()` returns a list with entries the lines of the file
(also with `\n`, since they are strings)

`Q = open("lol.txt", "r")`

`lines = Q.readlines(16)`

`print(lines)`

`linea = Q.readline(5)`

`print(linea)`

`linea = Q.readline(9)`

`print(linea)`

lol.txt

This is line 1
This is line 2
_____ 3

This is line 1

This

is line 2

Eg: If a file has

Pedro,	22,	1.78
Jorge,	25,	1.67
:	:	:

how to extract this data? Use `split(",")` !

~~Notes~~

~~Notes~~ Notes

D = open("datos.txt", "r")

Z = []

for line in D:

currentline = line.split(",")

(even better, line.rstrip().split(","))

Z.append(currentline)

Functions

inputs
↙ ↘
def add(a, b):
 mysum = a + b
 return mysum

def add_and_sub(a, b):
 mysum = a + b
 mydiff = a - b
 return (mysum, mydiff)

⚠ Warning: For mutable objects, a function ^{might} change its value!

def modify(mylist):
 mylist[0] = 10
 return (mylist)

L = [1, 3, 5, 7, 9]

M = modify(L)

M = [10, 3, 5, 7, 9]

L = [10, 3, 5, 7, 9]

!!!

Example: Function that creates a password:

import random

(random.choice([....]) selects randomly an elmt of the list)

def password(length):

pw = str()

chars = "abcd... 1234567890"

for i in range(length):

pw = pw + random.choice(chars)

Scope rules: It might be that at different points of a program there are several ~~more~~ functions or variables with the same name (eg inside a function and outside). Python follows

Local

Enclosing function

Global

Built-in

- Current function

- Function that called the current function, if any

- Module that the function was defined

- Python's built-in namespace

If Python cannot resolve a name, error!

Remark: Argument: Object that is passed to a function as its input when the function is called

Parameter: Variable that is used in the function definition to refer to that argument

(n = 360)

↑ ↑
variable object
(name of obj)

Example: Look how it works depending on mutable / immutable objects:

```
def update(n, x):
```

```
    n = 2
```

```
    x.append(4)
```

```
    print('update:', n, x)
```

```
def main():
```

```
    n = 1
```

```
    x = [0, 1, 2, 3]
```

```
    print('main:', n, x)
```

```
    update(n, x)
```

```
    print('main:', n, x)
```

```
main
```

Returns

```
main: 1 [0, 1, 2, 3]
```

```
update: 2 [0, 1, 2, 3, 4]
```

```
main: 1 [0, 1, 2, 3, 4]
```

n does not change its original value,
but x does

Object-oriented programming: Creating new object type ~~extending~~ extending another one.

Classes. This is analogous to define a function, but now this allows us to

define new methods, so new something().

Eg: Want to create a method ~~that~~ remove_min which removes the minimum element of a list:

```
class MyList(list):
```

```
    def remove_min(self):
```

```
        self.remove(min(self))
```

```
x = [2, 5, 10, 1]
```

```
y = MyList(x)
```

```
& dir(x) = (all methods for lists)
```

```
dir(y) = ( ) + remove_min
```

```
y.remove_min()
```

1 disappeared from y

Numpy

Provides new data type.

import numpy as np

$z_v = \text{np.zeros}(5)$ # zero vector of dim 5

$z_m = \text{np.zeros}((5, 3))$ # zero matrix of dim 5×3

np.empty

$x = \text{np.array}([7, -1])$

$x[1] = -1$

$A = \text{np.array}([[1, 2], [3, 4]])$

$A[0, 0] = 1$

$A.\text{transpose}()$

$A[:, 1] = \text{array}([2, 4])$

column

$A[0:] = \text{array}([1, 2])$

⚠ : $[1, 2] + [3, 4] = [1, 2, 3, 4]$

$\text{np.array}([1, 2]) + \text{np.array}([3, 4]) = \text{array}([4, 6])$
 ||
 z_1

$z_1 + 1 = \text{array}([2, 3])$

~~$z_1[0, 1] = 1$~~

$v = \text{np.array}([3, 5, 8, 16, 24])$

$v[\underbrace{[1, 2]}_{\text{list of indices}}] = \text{array}([5, 8])$

list of indices

I can use booleans as array in the following sense:

$z_1 = \text{np.array}([1, 3, 5, 7, 9])$

$z_1 > 6 \rightarrow$ gives me an array $([False, False, False, True, True])$

Then I can do

$z_1[z_1 > 6] \rightarrow$ gives me array $([7, 9])$, i.e., returns an array only with the True entries.

This is a "logical array"

Warning! : Slicing an array using colon $v[3:5]$ gives a view of the object, but does not create a new object.

Using ~~arr~~ indexing does create a copy:

$z_1 = \text{np.array}([1, 3, 5, 7, 9])$

$w = z_1[0:3]$

$w[0] = 3$

$w = \text{array}(3, 3, 5)$

$z_1 = (3, 3, 5, 7, 9) !!$

z_1 as before

$\text{ind} = \text{np.array}([0, 1, 2])$ # indexes

$w = z_1[\text{ind}]$

$w[0] = 3$

$w = (3, 3, 5)$

$z_1 = (1, 3, 5, 7, 9)$

Better: z_1 as before

$w = \text{np.copy}(z_1[0:3])$

$w[0] = 3$

(8)

`np.linspace(0, 100, 11)` gives an array of 11 points spaced equally between 0 & 100, ie, 0, 10, 20, ...

`np.logspace(1, 4, 4)` gives an array of 4 points spaced "log. equally" bet.
~~10¹ & 10⁴~~ ie, 10¹, 10², 10³, 10⁴.

`A = np.array([[123], [456]])`, `A.shape` : (2, 3)
`A.size` : 6

`x = np.random.random(10)` # randomly chooses 10 num in (0, 1)

`np.any(x > 0.9)` → True # is any of its elmts > 0.9?

`np.all(x >= 0.1)` → ~~Are~~ all Are all of its elmts ≥ 0.1?

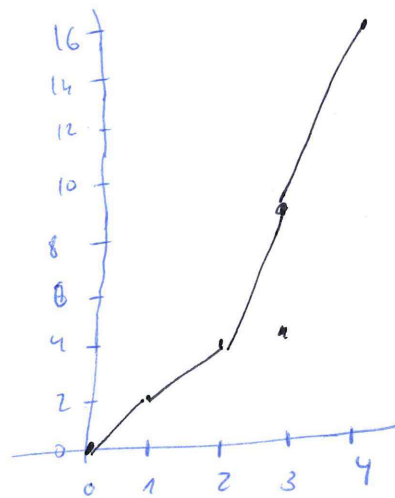
Matplotlib & Pyplot

`import matplotlib.pyplot as plt`

`plt.plot([0, 1, 4, 9, 16])`

(we ; at the end ~~has~~ to not show the object;
 ie the name of the object)

→



ie it plots the piecewise linear function
 passing through (0,0), (1,1), (2,4), (3,9), (4,16)

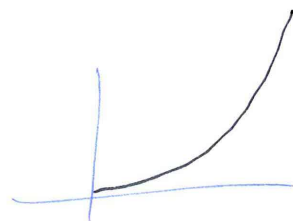
`x = np.linspace(0, 10, 20)`

`y = x**2`

`plt.plot(x, y)`

→

gives



smoother

x = (as before)

y1 = x ** 2.0

y2 = x ** 1.5

plt.plot(x, y1, "bo-")

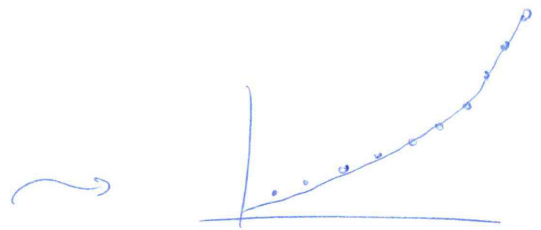
blue nodes
as circles
solid line

b = blue

g = green

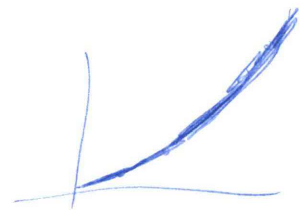
r = red, etc

o = circle
s = square



plt.plot(x, y1, "bo-", linewidth=2, markersize=4)

↑ line ↑ nodes



Example:

x = np.linspace(0, 10, 20)

y1 = x ** 2.0

y2 = x ** 1.5

plt.plot(x, y1, "bo-", linewidth=2, markersize=12, label="\$x^{2.0}\$")

plt.plot(x, y2, "gs-", label="\$x^{1.5}\$")

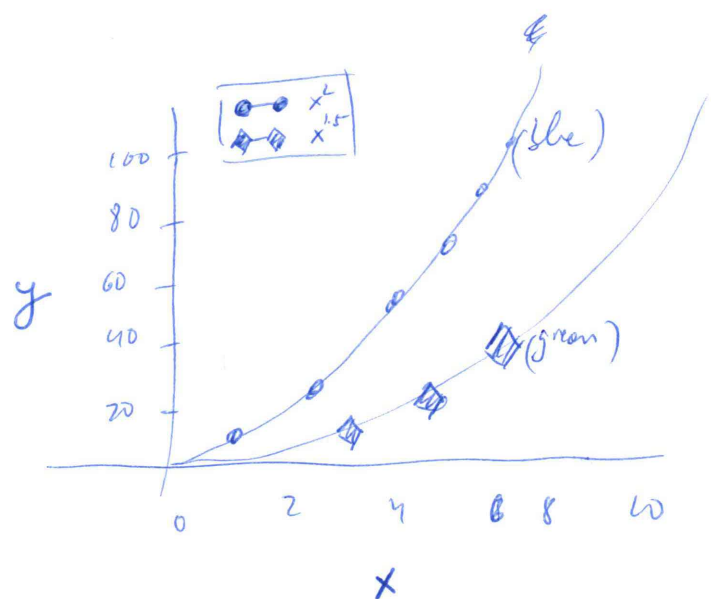
plt.xlabel("\$x\$")

plt.ylabel("\$y\$")

plt.axis([0.5, 10.5, -5, 105])

plt.legend(loc="upper left")

plt.savefig("myplot.pdf")



Remark : `plt. loglog (x, y, "bo-")` returns a plot with ^{both axes} ~~axis~~ in log scale.

`plt. semi log x ( )` _____ with log scale in x axes

`plt. semi log y ( )` _____ y

`np. logspace (-1, 1, 20)` gives log-eg. space from 10^{-1} to 10^1 .

Histograms :

`x = np. random. normal (size = 1000)`

generates an array of 1000 numbers
get from the standard normal distribution
 $N(0, 1)$.

`plt. hist (x);` # returns 

`plt. hist (x, normed = True);` # instead of plotting the number of observations, it plots
density the proportion of observations

`plt. hist (x, norm = True, bins = np. linspace (-5, 5, 21))` # set differences bet of
0.5, or 20 ~~bins~~ bins

`x = np. random. gamma (2, 3, 100000)` ^{$\sim \gamma(2, 3)$}

`plt. hist (x, bins = 30, cumulative = True)` # gives the cumulative histogram,
`(histtype = "step")`

↑ # appearance of the histogram

Can put many plots into one:

```
plt.figure()
```

```
plt.subplot(221)
```

22 indicates 2x2

```
plt.hist(x, ...)
```

```
plt.subplot(222)
```

```
{  
:  
}
```

223

224

Random processes

```
import random
```

```
random.choice([0,1]) # returns 0 or 1 randomly
```

{ or any list, or a range object, eg range(1,7) for 1,...,6.

```
random.choice(random.choice([range(1,7), range(1,9), range(1,11)])) # encodes the
```

situation where you have 3 dies of 6, 8 & 10 faces, you choose ~~one~~ at random and you roll the die.

Example: Roll the dice a number of times and plot the results. (say 100 times)

```
rolls = []
```

```
for k in range(100):
```

```
    rolls.append(random.choice([1,...,6]))
```

```
plt.hist(rolls, bins = np.linspace(0.5, 6.5, 7);
```

Example: What probability distribution follows the random variable consisting of the sum of rolling 10 dice?

`ys = []`

for rep in range(500000)

`j = 0`

for k in range(10):

`x = random.choice([1, ..., 6])`

`j = j + x`

`ys.append(j)`

`plt.hist(ys);` # normal distribution, (Central Limit Theorem)

`np.random.random(5)` creates a list w/ 5 numbers chosen from the uniform dist in $[0, 1]$
`(5, 3)` matrix w/ 5×3

`np.random.normal(0, 1)` # chooses a number from $N(0, 1)$
`(0, 1, (6, 8))` # ~~chooses~~ creates a 6×8 matrix of random numbers.

`np.random.randint(1, 7)` # chooses a random integer $1, \dots, 6$.

`A = np.array ( some matrix )`

`np.sum(A)` # sum all entries

`np.sum(A, axis=0)` # sum the elts in every column (so over the rows)

`np.sum(A, axis=1)` # sum the elts in every row (so over the columns)

Example (The previous example revisited):

`X = np.random.randint(1, 7, (100000, 10))`

`Y = np.sum(X, axis=1)`

`plt.hist(Y).`

How to measure time?

import time

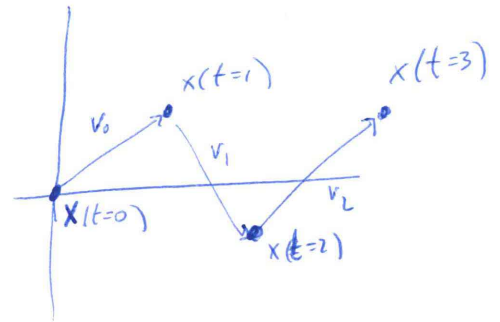
start_time = time.clock()

end_time = time.clock()

total_time = end_time - start_time

Random walk: A walker moves in $\mathbb{R}^2$ randomly at times $t=0, 1, 2, 3, \dots$.

$$x(t=k) = x(t=0) + \sum v_i$$



np.cumsum (matrix, axis = 0 or 1) # returns a matrix with partial sums in any column/row

np.concatenate (a1, a2) # concatenate two np arrays

Example: Plot a 5-iteration random walker starting from $(0,0) \in \mathbb{R}^2$:

X_0 = np.array ([0],[0])

delta-X = np.random.normal (0,1, (2,5))

X = np.concatenate ((X_0, np.cumsum (delta-X, axis=1), axis=1)

plt.plot (X[0], X[1], "ro-")

plt.savefig ("rw.pdf")

More about ~~the~~ string of characters

- `F.read()` creates a unique string with all the text, also `\n`. How to remove those? A easy form is to use `F.replace("\n", "")`.

⚠ Since strings are immutable, this does not change `F`, one has to do `F = F.replace(...)`

⚠ In some computers, there are `\r` also hidden. Better do `F = F.replace("r", "")` as well

- If you create a function and want to create a "help" entry, this can be done by creating a "docstring": after `def --- :`, write `""" This is the manual ... """`.

Remark. A safer way to read files is

with `open("dna.txt", "r")` as `F`:

`seq = F.read`

- Count words from a text

from collections import Counter

`text = "hola me llamo Pablo Casado y he perdido cinco elecciones en un año".`

`word_counts = Counter(text.split(" "))`

creates a ^{rehash} dict with keys the different words and values the frequency of the words.

Remark: To turn text into a string with no capital letters or punctuation, do

`text = text.lower()`

`skip = [".", ",", ";", ":", "!", "(", ")", "&"]`

for `ch` in `skip`: `text = text.replace(ch, " ")`

• Read a book and save result

with `open("book.txt", "r", encoding="utf8")` as `current_file`:

`text = current_file.read()`

`text = text.replace("\n", " ").replace("\r", " ")`

• Find fragment in a book

`ind = text.find("What's in a name?")`

`print(text[ind : ind + 100])` # prints the 100 char after the start of the sentence.

• Navigate file directories

`import os`

`os.listdir("./Descargas")` # list files in ./Descargas
↓
where you already are

what if I have many directories English, German, ... each of them w/ books of diff ^(a new folder) lang. authors
Shakespeare ... and inside those he added .txt's files with the books? Can do

for language in `os.listdir("./Books")`:

for author in `os.listdir("./Books/" + language)`:

for title in `os.listdir("./Books/" + language + "/" + author)`:

input file = `"./Books/" + language + "/" + author + "/" + title`
↙
(forgot the indentation)

• Create table of data with Pandas

```
import pandas as pd
```

```
table = pd.DataFrame ( columns = ("name", "age") )
```

```
table.loc[1] = "James", 22
```

```
table.loc[2] = "Pute", 30
```

Result:		
	name	age
1	James	22
2	Pute	30

```
table.head()
table.tail()
```

} # gives the first / last 5 lines

Remark: "hole".capitalize() returns "Hole"

table.name # gives a list with all entries of the column "name".

```
table = pd.read_csv ("data.csv", index_col = 0) # reads a csv table
```

More numpy

• Distance bet. pts:

```
import numpy as np
```

```
p1 = np.array ([1, 1])
```

```
p2 = np.array ([4, 4])
```

```
p2 - p1 → array ([3, 3])
```

```
np.power (p2 - p1, 2) # does  $(3^2, 3^2)$ , to return array ([9, 9])
```

```
np.sqrt (np.sum (np.power (p2 - p1, 2)))
```

- Find the mode of a list (most repeated value)

```
import scipy.stats as ss
```

```
votes = [1, 2, 3, 3, 1, 2, 1, 3, 3, 3]
```

```
mode, count = ss.mstats.mode(votes)
```

$\uparrow$
 value

$\uparrow$
 number of times
 it appears

- Submatrices

Given a np. array A (matrix), one can get a submatrix by doing $A[1:2, 0:2]$ eg.
 the first column is $A[:, 0]$, second column is $A[:, 1]$, etc.

$A.shape[0]$ returns # rows; $A.shape[1]$ returns # columns.

- Array

$\text{np.arange}(0, 10, 2, [\text{dtype} = \text{int}])$ # returns an array $[0, 2, 4, 6, 8]$.

$\text{np.arange}(10).shape$ (5, 2) gives the matrix

0	1	2	3	4
5	6	7	8	9

$\text{np.repeat}(n, k)$ gives an array $(n, n, n, \dots, n)$ k times

- Sort (Ordering) a list

```
a = np.array([2, 5, 8, -4, 3, 6])
```

$\text{np.sort}(a)$ returns the ordered list

$\text{np.argsort}(a)$ returns a list with the indices corresponding to the sorted elements, i.e.

if $a = (x_0, \dots, x_n)$, and $\text{np.sort}(a) = (x_4, x_3, x_0, x_2, \dots)$ then

$\text{np.argsort}(a)$ returns $(4, 3, 0, 2, \dots)$

• A list with the 3 largest numbers of a can be got by doing

ind = np.argsort(a)

~~data~~

a[ind[0:2]]

Example (KNN algorithm): Given two sets of pts $A = \{x_1, \dots, x_m\}$, $B = \{y_1, \dots, y_n\}$, and a pt $p \in \mathbb{R}^n$, a function which determines to what set the pt p is the closest, based on the nearest $K \in \mathbb{N}$ points in $A \cup B$.

def distance(p, q): (easy)

def find_nearest_neighbors(p, points, K)

Find the K nearest neighbors and return indices in points

distances = np.zeros(points.shape[0])

for i in range(len(distances)):

distances[i] = distance(p, points[i])

ind = np.argsort(distances)

return ind[:K]

def knn_predict(p, points, outcomes, K=5)

points is supposed to be A ∪ B and

ind = find_nearest_neighbors(p, points, K)

outcomes $\underbrace{[0, 0, \dots, 0]}_n, \underbrace{[1, \dots, 1]}_m$

~~mode~~ mode, count = ss.stats.mode(outcomes[ind])

return mode.

• $\text{ss.norm}(\mu, \sigma). \text{rvs}(n, m)$ returns a $n \times m$ matrix with entries chosen at random $\sim N(\mu, \sigma)$

• Enumerate

$\text{seasons} = [\text{"spring"}, \text{"autumn"}, \dots]$

$\text{list}(\text{enumerate}(\text{seasons}))$ returns $(0, \text{"spring"}), (1, \text{"autumn"}), \dots$

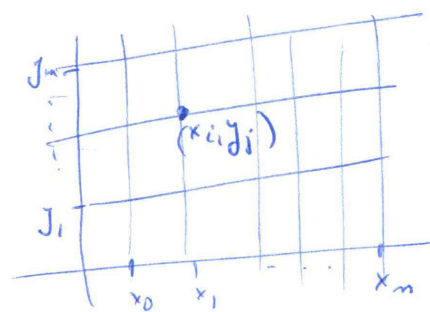
$\uparrow$ index $\uparrow$ element

• Meshgrid : Let $x = (x_1, \dots, x_m)$, $y = (y_1, \dots, y_m)$

$xx, yy = \text{np.meshgrid}(x, y)$ ^{np arrays} returns two matrices:

$$xx = \begin{pmatrix} x_i \end{pmatrix}_{\substack{i=1 \dots n \\ j=1 \dots m}}$$

$$yy = \begin{pmatrix} y_j \end{pmatrix}_{\substack{i=1 \dots n \\ j=1 \dots m}}$$



$$\text{let } A = \begin{pmatrix} (x_i, y_j) \end{pmatrix}_{\substack{i=1 \dots n \\ j=1 \dots m}}$$

(so xx has repeated rows, yy has repeated columns)

More on Pandas

import pandas as pd

(14)

Two main objects: series & DataFrame

1-dim array-like 2-dim array-like

Series

x = pd.Series([6, 8, 3, 6])

indices values
↓ ↓
0 6
1 8
2 3
3 6

Can define your own index labels:

x = pd.Series([6, 8, 3, 6], index=["f", "w", "e", "r"]) →

f 6
w 8
e 3
r 6

Can access to the value by doing x["w"] or x[["w", "r"]].

• If age = { "Tom": 50, ... } is a dict, get a pd series by doing

x = pd.Series(age)

Tom 50
: :
: :

• DataFrame: Before we saw one way. Can also turn a dict into DataFrame; by using lists in the values.

data = { 'name': ['Tom', ...], 'age': [50, ...], 'zip': ['06008', ...] }

x = pd.DataFrame(data, columns=["name", "age", "zip"])

→

	name	age	ZIP
0	Tom	50	06008
1			
2			
3			

Can recover a single column by doing `x["name"]` or `x.name`.

• How to ~~reset~~ change the indices 0, 1, ...? By doing

`x.reindex(["Person1", ...])`

array-like objects

• If `x, y` are pd Series of the same dim, then `x + y` sums the entries w/ the same labels, and returns NaN if a label appears only in one of `x` or `y`.

• Read files: Either is it in csv or a txt with data separated by ','s, do

`x = pd.read_csv("data.txt")`

To add a column from another file,

`x["new-column"] = pd.read_csv("new-data.txt")`

! `x["new-column"]` adds a column, `x.loc[7]` adds a new row

`x.iloc[0:10]` slices the first 10 rows

`x.iloc[0:10, 1:3]` ————— and 2nd, 3rd columns.

`x.columns` #lists the name of the columns.

• Compute correlations

Let x be a `pd.DataFrame` object w/ columns some attributes of same axis of data.

Would like to compare two such attributes

`corr_x = pd.DataFrame.corr(x)` # compute correlation matrix!

Graphical plotting:

import matplotlib.pyplot as plt

`plt.figure(figsize=(10,10))`

`plt.pcolor(x)` # actually plots the content

`plt.colorbar()` # the legend for the colors

`plt.savefig('...pdf')`

Remark. `x.transpose()` gives the actual transpose of the table.

• Spectral co-clustering: Attempt to find clusters of related data (eg words from books, and find clusters by language).

from sklearn.cluster.bicluster import SpectralCoClustering

`model = SpectralCoClustering(n_clusters=6, random_state=0)`

`model.fit(corr_x)`

`model.rows_` # returns a $(n_clusters \times \#rows\ in\ x)$ matrix with True/False
 ↑
 the data

`np.sum(model.rows == , axis = 1)` # returns an array with the # of obs in each cluster

Examples with Pandas

```
birddata = pd.read_csv("bird_tracking.csv")
```

```
ix = birddata.bird_name == "Eric"
```

returns the indices corresponding to the value "Eric" in the column "bird_name".

```
x, y = birddata.longitude[ix], birddata.latitude[ix]
```

returns arrays with the columns long, lat corresponding to ix, i.e. Eric.

```
plt.figure(figsize=(17,7))
```

```
plt.plot(x, y, "o")
```

```
bird_names = pd.unique(birddata.bird_name)
```

returns an array w/ the unique entries of the column bird_name (i.e. without repetition)

Example:

bird_names (as before)

```
for bird in bird_names:
```

```
    ix = birddata.bird_name == bird
```

```
    x, y = birddata.longitude[ix], birddata.latitude[ix]
```

```
    plt.plot(x, y, "o")
```

```
plt.xlabel("Longitude"); plt.ylabel("Latitude")
```

```
plt.legend(loc="lower right")
```

Useful trick : How to ignore the NaN's of a list? (This gives error if you want to plot): let x be an array with numbers and some NaN's. Then

```
ind = np.isnan(x) # ind returns an array with True if a value is NaN
plt.hist(x[~ind]) # ~ind is the complement of ind, so we only plot the "False" values in ind.
```

• Pandas can also plot histograms :

```
birddata.speed.az.plot(kind='hist', range=[0, 30])
plt.xlabel("2D speed")
```

Good point of this : Do not have to deal with NaN's.

Handling time with Datetime

```
import datetime
datetime.datetime.today() # returns a datetime-type obj datetime.datetime(2020, 6, 18, 11, 43, 50, 1)
                           Y M D H M S
time - 1 = datetime.datetime.today()
time - 2 = _____
time - 2 - time - 1 → returns the time it has passed (timedelta object)
```

• Say we have a string "2013-08-15 00:18:08+00". How to read off the date from this?

```
datetime.datetime.strptime(x[:10], "%Y-%m-%d %H:%M:%S")
```

How to add a list to a pandas DataFrame?

$\text{birddata}["\text{newcolumn}"] = \text{pd.Series}(x, \text{index}=\text{birddata.index})$

↑ ↑
the existed DF the list

• How many days have passed? I can divide the whole object: ~~can do~~

Perkins

$$\text{time}_2 - \text{time}_1 / \text{date time} - \text{time delta} (\text{days} = 1)$$

Graphs with NetworkX

```
import networkx as nx
```

`g = nx.Graph()` # creates an empty graph

g.add_node(1) # creates a node labeled w/ 1

g. add - nodes from $([2, 3, "u", "v"])$

$g_nodes()$ # show the nodes

$g.add_edge(1,2)$ # adds a node



g.add_edges_from([(1,5), ("u","w")]) $\neq$ 1 nodes are added automatically!

g. remove - node (2)
edge

g. remove - nodes - from $[4, 5]$
edges

g. number of nodes ()
edges

Plotting a graph

17

```
import matplotlib.pyplot as plt
```

```
nx.draw(G, with_labels=True, node_color="lightblue", edge_color="grey")
```

`G.degree()` # returns a dict-like object w/ the list of nodes & degrees

• Erdős-Rényi random graphs : Parameters (N, p) : N = #graphs, p = probab. that there is a graph for every pair of nodes

```
from scipy.stats import bernoulli
```

`bernoulli.rvs(p=0.2)` # bernoulli distr, returns 0 or 1 w/ prob p

```
G = nx.Graph(); N=20; p=0.2
```

```
G.add_nodes(range(N))
```

```
for nod1 in G.nodes():
```

```
    for nod2 in G.nodes:
```

can be omitted

```
        if bernoulli.rvs(p=p) == True:
```

```
            G.add_edge(nod1, nod2)
```

(creates a directed graph, if not directed should and (and $nod1 < nod2$))

Degree distribution

```
def deg_distr(G):
```

```
    deg_sequence = [d for n, d in G.degree]
```

```
    plt.hist(deg_sequence, histtype="step")
```

```
    plt.xlabel('...'); plt.ylabel('...'); plt.title('...')
```

Def: Let G be a graph. The adjacent matrix of G is $A = (a_{ij})$, $n \times n$ matrix,
 $n = \# \text{ nodes}$, and $a_{ij} = \begin{cases} 0, & \text{if } \nexists \text{ edge bet. } v_i \& v_j \\ 1, & \exists \end{cases}$

• Ofc G can be reconstructed from A . To do this w/ Python

$G = nx.to_network_graph(A)$

• Connected components of a graph

How to get the largest conn comp? Two ways:

(1) $g = nx.connected_component_subgraphs(G).next()$

(2) $g = \max(nx.connected_component_subgraphs(G), key = len)$

Statistics with Python

import numpy as np

——— scipy as ss

——— matplotlib.pyplot as plt

$x = 10 * ss.uniform.rvs(size = n)$

creates an array w/ n values chosen from the uniform dist'n $(0,1)$

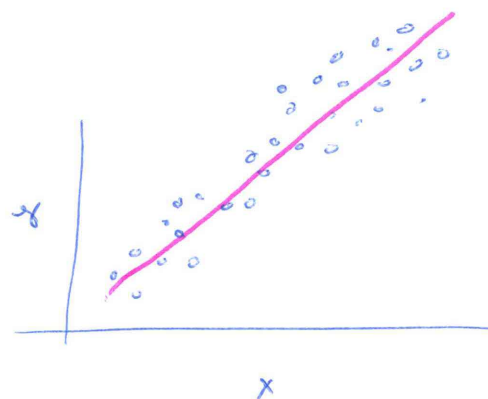
$y = 5 + 2 * x + \underbrace{ss.norm.rvs(loc = 0, scale = 1, size = 50)}_{\text{noise term} = N(0,1)}$

noise term $= N(0,1)$

Can plot this by doing

```
plt.figure()
plt.plot(x, y, "o", ms=5)
xx = np.array([0, 10])
plt.plot(xx, 5 + 2 * xx)
plt.xlabel("x"); plt.ylabel("y")
```

market size



• least squares problem: Given $Y = \beta_0 + \beta_1 X$ and data $\{(x_i, y_i)\}$, want to build $\hat{y} = \hat{\beta}_0 + \hat{\beta}_1 \hat{x}$ which minimizes $\sum e_i^2$, where $e_i = \hat{y}_i - y_i$.

import statsmodels.api as sm

mod = sm.OLS(y, x)

est = mod.fit()

print(est.summary())

for a estimation for $Y = \beta \cdot X$

X = sm.add_constant(x)

mod = sm.OLS(y, X)

est = mod.fit()

print(est.summary())

for $Y = \beta_0 + \beta_1 X$

• sciKit-learn for linear regression

$x1 = 10 * \text{ss.uniform.rvs}(\text{size}=500)$

$x2 = \underline{\hspace{10cm}}$

$y = 5 + 2 * x1 - 1 * x2 + \text{ss.norm.rvs}(\text{loc}=0, \text{scale}=1, \text{size}=500)$

$X = \text{np.stack}([x1, x2], \text{axis}=1)$

takes two arrays and stacks them together, i.e., concatenates them, so you get a matrix. In this case as columns ($\text{axis}=1$)

`from mpl_toolkits.mplot3d import Axes3D`

`fig = plt.figure()`

`ax = fig.add_subplot(111, projection='3d')`

`ax.scatter(X[:,0], X[:,1], y, c=y)`

} Produces a 3d graph

`from sklearn.linear_model import LinearRegression`

`lm = LinearRegression(fit_intercept=True)`

`lm.fit(X, y)`

`lm.intercept_` # returns β_0

`lm.coef_[0]` # — β_1

`lm.coef_[1]` # — β_2

• How to "predict" values?

`X0 = np.array([2, 4]).reshape(1, -1)`

`lm.predict(X0)`

- (Assessing model accuracy) : To evaluate the performance of the regression model, one has to quantify how much the predictions agree w/ the observed data, eg mean square error. Out of some data, we want to split it randomly ~~and~~, see how they perform differently and compare results.

```
from sklearn.model_selection import train_test_split
```

```
X_train, X_test, y_train, y_test = train_test_split(X, y, train_size=0.5, random_state=1)
```

```
lm = LinearRegression(fit_intercept=True) # for the  $\beta_0$  term
```

```
lm.fit(X_train, y_train) # get the model
```

```
lm.score(X_test, y_test)
```

