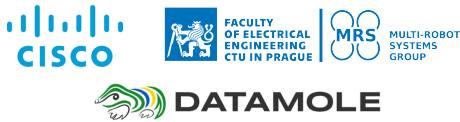


Classical methods and models in MLE with Scikit Learn

BECM33MLE — Machine Learning Engineering

Dr. Tomáš Báča

Multi-Robot Systems group, Faculty of Electrical Engineering
Czech Technical University in Prague



Lecture 2: Classical methods and models in MLE with Scikit Learn

2025-10-11

Classical methods and models in MLE with Scikit Learn
BECM33MLE — Machine Learning Engineering

Dr. Tomáš Báča

Multi-Robot Systems group, Faculty of Electrical Engineering
Czech Technical University in Prague



- always recommend to work in a **Virtual Environment**
- it stores copies of the particular dependencies locally
- one solution to a **dependency hell**

Dependencies (requirements.txt)

```
1 numpy==2.3.3
2 matplotlib==3.10.6
3 scikit-learn==1.7.2
4 pandas==2.3.2
5 jupyter
6 jupyterlab-vim
```

- the version is fixed for some of the packages to make the code within the examples work in the future

Setting up python environment (Ubuntu 24.04)

```
1 # install the python3's virtual environment
2 sudo apt get install python3-venv
3
4 # create the virtual environment
5 python3 -m venv python-env
6
7 # activate the environment
8 source ./python-env/bin/activate
9
10 # install dependencies manually
11 pip install numpy ...
```

or install dependencies in bulk

```
1 # install deps from requirements.txt
2 python3 -m pip install -r requirements.txt
```

Lecture 2: Classical methods and models in MLE with Scikit Learn

Python Environment

Python Virtual Environment

- there are other options to make sure you dependencies are met and are not colliding with other projects:
 - Conda package manager,
 - Docker,
 - Singularity / Apptainer container system,
 - Nix package manager.

- web-interface with an editor
- allows execution of portions of your code
- remembers the state of variables
- remembers where you left it
- embedded visualization
- well-suited for learning and prototyping
- (bonus) vim-like key bindings
- **Jupyter Notebook**
 - single document, simple
- **Jupyter Lab**
 - multiple documents at once
 - similar to an IDE
 - better file management

Running Jupyter Notebook

```
1 # 1. activate your python env
2 source ./python-env/bin/activate
3
4 # 2. run
5 jupyter notebook
```

3. check your browser (localhost:8888)

Running Jupyter Lab

```
1 # 1. activate your python env
2 source ./python-env/bin/activate
3
4 # 2. run
5 jupyter lab
```

3. check your browser (localhost:8888)

Lecture 2: Classical methods and models in MLE with Scikit Learn

Python Environment

Jupyter Notebook & Jupyter Lab

Jupyter Notebook & Jupyter Lab

- web-interface with an editor
- allows execution of portions of your code
- remembers the state of variables
- remembers where you left it
- embedded visualization
- well-suited for learning and prototyping
- (bonus) vim-like key bindings
- **Jupyter Notebook**
 - single document, simple
- **Jupyter Lab**
 - multiple documents at once
 - similar to an IDE
 - better file management

Running Jupyter Notebook

```
1 # 1. activate your python env
2 source ./python-env/bin/activate
3
4 # 2. run
5 jupyter notebook
```

Running Jupyter Lab

```
1 # 1. activate your python env
2 source ./python-env/bin/activate
3
4 # 2. run
5 jupyter lab
```

- You can, of course, run your python scripts manually. However, this makes development and tweaking of the algorithms tedious.
- In practice, you can design the first part of the architecture in Jupyter and then move to more traditional program architecture.

Jupyter Notebook & Jupyter Lab

Lecture 2:
Classical
methods
and
models in
MLE with
Scikit
Learn

Tomáš
Báča

Python
Environ-
ment

Pandas

Scikit
Learn

Toy dataset

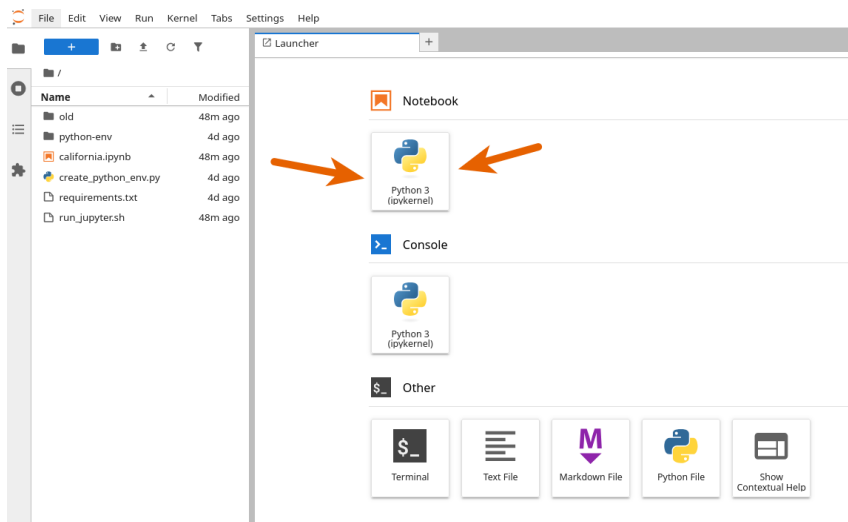
Regression
pipeline

Preprocessing

Metrics

Realworld
example

References



Tomáš Báča (CTU in Prague)

Lecture 2: Classical methods and models in MLE with Scikit Learn

September 30th, 2025

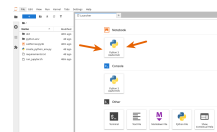
4 / 48

Lecture 2: Classical methods and models in MLE with Scikit Learn

Python Environment

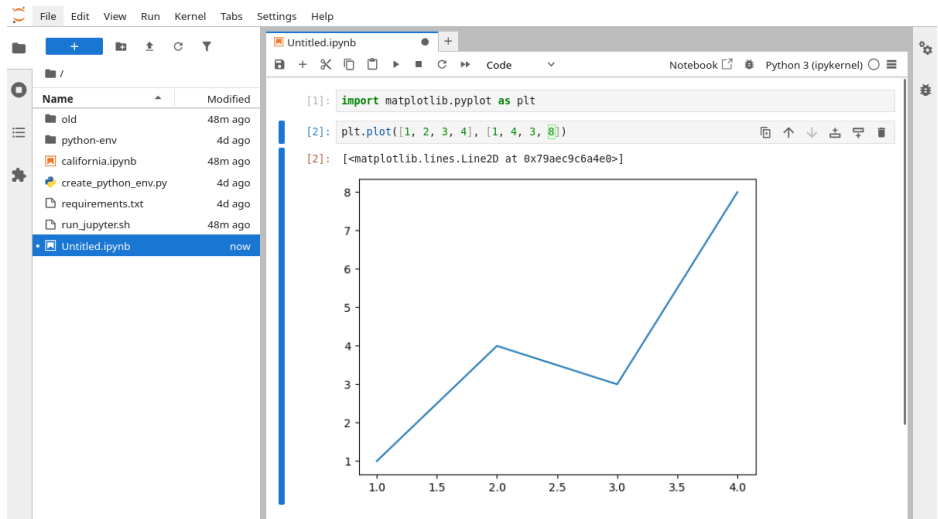
Jupyter Notebook & Jupyter Lab

Jupyter Notebook & Jupyter Lab



2025-10-11

Jupyter Notebook & Jupyter Lab

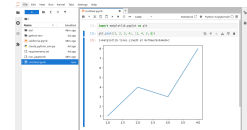


Lecture 2: Classical methods and models in MLE with Scikit Learn

Python Environment

Jupyter Notebook & Jupyter Lab

Jupyter Notebook & Jupyter Lab



- use <tab> for code completion
- use shift+<tab> for explanation of objects and their properties
- use ctrl+enter to evaluate current cell

- Python library
- manipulation of tabular data
- "Excel in Python"
- can load common formats (CSV, Excel, SQL database, JSON, ...)
- suitable for:
 - data cleaning and processing
 - merging and joining of datasets
 - visualization (Matplotlib integration)
 - handles temporal data
 - integrates with *Numpy* and *SkLearn*

Pandas dataframe

- the main workhorse of Pandas
- 2D, size-mutable data structure

	MedInc	HouseAge	AveRooms	AveBedrms	Population	AveOccup	Latitude	Longitude
0	8.3252	41.0	6.984127	1.023810	322.0	2.555556	37.88	-122.23
1	8.3014	21.0	6.238137	0.971880	2401.0	2.109842	37.86	-122.22
2	7.2574	52.0	8.288136	1.073446	496.0	2.802260	37.85	-122.24
3	5.6431	52.0	5.817352	1.073059	558.0	2.547945	37.85	-122.25
4	3.8462	52.0	6.281853	1.081081	565.0	2.181467	37.85	-122.25
...	...	...	...	...	...	...	...	...
20635	1.5603	25.0	5.045455	1.133333	845.0	2.560606	39.48	-121.09
20636	2.5568	18.0	6.114035	1.315789	356.0	3.122807	39.49	-121.21
20637	1.7000	17.0	5.205543	1.120092	1007.0	2.325635	39.43	-121.22
20638	1.8672	18.0	5.329513	1.171920	741.0	2.123209	39.43	-121.32
20639	2.3886	16.0	5.254717	1.162264	1387.0	2.616981	39.37	-121.24

20640 rows × 8 columns

Lecture 2: Classical methods and models in MLE with Scikit Learn

Pandas

Pandas — manipulating data in Python

Pandas — manipulating data in Python

Pandas dataframe

- Python library
- manipulation of tabular data
- "Excel in Python"
- can load common formats (CSV, Excel, SQL database, JSON, ...)
- suitable for:
 - data cleaning and processing
 - merging and joining of datasets
 - visualization (Matplotlib integration)
 - handles temporal data
 - integrates with *Numpy* and *SkLearn*

Size-mutable data structure

- 2D, size-mutable data structure
- MedInc, HouseAge, AveRooms, AveBedrms, Population, AveOccup, Latitude, Longitude
- 0 8.3252 41.0 6.984127 1.023810 322.0 2.555556 37.88 -122.23
- 1 8.3014 21.0 6.238137 0.971880 2401.0 2.109842 37.86 -122.22
- 2 7.2574 52.0 8.288136 1.073446 496.0 2.802260 37.85 -122.24
- 3 5.6431 52.0 5.817352 1.073059 558.0 2.547945 37.85 -122.25
- 4 3.8462 52.0 6.281853 1.081081 565.0 2.181467 37.85 -122.25
- ...
- 20635 1.5603 25.0 5.045455 1.133333 845.0 2.560606 39.48 -121.09
- 20636 2.5568 18.0 6.114035 1.315789 356.0 3.122807 39.49 -121.21
- 20637 1.7000 17.0 5.205543 1.120092 1007.0 2.325635 39.43 -121.22
- 20638 1.8672 18.0 5.329513 1.171920 741.0 2.123209 39.43 -121.32
- 20639 2.3886 16.0 5.254717 1.162264 1387.0 2.616981 39.37 -121.24

- free and open source **Python library**
- implements many classical ML methods suitable for work with *small data*
- perfect for introduction into practical ML
- elegant interface that leads to clean and simple code
- plenty of examples and community support
- https://scikit-learn.org/stable/user_guide.html

Scikit learn can do

- Classification
 - NN, SVM, random forests, logistic regression, ...
- Regression
 - gradient boosting, random forests, logistic regression, ...
- Clustering
 - k-Means, hierarchical clustering, ...
- Dimensionality reduction
 - PCA, linear & nonlinear manifold learning
- Model selection, hyper parameter optimization
 - grid search, cross validation, metrics, ...
- Data preprocessing
 - data preprocessing, feature extraction, ...

Lecture 2: Classical methods and models in MLE with Scikit Learn

Scikit Learn

Scikit Learn

2025-10-11

Scikit Learn

- Scikit Learn can do
- free and open source **Python library**
 - implements many classical ML methods suitable for work with *small data*
 - perfect for introduction into practical ML
 - elegant interface that leads to clean and simple code
 - plenty of examples and community support
 - https://scikit-learn.org/stable/user_guide.html
- Scikit Learn can do
- Classification
 - NN, SVM, random forests, logistic regression, ...
 - Regression
 - gradient boosting, random forests, logistic regression, ...
 - Clustering
 - k-Means, hierarchical clustering, ...
 - Dimensionality reduction
 - PCA, linear & nonlinear manifold learning
 - Model selection, hyper parameter optimization
 - grid search, cross validation, metrics, ...
 - Data preprocessing
 - data preprocessing, feature extraction, ...

Importing a toy dataset — California housing

Lecture 2:
Classical methods
and
models in
MLE with
Scikit
Learn

Tomáš
Báča

Python
Environ-
ment

Pandas

Scikit
Learn

Toy dataset

Regression

pipeline

Preprocessing

Metrics

Realworld

example

References

- data matrix:
 - median house age
 - median income
 - avg number of rooms
 - avg number of bedrooms
 - population in census block
 - occupancy
 - latitude
 - longitude
- target value:
 - \$ $\frac{\text{house price}}{100000}$ USD

Importing a dataset

```
1 from sklearn.datasets import  
  fetch_california_housing  
2 from pandas import pd  
3  
4 housing = fetch_california_housing(as_frame=True)
```

What is in the dataset?

```
1 # plaintext print  
2 print(housing.data)  
3 print(housing.target)  
4  
5 # print in jupyter notebook  
6 housing.data  
7 housing.target
```

Tomáš Báča (CTU in Prague)

Lecture 2: Classical methods and models in MLE with Scikit Learn

September 30th, 2025

8 / 48

Lecture 2: Classical methods and models in MLE with Scikit Learn

Scikit Learn

Toy dataset

Importing a toy dataset — California housing

Importing a toy dataset — California housing

```
Importing a dataset  
1 from sklearn.datasets import  
  fetch_california_housing  
2 from pandas import pd  
3 housing = fetch_california_housing(as_frame=True)  
4  
What is in the dataset?  
5 # plaintext print  
6 print(housing.data)  
7 print(housing.target)  
8  
9 # print in jupyter notebook  
10 housing.data  
11 housing.target
```

- some other datasets: https://scikit-learn.org/stable/datasets/toy_dataset.html

2025-10-11

Data matrix

	MedInc	HouseAge	AveRooms	AveBedrms	Population	AveOccup	Latitude	Longitude
0	8.3252	41.0	6.984127	1.023810	322.0	2.555556	37.88	-122.23
1	8.3014	21.0	6.238137	0.971880	2401.0	2.109842	37.86	-122.22
2	7.2574	52.0	8.288136	1.073446	496.0	2.802260	37.85	-122.24
3	5.6431	52.0	5.817352	1.073059	558.0	2.547945	37.85	-122.25
4	3.8462	52.0	6.281853	1.081081	565.0	2.181467	37.85	-122.25
...	...	...	...	...	...	...	...	...
20635	1.5603	25.0	5.045455	1.133333	845.0	2.560606	39.48	-121.09
20636	2.5568	18.0	6.114035	1.315789	356.0	3.122807	39.49	-121.21
20637	1.7000	17.0	5.205543	1.120092	1007.0	2.325635	39.43	-121.22
20638	1.8672	18.0	5.329513	1.171920	741.0	2.123209	39.43	-121.32
20639	2.3886	16.0	5.254717	1.162264	1387.0	2.616981	39.37	-121.24

20640 rows × 8 columns

Target value

0	4.526
1	3.585
2	3.521
3	3.413
4	3.422
...	...
20635	0.781
20636	0.771
20637	0.923
20638	0.847
20639	0.894

Lecture 2: Classical methods and models in MLE with Scikit Learn

Scikit Learn

Toy dataset

California housing dataset

2025-10-11

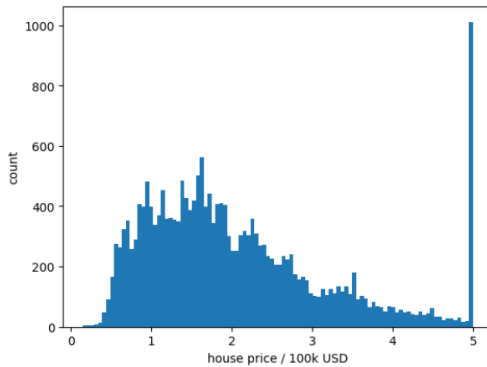
- this dataset is small enough for quick testing and tuning

California housing dataset

Data matrix								
	MedInc	HouseAge	AveRooms	AveBedrms	Population	AveOccup	Latitude	Longitude
0	8.3252	41.0	6.984127	1.023810	322.0	2.555556	37.88	-122.23
1	8.3014	21.0	6.238137	0.971880	2401.0	2.109842	37.86	-122.22
2	7.2574	52.0	8.288136	1.073446	496.0	2.802260	37.85	-122.24
3	5.6431	52.0	5.817352	1.073059	558.0	2.547945	37.85	-122.25
4	3.8462	52.0	6.281853	1.081081	565.0	2.181467	37.85	-122.25
...	...	...	...	...	...	...	...	...
20635	1.5603	25.0	5.045455	1.133333	845.0	2.560606	39.48	-121.09
20636	2.5568	18.0	6.114035	1.315789	356.0	3.122807	39.49	-121.21
20637	1.7000	17.0	5.205543	1.120092	1007.0	2.325635	39.43	-121.22
20638	1.8672	18.0	5.329513	1.171920	741.0	2.123209	39.43	-121.32
20639	2.3886	16.0	5.254717	1.162264	1387.0	2.616981	39.37	-121.24

Target value	
0	4.526
1	3.585
2	3.521
3	3.413
4	3.422
...	...
20635	0.781
20636	0.771
20637	0.923
20638	0.847
20639	0.894

histogram of the house price



Plotting the histogram

```
1 plt.hist(housing.target, bins=100)
2 plt.xlabel("house price / 100k USD")
3 plt.ylabel("count")
```

Observation?

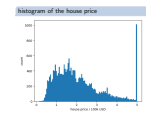
- note that the maximum price is 5k USD
- the histogram has a clear spike in that price

Lecture 2: Classical methods and models in MLE with Scikit Learn

Scikit Learn

Toy dataset

California housing dataset — targets



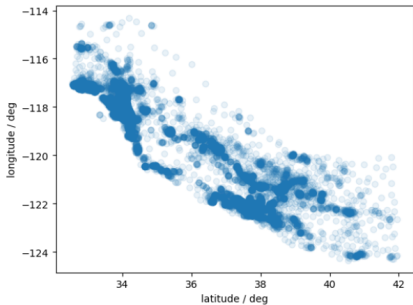
Plotting the histogram

```
plt.hist(housing.target, bins=100)
plt.xlabel("house price / 100k USD")
plt.ylabel("count")
```

Observation?

- note that the maximum price is 5k USD
- the histogram has a clear spike in that price

Location density of data samples



Plotting the histogram

```
1 plt.scatter(housing.data["Latitude"],  
             housing.data["Longitude"], alpha=0.1)  
2 plt.xlabel("latitude / deg")  
3 plt.ylabel("longitude / deg")
```

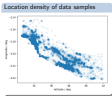
Lecture 2: Classical methods and models in MLE with Scikit Learn

Scikit Learn

Toy dataset

California housing dataset — locations

California housing dataset — locations



Plotting the histogram

```
plt.scatter(housing.data["Latitude"],  
            housing.data["Longitude"], alpha=0.1)  
plt.xlabel("latitude / deg")  
plt.ylabel("longitude / deg")
```

Loading X and y of the dataset

`return_X_y=True`

```
1 from sklearn.datasets import fetch_california_housing
2
3 X, y = fetch_california_housing(return_X_y=True)
```

X

```
[[ 8.3252  41.  6.98412698 ... 2.55555556
 37.88 -122.23 ] 6.23813708 ... 2.10984183
 [ 8.3814  21.  6.23813708 ... 2.10984183
 37.86 -122.22 ] 8.28813559 ... 2.80225989
 [ 7.2574  52.  8.28813559 ... 2.80225989
 37.85 -122.24 ] ...
 [ 1.7  17.  5.20554273 ... 2.3256351
 39.43 -121.22 ] 5.32951289 ... 2.12320917
 [ 1.8672  18.  5.32951289 ... 2.12320917
 39.43 -121.32 ] 5.25471698 ... 2.61698113
 [ 2.3886  16.  5.25471698 ... 2.61698113
 39.37 -121.24 ]]
```

y

[4.526 3.585 3.521 ... 0.923 0.847 0.894]

Lecture 2: Classical methods and models in MLE with Scikit Learn

Scikit Learn

Toy dataset

California housing dataset — loading for learning

2025-10-11

- X is a 2D matrix with individual attributes in the columns
- y is a 1D vector
- the scikit API is build to work with corresponding X and y

California housing dataset — loading for learning

Loading X and y of the dataset

return_X_y=True

```
from sklearn.datasets import fetch_california_housing
X, y = fetch_california_housing(return_X_y=True)
```

X

```
[[ 8.3252  41.  6.98412698 ... 2.55555556
 37.88 -122.23 ] 6.23813708 ... 2.10984183
 [ 8.3814  21.  6.23813708 ... 2.10984183
 37.86 -122.22 ] 8.28813559 ... 2.80225989
 [ 7.2574  52.  8.28813559 ... 2.80225989
 37.85 -122.24 ] ...
 [ 1.7  17.  5.20554273 ... 2.3256351
 39.43 -121.22 ] 5.32951289 ... 2.12320917
 [ 1.8672  18.  5.32951289 ... 2.12320917
 39.43 -121.32 ] 5.25471698 ... 2.61698113
 [ 2.3886  16.  5.25471698 ... 2.61698113
 39.37 -121.24 ]]
```

y

```
[4.526 3.585 3.521 ... 0.923 0.847 0.894]
```

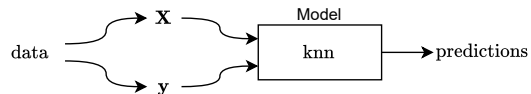
The *simplest* regression model

- the minimal example that *does something*
- using KNN to predict the price by finding nearest neighbors in the dataset and averaging their price

KNN regression pipeline

```
1 from sklearn.datasets import fetch_california_housing
2 from sklearn.neighbors import KNeighborsRegressor
3
4 X, y = fetch_california_housing(return_X_y=True)
5
6 model = KNeighborsRegressor()
7
8 # training
9 model.fit(X, y)
10
11 # making predictions
12 prediction = model.predict(X)
```

"The ML model"



Lecture 2: Classical methods and models in MLE with Scikit Learn

Scikit Learn

Regression pipeline

The *simplest* regression model

2025-10-11

- The *model* can be easily swapped for other model, e.g., the `LinearRegression`. All the sklearn models have the same interface.
- This approach has may other flaws besides the ones mentioned here. We are going to get to some of them :-).

The simplest regression model

```
1 # the minimal example
2 # that does something
3
4 from sklearn.datasets import fetch_california_housing
5 from sklearn.neighbors import KNeighborsRegressor
6
7 X, y = fetch_california_housing(return_X_y=True)
8
9 model = KNeighborsRegressor()
10
11 # training
12 model.fit(X, y)
13
14 # making predictions
15 prediction = model.predict(X)
```

The ML model



The *simplest* regression model

- the minimal example that *does something*
- using KNN to predict the price by finding nearest neighbors in the dataset and averaging their price

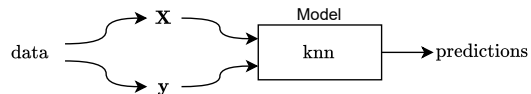
Possible flaws?

- we are using all the data to train the model at once
- we are doing our predictions on the same data we used for training

KNN regression pipeline

```
1 from sklearn.datasets import fetch_california_housing
2 from sklearn.neighbors import KNeighborsRegressor
3
4 X, y = fetch_california_housing(return_X_y=True)
5
6 model = KNeighborsRegressor()
7
8 # training
9 model.fit(X, y)
10
11 # making predictions
12 prediction = model.predict(X)
```

"The ML model"



Lecture 2: Classical methods and models in MLE with Scikit Learn

Scikit Learn

Regression pipeline

The *simplest* regression model

2025-10-11

The simplest regression model

```
1 # the minimal example
2 # that does something
3 from sklearn.datasets import fetch_california_housing
4 from sklearn.neighbors import KNeighborsRegressor
5 X, y = fetch_california_housing(return_X_y=True)
6 model = KNeighborsRegressor()
7 model.fit(X, y)
8
9 # making predictions
10 prediction = model.predict(X)
```

Possible flaws?

• we are using all the data to train the model at once

• we are doing our predictions on the same data we used for training

"The ML model"

```
graph LR
    data --> X
    data --> y
    X --> Model[knn]
    y --> Model
    Model --> predictions
```

- The *model* can be easily swapped for other model, e.g., the `LinearRegression`. All the sklearn models have the same interface.
- This approach has may other flaws besides the ones mentioned here. We are going to get to some of them :-).

The *simplest* regression model — results

Lecture 2:
Classical
methods
and
models in
MLE with
Scikit
Learn

Tomáš
Báča

Python
Environ-
ment

Pandas

Scikit
Learn

Toy dataset

Regression
pipeline

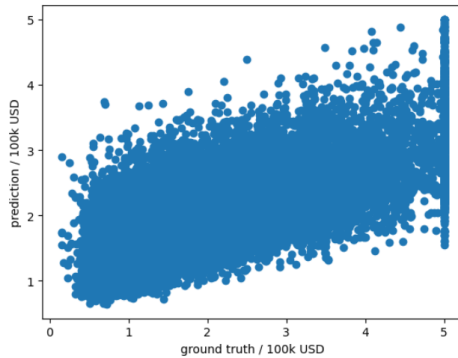
Preprocessing

Metrics

Realworld
example

References

Predictions vs. ground truth



Showing the predictions

```
1 plt.scatter(y, pred)
2 plt.xlabel("ground truth / 100k USD")
3 plt.ylabel("prediction / 100k USD")
```

Observations

- resemblance of a predictive behavior is there
- ideally, the data would form a line with slope = 1
- see the pattern caused by the capped house price at 500k USD?
- far from ideal, but this is a good start

Tomáš Báča (CTU in Prague)

Lecture 2: Classical methods and models in MLE with Scikit Learn

September 30th, 2025

14 / 48

Lecture 2: Classical methods and models in MLE with Scikit Learn

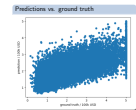
Scikit Learn

Regression pipeline

The *simplest* regression model — results

2025-10-11

The simplest regression model — results



Showing the predictions

```
plt.scatter(y, pred)
plt.xlabel("ground truth / 100k USD")
plt.ylabel("prediction / 100k USD")
```

Observations

- resemblance of a predictive behavior is there
- ideally, the data would form a line with slope = 1
- see the pattern caused by the capped house price at 500k USD?
- far from ideal, but this is a good start

The *simplest* regression model — improving?

Lecture 2:
Classical methods
and models in
MLE with
Scikit Learn

Tomáš
Báča

Python
Environ-
ment

Pandas

Scikit
Learn

Toy dataset
Regression
pipeline
Preprocessing

Metrics

Realworld
example

References

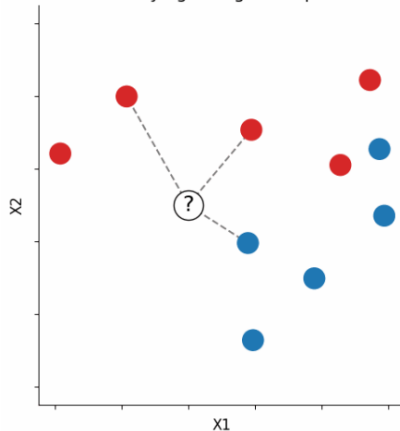
- our data has varying physical quantities in different dimensions
- one axis is in feet², other in degrees, other in multiples of 100k USD

Solution?

- scaling data such that their distribution is *similar*
- many ways how to do that, but let's start with a simple practical example

KNN

Classifying a single test point



Tomáš Báča (CTU in Prague)

Lecture 2: Classical methods and models in MLE with Scikit Learn

September 30th, 2025

15 / 48

Lecture 2: Classical methods and models in MLE with Scikit Learn

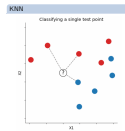
Scikit Learn

Regression pipeline

The *simplest* regression model — improving?

The simplest regression model — improving?

- our data has varying physical quantities in different dimensions
 - one axis is in feet², other in degrees, other in multiples of 100k USD
- Solution?**
- scaling data such that their distribution is *similar*
 - many ways how to do that, but let's start with a simple practical example



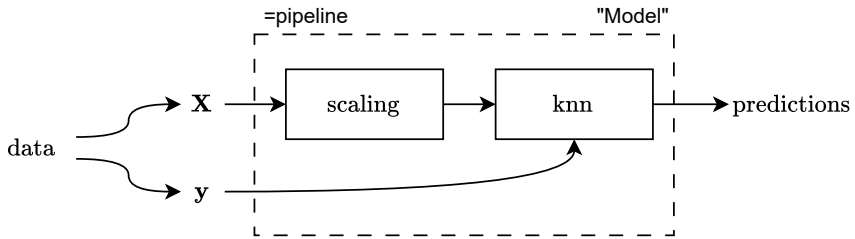
- KNN uses Minkowski metric with power = 2 (= l2 norm). If the attributes have different scales, then the metric will not be influenced evenly by the different attributes.

2025-10-11

The *simplest* regression model — creating a pipeline

- let's add a scaling block in front of our KNN regressor
- with two block, we are already forming a *pipeline*

The *simplest* pipeline diagram



Lecture 2: Classical methods and models in MLE with Scikit Learn

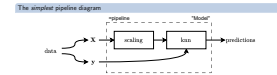
Scikit Learn

Regression pipeline

The *simplest* regression model — creating a pipeline

The simplest regression model — creating a pipeline

- let's add a scaling block in front of our KNN regressor
- with two block, we are already forming a pipeline



- the Pipeline has the same interface as the KNN regressor
 - .fit(X, y)
 - .predict(X)

KNN regression pipeline

```
1 from sklearn.datasets import fetch_california_housing
2 from sklearn.neighbors import KNeighborsRegressor
3 from sklearn.preprocessing import StandardScaler
4 from sklearn.pipeline import Pipeline
5
6 X, y = fetch_california_housing(return_X_y=True)
7
8 model = Pipeline([
9     ("scaler", StandardScaler()),
10    ("predictor", KNeighborsRegressor()),
11 ])
12
13 # training
14 model.fit(X, y)
15
16 # making predictions
17 prediction = model.predict(X)
```

Lecture 2: Classical methods and models in MLE with Scikit Learn

Scikit Learn

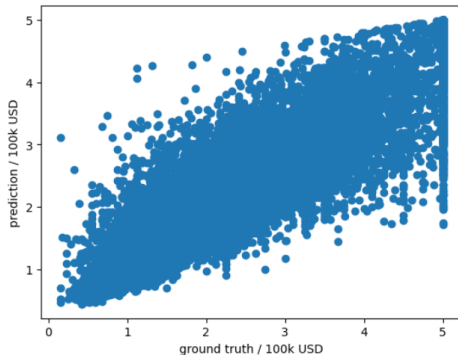
Regression pipeline

Adding preprocessing and pipeline

```
[KNN regression pipeline]
1 from sklearn.datasets import fetch_california_housing
2 from sklearn.neighbors import KNeighborsRegressor
3 from sklearn.preprocessing import StandardScaler
4 from sklearn.pipeline import Pipeline
5
6 X, y = fetch_california_housing(return_X_y=True)
7
8 # The Pipeline has the
9 # same interface as the
10 # KNN regressor
11
12 model = Pipeline([
13     ("scaler", StandardScaler()),
14     ("predictor", KNeighborsRegressor()),
15 ])
16
17 # training
18 model.fit(X, y)
19
20 # making predictions
21 prediction = model.predict(X)
```

The *simplest* regression model — results

Predictions vs. ground truth now



Showing the predictions

```
1 plt.scatter(y, pred)
2 plt.xlabel("ground truth / 100k USD")
3 plt.ylabel("prediction / 100k USD")
```

Observations

- better performance than **before**
- as simple to use as **before**

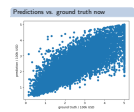
Lecture 2: Classical methods and models in MLE with Scikit Learn

Scikit Learn

Regression pipeline

The *simplest* regression model — results

The simplest regression model — results



Showing the predictions

```
plt.scatter(y, pred)
plt.xlabel("ground truth / 100k USD")
plt.ylabel("prediction / 100k USD")
```

Observations

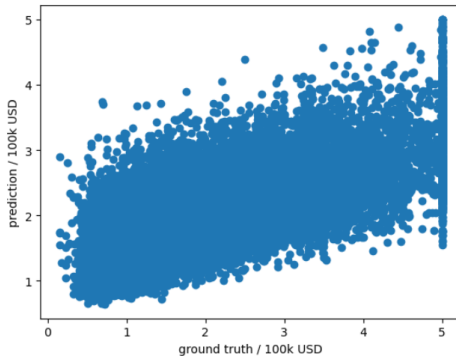
- better performance than **before**
- as simple to use as **before**

- The results already looks better with just the *StandardScaler*.
- But what does the standard scaler do?
 - more on that later, but in the meantime:
 - the standard scaler subtracts the mean of the data and then divides the data by the standard deviation.

The *simplest* regression model — results

Lecture 2:
Classical methods
and models in
MLE with
Scikit Learn
Tomas
Báča
Python
Environ-
ment
Pandas
Scikit
Learn
Toy dataset
Regression
pipeline
Preprocessing
Metrics
Realworld
example
References

Predictions vs. ground truth *before*



Showing the predictions

```
1 plt.scatter(y, pred)
2 plt.xlabel("ground truth / 100k USD")
3 plt.ylabel("prediction / 100k USD")
```

Observations

- better performance than *before*
- as simple to use as *before*

Tomas Báča (CTU in Prague)

Lecture 2: Classical methods and models in MLE with Scikit Learn

September 30th, 2025

18 / 48

Lecture 2: Classical methods and models in MLE with Scikit Learn

Scikit Learn

Regression pipeline

The *simplest* regression model — results

The simplest regression model — results



Showing the predictions

```
1 plt.scatter(y, pred)
2 plt.xlabel("ground truth / 100k USD")
3 plt.ylabel("prediction / 100k USD")
```

Observations

- better performance than *before*
- as simple to use as *before*

- The results already looks better with just the *StandardScaler*.
- But what does the standard scaler do?
 - more on that later, but in the meantime:
 - the standard scaler subtracts the mean of the data and then divides the data by the standard deviation.

2025-10-11

Problems with current approach?

1. model parameters

- our model might have parameters that might need changing to improve “our *performance*”

Let's add `n_neighbors=1`

```
1 ...  
2  
3 model = Pipeline([  
4     ("scaler", StandardScaler()),  
5     ("predictor", KNeighborsRegressor(  
6         n_neighbors=1)),  
7     ])  
8 ...
```

Lecture 2: Classical methods and models in MLE with Scikit Learn

Scikit Learn

Regression pipeline

Problems with current approach?

Problems with current approach?

1. model parameters

• our model might have parameters that might need changing to improve “our performance”

Let's add `n_neighbors=1`

```
model = Pipeline([  
    ("scaler", StandardScaler()),  
    ("predictor", KNeighborsRegressor(  
        n_neighbors=1)),  
])
```

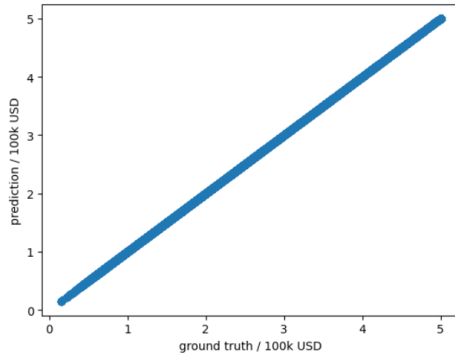
- When setting the KNN to use only 1 neighbor, the methods suddenly produces perfect predictions.
- That is caused by the KNN being asked to predict using the same samples it remembered during the training.
- So now we know we have a problem:
 - There might be hyper parameters that influence the performance we don't know which values to pick.
 - We can not use the same dataset for training as well as for evaluation, that gives us false sense of success.

Problems with current approach?

1. model parameters

- our model might have parameters that might need changing to improve “our *performance*”

Predictions vs. ground truth for 1 neighbor



Let's add `n_neighbors=1`

```
1 ...  
2  
3 model = Pipeline([  
4     ("scaler", StandardScaler()),  
5     ("predictor", KNeighborsRegressor(  
6         n_neighbors=1)),  
7 ])  
8 ...
```

This leads to problem #2

- our performance evaluation method is not very good

Lecture 2: Classical methods and models in MLE with Scikit Learn

Scikit Learn

Regression pipeline

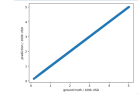
Problems with current approach?

Problems with current approach?

1. model parameters

- our model might have parameters that might need changing to improve “our performance”

Predictions vs. ground truth for 1 neighbor



Let's add `n_neighbors=1`

```
model = Pipeline([  
    ("scaler", StandardScaler()),  
    ("predictor", KNeighborsRegressor(  
        n_neighbors=1)),  
])
```

This leads to problem #2

- our performance evaluation method is not very good

- When setting the KNN to use only 1 neighbor, the methods suddenly produces perfect predictions.
- That is caused by the KNN being asked to predict using the same samples it remembered during the training.
- So now we know we have a problem:
 - There might be hyper parameters that influence the performance we don't know which values to pick.
 - We can not use the same dataset for training as well as for evaluation, that gives us false sense of success.

Problems with current approach?

2. our evaluation method is wrong

- we are checking performance of the model using the same data we trained on
- this makes selecting the right parameters impossible (for other data than the one we are training on)

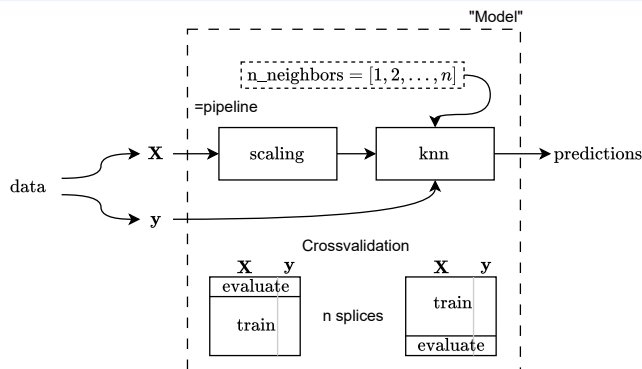
Grid search / hyperparameter optimization

- selecting parameter of *block* that can not be trained by the block itself
- iterates over all combinations of parameters

Crossvalidation

- splicing data to several pieces
- training on majority of the data while testing on the rest
- then swapping the splice

Using grid search with crossvalidation



Lecture 2: Classical methods and models in MLE with Scikit Learn

Scikit Learn

Regression pipeline

Problems with current approach?

2025-10-11

Problems with current approach?

2. our evaluation method is wrong

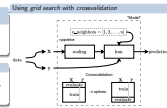
- we are checking performance of this model using the same data we trained on
- this makes selecting the right parameters impossible (for other data than the one we are training on)

Grid search / hyperparameter optimization

- selecting parameter of block that can not be trained by the block itself
- iterates over all combinations of parameters

Crossvalidation

- splicing data to several pieces
- training on majority of the data while testing on the rest
- then swapping the splice



Observations

- GridSearchCV has the same interface
 - `.fit(X, y)`
 - `.predict(X)`
- whole parts of the pipeline can be disabled/enabled

Checking the results

- in Jupyter notebook

```
1 pd.DataFrame(model.cv_results_)
```

Introducing GridSearchCV

```

1 ...
2
3 from sklearn.model_selection import
   GridSearchCV
4
5 pipeline = Pipeline([
6     ("scaler", StandardScaler()),
7     ("predictor", KNeighborsRegressor()),
8 ])
9
10 model = GridSearchCV(
11     pipeline,
12     param_grid={"predictor__n_neighbors":
13                 [1, 2, 3, 4, 5, 6]
14                 },
15     cv=3
16 )
17
18 model.fit(X, y)
19
20 ...

```

Lecture 2: Classical methods and models in MLE with Scikit Learn

Scikit Learn

Regression pipeline

GridSearch with Crossvalidation

GridSearch with Crossvalidation

Observations

• GridSearchCV has the same interface
• `.fit(X, y)`
• `.predict(X)`
• whole parts of the pipeline can be disabled/enabled

Checking the results

• in Jupyter notebook

```
pd.DataFrame(model.cv_results_)
```

Introducing GridSearchCV

```

from sklearn.model_selection import
GridSearchCV

pipeline = Pipeline([
    ("scaler", StandardScaler()),
    ("predictor", KNeighborsRegressor()),
])

model = GridSearchCV(
    pipeline,
    param_grid={"predictor__n_neighbors":
                [1, 2, 3, 4, 5, 6]
                },
    cv=3
)

model.fit(X, y)

```

- whole parts of the pipeline can be skipped by replacing them using the *passthrough* step in the `param_grid`:
'scaler': [StandardScaler(), 'passthrough']

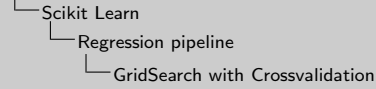
Grid search results

	mean_fit_time	std_fit_time	mean_score_time	std_score_time	param_predictor_n_neighbors	params	split0_test_score	split1_test_score	split2_test_score	mean_test_score	std_test_score	rank_test_score
0	0.009950	0.000354	0.218222	0.016866	1	(predictor_n_neighbors: 1)	0.324068	0.334830	0.323371	0.327423	0.005245	5
1	0.009837	0.000220	0.246559	0.021540	2	(predictor_n_neighbors: 2)	0.468788	0.503457	0.424388	0.465544	0.032361	6
2	0.009659	0.000020	0.269420	0.019367	3	(predictor_n_neighbors: 3)	0.518547	0.543340	0.473595	0.511827	0.028867	4
3	0.009642	0.000061	0.275874	0.021485	4	(predictor_n_neighbors: 4)	0.540323	0.564974	0.499827	0.535041	0.026857	3
4	0.010023	0.000212	0.289105	0.020307	5	(predictor_n_neighbors: 5)	0.551149	0.579313	0.511781	0.547414	0.027696	2
5	0.009678	0.000049	0.305267	0.028973	6	(predictor_n_neighbors: 6)	0.558435	0.586185	0.521134	0.555251	0.026652	1

Observations

- the best KNN model is the one with 6 neighbors (makes sense)
- the model won with a *mean_test_score* of 0.555
 - what score is it
 - how is it calculated?
 - how do we change it?
 - we will get back to this

Lecture 2: Classical methods and models in MLE with Scikit Learn



GridSearch with Crossvalidation

Grid search results

	mean_fit_time	std_fit_time	mean_score_time	std_score_time	param_predictor_n_neighbors	params	split0_test_score	split1_test_score	split2_test_score	mean_test_score	std_test_score	rank_test_score
0	0.009950	0.000354	0.218222	0.016866	1	(predictor_n_neighbors: 1)	0.324068	0.334830	0.323371	0.327423	0.005245	5
1	0.009837	0.000220	0.246559	0.021540	2	(predictor_n_neighbors: 2)	0.468788	0.503457	0.424388	0.465544	0.032361	6
2	0.009659	0.000020	0.269420	0.019367	3	(predictor_n_neighbors: 3)	0.518547	0.543340	0.473595	0.511827	0.028867	4
3	0.009642	0.000061	0.275874	0.021485	4	(predictor_n_neighbors: 4)	0.540323	0.564974	0.499827	0.535041	0.026857	3
4	0.010023	0.000212	0.289105	0.020307	5	(predictor_n_neighbors: 5)	0.551149	0.579313	0.511781	0.547414	0.027696	2
5	0.009678	0.000049	0.305267	0.028973	6	(predictor_n_neighbors: 6)	0.558435	0.586185	0.521134	0.555251	0.026652	1

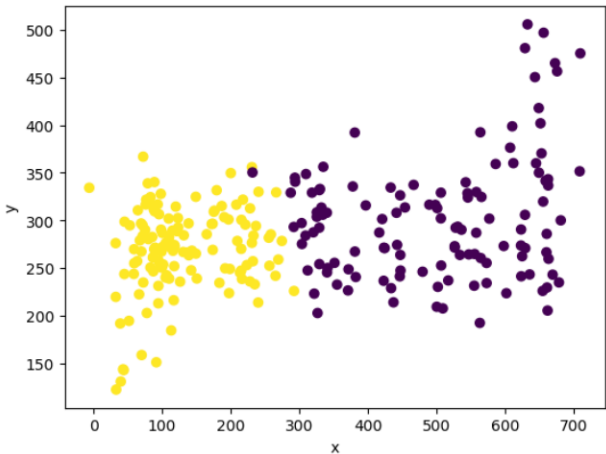
Observations

- the best KNN model is the one with 6 neighbors (makes sense)
- the model won with a *mean_test_score* of 0.555
 - what score is it
 - how is it calculated?
 - how do we change it?
 - we will get back to this

Observations

- two classes
- should be nicely separable
- outliers near the edges
- data in arbitrary range on both axes

A dataset with outliers

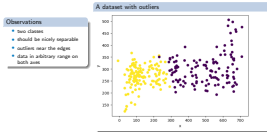


Lecture 2: Classical methods and models in MLE with Scikit Learn

Scikit Learn

Preprocessing

Data preprocessing — more about sklearn's transformers



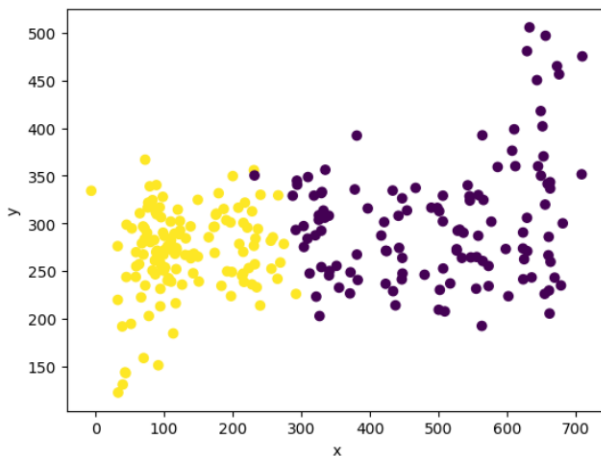
Observations

- two classes
- should be nicely separable
- outliers near the edges
- data in arbitrary range on both axes

Let's try standard scaler

- what does it even do?
- how will it perform?
- will it be influenced by the outliers?

A dataset with outliers

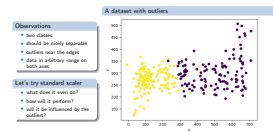


Lecture 2: Classical methods and models in MLE with Scikit Learn

Scikit Learn

Preprocessing

Data preprocessing — more about sklearn's transformers



Standard scaler

$$z = \frac{x - \mu}{\sigma},$$

where

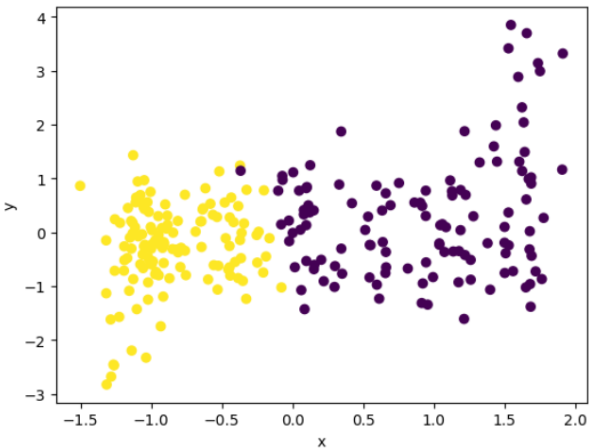
$$\mu = \frac{1}{N} \sum_{i=1}^N x_i$$

$$\sigma = \sqrt{\frac{1}{N} \sum_{i=1}^N (x_i - \mu)^2}$$

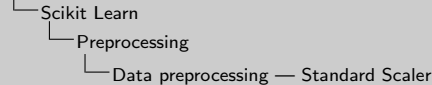
Observations

- outliers still visible
- axes' range is *near zero*
- the shape of the data is similar

Dataset after transformation



Lecture 2: Classical methods and models in MLE with Scikit Learn



Data preprocessing — Standard Scaler

Standard scaler:

$$z = \frac{x - \mu}{\sigma},$$

where

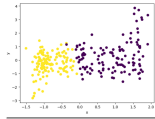
$$\mu = \frac{1}{N} \sum_{i=1}^N x_i$$

$$\sigma = \sqrt{\frac{1}{N} \sum_{i=1}^N (x_i - \mu)^2}$$

Observations

- outliers still visible
- axes' range is near zero
- the shape of the data is similar

Dataset after transformation



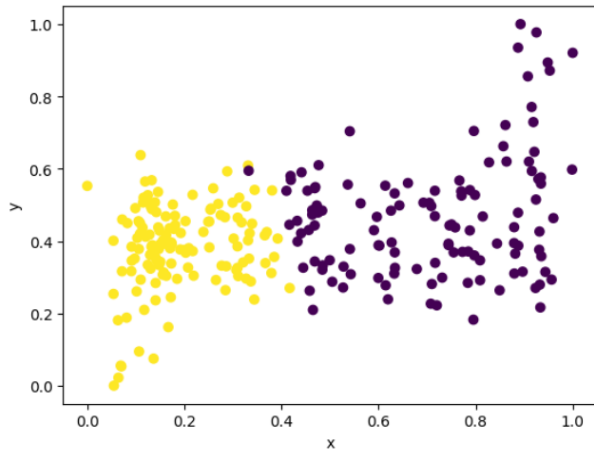
MinMax scaler

$$z = \frac{x - \min}{\max - \min},$$

Observations

- outliers still visible
- axes' range is between 0 and 1
- distribution is preserved

Dataset after transformation

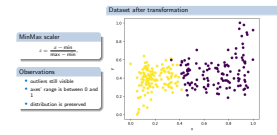


Lecture 2: Classical methods and models in MLE with Scikit Learn

Scikit Learn

Preprocessing

Data preprocessing — MinMax scaler



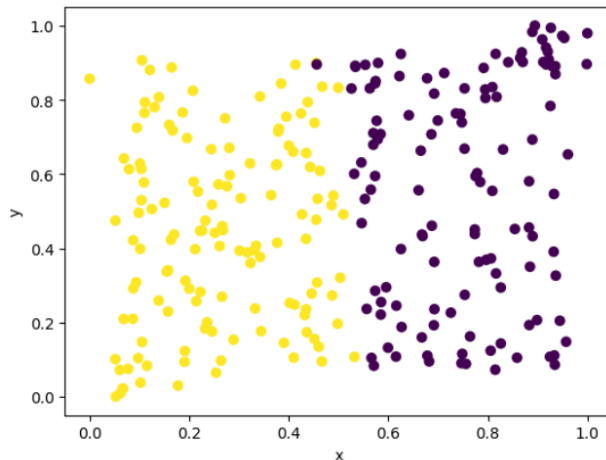
Quantile transformer

- reshapes the data to have uniform (or normal) distribution
- robust scaler

Observations

- uniform output (default)
- outliers are not prominent anymore
- non-linear transformation
 - will break linear correlations between features

Dataset after transformation (uniform, 10 quantiles)

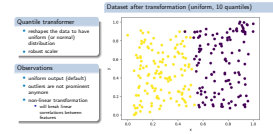


Lecture 2: Classical methods and models in MLE with Scikit Learn

Scikit Learn

Preprocessing

Data preprocessing — Quantile transformer



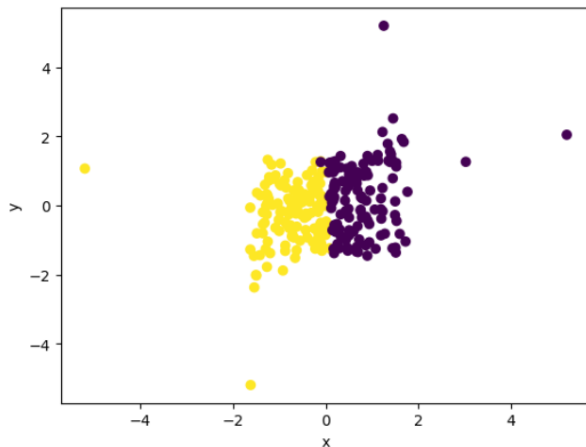
Quantile transformer

- reshapes the data to have uniform (or normal) distribution
- robust scaler

Observations

- uniform output (default)
- outliers are not prominent anymore
- non-linear transformation
 - will break linear correlations between features

Dataset after transformation (normal, 10 quantiles)

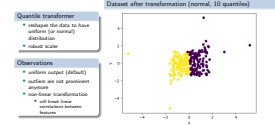


Lecture 2: Classical methods and models in MLE with Scikit Learn

Scikit Learn

Preprocessing

Data preprocessing — Quantile transformer



Linear separability

Linearly non-separable two-class dataset

Problem?

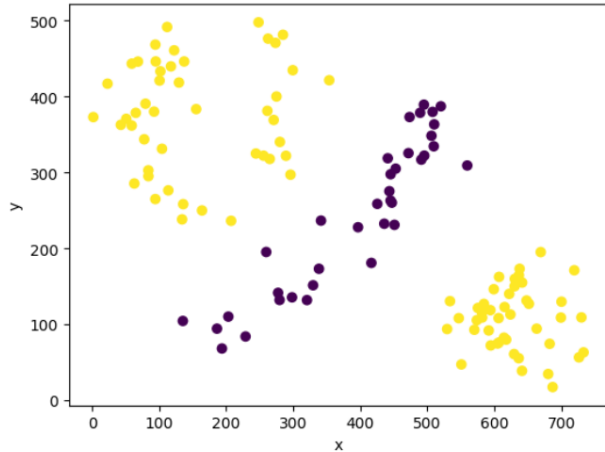
- classes are not linearly separable
- a.k.a., a single separating hyperplane can not be found

Solution?

- use more complex classifier
 - nonlinear
 - e.g., boosting
 - decision tree

Easier solution?

- use more complex classifier
- lift feature to new dimension



Lecture 2: Classical methods and models in MLE with Scikit Learn

- Scikit Learn
 - Preprocessing
 - Linear separability

Linear separability

Problem?

- classes are not linearly separable
- a.k.a., a single separating hyperplane can not be found

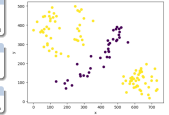
Solution?

- use more complex classifier
 - nonlinear
 - e.g., boosting
 - decision tree

Easier solution?

- use more complex classifier
- lift feature to new dimension

Linearly non-separable two-class dataset



Linear separability

Lecture 2:
Classical methods
and models in
MLE with
Scikit
Learn

Tomáš
Báča

Python
Environ-
ment

Pandas

Scikit
Learn

Toy dataset
Regression
pipeline
Preprocessing

Metrics

Realworld
example

References

Problem?

- classes are not linearly separable
- a.k.a., a single separating hyperplane can not be found

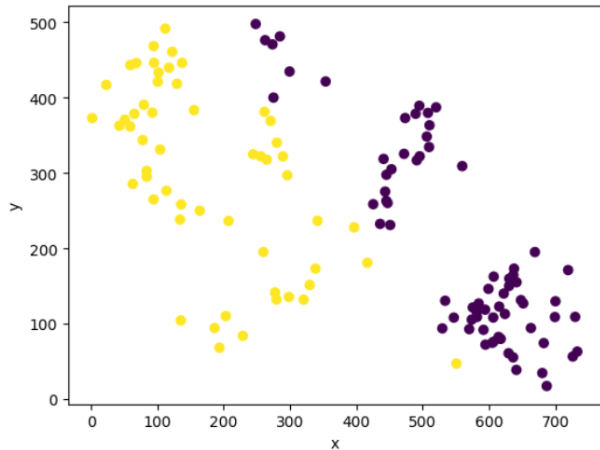
Solution?

- use more complex classifier
 - nonlinear
 - e.g., boosting
 - decision tree

Easier solution?

- use more complex classifier
- lift feature to new dimension

Badly-classified by linear classifier



Tomáš Báča (CTU in Prague)

Lecture 2: Classical methods and models in MLE with Scikit Learn

September 30th, 2025

27 / 48

Lecture 2: Classical methods and models in MLE with Scikit Learn

Scikit Learn

Preprocessing

Linear separability

2025-10-11

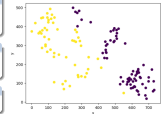
Linear separability

Problem?
• classes are not linearly separable
• a.k.a., a single separating hyperplane can not be found

Solution?
• use more complex classifier
• nonlinear
• e.g., boosting
• decision tree

Easier solution?
• use more complex classifier
• lift feature to new dimension

Badly-classified by linear classifier



Lifting features using polynomial interactions

- example for order = 2

$$\mathbf{x} = [x_1, x_2]^T$$

$$\mathbf{x}_{\text{new}} = \Phi(\mathbf{x}) = [1, x_1, x_2, x_1^2, x_1x_2, x_2^2]^T$$

Observations

- leads to simple classifier, which can be robust
- e.g., all favourite SVM
- makes computations more expensive during production
- can be simplified with *the kernel trick*

Using polynomial features

```

1  ...
2
3  from sklearn.preprocessing import
   PolynomialFeatures
4  from sklearn.linear_model import
   LogisticRegression
5
6  pipeline = Pipeline([
7      ("scale", PolynomialFeatures()),
8      ("classifier", LogisticRegression(
9          max_iter=1000, class_weight="balanced"
10         )),
11 ])
12
13 pred = pipeline.fit(X, y).predict(X)
14
15 ...

```

Lecture 2: Classical methods and models in MLE with Scikit Learn

Scikit Learn

Preprocessing

Polynomial features

2025-10-11

Polynomial features

Lifting features using polynomial interactions

• example for order = 2
 $\mathbf{x} = [x_1, x_2]^T$
 $\mathbf{x}_{\text{new}} = \Phi(\mathbf{x}) = [1, x_1, x_2, x_1^2, x_1x_2, x_2^2]^T$

Observations

- leads to simple classifier, which can be robust
- e.g. all favourite SVM
- makes computations more expensive during production
- can be simplified with the kernel trick

Using polynomial features

```

from sklearn.preprocessing import
PolynomialFeatures
from sklearn.linear_model import
LogisticRegression

pipeline = Pipeline([
    ("scale", PolynomialFeatures()),
    ("classifier", LogisticRegression(
        max_iter=1000, class_weight="balanced"
    )),
])

pred = pipeline.fit(X, y).predict(X)

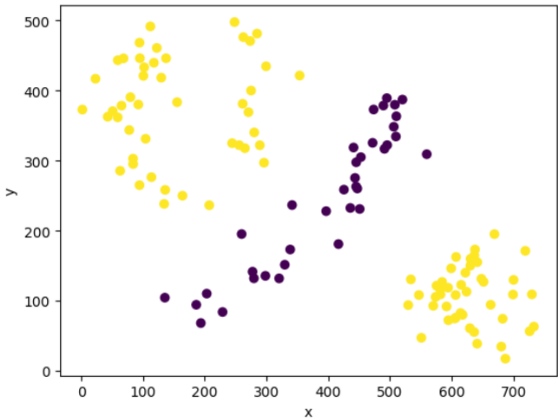
```

- The *Kernel trick* with the SVM (or other classifier that uses a dot product between the attributes and its weights) is about doing just the dot product using the lifted dimensions but training the classified with the original dimension.

Observations

- simple to use
- just like any other Sklearn transformer

Logistic regression with polynomial features



Lecture 2: Classical methods and models in MLE with Scikit Learn

Scikit Learn

Preprocessing

Polynomial features

Polynomial features

- Observations
- simple to use
 - just like any other Sklearn transformer

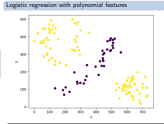
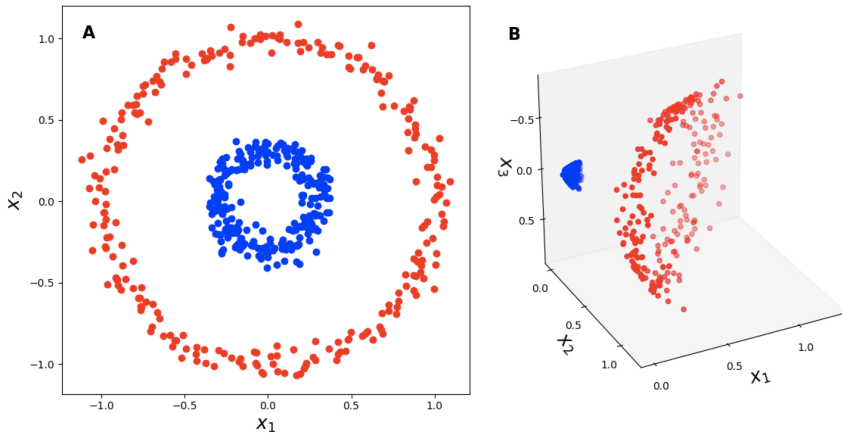


Illustration of lifting from 2D to 3D



Source: <https://gregorygundersen.com/blog/2019/12/10/kernel-trick/>

Lecture 2: Classical methods and models in MLE with Scikit Learn

Scikit Learn

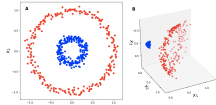
Preprocessing

Polynomial features

2025-10-11

Polynomial features

Illustration of lifting from 2D to 3D



Source: <https://gregorygundersen.com/blog/2019/12/10/kernel-trick/>

Encoding text and categorical attributes

Using OneHot encoder

How to encode text?

- added unique ID in 1D space might cause problems
- better way?
 - give each class a unique unit vector in its own dimension
 - this this creates as many dimensions (+1) as the # of classes

Even better way?

- learn unique vectors in lower dimensional space
- called **embedding**
- heart (mouth) of transformers

```
1 from sklearn.preprocessing import OneHotEncoder
2
3 X = [['red'], ['green'], ['blue'], ['red']]
4
5 encoder = OneHotEncoder(handle_unknown="ignore")
6 encoder.fit(X)
7
8 new_X = encoder.transform(X).todense()
```

```
1 matrix([[0., 0., 1.],
2         [0., 1., 0.],
3         [1., 0., 0.],
4         [0., 0., 1.]])
```

Inverse:

```
1 encoder.inverse_transform([[0, 1, 0]])
```

```
1 [['green']]
```

Lecture 2: Classical methods and models in MLE with Scikit Learn

Scikit Learn

Preprocessing

Encoding text and categorical attributes

How to encode text?

- added unique ID in 1D space might cause problems
- better way?
 - give each class a unique unit vector in its own dimension
 - this this creates as many dimensions (+1) as the # of classes

Even better way?

- learn unique vectors in lower dimensional space
- called **embedding**
- heart (mouth) of transformers

Using Onehot encoder

```
1 from sklearn.preprocessing import OneHotEncoder
2
3 X = [['red'], ['green'], ['blue'], ['red']]
4
5 encoder = OneHotEncoder(handle_unknown="ignore")
6 encoder.fit(X)
7
8 new_X = encoder.transform(X).todense()
```

```
1 matrix([[0., 0., 1.],
2         [0., 1., 0.],
3         [1., 0., 0.],
4         [0., 0., 1.]])
```

Inverse:

```
1 encoder.inverse_transform([[0, 1, 0]])
2
3 [['green']]
```

Objects in Scikit

Lecture 2:
Classical methods
and models in
MLE with
Scikit
Learn

Tomáš
Báča

Python
Environ-
ment

Pandas

Scikit
Learn

Toy dataset

Regression
pipeline

Preprocessing

Metrics

Realworld
example

References

Estimator

- implements `.fit()`
- only consumes data
- e.g., some statistics collector

Transformer

- implements `.fit()`, `.transform()`
- scalers
- dimensionality reducers (PCA)
- imputers (handle missing values)

Predictor

- implements `.fit()`, `.predict()`
- regressors, classifiers

Pipeline

- is a Predictor

GridSearchCV

- is a Predictor

Tomáš Báča (CTU in Prague)

Lecture 2: Classical methods and models in MLE with Scikit Learn

September 30th, 2025

31 / 48

Lecture 2: Classical methods and models in MLE with Scikit Learn

Scikit Learn

Preprocessing

Objects in Scikit

2025-10-11

Objects in Scikit

Estimator

- implements `.fit()`
- only consumes data
- e.g., some statistics collector

Transformer

- implements `.fit()`, `.transform()`
- scalers
- dimensionality reducers (PCA)
- imputers (handle missing values)

Predictor

- implements `.fit()`, `.predict()`
- regressors, classifiers

Pipeline

- is a Predictor

GridSearchCV

- is a Predictor

Precision

$$\text{precision} = \frac{TP}{TP + FP}$$

- given I mark the sample as relevant, high often am I right

Recall

$$\text{recall} = \frac{TP}{TP + FN}$$

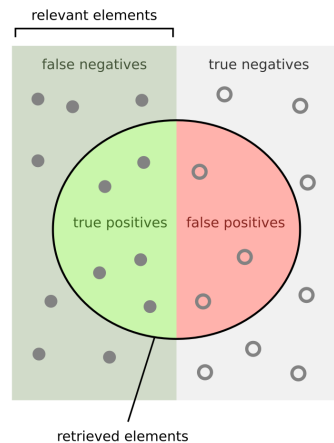
- did I marked all the relevant instances?

Accuracy

$$\text{accuracy} = \frac{TP + TN}{TP + TN + FP + FN}$$

- how precise am I in marking correctly both instances

Diagram of possible results



Metrics example

Lecture 2:
Classical methods
and models in
MLE with
Scikit
Learn

Tomáš
Báča

Python
Environ-
ment

Pandas

Scikit
Learn

Toy dataset
Regression
pipeline
Preprocessing

Metrics

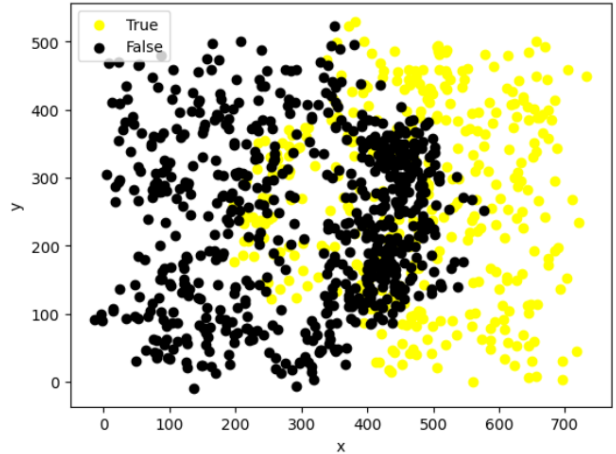
Realworld
example

References

Two non-separable classes

Is there a single best classifier?

- no, it depends on our metrics
- do we care more about
 - precision?
 - accuracy?
 - recall?
- or some combination of these?



Tomáš Báča (CTU in Prague)

Lecture 2: Classical methods and models in MLE with Scikit Learn

September 30th, 2025

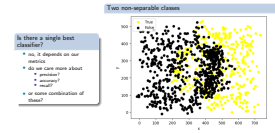
33 / 48

Lecture 2: Classical methods and models in MLE with Scikit Learn

Metrics

Metrics example

Metrics example



2025-10-11

Observations

- we are trying Logistic Regression
- we are varying the class weight
- we are searching for the *best-performing* predictor

Let's run this

- let's run it as it is
- the default score for logistic regression is *Accuracy*

Trying logistic regression

```

1  ...
2  from sklearn.linear_model import LogisticRegression
3  from sklearn.model_selection import GridSearchCV
4
5  pipeline = Pipeline([
6      ("classifier", LogisticRegression(max_iter=1000,
7      verbose=0)),
8  ])
9
10 model = GridSearchCV(
11     estimator=pipeline, cv=4, param_grid={
12         'classifier__class_weight': [{0: 5, 1: v} for v in
13         range(1, 10)]},
14 )
15
16 pred = model.fit(X, y).predict(X)
17 ...

```

Lecture 2: Classical methods and models in MLE with Scikit Learn

Metrics

Logistic regression with default metrics

Observations

- we are trying Logistic Regression
- we are varying the class weight
- we are searching for the *best-performing* predictor

Let's run this

- let's run it as it is
- the default score for logistic regression is *Accuracy*

Trying logistic regression

```

1  from sklearn.linear_model import LogisticRegression
2  from sklearn.model_selection import GridSearchCV
3
4  pipeline = Pipeline([
5      ("classifier", LogisticRegression(max_iter=1000,
6      verbose=0)),
7  ])
8
9  model = GridSearchCV(
10     estimator=pipeline, cv=4, param_grid={
11         'classifier__class_weight': [{0: 5, 1: v} for v in
12         range(1, 10)]},
13 )
14
15 pred = model.fit(X, y).predict(X)
16 ...

```

Metrics example — Accuracy

Lecture 2:
Classical methods
and models in
MLE with
Scikit
Learn

Tomáš
Báča

Python
Environ-
ment

Pandas

Scikit
Learn

Toy dataset
Regression
pipeline
Preprocessing

Metrics

Realworld
example

References

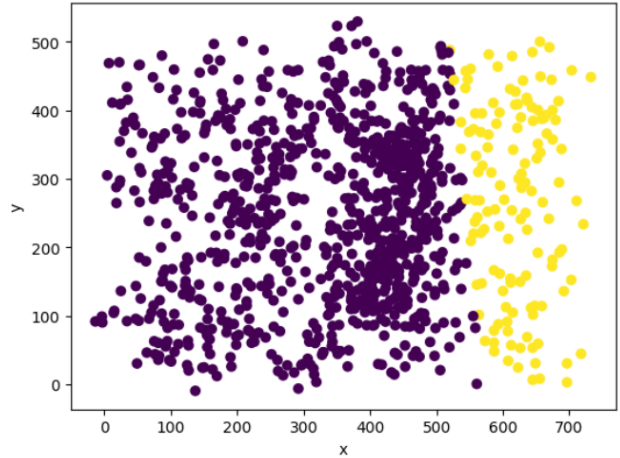
Logistic Regression selected by the Accuracy score

Observations

- the decision is *cut* at a reasonable compromise
- many FP as well as FN

Metrics on the dataset

- **accuracy** = 0.77
- **precision** = 0.97
- **recall** = 0.35



Tomáš Báča (CTU in Prague)

Lecture 2: Classical methods and models in MLE with Scikit Learn

September 30th, 2025

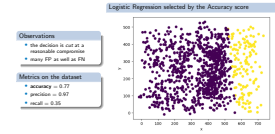
35 / 48

Lecture 2: Classical methods and models in MLE with Scikit Learn

Metrics

Metrics example — Accuracy

Metrics example — Accuracy



2025-10-11

Observations

- we are trying Logistic Regression
- we are varying the class weight
- we are searching for the *best-performing* predictor

Scoring

- new scoring argument
 - we can add arbitrary scoring functions
 - all will be evaluated during the grid search
- new refit argument
 - the pipeline will be re-trained using parameters that won using the pre-defined score

Logistic regression precision

```

1 from sklearn.linear_model import LogisticRegression
2 from sklearn.model_selection import GridSearchCV
3 from sklearn.metrics import precision_score,
  accuracy_score, recall_score
4
5 pipeline = Pipeline([
6     ("classifier", LogisticRegression(max_iter=1000,
7         verbose=0)),
8 ])
9
10 model = GridSearchCV(
11     estimator=pipeline, cv=4, param_grid={
12         'classifier__class_weight': [{0: 5, 1: v} for v in
13             range(1, 10)]},
14     scoring={'precision': make_scorer(precision_score),
15         'recall': make_scorer(recall_score), '
16         accuracy': make_scorer(accuracy_score)},
17     refit='precision',
18 )
19
20 pred = model.fit(X, y).predict(X)
    
```

Lecture 2: Classical methods and models in MLE with Scikit Learn

Metrics

Logistic regression with Precision metrics

Logistic regression with Precision metrics

Observations	Logistic regression precision
• we are trying Logistic Regression • we are varying the class weight • we are searching for the <i>best-performing</i> predictor	<pre> from sklearn.linear_model import LogisticRegression from sklearn.model_selection import GridSearchCV from sklearn.metrics import precision_score, accuracy_score, recall_score pipeline = Pipeline([("classifier", LogisticRegression(max_iter=1000, verbose=0)),]) model = GridSearchCV(estimator=pipeline, cv=4, param_grid={ 'classifier__class_weight': [{0: 5, 1: v} for v in range(1, 10)]}, scoring={'precision': make_scorer(precision_score), 'recall': make_scorer(recall_score), ' accuracy': make_scorer(accuracy_score)}, refit='precision',) pred = model.fit(X, y).predict(X) </pre>

Metrics example — Precision

Lecture 2:
Classical methods
and models in
MLE with
Scikit
Learn
Tomas
Bacha
Python
Environ-
ment
Pandas
Scikit
Learn
Toy dataset
Regression
pipeline
Preprocessing
Metrics
Realworld
example
References

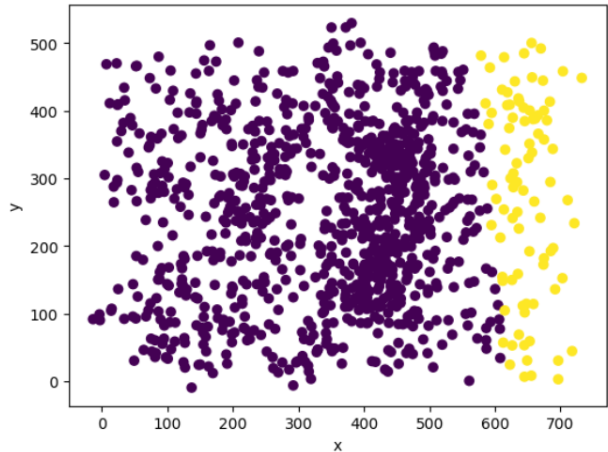
Observations

- the decision is *cut* at the very right
- minimizes false positives
- lot of missed *true* sample

Metrics on the dataset

- accuracy = 0.73
- **precision = 1.0**
- recall = 0.23

Logistic Regression selected by the Precision score



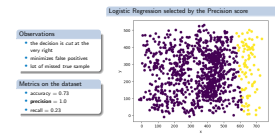
Lecture 2: Classical methods and models in MLE with Scikit Learn

Metrics

Metrics example — Precision

2025-10-11

Metrics example — Precision



Observations

- we are trying Logistic Regression
- we are varying the class weight
- we are searching for the *best-performing* predictor

Scoring

- new scoring argument
 - we can add arbitrary scoring functions
 - all will be evaluated during the grid search
- new refit argument
 - the pipeline will be re-trained using parameters that won using the pre-defined score

Logistic regression recall

```

1 from sklearn.linear_model import LogisticRegression
2 from sklearn.model_selection import GridSearchCV
3 from sklearn.metrics import precision_score,
  accuracy_score, recall_score
4
5 pipeline = Pipeline([
6     ("classifier", LogisticRegression(max_iter=1000,
7         verbose=0)),
8 ])
9
10 model = GridSearchCV(
11     estimator=pipeline, cv=4, param_grid={
12         'classifier__class_weight': [{0: 5, 1: v} for v in
13             range(1, 10)]},
14     scoring={'precision': make_scorer(precision_score),
15         'recall': make_scorer(recall_score), '
16         accuracy': make_scorer(accuracy_score)},
17     refit='recall',
18 )
19
20 pred = model.fit(X, y).predict(X)
    
```

Lecture 2: Classical methods and models in MLE with Scikit Learn

Metrics

Logistic regression with Recall metrics

Logistic regression with Recall metrics

Observations	Logistic regression recall
• we are trying Logistic Regression • we are varying the class weight • we are searching for the <i>best-performing</i> predictor	<pre> from sklearn.linear_model import LogisticRegression from sklearn.model_selection import GridSearchCV from sklearn.metrics import precision_score, accuracy_score, recall_score pipeline = Pipeline([("classifier", LogisticRegression(max_iter=1000, verbose=0)),]) model = GridSearchCV(estimator=pipeline, cv=4, param_grid={ 'classifier__class_weight': [{0: 5, 1: v} for v in range(1, 10)]}, scoring={'precision': make_scorer(precision_score), 'recall': make_scorer(recall_score), ' accuracy': make_scorer(accuracy_score)}, refit='recall',) pred = model.fit(X, y).predict(X) </pre>
Scoring	
• new scoring argument • we can add arbitrary scoring functions • all will be evaluated during the grid search • new refit argument • the pipeline will be re-trained using parameters that won using the pre-defined score	

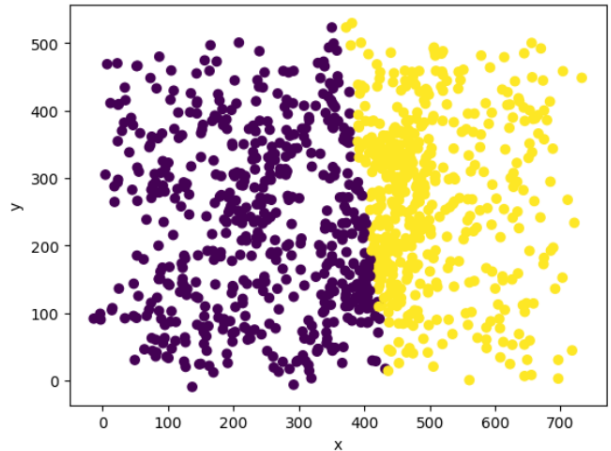
Observations

- the decision is *cut* at the very right
- minimizes false positives
- lot of missed *true* sample

Metrics on the dataset

- accuracy = 0.67
- precision = 0.52
- **recall = 0.72**

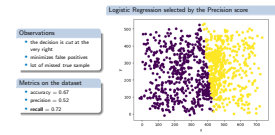
Logistic Regression selected by the Precision score



Lecture 2: Classical methods and models in MLE with Scikit Learn

Metrics

Metrics example — Recall



- as simple as dumping the object into a file
- alternatively, you can use pickle

Storing a model

```
1 ...  
2  
3 from joblib import dump  
4  
5 dump(model, "my_model.joblib")  
6  
7 ...
```

Loading a model

```
1 ...  
2  
3 from joblib import load  
4  
5 model = load("my_model.joblib")  
6  
7 ...
```

Lecture 2: Classical methods and models in MLE with Scikit Learn

Metrics

Storing a trained sklearn model

Storing a trained sklearn model

- * as simple as dumping the object into a file
- * alternatively, you can use pickle

Storing a model

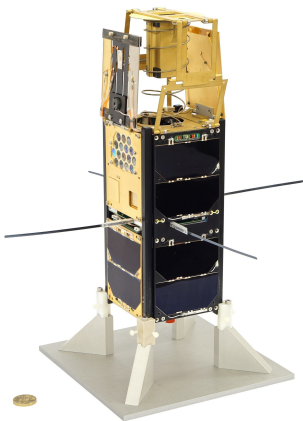
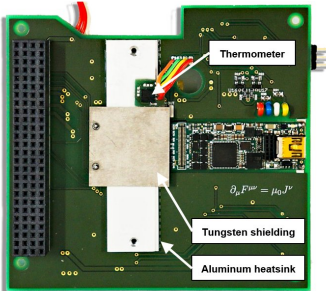
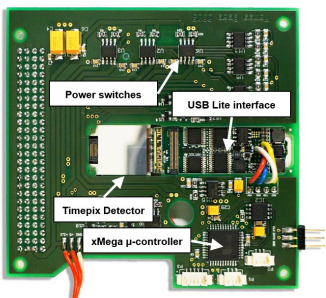
```
from joblib import dump  
dump(model, "my_model.joblib")
```

Loading a model

```
from joblib import load  
model = load("my_model.joblib")
```

VZLUSAT-1 — The first Czech CubeSat [1], [2]

- Custom-designed embedded Timepix board and software [2]
- Onboard image processing, filtering, automated acquisition
- The longest-operating Czech(-oslovakian) sat. (2017–2023)



Lecture 2: Classical methods and models in MLE with Scikit Learn

Realworld example

Realworld example — Realtime classification of ionizing particles

VZLUSAT-1 — The first Czech CubeSat [1], [2]
• Custom-designed embedded Timepix board and software [2]
• Onboard image processing, filtering, automated acquisition
• The longest operating Czech(-oslovakian) sat. (2017–2023)



Timepix - Ionizing radiation dosimetry and imaging

Lecture 2:
Classical
methods
and
models in
MLE with
Scikit
Learn

Tomáš
Báča

Python
Environ-
ment

Pandas

Scikit
Learn

Toy dataset
Regression
pipeline
Preprocessing

Metrics

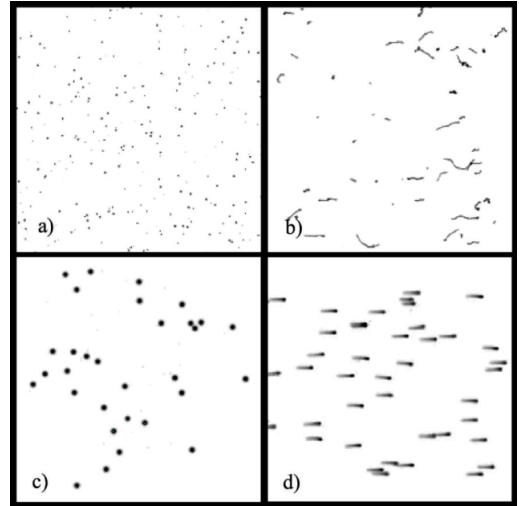
Realworld
example

References

- real-time recording of the interaction of an incoming particle and the sensor
- machine learning was applied to distinguish particle types [3]
- no dark-current noise (the images are spotless besides the actual data)

Examples in the Figure

- photons (gamma)
- electrons (beta)
- Helium nuclei (alpha)
- protons



Tomáš Báča (CTU in Prague)

Lecture 2: Classical methods and models in MLE with Scikit Learn

September 30th, 2025

42 / 48

Lecture 2: Classical methods and models in MLE with Scikit Learn

Realworld example

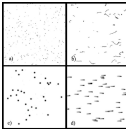
Timepix - Ionizing radiation dosimetry and imaging

Timepix - Ionizing radiation dosimetry and imaging

- real-time recording of the interaction of an incoming particle and the sensor
- machine learning was applied to distinguish particle types [3]
- no dark-current noise (the images are spotless besides the actual data)

Examples in the Figure

- photons (gamma)
- electrons (beta)
- Helium nuclei (alpha)
- protons



- The sensor is a digital counterpart of the *Cloud chamber* (https://en.wikipedia.org/wiki/Cloud_chamber)

2025-10-11

Random forests for particle track classification

Lecture 2:
Classical methods
and models in
MLE with
Scikit
Learn

Tomáš
Báča

Python
Environ-
ment

Pandas

Scikit
Learn

Toy dataset

Regression
pipeline

Preprocessing

Metrics

Realworld
example

References

Particle classification in a single image [4]

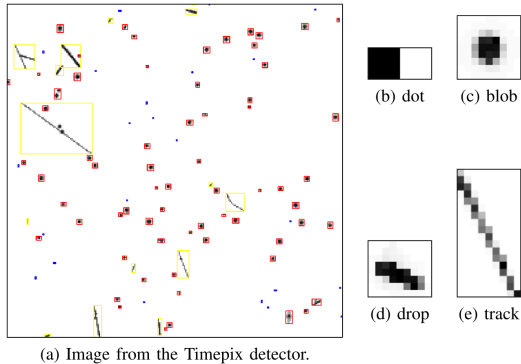
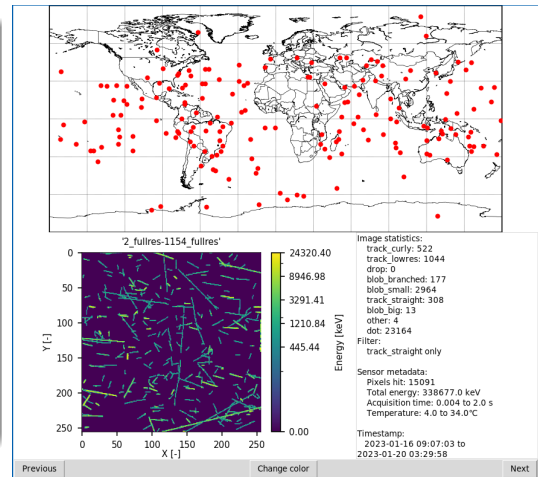


Image browser



Tomáš Báča (CTU in Prague)

Lecture 2: Classical methods and models in MLE with Scikit Learn

September 30th, 2025

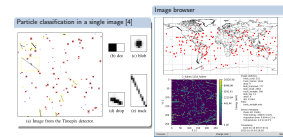
43 / 48

Lecture 2: Classical methods and models in MLE with Scikit Learn

Realworld example

Random forests for particle track classification

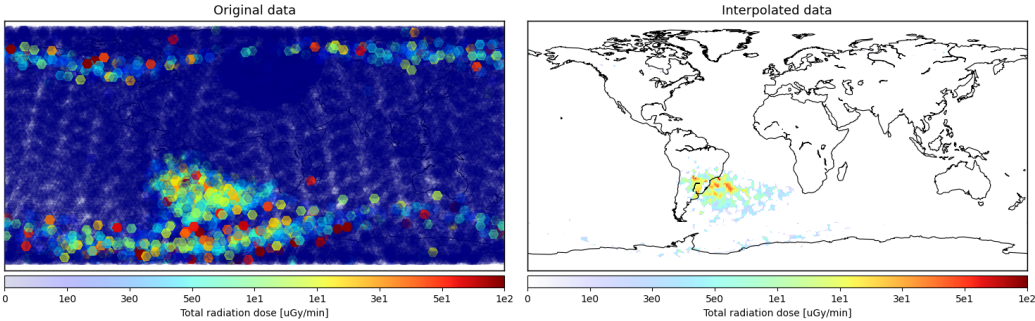
- We have built a classified dataset of the data from the 6-years of deployment in the Low-Earth Orbit:
<https://github.com/vzlsat/vzlsat1-timepix-data>



2025-10-11

Plot for Beta particles

VZLUSAT-1 Timepix radiation dose, 2017-07-13 10:17:57 to 2023-05-15 09:47:37

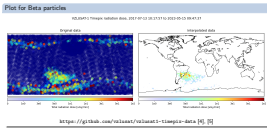


<https://github.com/vzlusat/vzlusat1-timepix-data> [4], [5]

Lecture 2: Classical methods and models in MLE with Scikit Learn

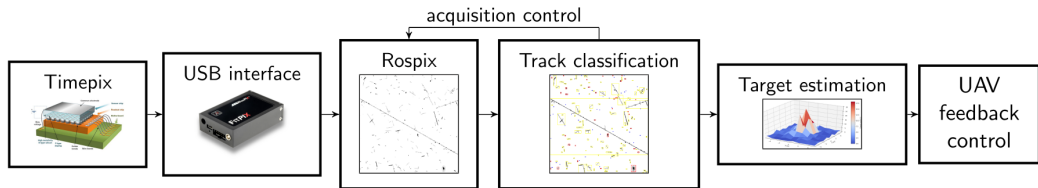
Realworld example

Low-Earth orbit particle classification



Proof of concept for UAVs

- Real-time particle track classifier for mobile robots and drones.
- We have finished well-evaluated TACR grant to develop this tech.



- [4] T. Baca, M. Jilek, P. Manek, P. Stibinger, V. Linhart, J. Jakubek, *et al.*, "Timepix Radiation Detector for Autonomous Radiation Localization and Mapping by Micro Unmanned Vehicles," in *2019 IEEE/RSJ International Conference on Intelligent Robots and Systems*, IEEE, 2019, pp. 1–8

Lecture 2: Classical methods and models in MLE with Scikit Learn

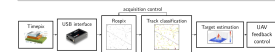
Realworld example

Realtime classification onboard UAVs

Realtime classification onboard UAVs

Proof of concept for UAVs

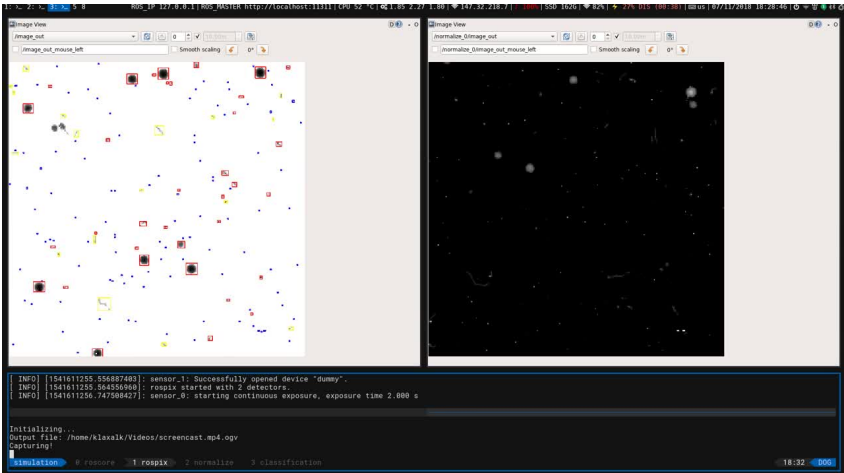
- Real-time particle track classifier for mobile robots and drones.
- We have finished well-evaluated TACR grant to develop this tech.



[4] T. Baca, M. Jilek, P. Manek, P. Stibinger, V. Linhart, J. Jakubek, *et al.*, "Timepix Radiation Detector for Autonomous Radiation Localization and Mapping by Micro Unmanned Vehicles," in *2019 IEEE/RSJ International Conference on Intelligent Robots and Systems*, IEEE, 2019, pp. 1–8.

- The particle classifier is a necessary tool for real-time control of the acquisition parameters of the detector (acquisition time, bias, etc.).
- The particle spectrum measurement is also needed for educated decision about what source of radiation we are observing, this leads to better localization with some advanced techniques, such as the Compton camera localization [baca2021gamm](#).

Realtime classification of ionizing particles in ROS

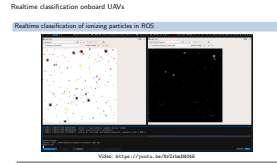


Video: <https://youtu.be/8rZrbmDHG4E>

Lecture 2: Classical methods and models in MLE with Scikit Learn

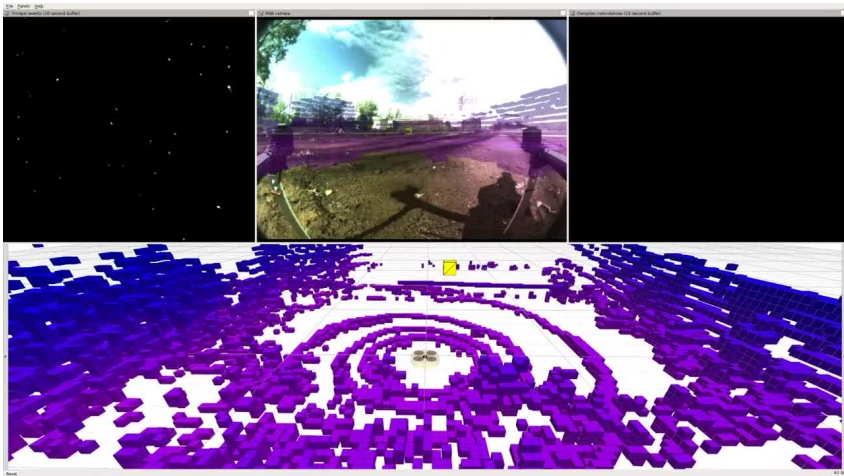
Realworld example

Realtime classification onboard UAVs



- The classification needs to run for hundreds of particles per second on a limited hardware of a drone.

Endgame: realtime search for ionizing radiation sources



Video: <https://youtu.be/oH4jMMHfGVA>

Tomas Bacea (CTU in Prague)

Lecture 2: Classical methods and models in MLE with Scikit Learn

September 30th, 2025

46 / 48

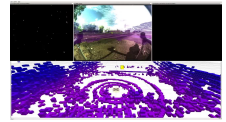
Lecture 2: Classical methods and models in MLE with Scikit Learn

Realworld example

Realtime classification onboard UAVs

Realtime classification onboard UAVs

Endgame: realtime search for ionizing radiation sources



- The classification needs to run for hundreds of particles per second on a limited hardware of a drone.

Thanks for your attention

Lecture 2: Classical methods and models in MLE with Scikit Learn

Conclusion

Realworld example

Thanks for your attention

Conclusion

[1] M. Urban, O. Nentvich, V. Stehlikova, T. Baca, V. Daniel, and R. Hudec, "VZLUSAT-1: Nanosatellite with miniature lobster eye X-ray telescope and qualification of the radiation shielding composite for space application," *Acta Astronautica*, vol. 140, pp. 96–104, 2017.

[2] T. Baca, M. Platkevici, J. Jakubek, et al., "Miniaturized X-ray telescope for VZLUSAT-1 nanosatellite with Timepix detector," *Journal of Instrumentation*, vol. 11, no. 10, p. C10007, 2016.

[3] T. Baca, M. Jilek, I. Vertat, et al., "Timepix in LEO Orbit onboard the VZLUSAT-1 Nanosatellite: 1-year of Space Radiation Dosimetry Measurements," *Journal of Instrumentation*, vol. 13, no. 11, p. C11010, 2018.

[4] T. Baca, M. Jilek, P. Manek, et al., "Timepix Radiation Detector for Autonomous Radiation Localization and Mapping by Micro Unmanned Vehicles," in *2019 IEEE/RSJ International Conference on Intelligent Robots and Systems*, IEEE, 2019, pp. 1–8.

[5] T. Baca, P. Stibinger, D. Doubravova, et al., "Gamma Radiation Source Localization for Micro Aerial Vehicles with a Miniature Single-Detector Compton Event Camera," in *2021 International Conference on Unmanned Aircraft Systems (ICUAS)*, IEEE, 2021, pp. 1–9.

Lecture 2: Classical methods and models in MLE with Scikit Learn

References

References

References I

[1] M. Urban, O. Nentvich, V. Stehlikova, T. Baca, V. Daniel, and R. Hudec, "VZLUSAT-1: Nanosatellite with miniature lobster eye X-ray telescope and qualification of the radiation shielding composite for space application," *Acta Astronautica*, vol. 140, pp. 96–104, 2017.

[2] T. Baca, M. Platkevici, J. Jakubek, et al., "Miniaturized X-ray telescope for VZLUSAT-1 nanosatellite with Timepix detector," *Journal of Instrumentation*, vol. 11, no. 10, p. C10007, 2016.

[3] T. Baca, M. Jilek, I. Vertat, et al., "Timepix in LEO Orbit onboard the VZLUSAT-1 Nanosatellite: 1-year of Space Radiation Dosimetry Measurements," *Journal of Instrumentation*, vol. 13, no. 11, p. C11010, 2018.

[4] T. Baca, M. Jilek, P. Manek, et al., "Timepix Radiation Detector for Autonomous Radiation Localization and Mapping by Micro Unmanned Vehicles," in *2019 IEEE/RSJ International Conference on Intelligent Robots and Systems*, IEEE, 2019, pp. 1–8.

[5] T. Baca, P. Stibinger, D. Doubravova, et al., "Gamma Radiation Source Localization for Micro Aerial Vehicles with a Miniature Single-Detector Compton Event Camera," in *2021 International Conference on Unmanned Aircraft Systems (ICUAS)*, IEEE, 2021, pp. 1–9.