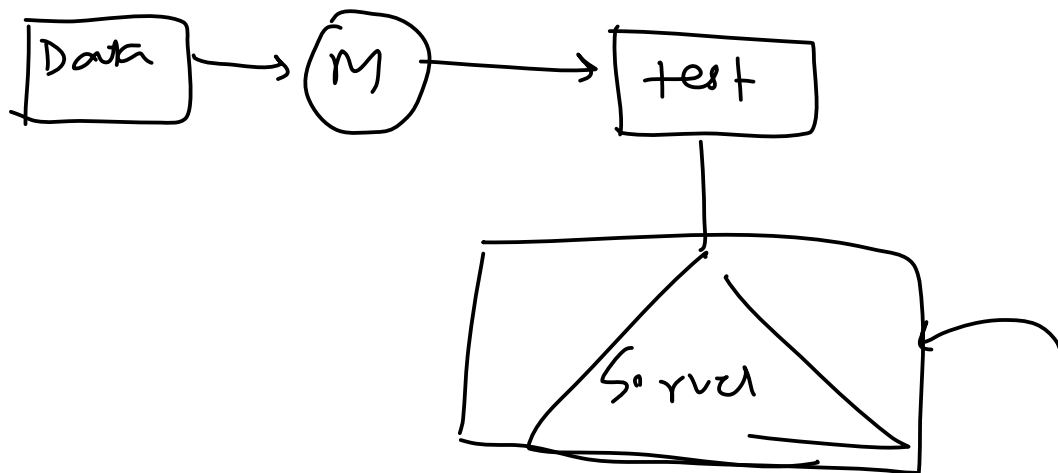


1. Batch Vs Online ML

Wednesday, March 17, 2021 5:30 PM

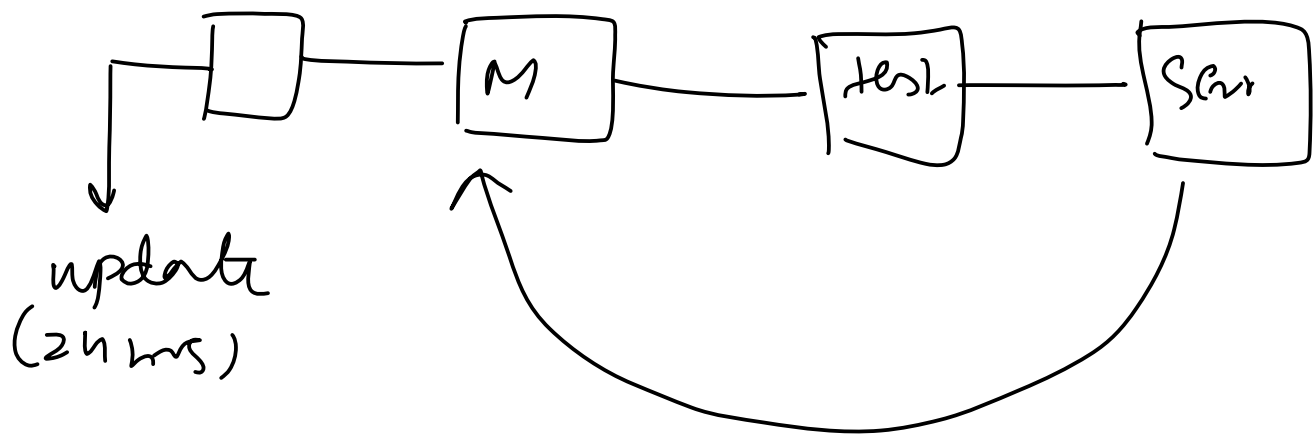
2. Batch/Offline ML

Wednesday, March 17, 2021 5:31 PM



3. The problem with Batch Learning

Wednesday, March 17, 2021 5:47 PM



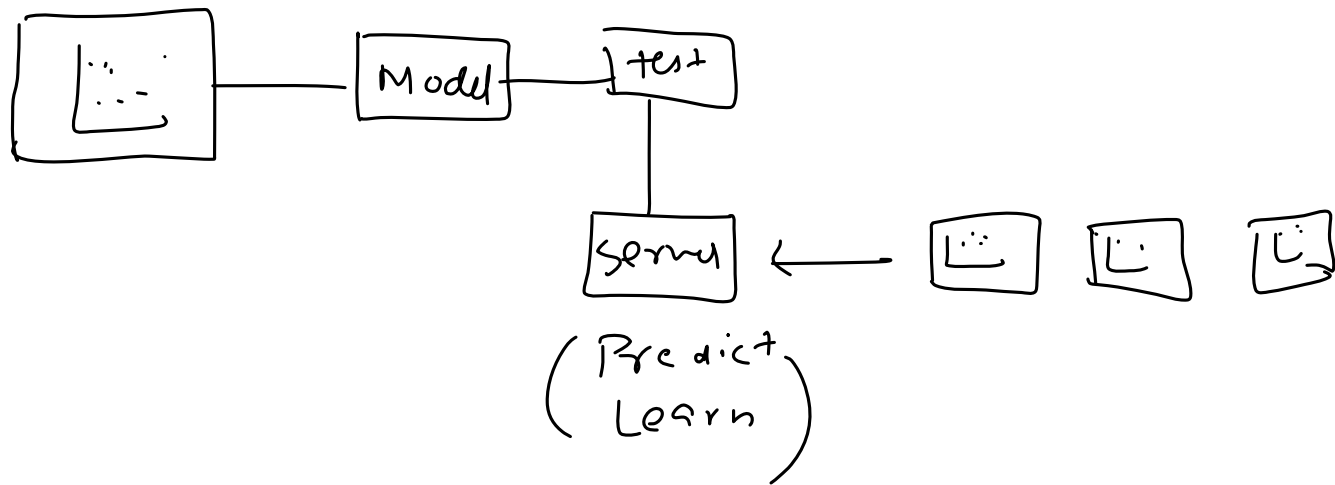
4. Disadvantages of Batch ML

Wednesday, March 17, 2021 5:32 PM

1. Lots of Data
2. Hardware Limitation
3. Availability

1. Online Machine Learning

Thursday, March 18, 2021 4:27 PM



2. When to use?

Thursday, March 18, 2021 4:33 PM

1. Where there is a concept drift
2. Cost Effective
3. Faster solution

3. How to implement?

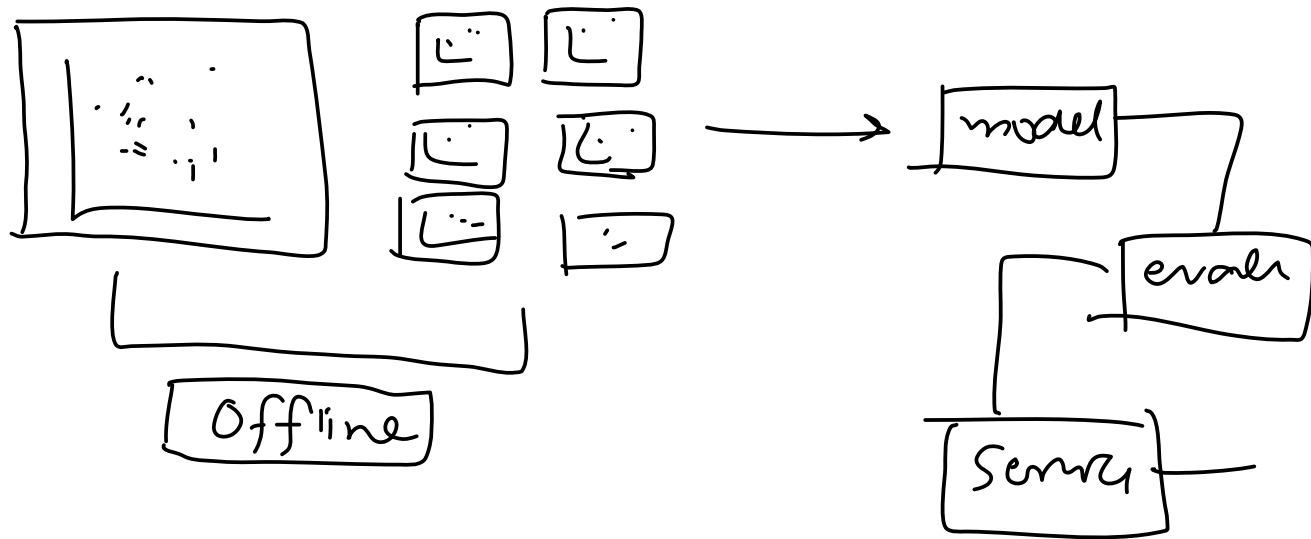
Thursday, March 18, 2021 4:28 PM

4. Learning Rate

Thursday, March 18, 2021 4:28 PM

5. Out of Core Learning

Thursday, March 18, 2021 4:28 PM



6. Disadvantage

Thursday, March 18, 2021 4:29 PM

1. Tricky to use
2. Risky

7. Batch Vs Online Learning

Thursday, March 18, 2021 4:29 PM

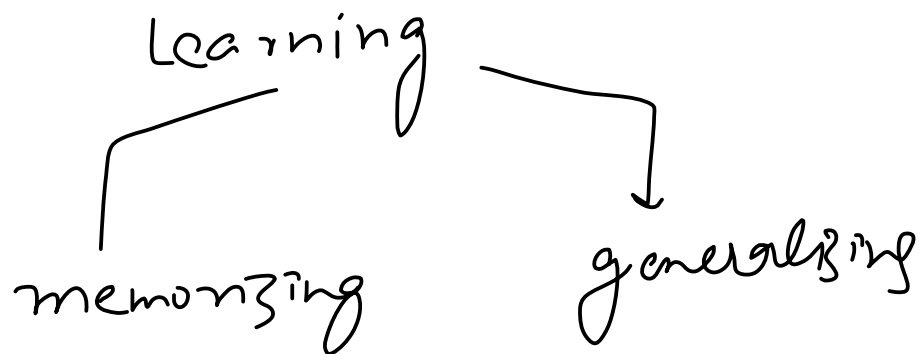
Offline Learning	Features	Online Learning
Less complex as model is constant	Complexity	Dynamic complexity as the model keeps evolving over time
Fewer computations, single time batch-based training	Computational Power	Continuous data ingestions result in consequent model refinement computations
Easier to implement	Use in Production	Difficult to implement and manage
Image Classification or anything related to Machine Learning - where data patterns remains constant without sudden concept drifts	Applications	Used in finance, economics, health where new data patterns are constantly emerging
Industry proven tools. E.g. Sci-kit, TensorFlow, Pytorch, Keras, Spark Mlib	Tools	Active research/New project tools: E.g. MOA, SAMOA, scikit-multiflow, streamDM



Image courtesy - <https://www.iunera.com/kraken/fabric/simple-introduction-to-online-learning-in-machine-learning/>

1. Instance Vs Model Based Learning

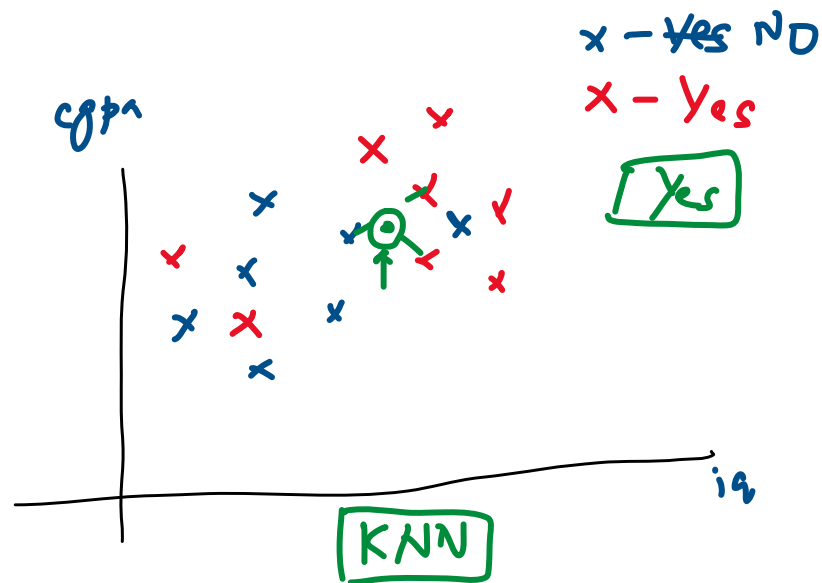
Friday, March 19, 2021 4:05 PM



2. Instance Based

Friday, March 19, 2021 4:06 PM

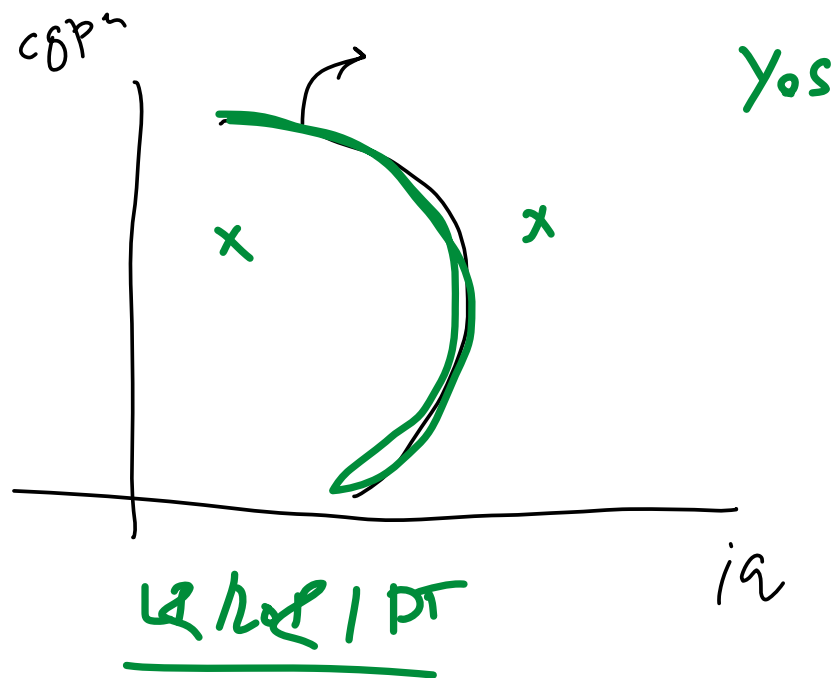
iq	cgpa	placement x/n
80	8	y
70	7	n
7.5, 103		



3. Model Based

Friday, March 19, 2021 4:06 PM

iq | cgpa | place



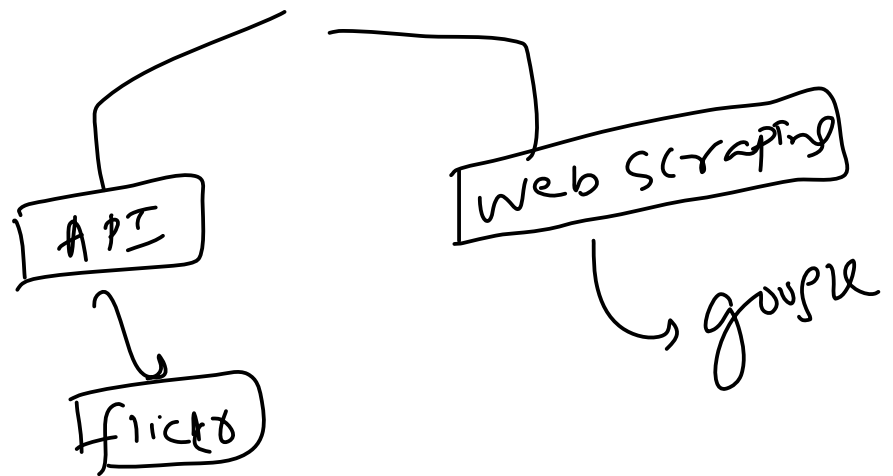
4. Differences

Friday, March 19, 2021 4:06 PM

Usual/Conventional Machine Learning	Instance Based Learning
Prepare the data for model training	Prepare the data for model training. No difference here
Train model from training data to estimate model parameters i.e. discover patterns	Do not train model. Pattern discovery postponed until scoring query received
Store the model in suitable form	There is no model to store
Generalize the rules in form of model, even before scoring instance is seen	No generalization before scoring. Only generalize for each scoring instance individually as and when seen
Predict for unseen scoring instance using model	Predict for unseen scoring instance using training data directly
Can throw away input/training data after model training	Input/training data must be kept since each query uses part or full set of training observations
Requires a known model form	May not have explicit model form
Storing models generally requires less storage	Storing training data generally requires more storage

1. Data Collection

Saturday, March 20, 2021 5:59 PM



2. Insufficient Data/Labelled Data

Saturday, March 20, 2021 6:00 PM

✓
A
(100)
M1

B
(10⁶)
M2 ✓

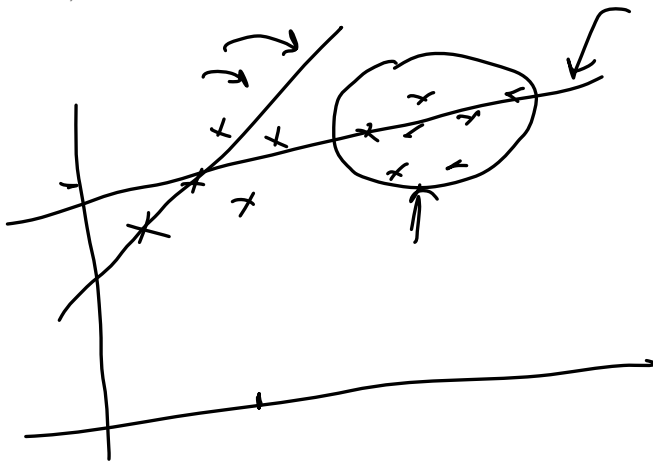
t₀, t₁₀,
↙ ↘

NLP



3. Non Representative Data

Saturday, March 20, 2021 6:00 PM



Sampling Noise

Sampling bias

4. Poor Quality Data

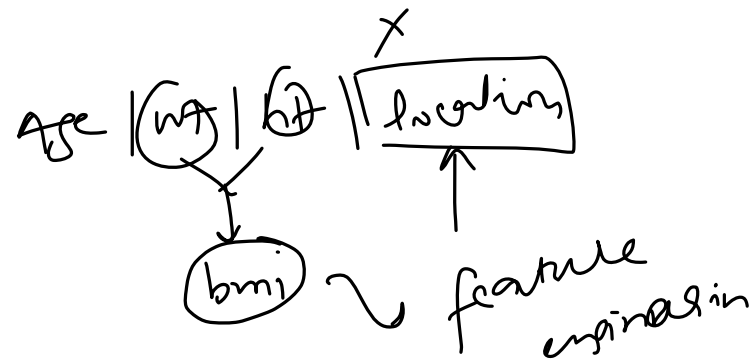
Saturday, March 20, 2021 6:00 PM

60%

5. Irrelevant Features

Saturday, March 20, 2021 6:00 PM

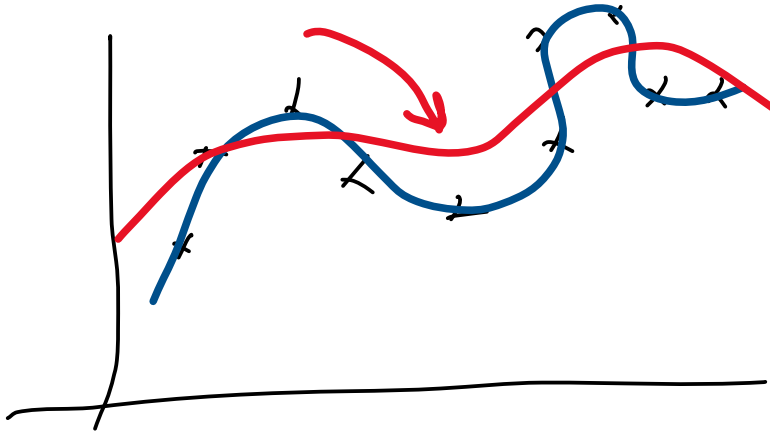
Garbage In
Garbage out



6. Overfitting

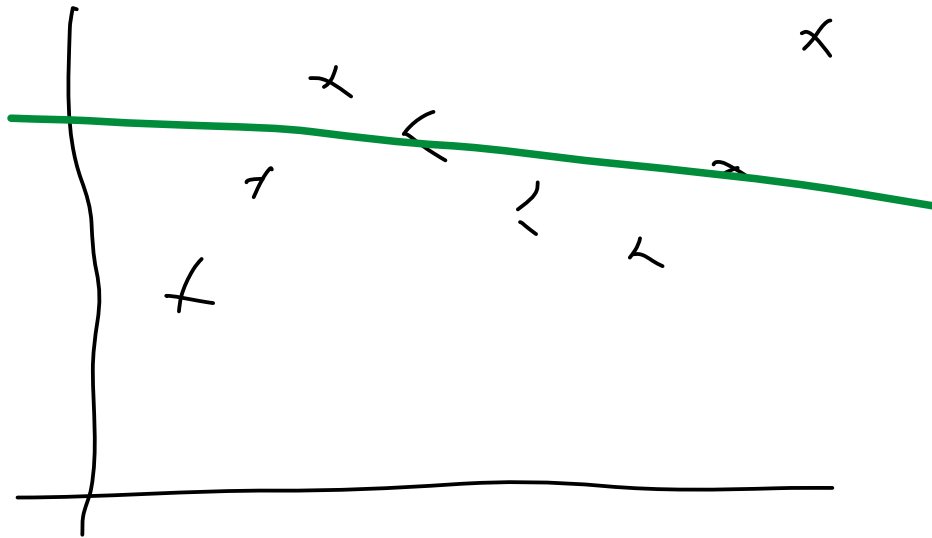
Saturday, March 20, 2021

6:01 PM



7. Underfitting

Saturday, March 20, 2021 6:01 PM

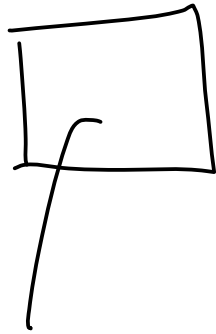


8. Software Integration

Saturday, March 20, 2021 6:01 PM

9. Offline Learning/ Deployment

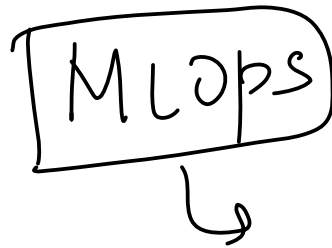
Saturday, March 20, 2021 6:01 PM



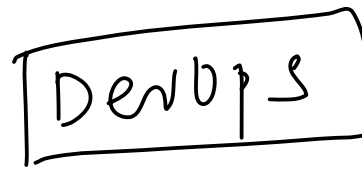
10. Cost Involved

Saturday, March 20, 2021 6:01 PM

MLops



DevOps



1. Retail - Amazon/Big Bazaar

Monday, March 22, 2021 6:07 PM



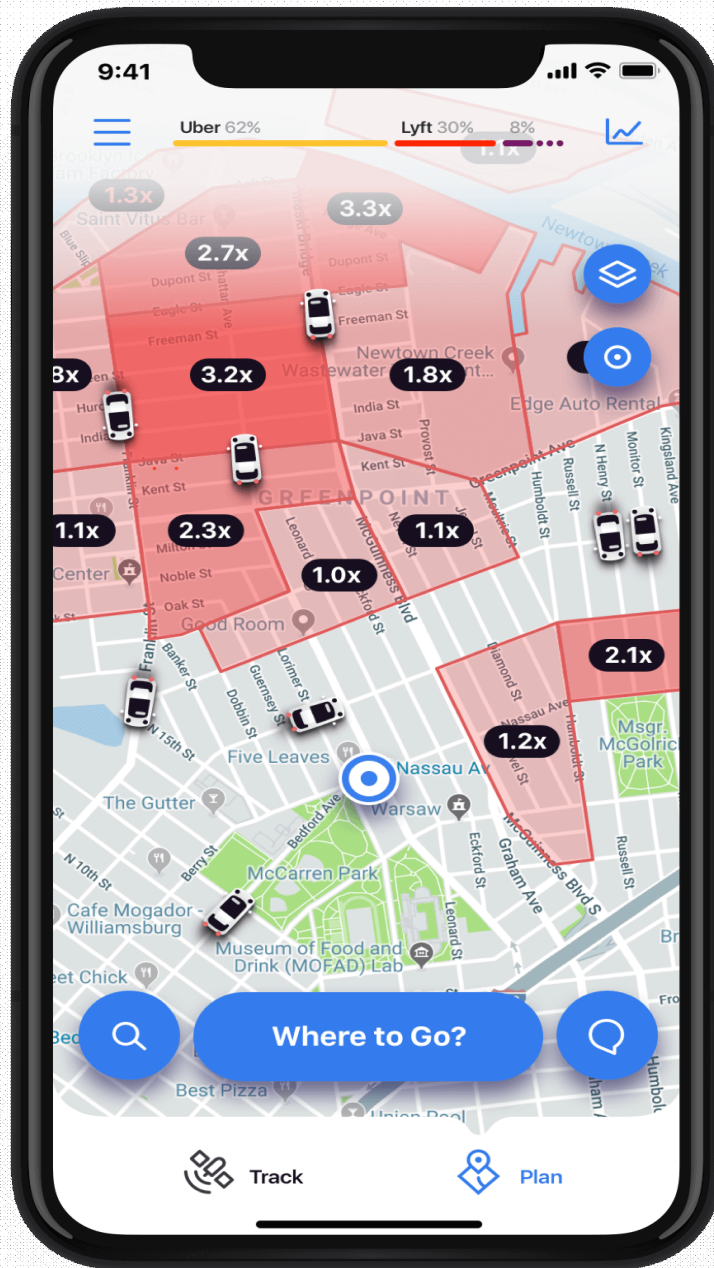
2. Banking and Finance

Monday, March 22, 2021 6:07 PM



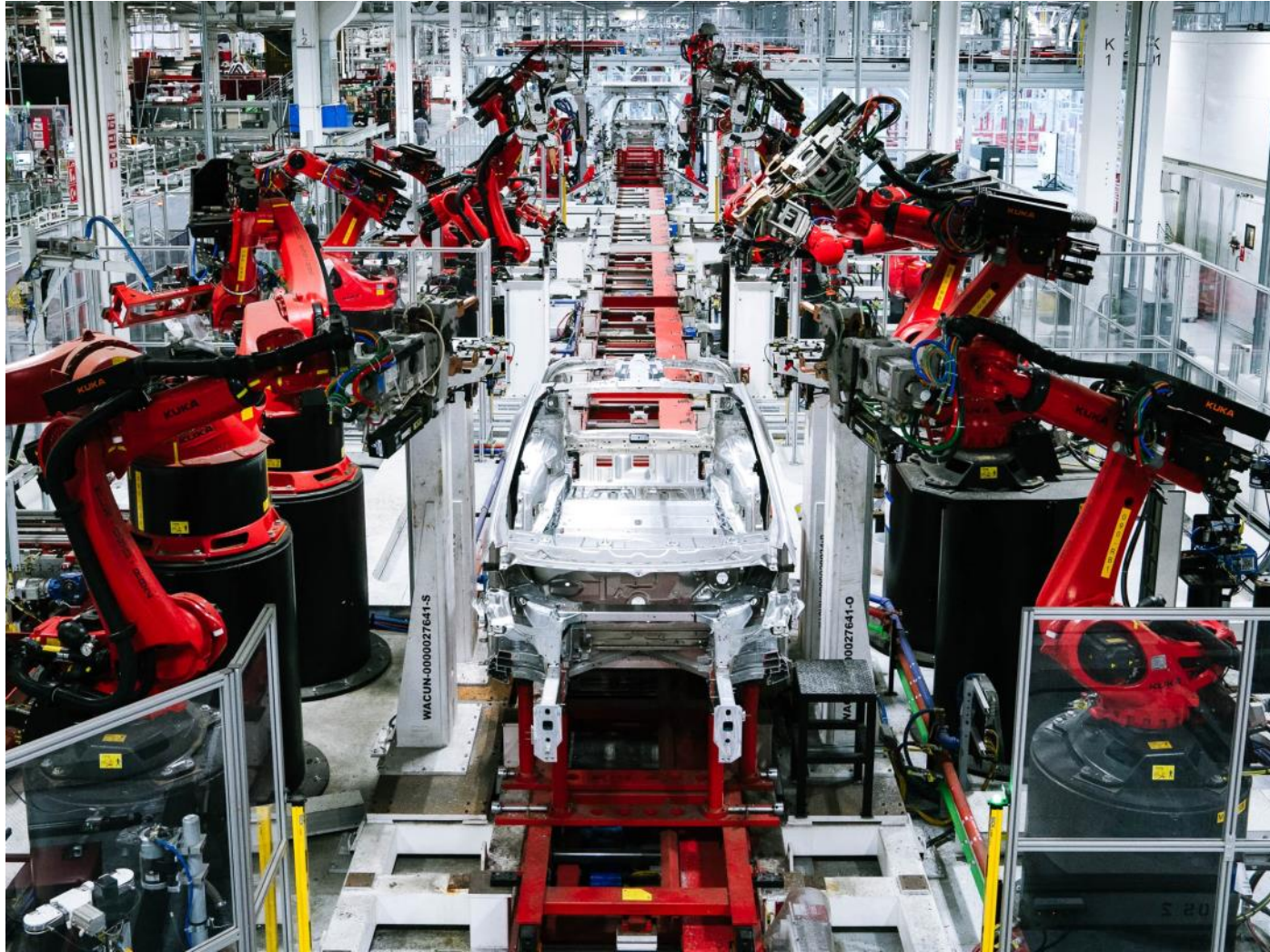
3. Transport - OLA

Monday, March 22, 2021 6:07 PM



4. Manufacturing - Tesla

Monday, March 22, 2021 6:08 PM



5. Consumer Internet - Twitter

Monday, March 22, 2021 6:08 PM



Machine Learning Development Life Cycle(MLDLC/MLDC)

Tuesday, March 23, 2021 12:09 PM

SDLC

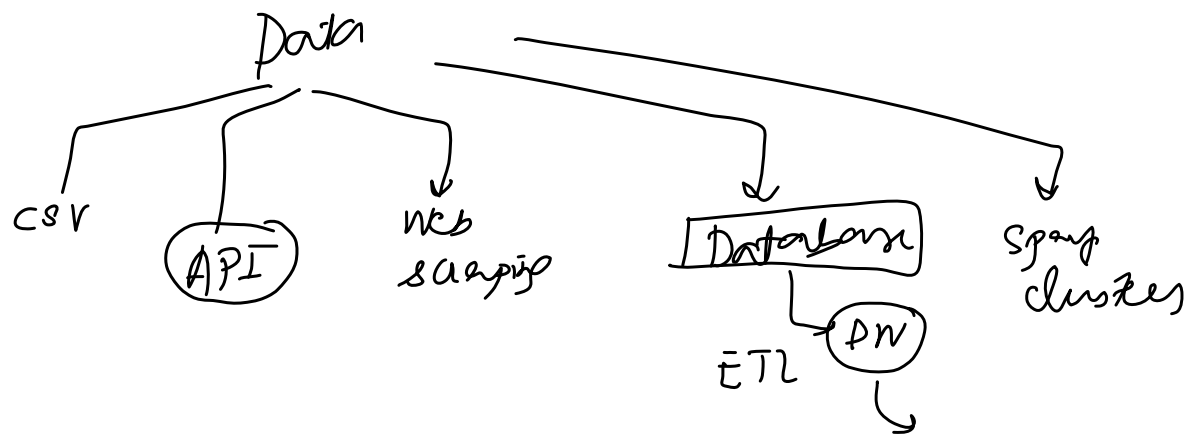
ML DLC

1. Frame the Problem

Tuesday, March 23, 2021 12:10 PM

2. Gathering Data

Tuesday, March 23, 2021 12:11 PM



3. Data Preprocessing

Tuesday, March 23, 2021 12:11 PM

- Remove duplicates
- Remove missing val
- Outliers
- Scale

4. Exploratory Data Analysis

Tuesday, March 23, 2021 12:11 PM

Vizs

Univariate/Bivariate

Outlier detection

Imbalance →

5. Feature Engineering and Selection

Tuesday, March 23, 2021 12:12 PM

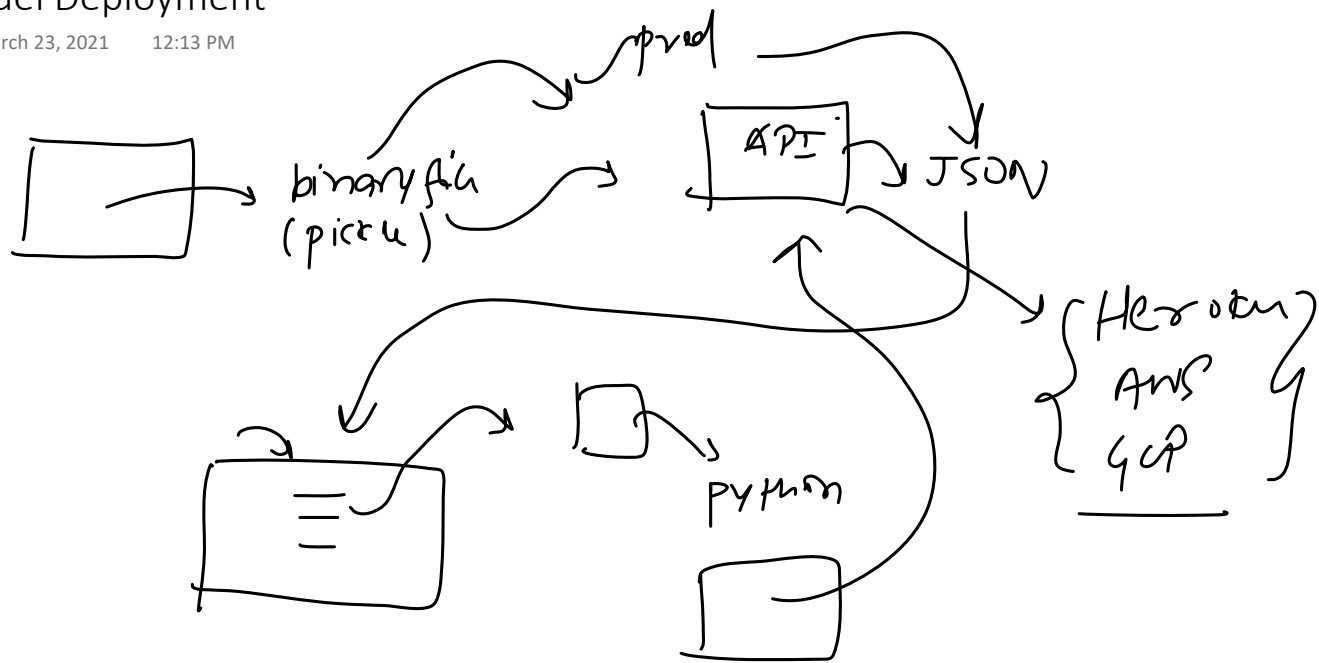
6. Model Training, Evaluation and Selection

Tuesday, March 23, 2021 12:12 PM

Ensemble
learning

7. Model Deployment

Tuesday, March 23, 2021 12:13 PM



8. Testing

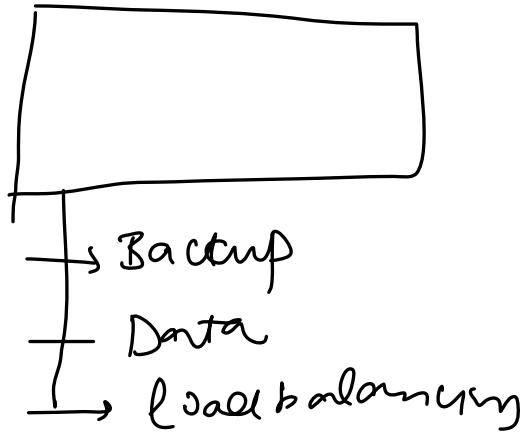
Tuesday, March 23, 2021 12:14 PM

A/B testing

9. Optimize

Tuesday, March 23, 2021

12:15 PM



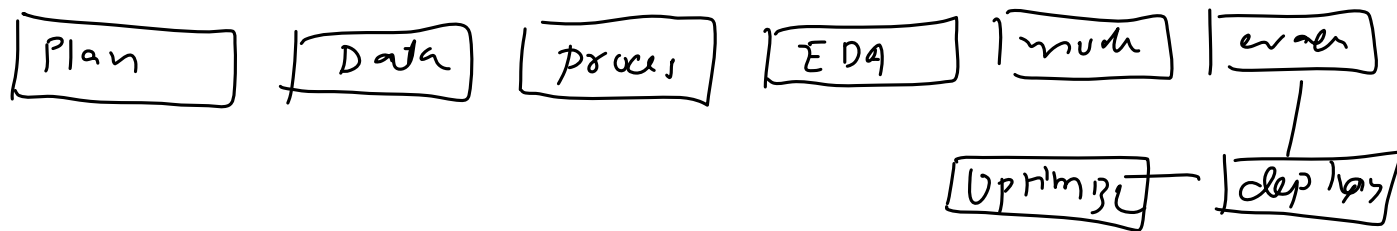
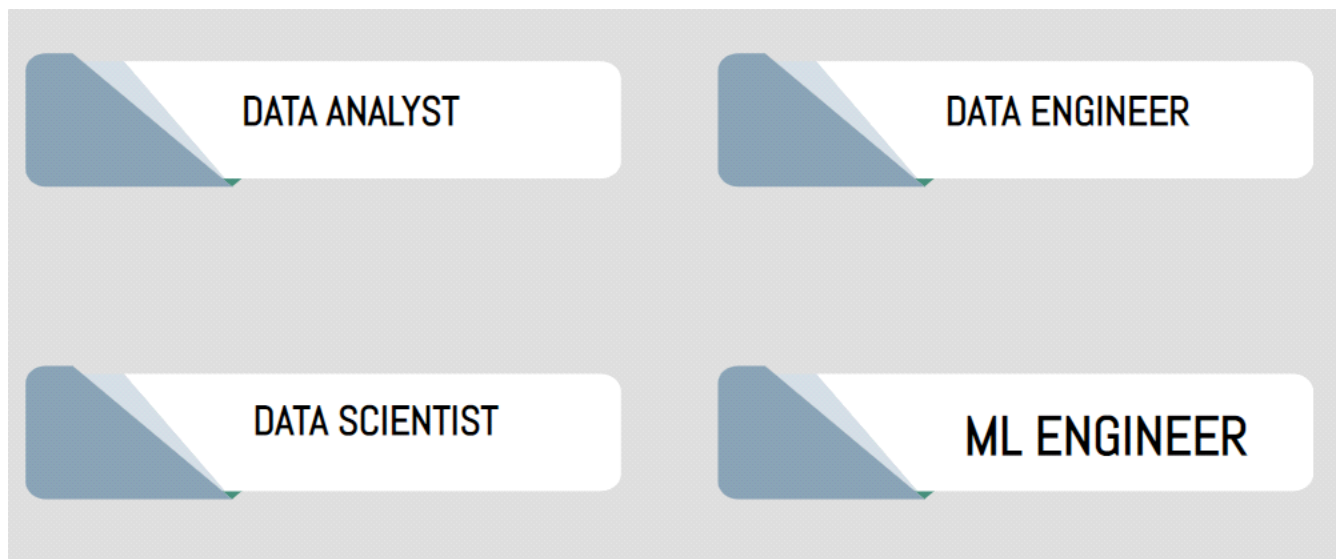
Retrain

Rolling

1. Various Data Based Job Roles

Wednesday, March 24, 2021

1:25 PM

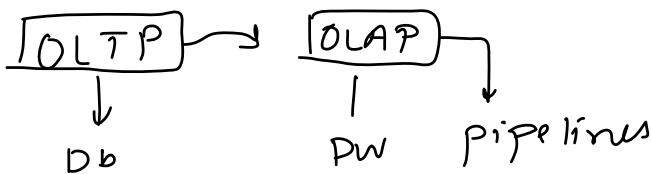


1. Data Engineer

Wednesday, March 24, 2021 1:25 PM

Job Roles

- Scrape Data from the given sources.
- Move/Store the data in optimal servers/warehouses.
- Build data pipelines/APIs for easy access to the data.
- Handle databases/data warehouses.



Skills Required

- Strong grasp of algorithms and data structures
- Programming Languages (Java/R/Python/Scala) and script writing
- Advanced DBMS's
- BIG DATA Tools (Apache Spark, Hadoop, Apache Kafka, Apache Hive)
- Cloud Platforms (Amazon Web Services, Google Cloud Platform)
- Distributed Systems
- Data Pipelines

2. Data Analyst

Wednesday, March 24, 2021 1:26 PM

Responsibilities of a Data Analyst

- *Cleaning and organizing Raw data.*
- *Analyzing data to derive insights.*
- *Creating data visualizations.*
- *Producing and maintaining reports.*
- *Collaborating with teams/colleagues based on the insight gained.*
- *Optimizing data collection procedures*

Skills

- *Statistical Programming*
- *Programming Languages (R/SAS/Python)*
- *Creative and Analytical Thinking*
- *Business Acumen — Medium to High preferred*
- *Strong Communication Skills.*
- *Data Mining, Cleaning, and Munging*
- *Data Visualization*
- *Data Story Telling*
- *SQL*
- *Advanced Microsoft Excel*

3. Data Scientist

Wednesday, March 24, 2021 1:26 PM

“A data scientist is someone who is better at statistics than any software engineer and better at software engineering than any statistician”.

4. ML Engineer

Wednesday, March 24, 2021 1:26 PM

Responsibilities

- Deploying machine learning models to production ready environment
- Scaling and optimizing the model for production
- Monitoring and maintenance of deployed models

Skills

- Mathematics
- Programming Languages (R/Python/Java/Scala mainly)
- Distributed Systems
- Data model and evaluation
- Machine Learning models
- Software Engineering & Systems design

5. Comparison

Wednesday, March 24, 2021 1:26 PM

	ANALYTICAL SKILLS	BUSINESS ACUMEN	DATA STORYTELLING	SOFT SKILLS	SOFTWARE SKILLS
DATA ANALYST	HIGH	MEDIUM TO HIGH	HIGH	MEDIUM TO HIGH	MEDIUM
DATA ENGINEER	MEDIUM	LOW	LOW	MEDIUM	HIGH
DATA SCIENTIST	HIGH	HIGH	HIGH	HIGH	MEDIUM
ML ENGINEER	MEDIUM TO HIGH	MEDIUM	LOW	HIGH	HIGH

1. What are Tensors

Thursday, March 25, 2021 4:44 PM

2. 0D Tensor/Scalar

Thursday, March 25, 2021 4:44 PM

(2) (3)
✓
Scalars → 0D Tensor 0

3. 1D Tensor/Vector

Thursday, March 25, 2021 4:45 PM

[1, 2, 3, 4] → 1D Tensor

Vector

1D array / array

ndim → 1

Axis
2 Dim

No. of axes = rank
= dim

1D Tensor / Vector

4D

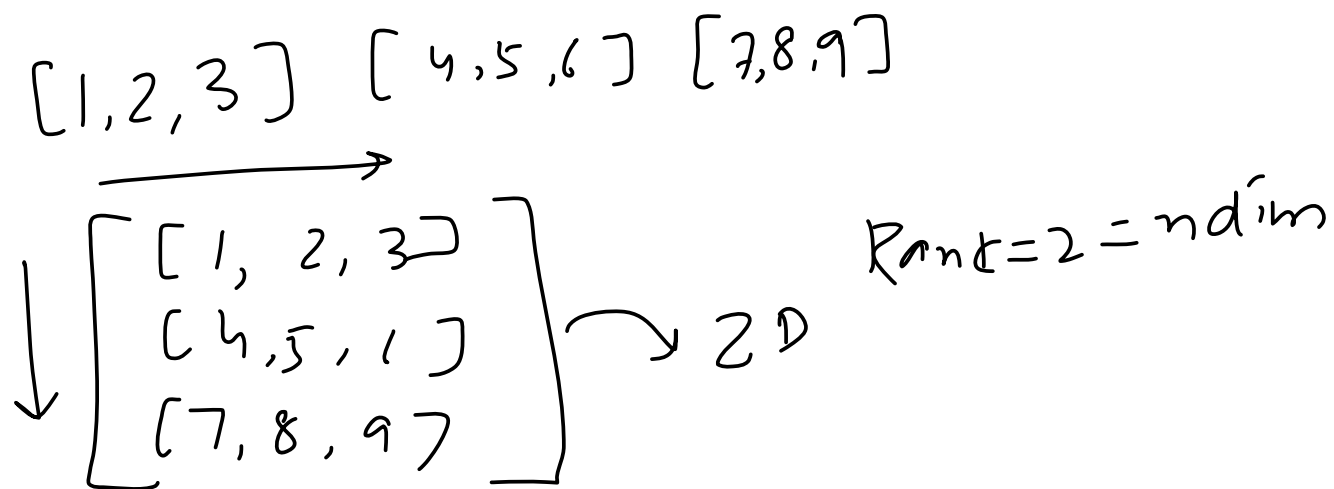
[1, 2] → vector (2)
↳ 1D Tensor

[0, 1, 2, 3]

4 scalars → vector

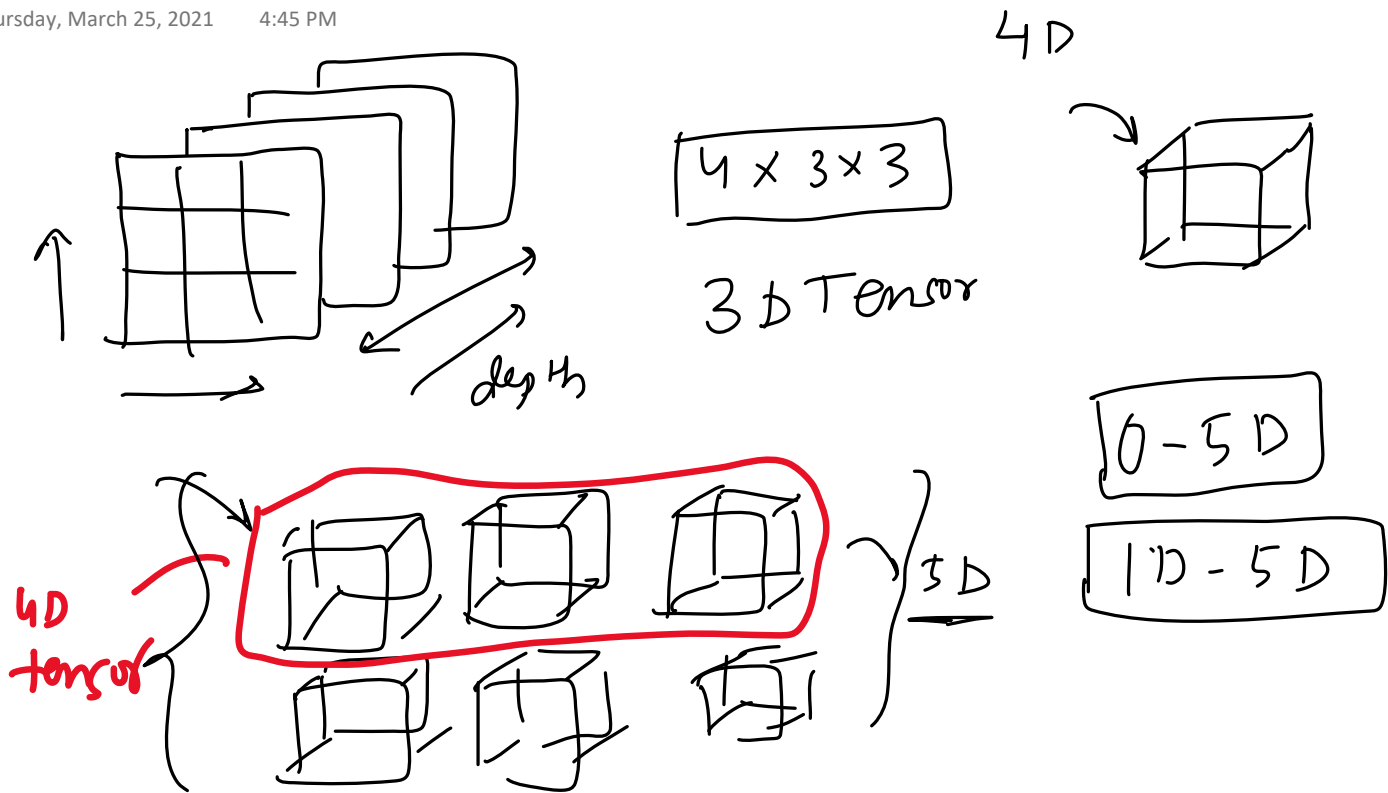
4. 2D Tensor/Matrices

Thursday, March 25, 2021 4:45 PM



5. ND Tensors

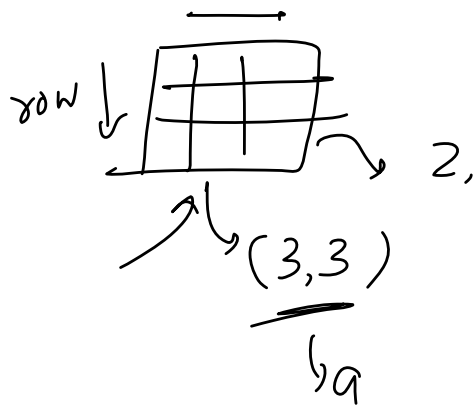
Thursday, March 25, 2021 4:45 PM



6. Rank, Axes and Shape

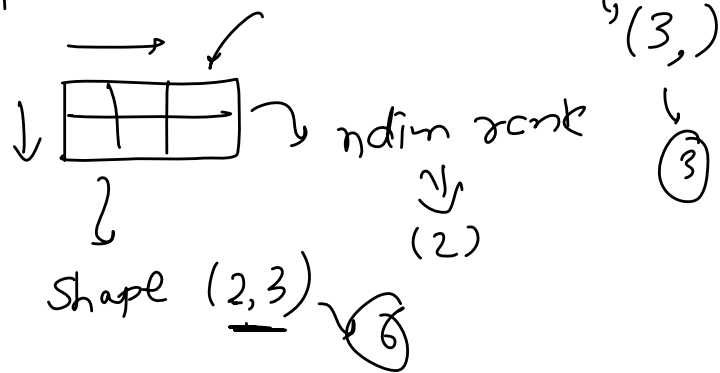
Thursday, March 25, 2021 4:45 PM

No. of axis = Rank = No. of dim

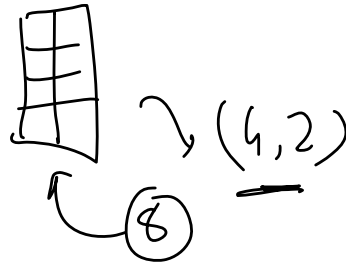


Size = 1

Shape



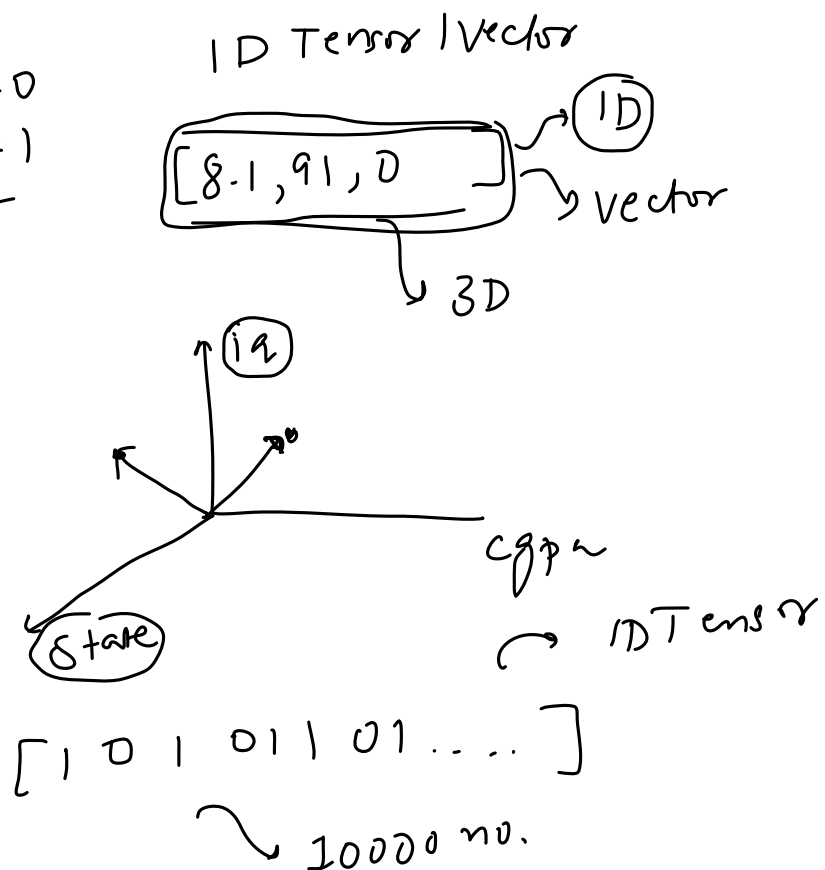
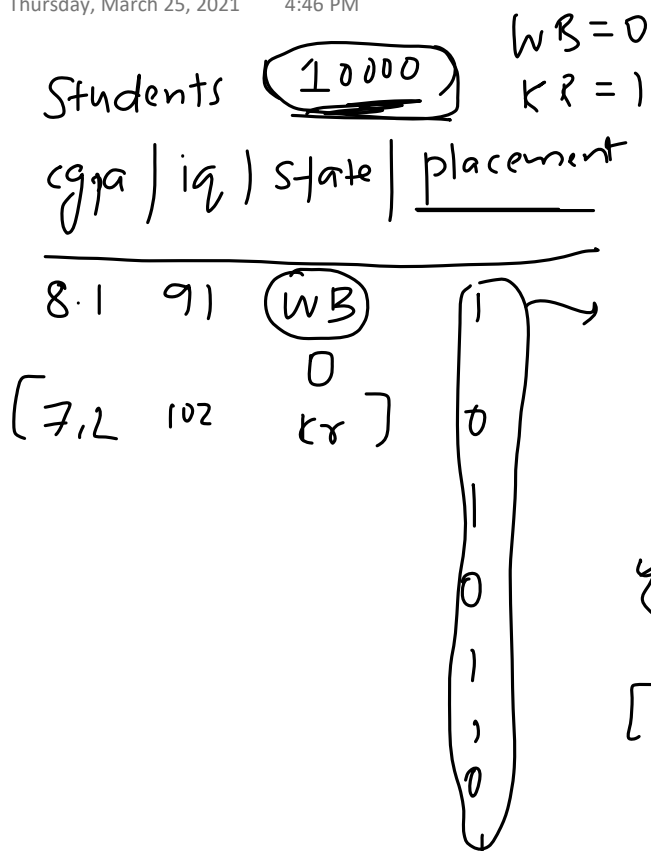
Size of tensor



7. Example of 1D Tensors

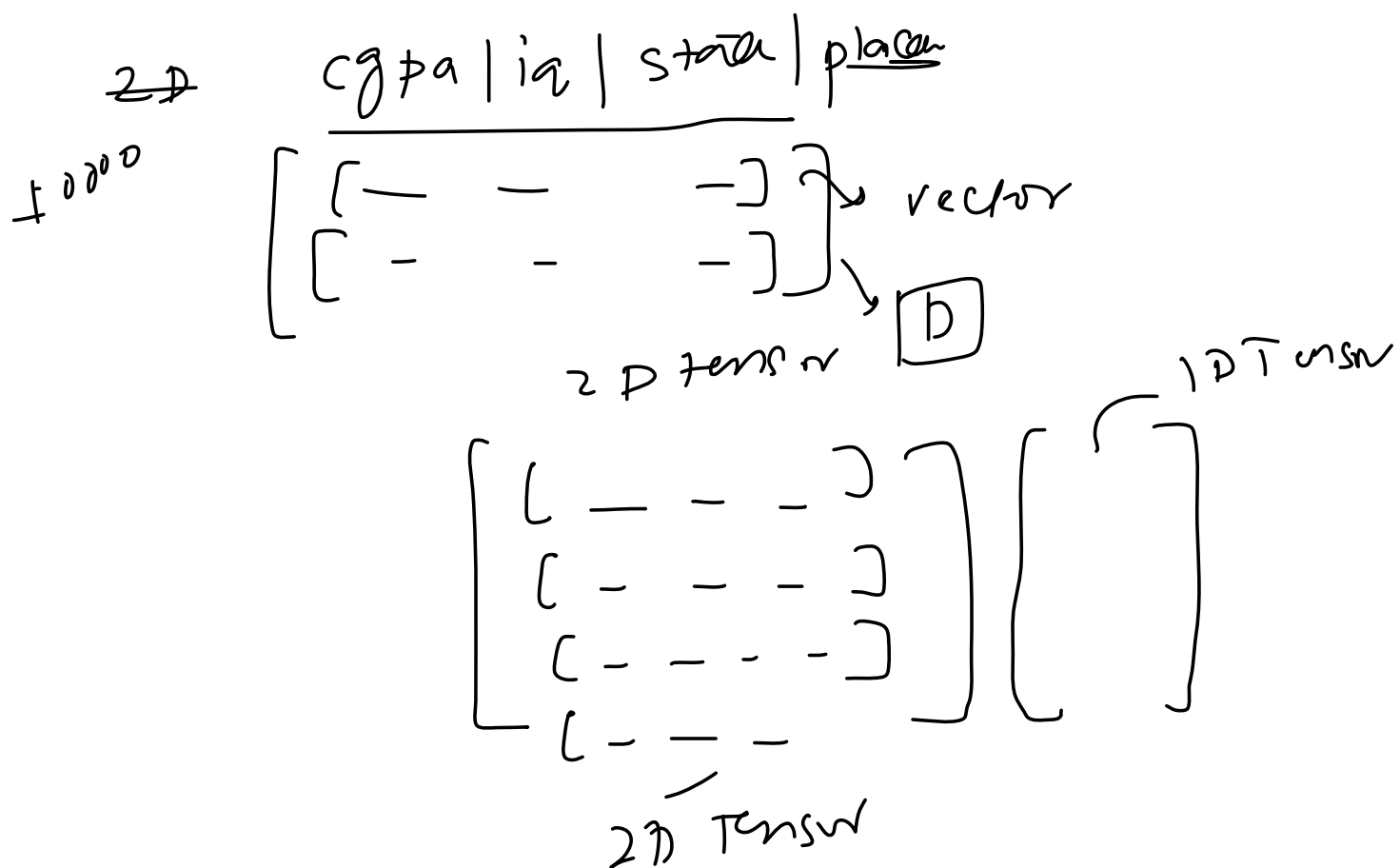
Thursday, March 25, 2021

4:46 PM



8. Example of 2D Tensors

Thursday, March 25, 2021 4:46 PM



9. Example of 3D Tensors

Thursday, March 25, 2021 4:46 PM

NLP

Hi Nitish

Hi Rahul

Hi Ankit

Hi	Nitish	Rahul	Ankit
1	0	0	0
0	1	0	0
0	0	1	0
0	0	0	1

$\left[\left[[1, 0, 0, 0], [0, 1, 0, 0] \right] \right]$ (2)
 $\left[[1, 0, 0, 0], [0, 0, 1, 0] \right]$ $\rightarrow$ 3D Tensor
 $\left[[1, 0, 0, 0], [0, 0, 0, 1] \right]$

Timeseries Data

(2)
Highest | Lowest

$\rightarrow$ 10 years
 $(365, 2)$

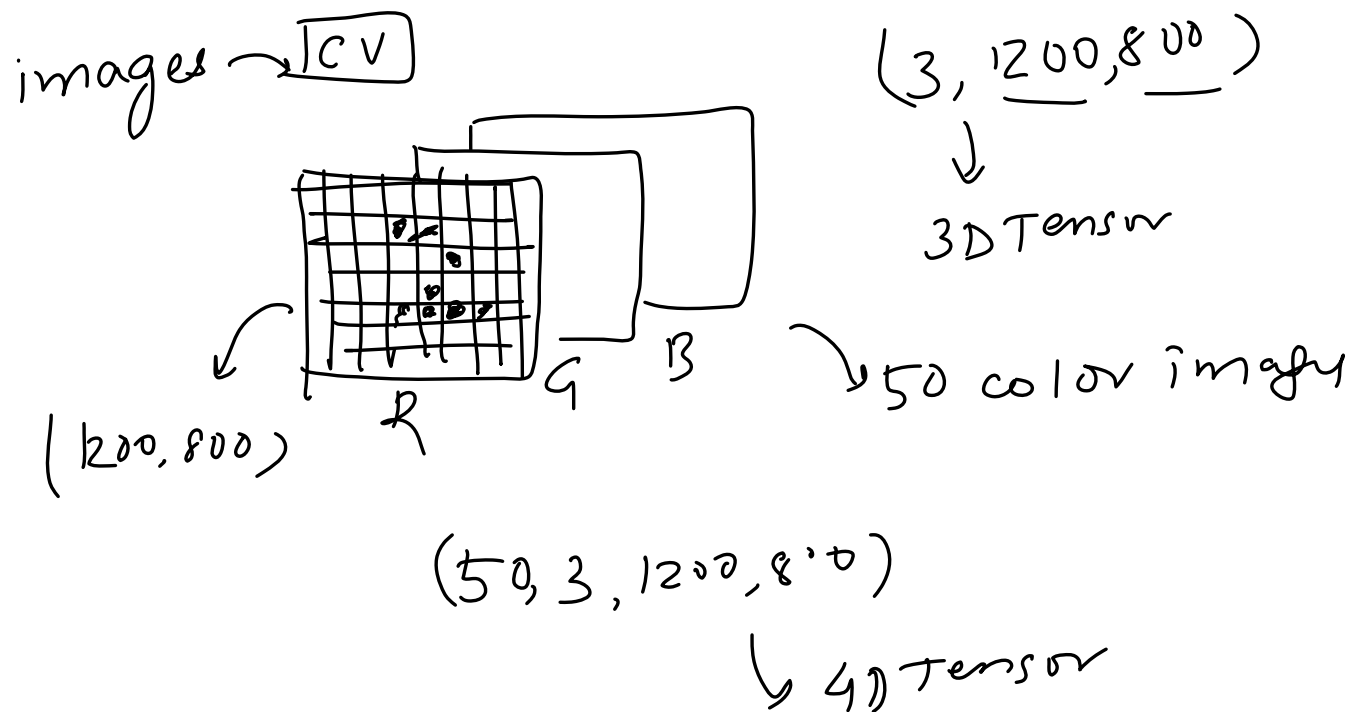
Day 1 — —
 Day 2 — — $\rightarrow$ 2D $\rightarrow$ 10
 Day 365 — —
 365

$(10, 365, 2)$ $\rightarrow$ 3D Tensor
 time axis

10. Example of 4D Tensors

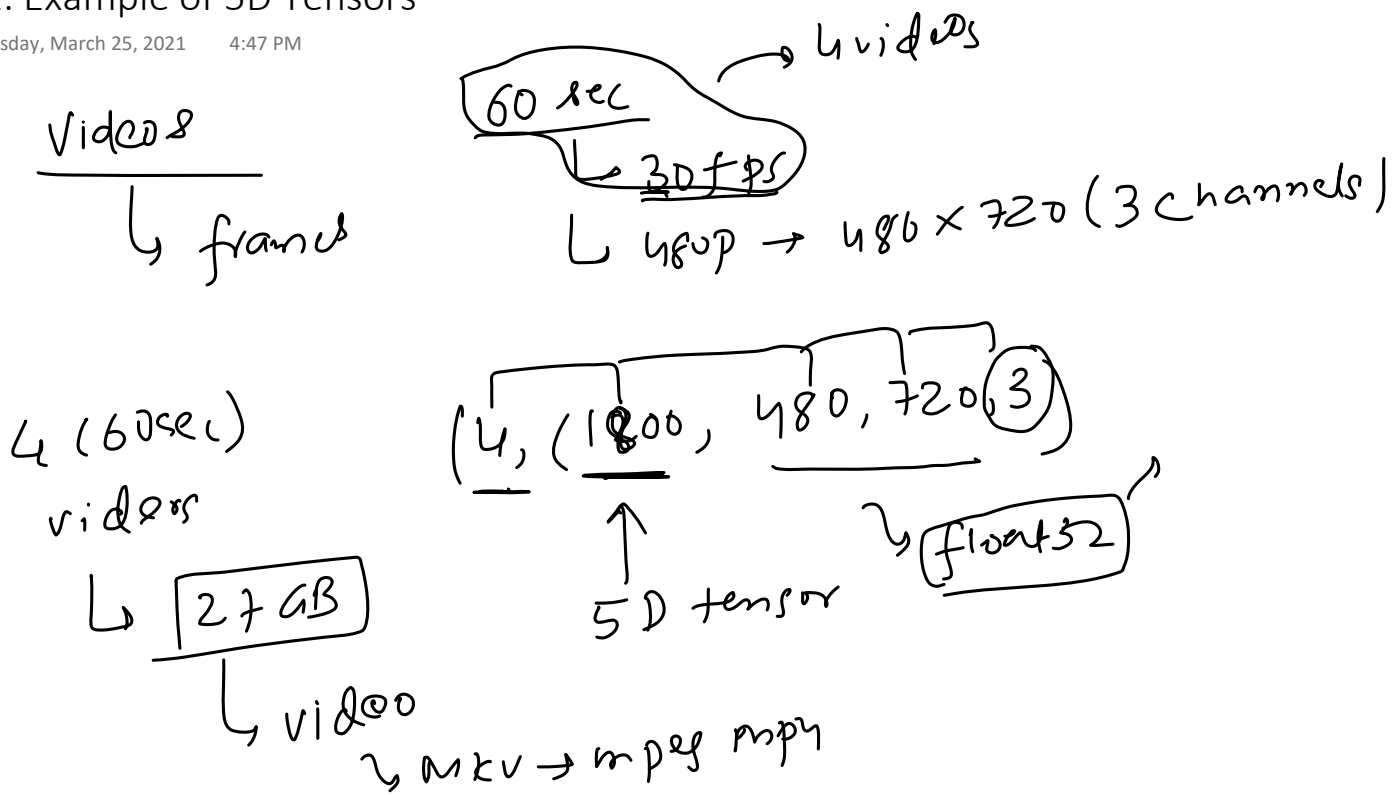
Thursday, March 25, 2021

4:47 PM



12. Example of 5D Tensors

Thursday, March 25, 2021 4:47 PM



1. Installing Anaconda

Friday, March 26, 2021 5:40 PM

2. Jupyter Notebook Intro

Friday, March 26, 2021 5:40 PM

3. Virtual Env

Friday, March 26, 2021 5:40 PM

4. Using Kaggle

Friday, March 26, 2021 5:41 PM

5. Using Google Colab

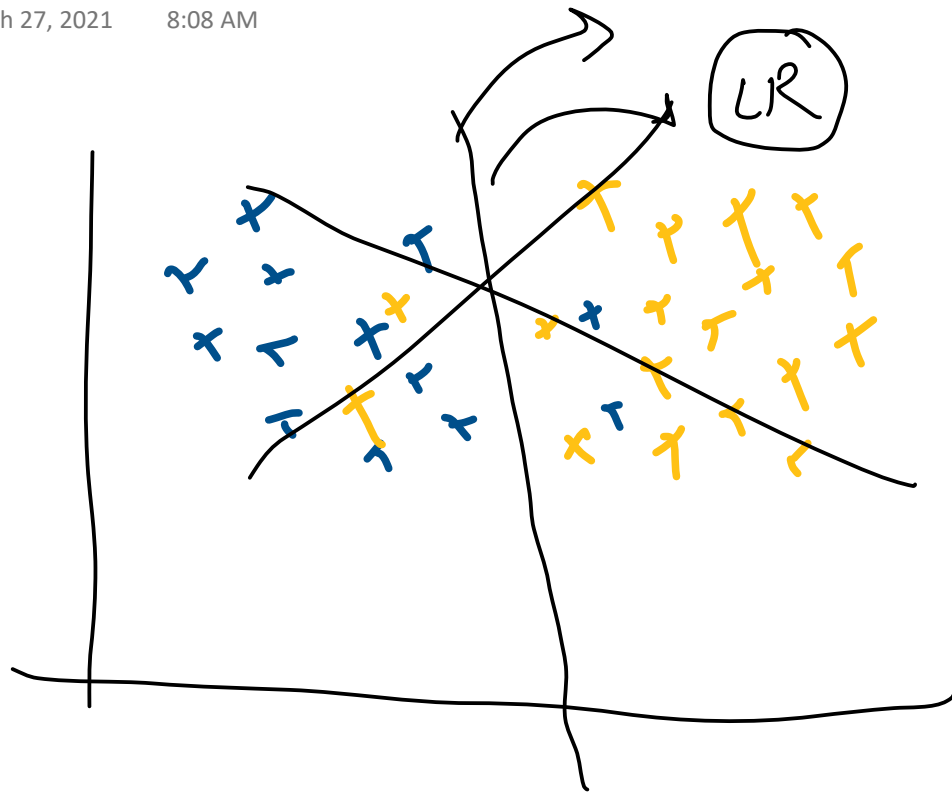
Friday, March 26, 2021 5:41 PM

6. Running Kaggle Data on Google Colab

Friday, March 26, 2021 5:41 PM

End to End Example

Saturday, March 27, 2021 8:08 AM



1. Business Problem to ML Problem

Monday, March 29, 2021 7:29 AM

Netflix

Churn rate ↓

Increase revenue

↳ 4% → 3.75%

↳ 4.1%

↳ 3.75%

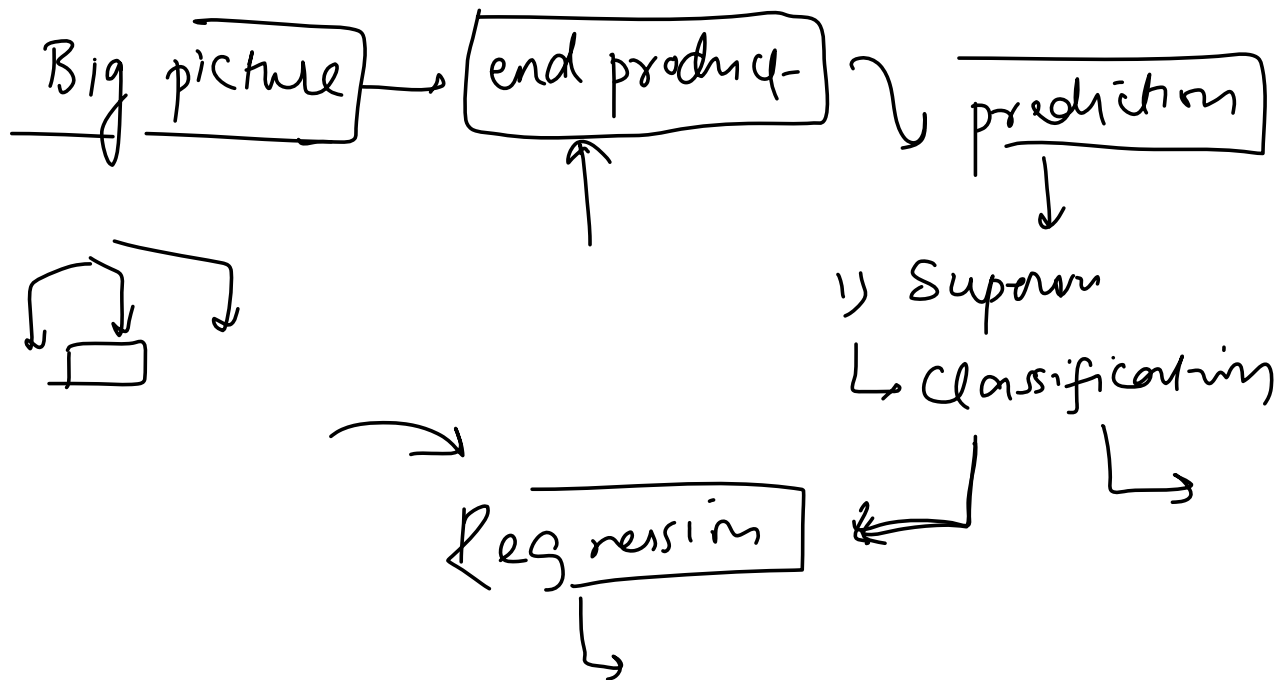
↑

↳ 3.5%

↳ 100 → 98 → 2% → 2.1%

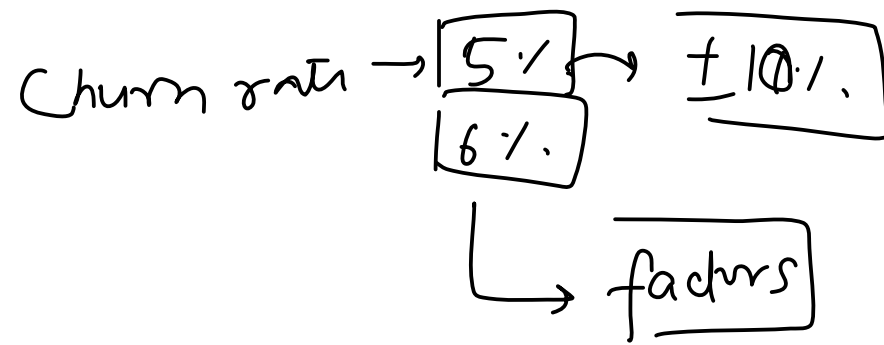
2. Type of Problem

Monday, March 29, 2021 7:30 AM



3. Current Solution

Monday, March 29, 2021 7:30 AM

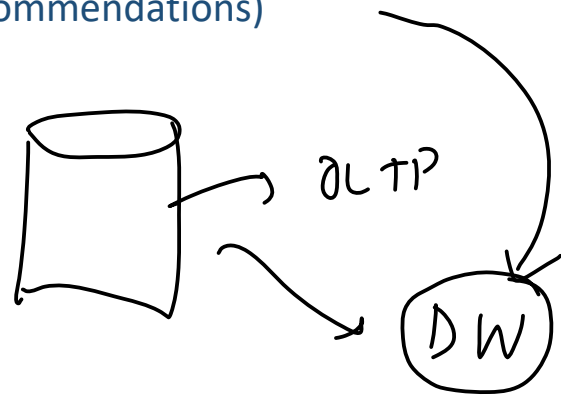


4. Getting Data

Monday, March 29, 2021 7:30 AM

1. Watch time
2. Search but did not find
3. Content left in the middle
4. Clicked on recommendations(order of recommendations)

Data engineer

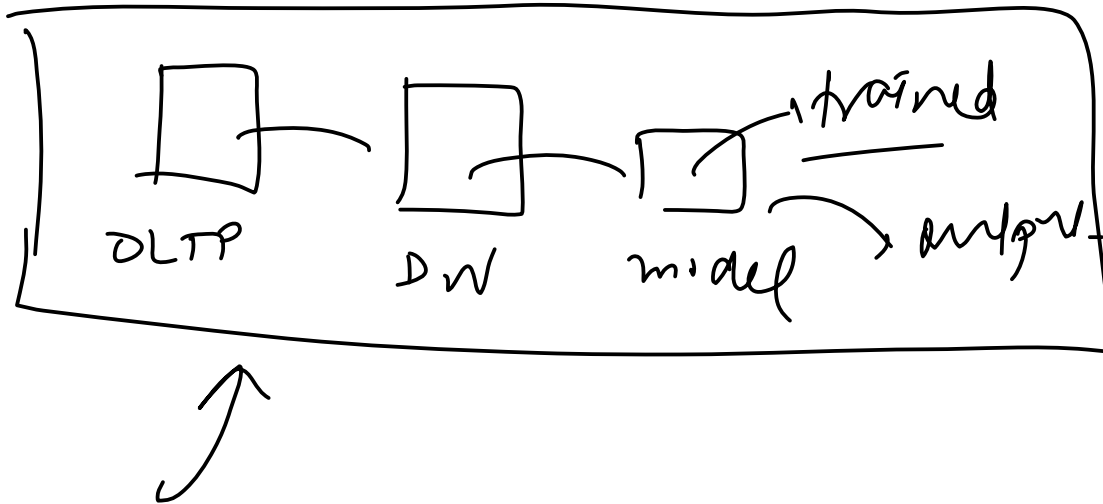


5. Metrics to measure

Monday, March 29, 2021 7:30 AM

6. Online Vs Batch?

Monday, March 29, 2021 7:31 AM

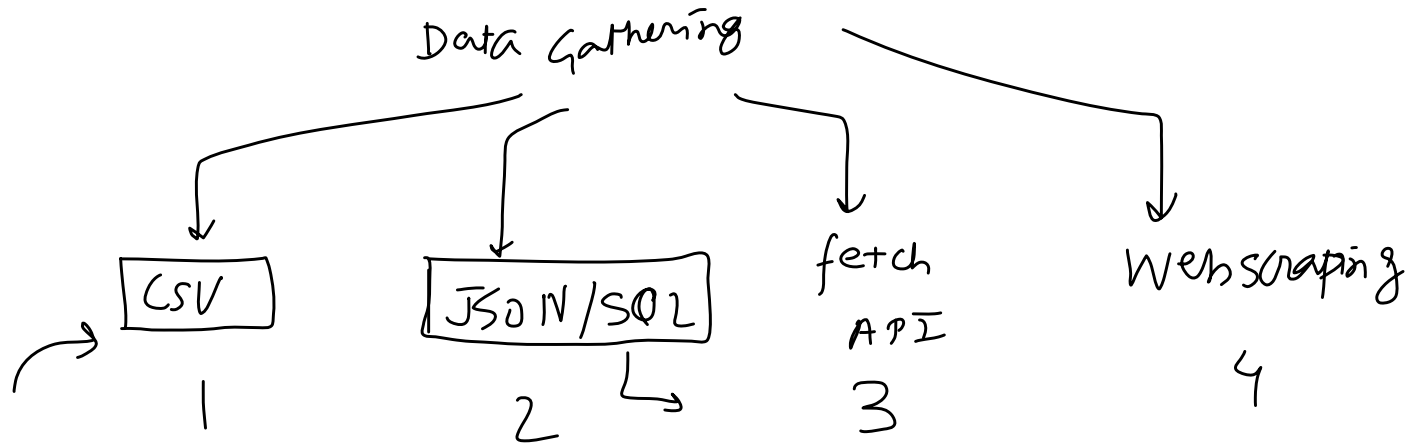


7. Check Assumptions

Monday, March 29, 2021 7:31 AM

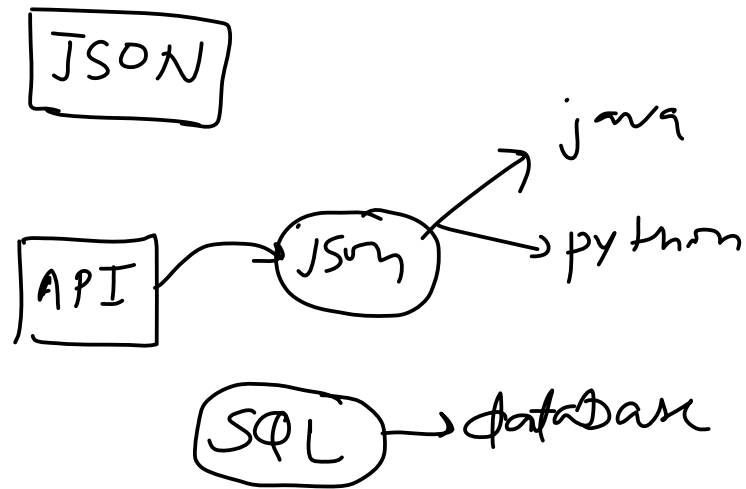
Gathering Data

Tuesday, March 30, 2021 5:35 PM



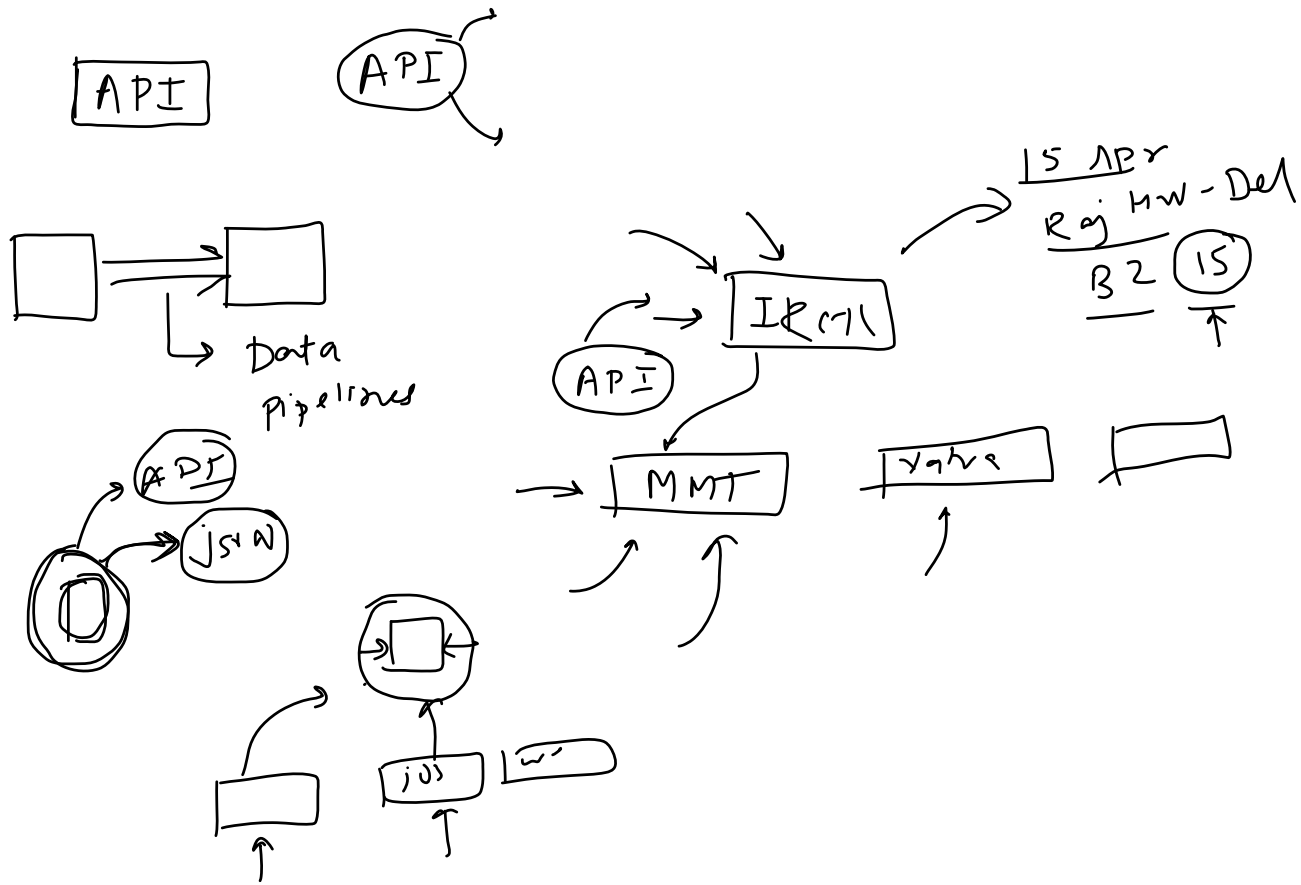
JSON/SQL

Wednesday, March 31, 2021 8:10 AM



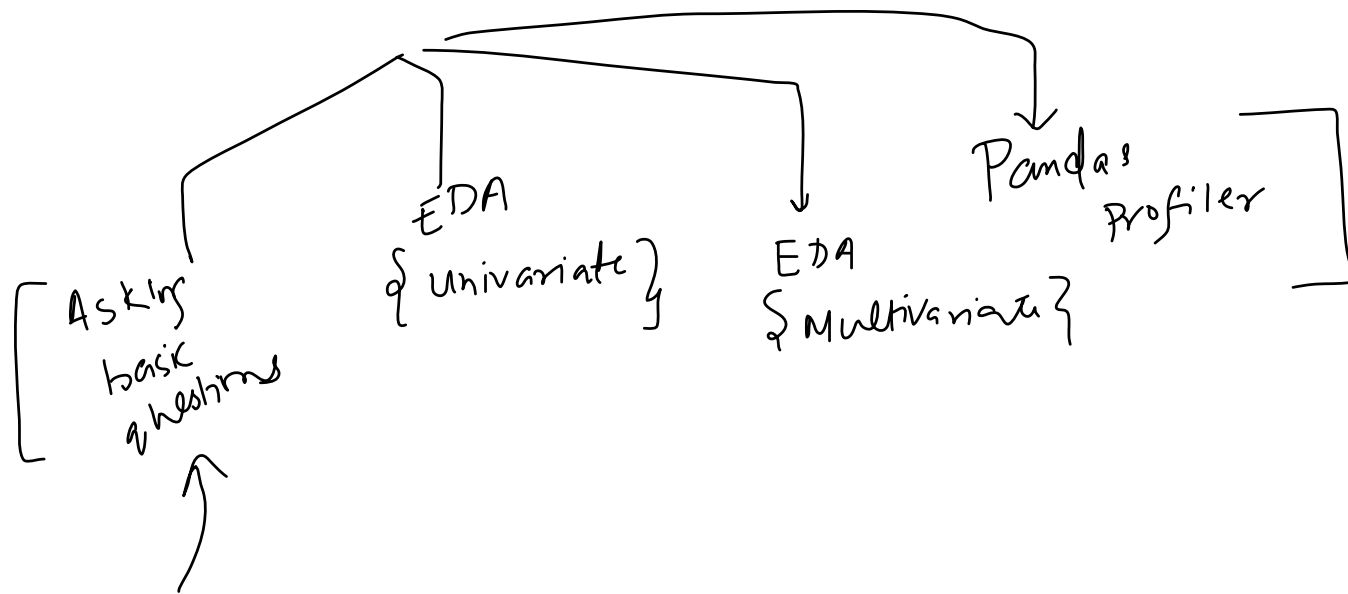
What is an API

Thursday, April 1, 2021 6:55 AM



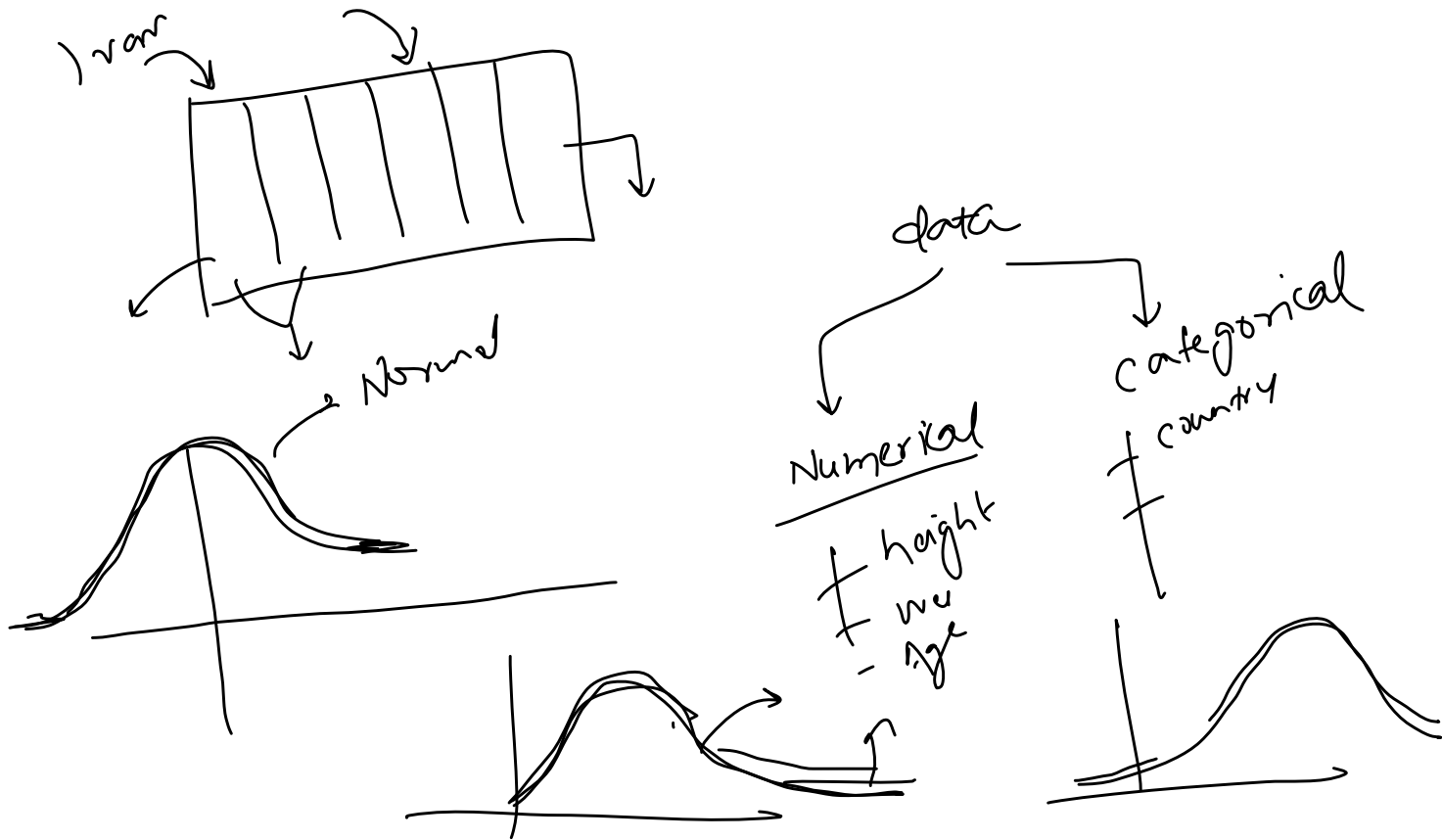
Understanding Your Data

Saturday, April 3, 2021 9:07 AM



Univariate Analysis

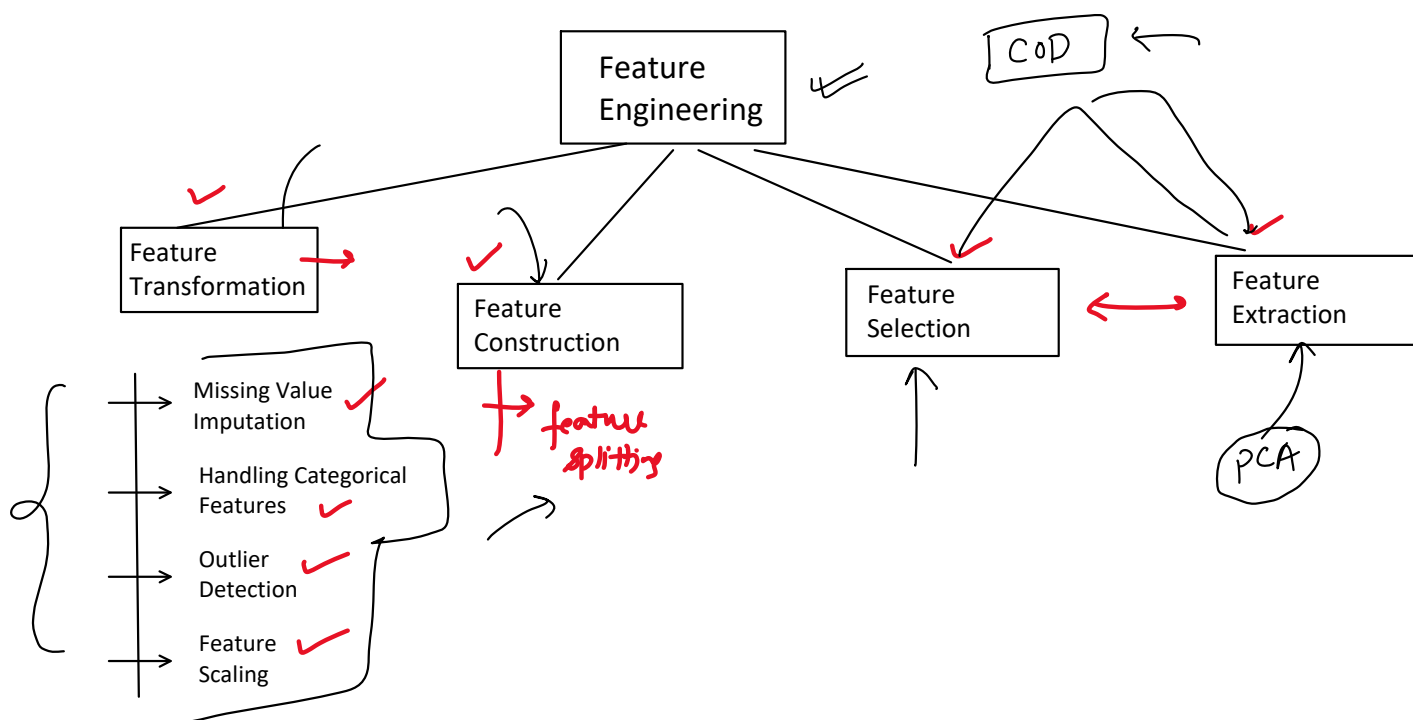
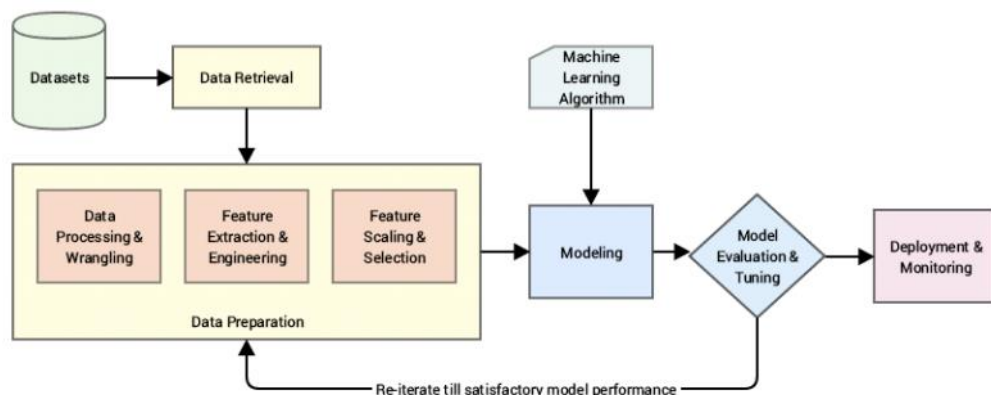
Sunday, April 4, 2021 5:39 PM



Feature Engineering

Thursday, April 8, 2021 6:02 PM

Feature engineering is the process of using domain knowledge to extract features from raw data. These features can be used to improve the performance of machine learning algorithms.



1.1 Missing Values Imputation

Thursday, April 8, 2021 7:05 PM

Average_Age = 26.0

ID	City	Age	Married ?
1	Lisbon	25	0
2	Berlin	25	1
3	Lisbon	30	1
4	Lisbon	30	1
5	Berlin	18	0
6	Lisbon	NaN	0
7	Berlin	30	1
8	Berlin	NaN	0
9	Berlin	25	1
10	Madrid	25	1



ID	City	Age	Married ?
1	Lisbon	25	0
2	Berlin	25	1
3	Lisbon	30	1
4	Lisbon	30	1
5	Berlin	18	0
6	Lisbon	26	0
7	Berlin	30	1
8	Berlin	26	0
9	Berlin	25	1
10	Madrid	25	1

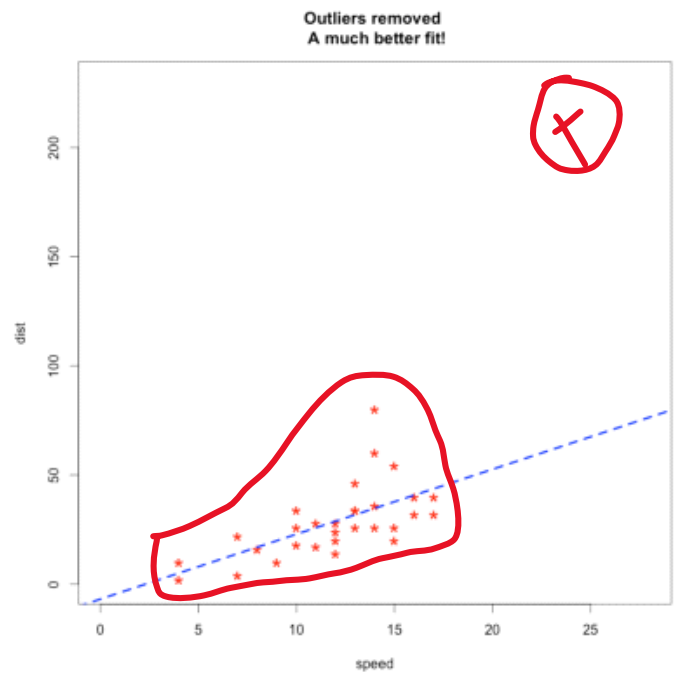
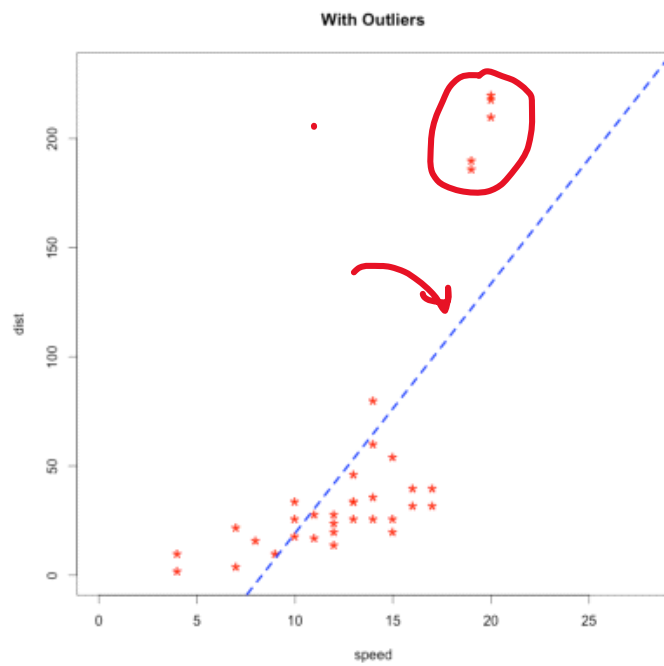
1.2 Handling Categorical Values

Thursday, April 8, 2021 7:06 PM

Index	Animal	One-Hot code 	Index	Dog	Cat	Sheep	Lion	Horse
0	Dog		0	1	0	0	0	0
1	Cat		1	0	1	0	0	0
2	Sheep		2	0	0	1	0	0
3	Horse		3	0	0	0	0	1
4	Lion		4	0	0	0	1	0

1.3 Outlier Detection

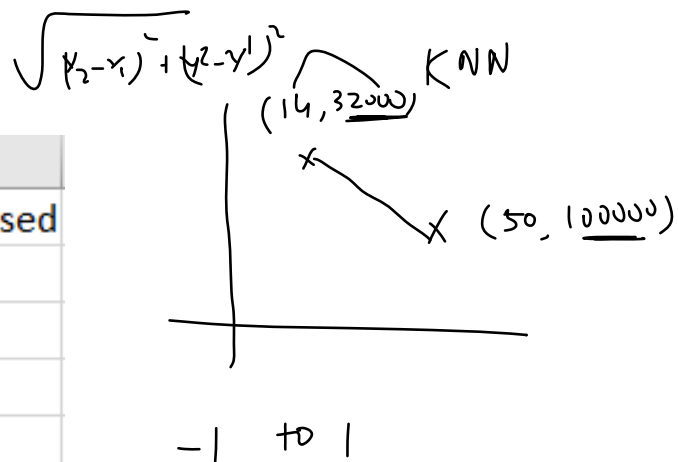
Thursday, April 8, 2021 7:07 PM



1.4 Feature Scaling

Thursday, April 8, 2021 7:08 PM

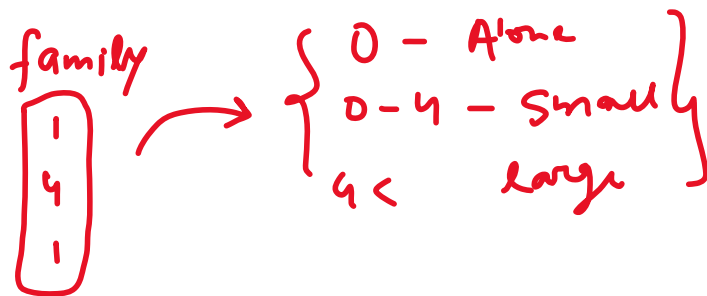
	A	B	C	D
1	Country	Age	Salary	Purchased
2	France	44	72000	No
3	Spain	27	48000	Yes
4	Germany	30	54000	No
5	Spain	38	61000	No
6	Germany	40		Yes
7	France	35	58000	Yes
8	Spain		52000	No
9	France	48	79000	Yes
10	Germany	50	83000	No
11	France	37	67000	Yes



2. Feature Construction

Thursday, April 8, 2021 7:09 PM

P. Id	Survived	Pclass	Name	Sex	Age	SibSp	Parch	Ticket	Fare	Cabin	Embarked
1	0	3	Braund, Mr. Owen Harris	male	22	- 1	- 0	A/5 21171	7.25		S
2	1	1	Cumings, Mrs. John Bradley (Florence Briggs)	female	38	1	- 0	PC 17599	71.2833	C85	C
3	1	3	Heikkinen, Miss. Laina	female	26	0	- 0	STON/O2.	7.925		S
4	1	1	Futrelle, Mrs. Jacques Heath (Lily May Peel)	female	35	- 1	0	113803	53.1	C123	S
5	0	3	Allen, Mr. William Henry	male	35	0	0	373450	8.05		S
6	0	3	Moran, Mr. James	male		0	0	330877	8.4583		Q
7	0	1	McCarthy, Mr. Timothy J	male	54	0	0	17463	51.8625	E46	S
8	0	3	Palsson, Master. Gosta Leonard	male	2	- 3	1	349909	21.075		S
9	1	3	Johnson, Mrs. Oscar W (Elisabeth Vilhelmin)	female	27	0	2	347742	11.1333		S
10	1	2	Nasser, Mrs. Nicholas (Adele Achem)	female	14	1	0	237736	30.0708		C



3. Feature Selection ✓

Thursday, April 8, 2021 7:11 PM

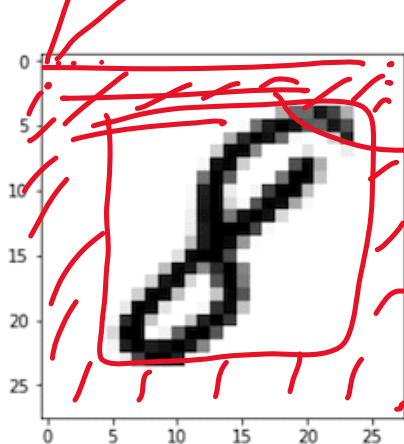
```
In [5]: train_df = pd.read_csv(f'{PATH}train.csv', header = 0)
```

```
In [19]: train_df
```

```
Out[19]:
```

784 features

	label	pixel0	pixel1	pixel2	pixel3	pixel4	pixel5	pixel6	pixel7	pixel8	...	pixel774
0	1	0	0	0	0	0	0	0	0	0	...	0
1	0	0	0	0	0	0	0	0	0	0	...	0
2	1	0	0	0	0	0	0	0	0	0	...	0
3	4	0	0	0	0	0	0	0	0	0	...	0



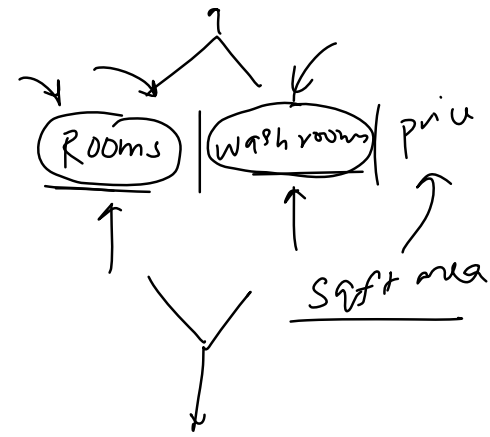
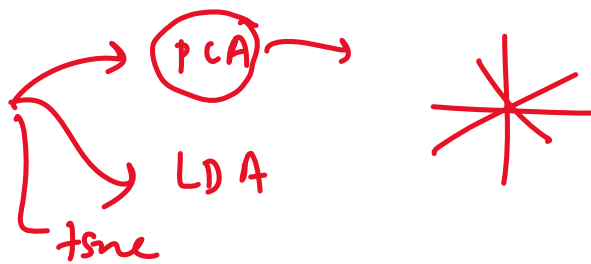
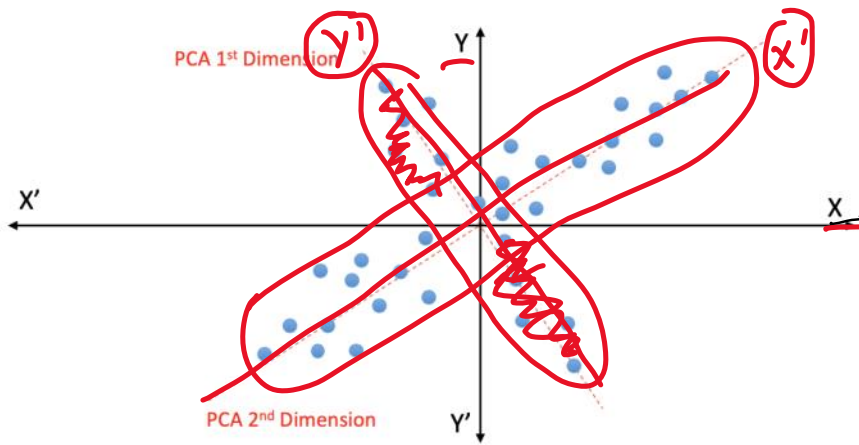
→ MNIST

50000
imgs

$$28 \times 28 = 784$$

4. Feature Extraction ✓

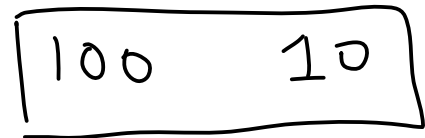
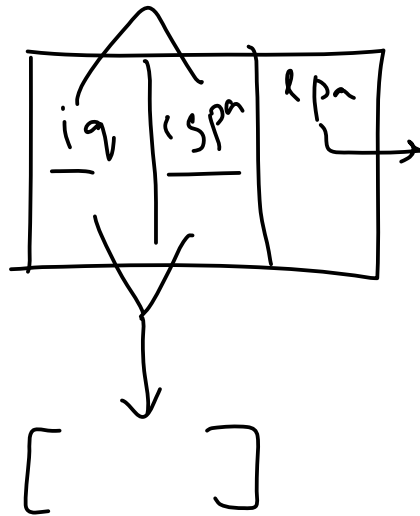
Thursday, April 8, 2021 7:13 PM



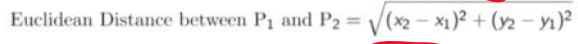
What is Feature Scaling?

Friday, April 9, 2021 4:21 PM

Feature Scaling is a technique to standardize the independent features present in the data in a fixed range



Friday, April 9, 2021 4:27 PM

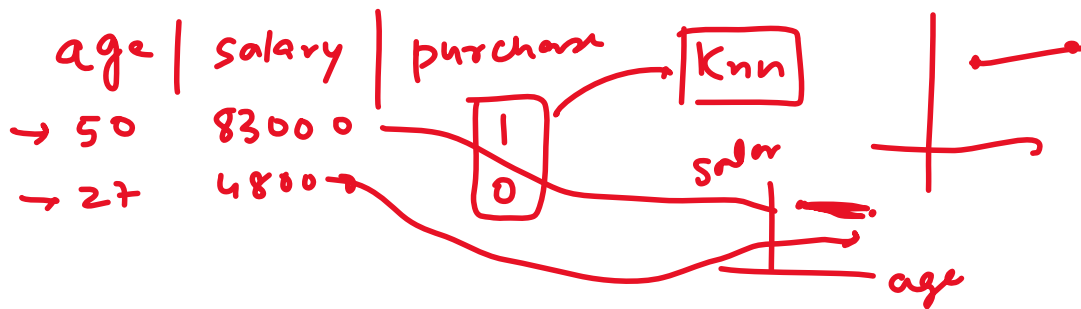


Example: x_1 & y_1 are in row 2, x_2 & y_2 are in row 9

$= 1225000000$ ✓

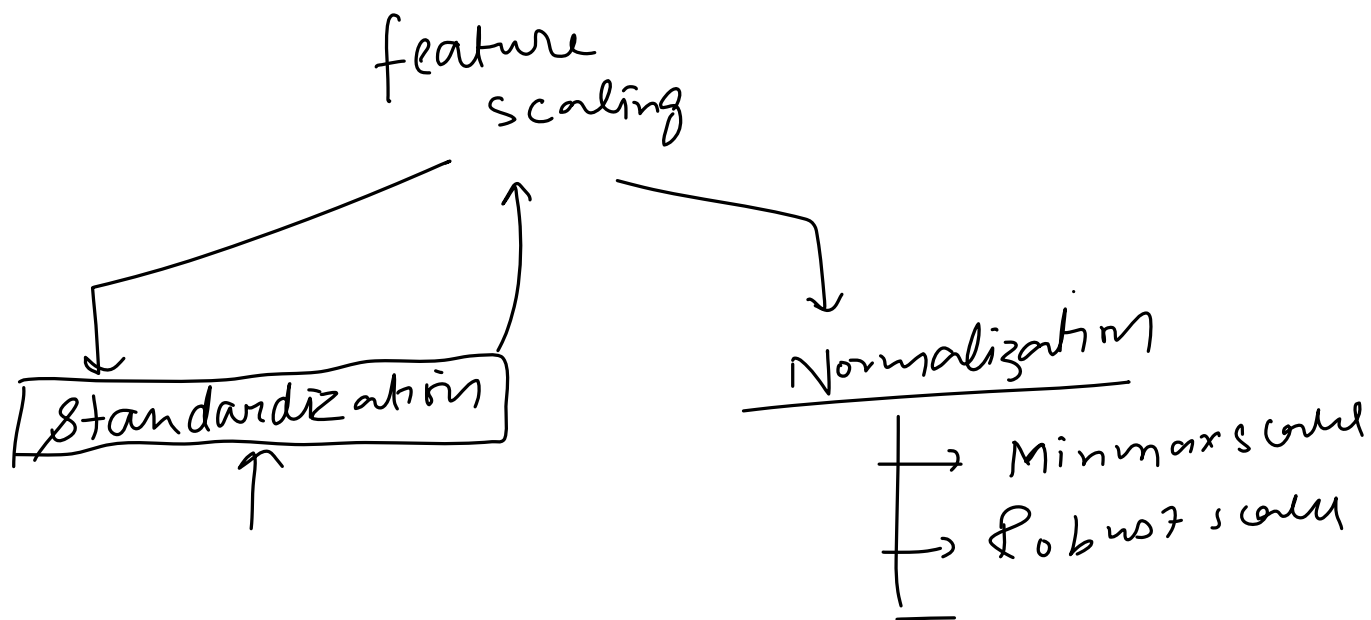
K_{nn}

$T = 529$



Types of Feature Scaling

Friday, April 9, 2021 4:27 PM

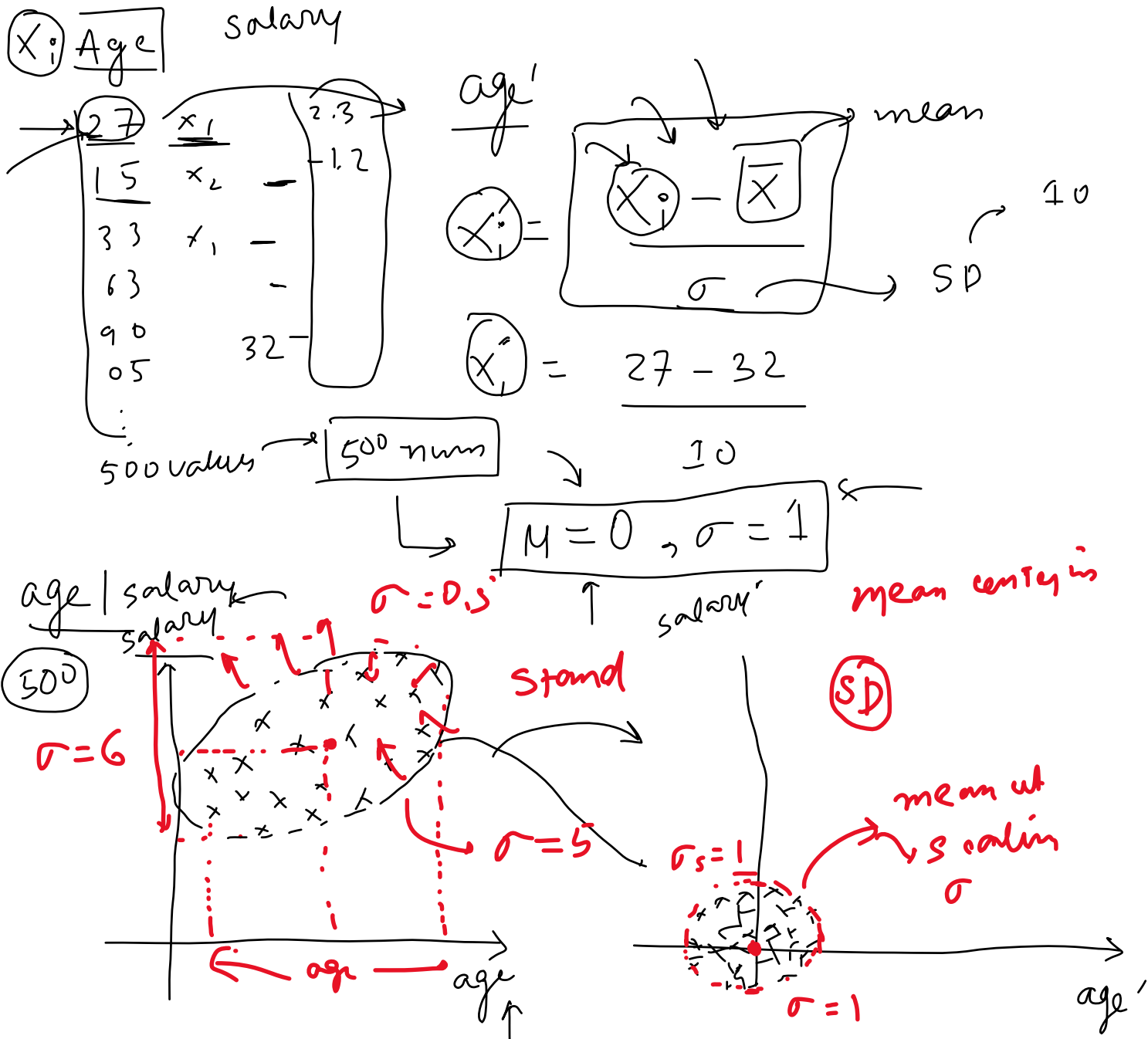


Standardization - Intuition

Friday, April 9, 2021 4:27 PM

$$\frac{15 - 32}{10}$$

Also called as Z-score Normalization



Example

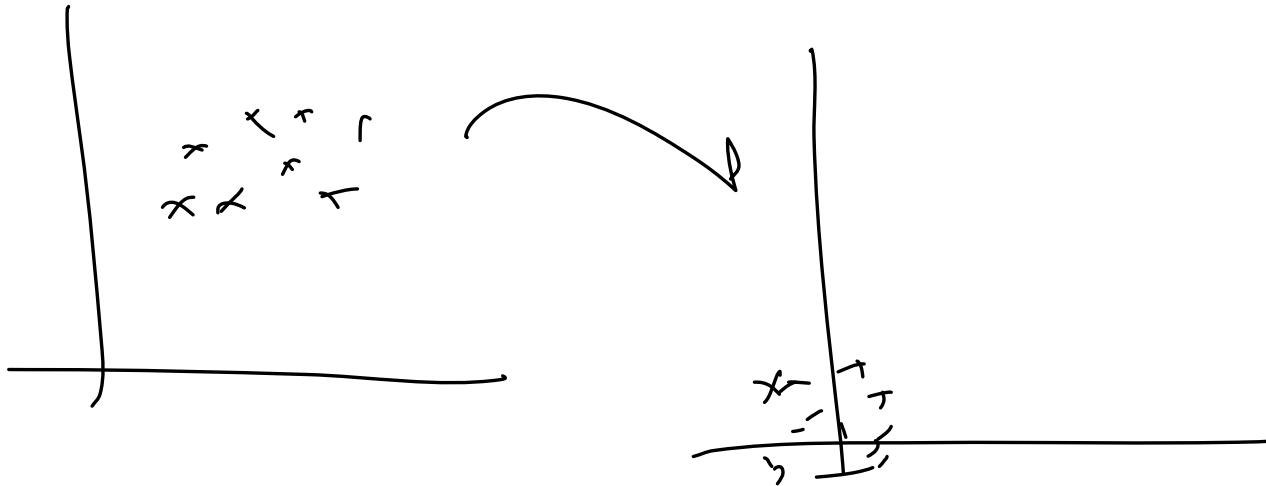
Friday, April 9, 2021

4:28 PM

Impact of Outliers

Friday, April 9, 2021 4:28 PM

(x) ?



Friday, April 9, 2021 4:28 PM

Algorithm(s)	Reason of applying feature scaling
1. <u>K-Means</u>	Use the Euclidean distance measure.
2. <u>K-Nearest-Neighbours</u>	Measure the distances between pairs of samples and these distances are influenced by the measurement units
3. <u>Principal Component Analysis (PCA)</u>	Try to get the feature with <u>maximum variance</u>
4. <u>Artificial Neural Network</u>	Apply Gradient Descent
5. <u>Gradient Descent</u>	Theta calculation becomes faster after feature scaling and the learning rate in the update equation of Stochastic Gradient Descent is the same for every parameter

Normalization

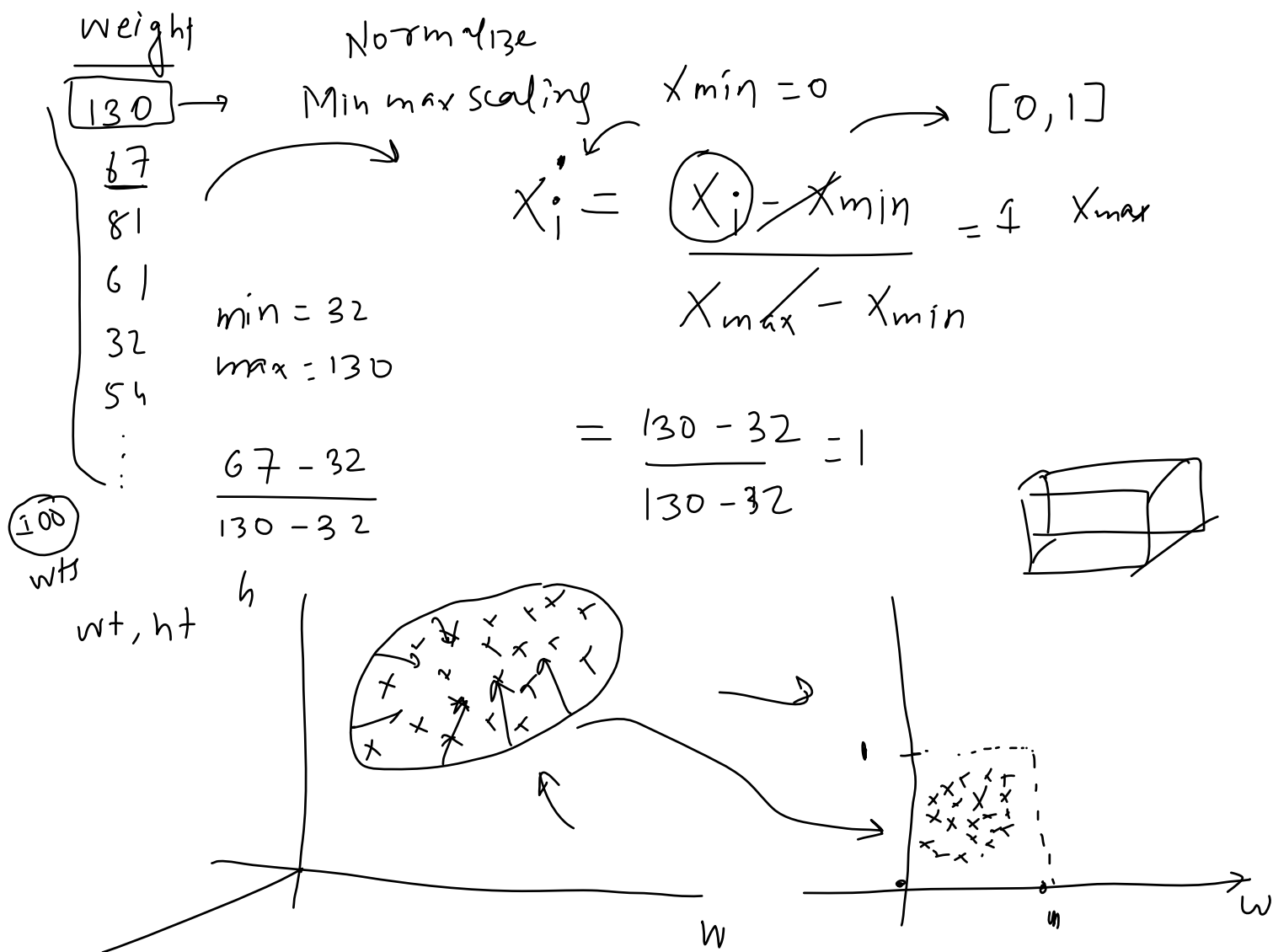
Saturday, April 10, 2021 1:15 PM

Normalization is a technique often applied as part of data preparation for machine learning. The goal of normalization is to change the values of numeric columns in the dataset to use a common scale, without distorting differences in the ranges of values or losing information

- + MinMaxScaling
 - + Mean normalization
 - + Max Absorb
 - Robust scaling
- 

MinMaxScaling - Intuition

Saturday, April 10, 2021 1:16 PM



Code Example

Saturday, April 10, 2021 1:17 PM

wine_data →

↗

Mean Normalization

Saturday, April 10, 2021 1:16 PM

wt
200
100

normal

val < mean

-val | 1 - val

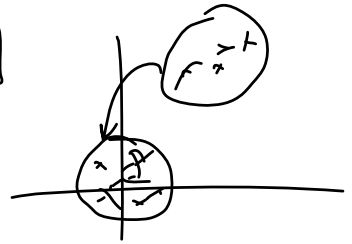
$$x'_i = \frac{x_i - x_{\text{mean}}}{x_{\text{max}} - x_{\text{min}}} [-1 \text{ to } 1]$$

mean
center

sklearn

centered
data

Standardization



MaxAbsScaling

Saturday, April 10, 2021 1:17 PM

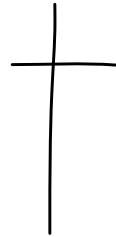
Wt
200
100
300



$$X'_i = \frac{X_i}{|X_{\max}|}$$



MaxAbsScaler



sparse data



0's

Robust Scaling

Saturday, April 10, 2021 1:17 PM

wt

200

300

100

$$X_i' = \frac{X_i - X_{\text{median}}}{\text{IQR} \{ 75^{\text{th}} \text{ pt} - 25^{\text{th}} \text{ pt} \}}$$

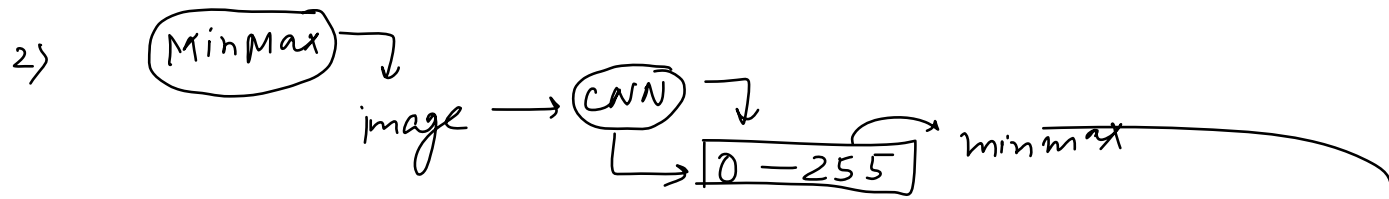
Robust scale

→ Robust to Outliers

Normalization Vs Standardization

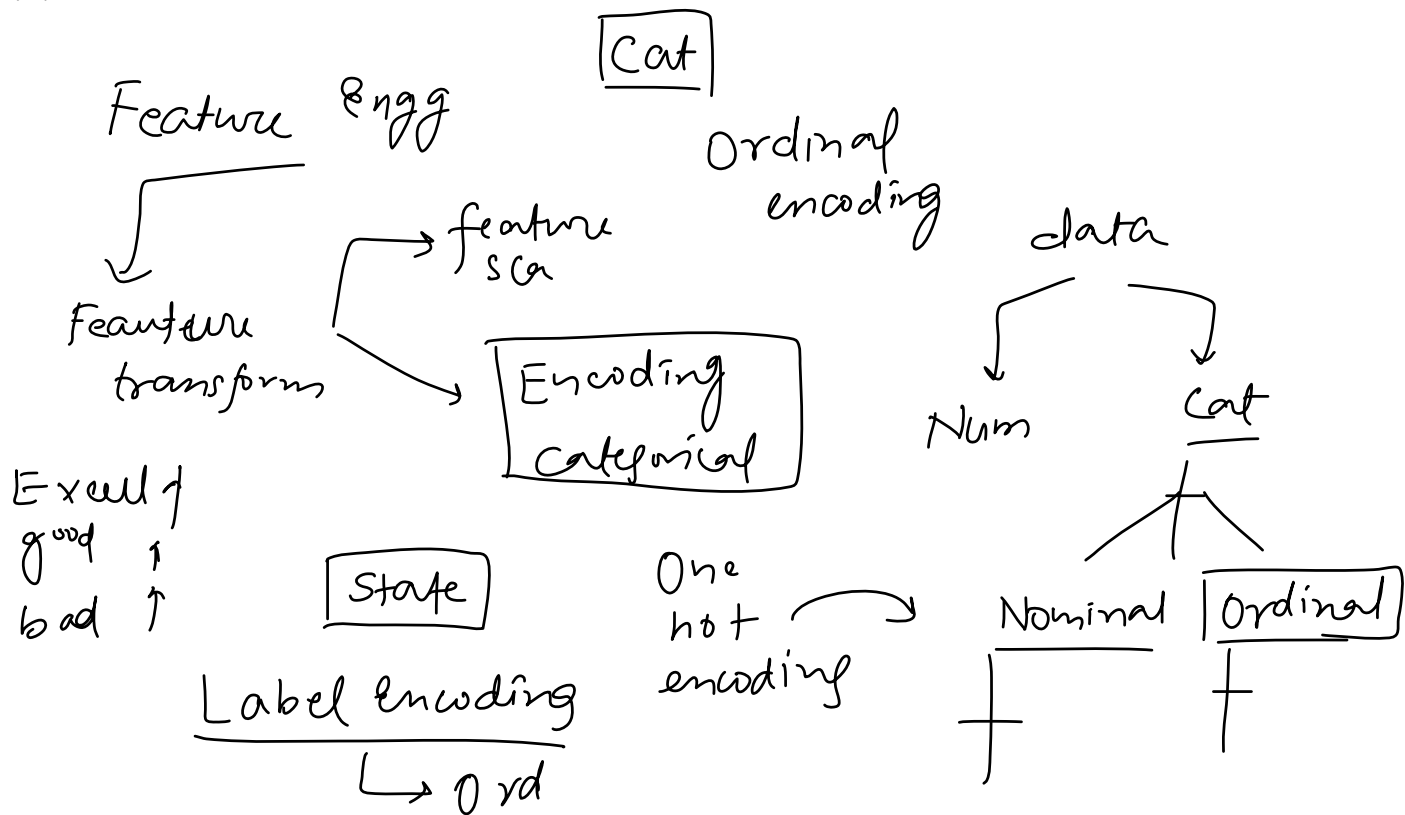
Saturday, April 10, 2021 1:17 PM

1) Is feature scaling required?



Encoding Categorical Variables

Monday, April 12, 2021 3:53 PM



Ordinal Encoding

Monday, April 12, 2021 3:54 PM

Nominal (OME)

ordinal
encoding

$PG > UG > HS$

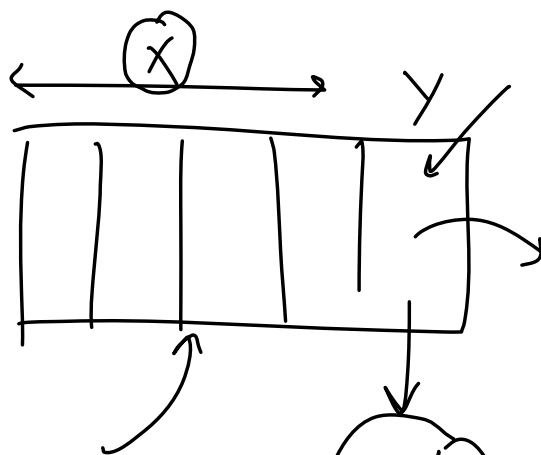
<u>Education</u>		
HS	→	0
UG	—	1
PG	—	2
PG	—	2
UG	—	1
HS	—	0
UG	—	2

$\left\{ \begin{array}{l} PG \rightarrow 2 \\ UG \rightarrow 1 \\ HS \rightarrow 0 \end{array} \right\}$

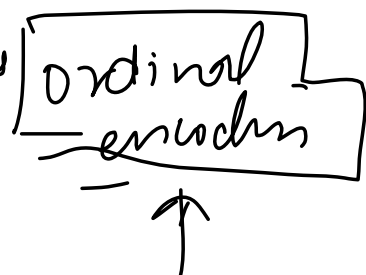
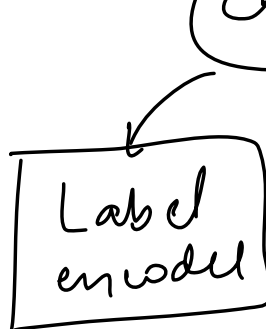
Label Encoding

Monday, April 12, 2021 3:54 PM

Ordinal
encoder



Categorical
classific

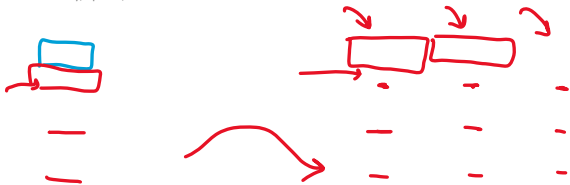


Example

Monday, April 12, 2021 3:54 PM

OneHotEncoding

Tuesday, April 13, 2021 12:24 PM



$[Y, B, R]$
0 1 2

$[1, 0, 0] \rightarrow \text{Yellow}$
 $[0, 1, 0] \rightarrow \text{blue}$
 $[0, 0, 1] \rightarrow \text{red}$

50 diff

Color	Target
Yellow	0
Yellow	1
Blue	1
Yellow	1
Red	1
Yellow	0
Red	1
Red	0
Yellow	1
Blue	0



Color_Y	Color_B	Color_R	Target
1	0	0	0
1	0	0	1
0	1	0	1
1	0	0	1
0	0	1	1
1	0	0	0
0	0	1	1
0	0	1	0
1	0	0	1
0	1	0	0

One Hot Encoding

Dummy Variable Trap

Tuesday, April 13, 2021 12:26 PM

Color	Target
Yellow	0
Yellow	1
Blue	1
Yellow	1
Red	1
Yellow	0
Red	1
Red	0
Yellow	1
Blue	0

① → **Multic**

n cat
($n-1$) ✓
cols →

Color_Y	Color_B	Color_R	Target
1	0	0	0
1	0	0	1
0	1	0	1
1	0	0	1
0	0	1	1
1	0	0	0
0	0	1	1
0	0	1	0
1	0	0	1
0	1	0	0

① n
 n cols
① $n-1$ cols

One Hot Encoding

Agneticollinearity →

Y - $\begin{bmatrix} 0 & 0 \end{bmatrix}$ ✓
B - $\begin{bmatrix} 0 & 1 & 0 \end{bmatrix}$ ✓
R - $\begin{bmatrix} 0 & 0 & 1 \end{bmatrix}$ ✓

$\sum = 1$

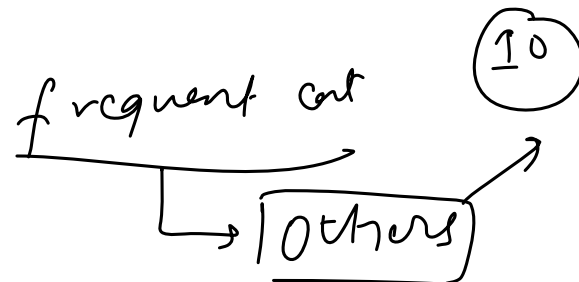
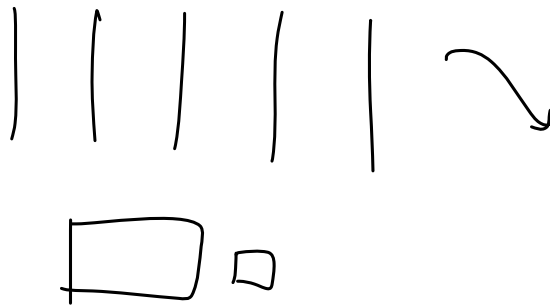
$\sum = 1$ → linear model

Linear Reg
logistic

OHE using most frequent variables

Tuesday, April 13, 2021 12:31 PM

brand → nominal → 40 brands
cat → OHE



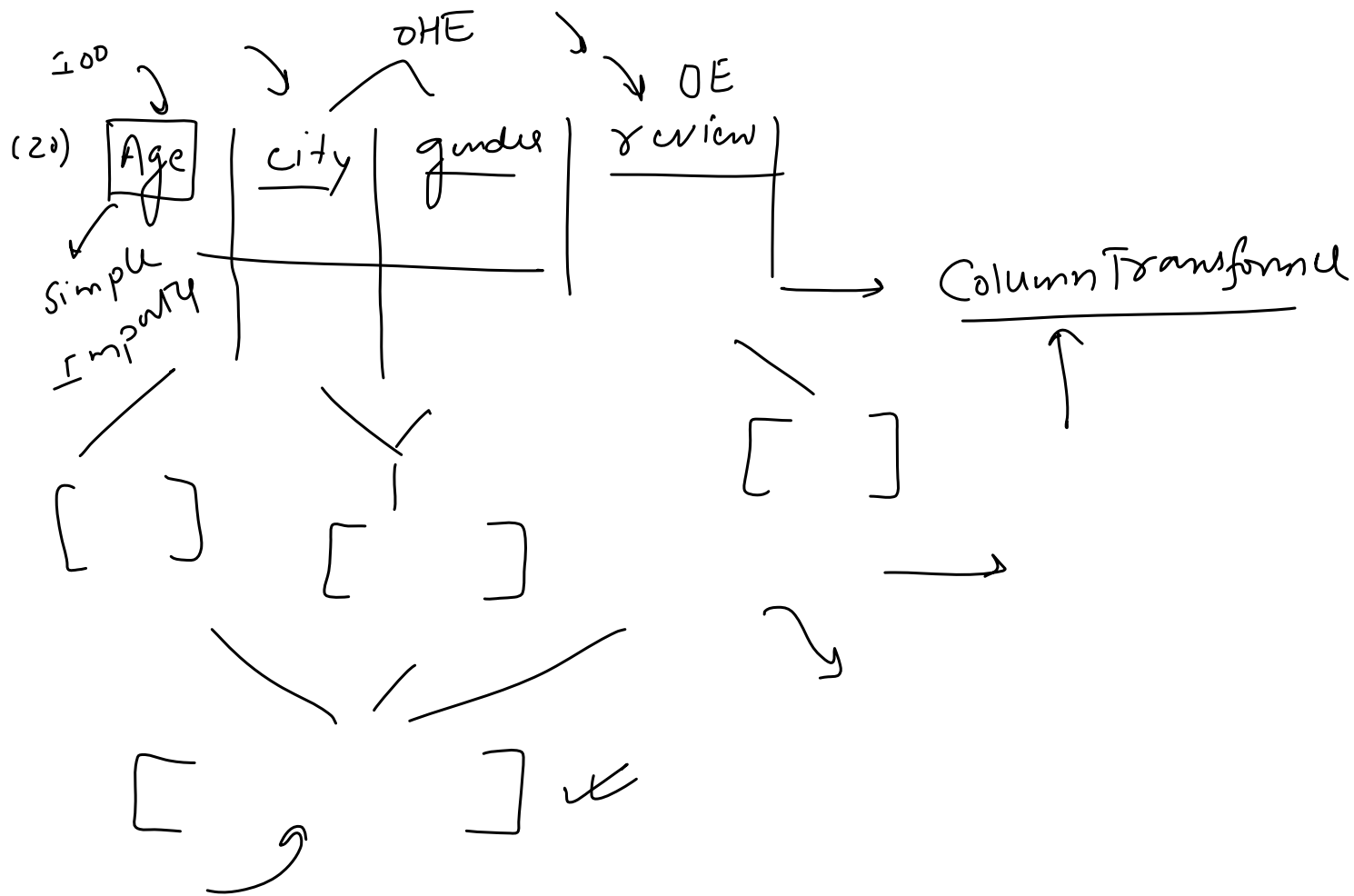
Example

Tuesday, April 13, 2021 12:26 PM

	brand	km_driven	fuel	owner	selling_price
0	Maruti	145500	Diesel	First Owner	450000
1	Skoda	120000	Diesel	Second Owner	370000
2	Honda	140000	Petrol	Third Owner	158000
3	Hyundai	127000	Diesel	First Owner	225000
4	Maruti	120000	Petrol	First Owner	130000

Column Transformer

Wednesday, April 14, 2021 5:16 PM

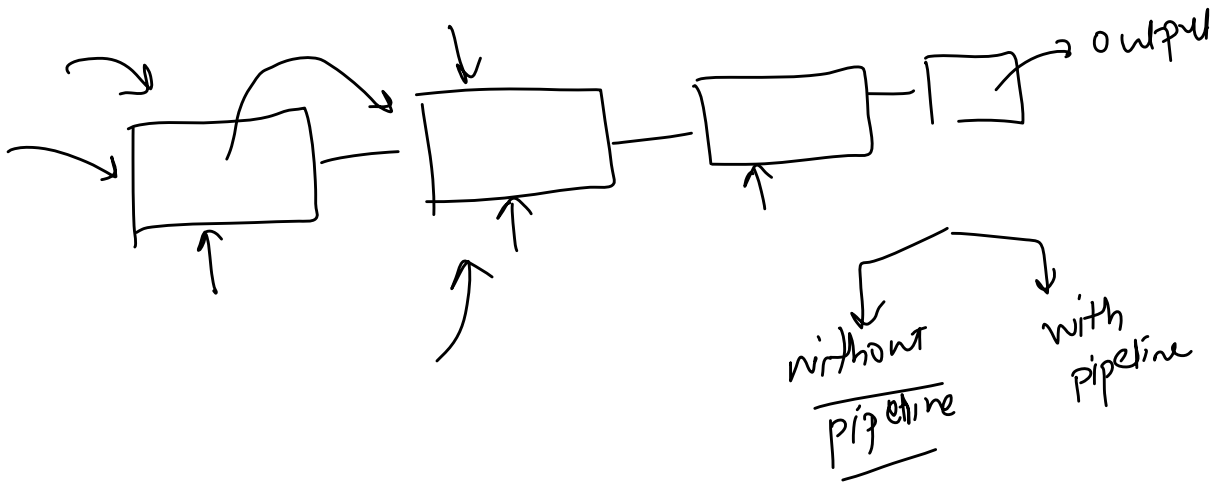


Internship Project Discussion

Wednesday, April 14, 2021 6:41 PM

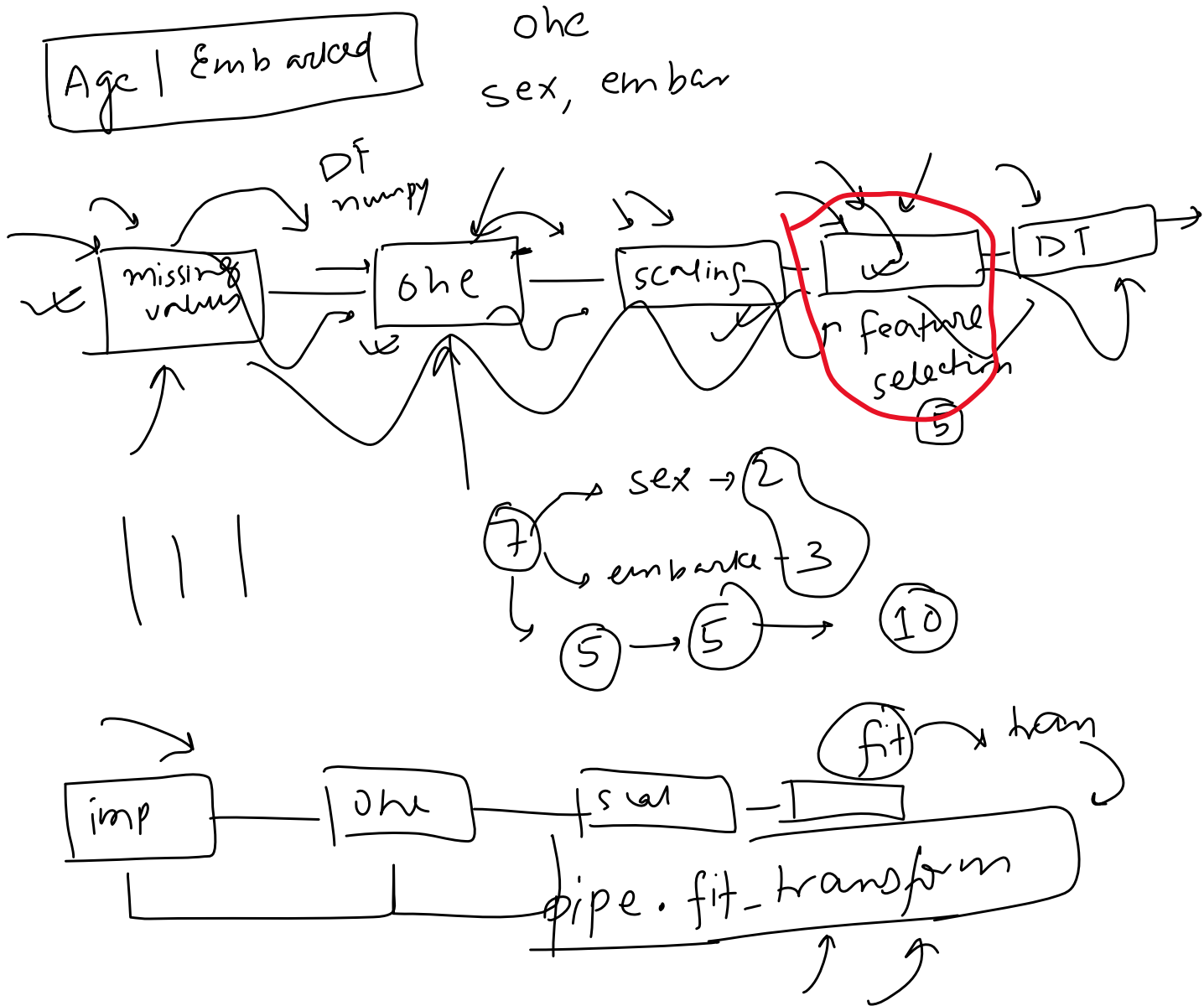
Pipelines chains together multiple steps so that the output of each step is used as input to the next step.

Pipelines makes it easy to apply the same preprocessing to train and test!



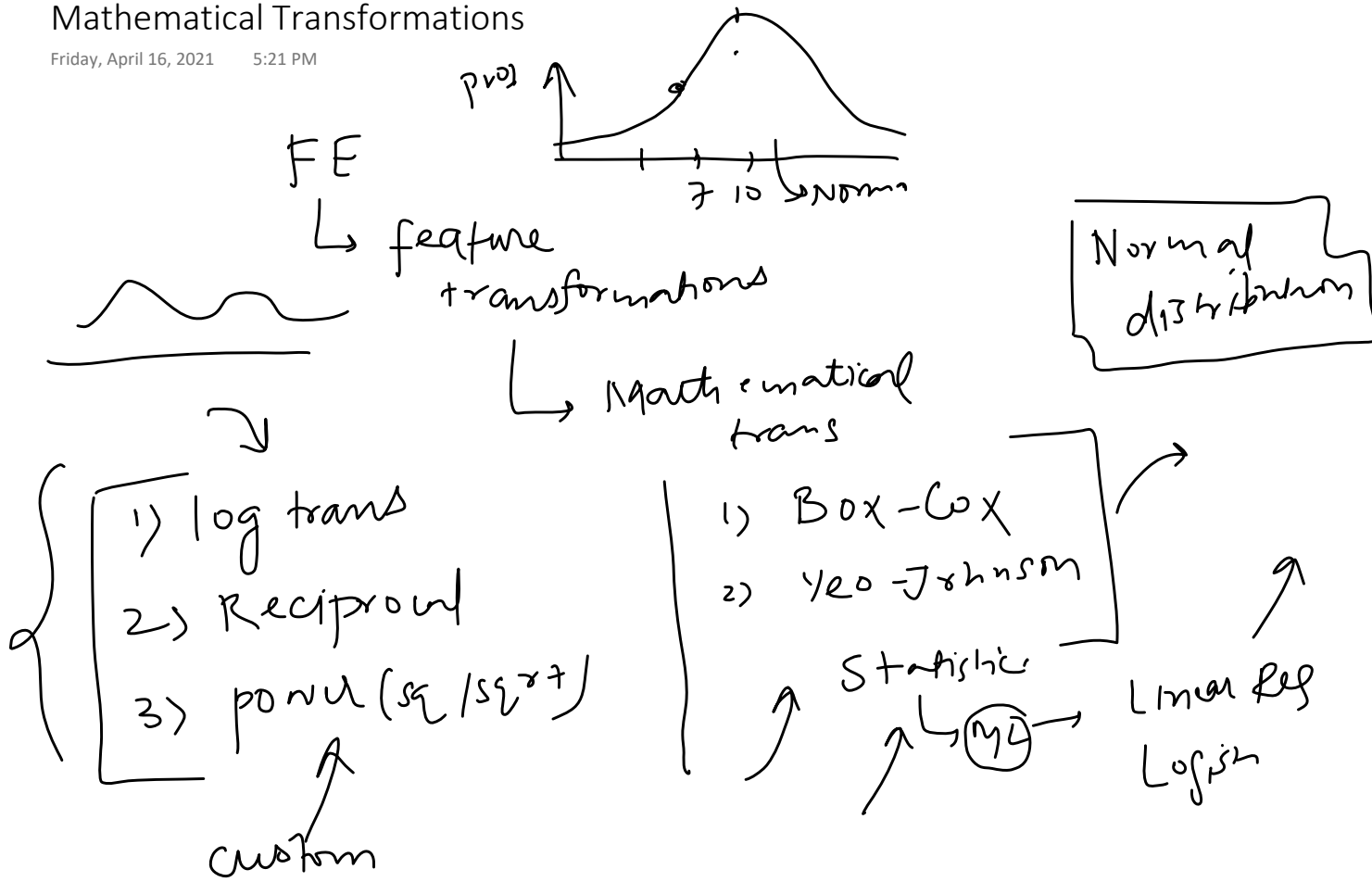
Strategy

Thursday, April 15, 2021 6:58 PM



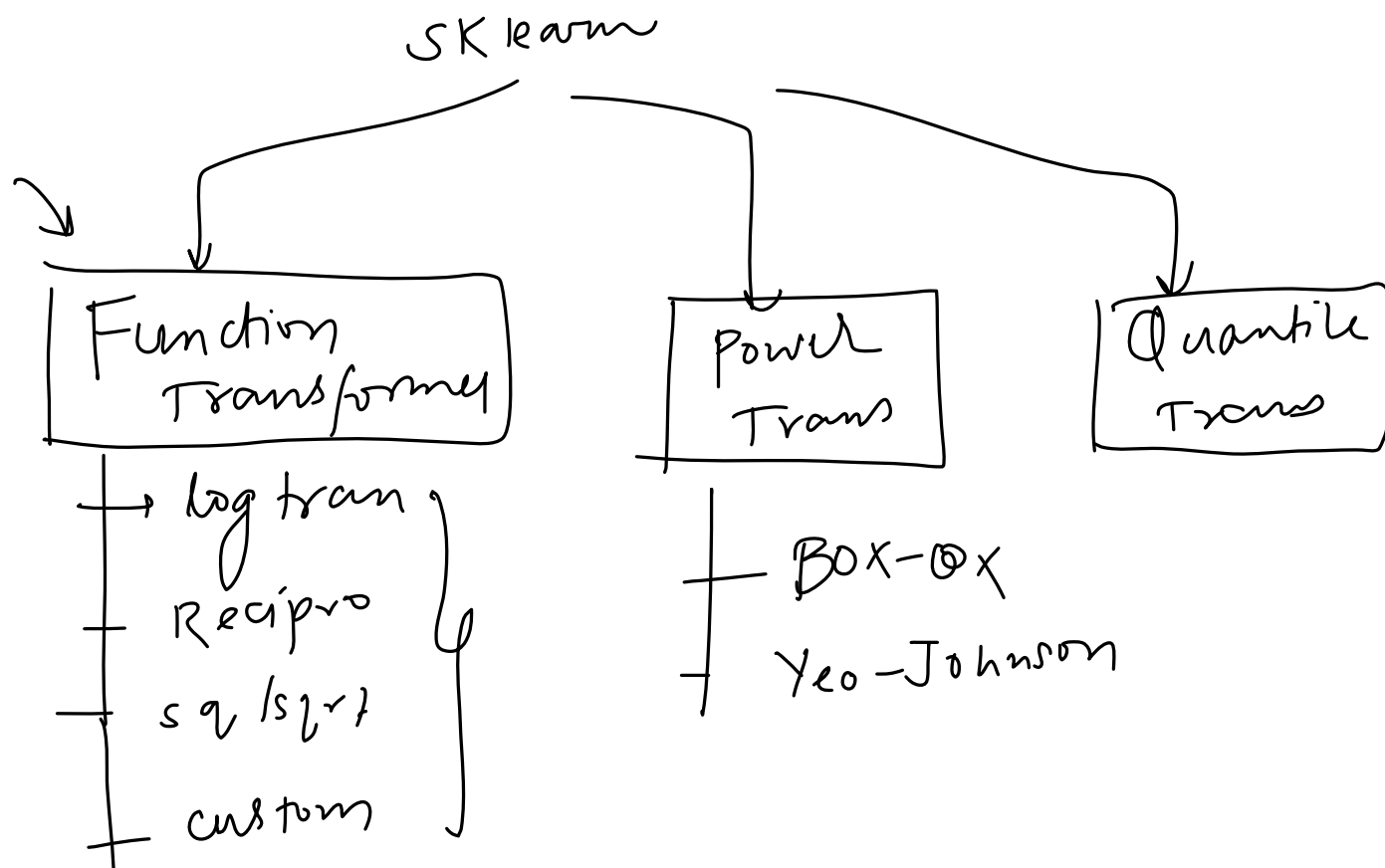
Mathematical Transformations

Friday, April 16, 2021 5:21 PM



Function Transformer

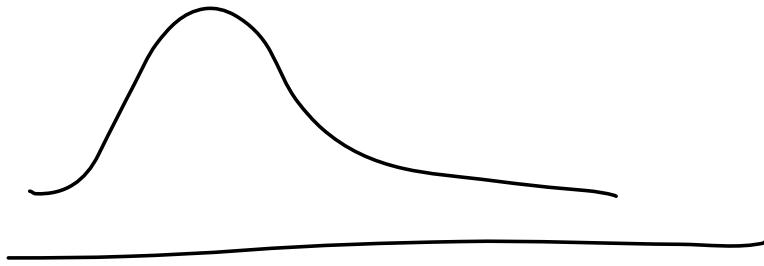
Friday, April 16, 2021 5:22 PM



How to find if data is normal?

Friday, April 16, 2021 6:30 PM

`sns.distplot`



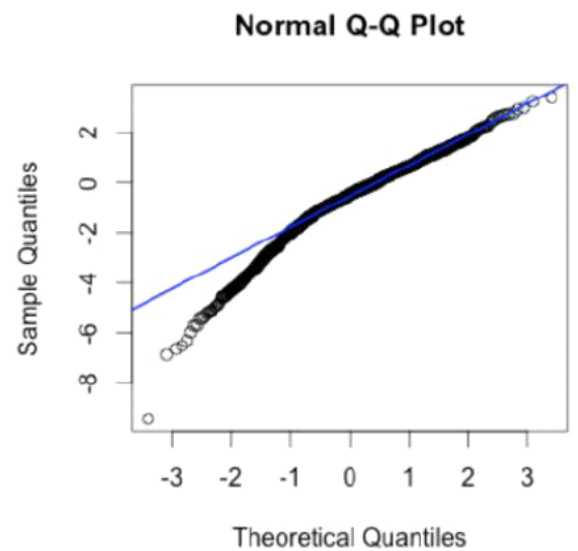
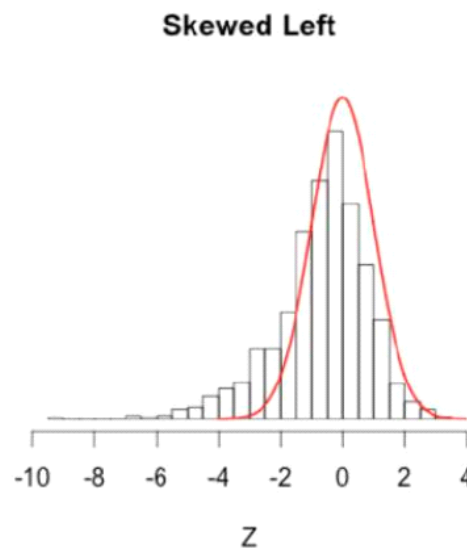
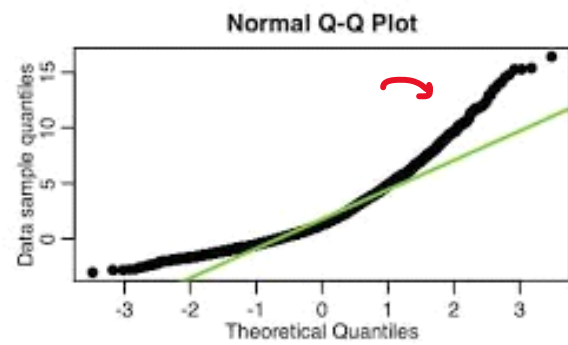
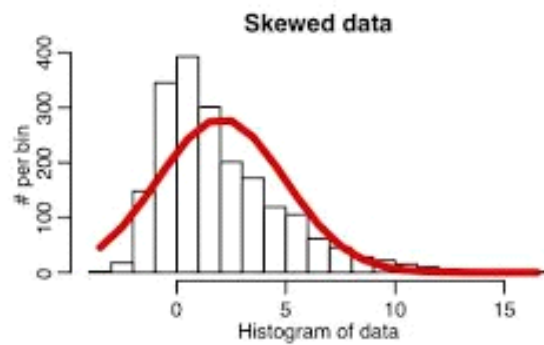
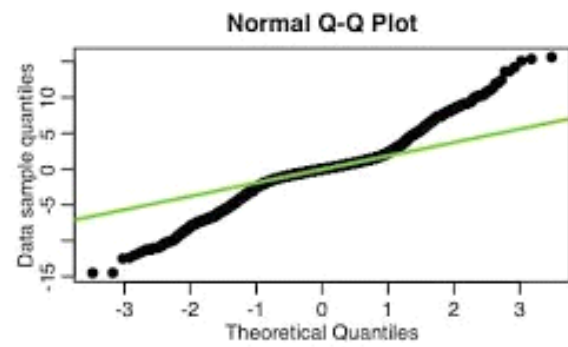
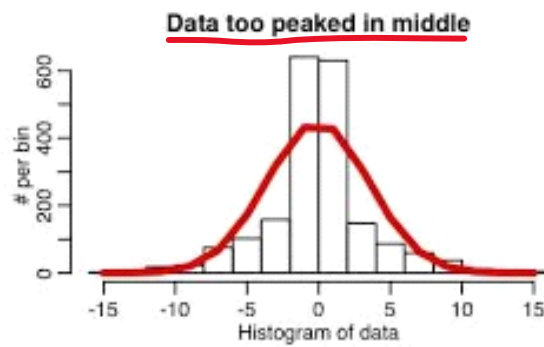
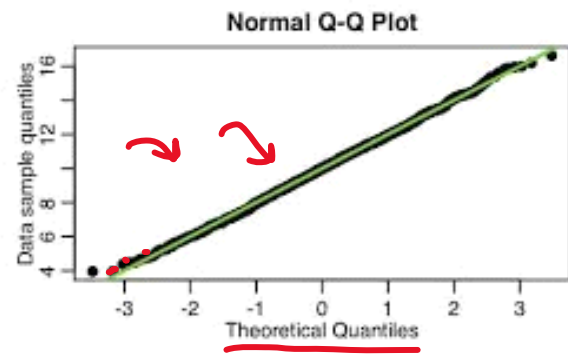
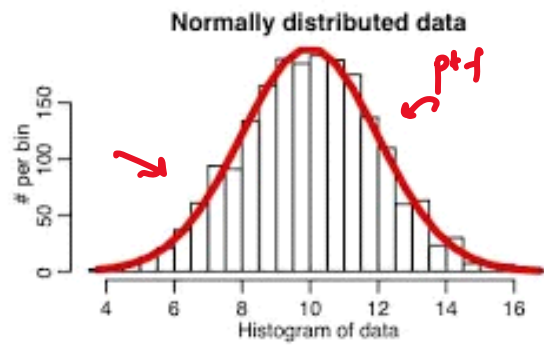
`QQ plot`

`pd.skew()` = 0

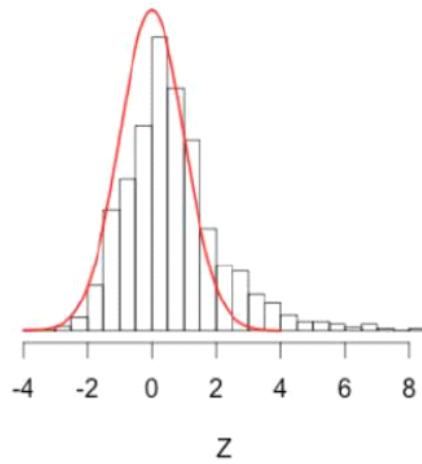
statistics

QQ Plots

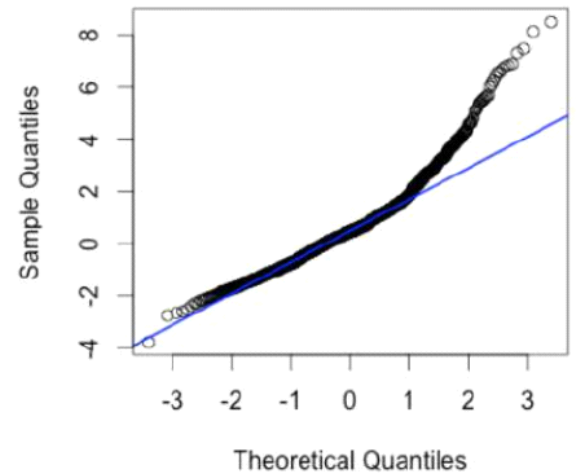
Friday, April 16, 2021 5:23 PM



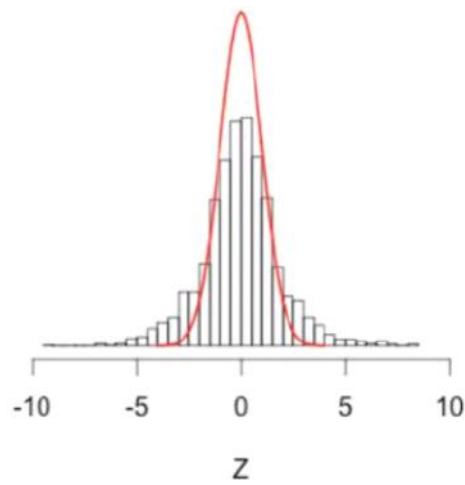
Skewed Right



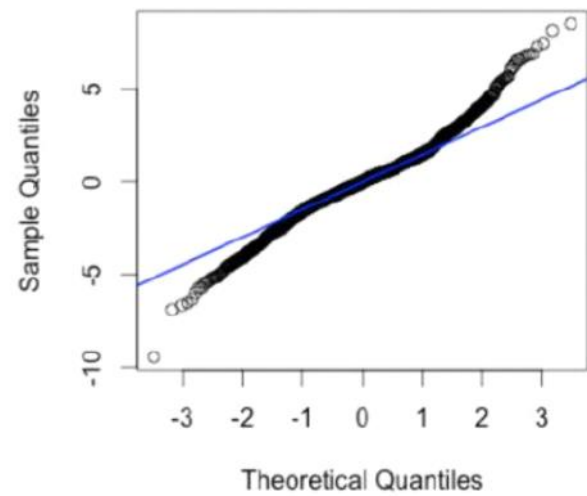
Normal Q-Q Plot



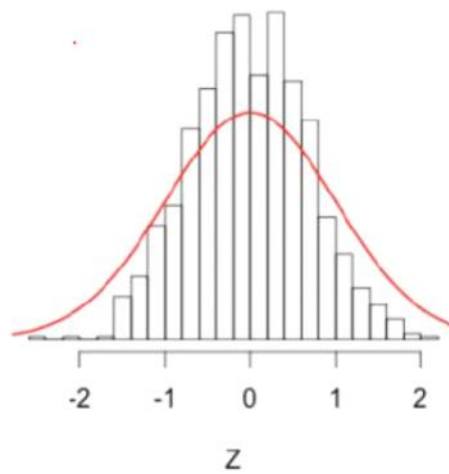
Fat Tails



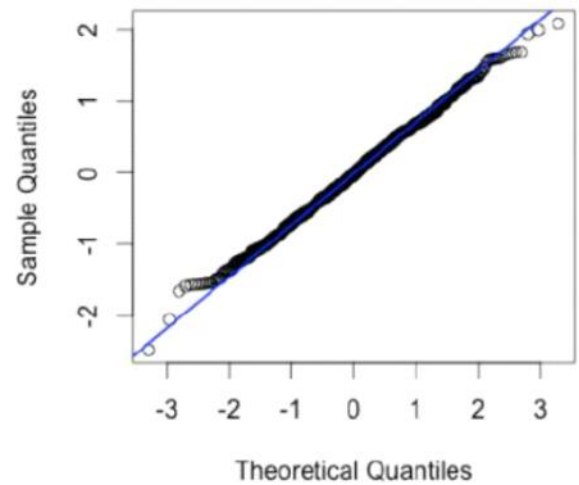
Normal Q-Q Plot



Thin Tails

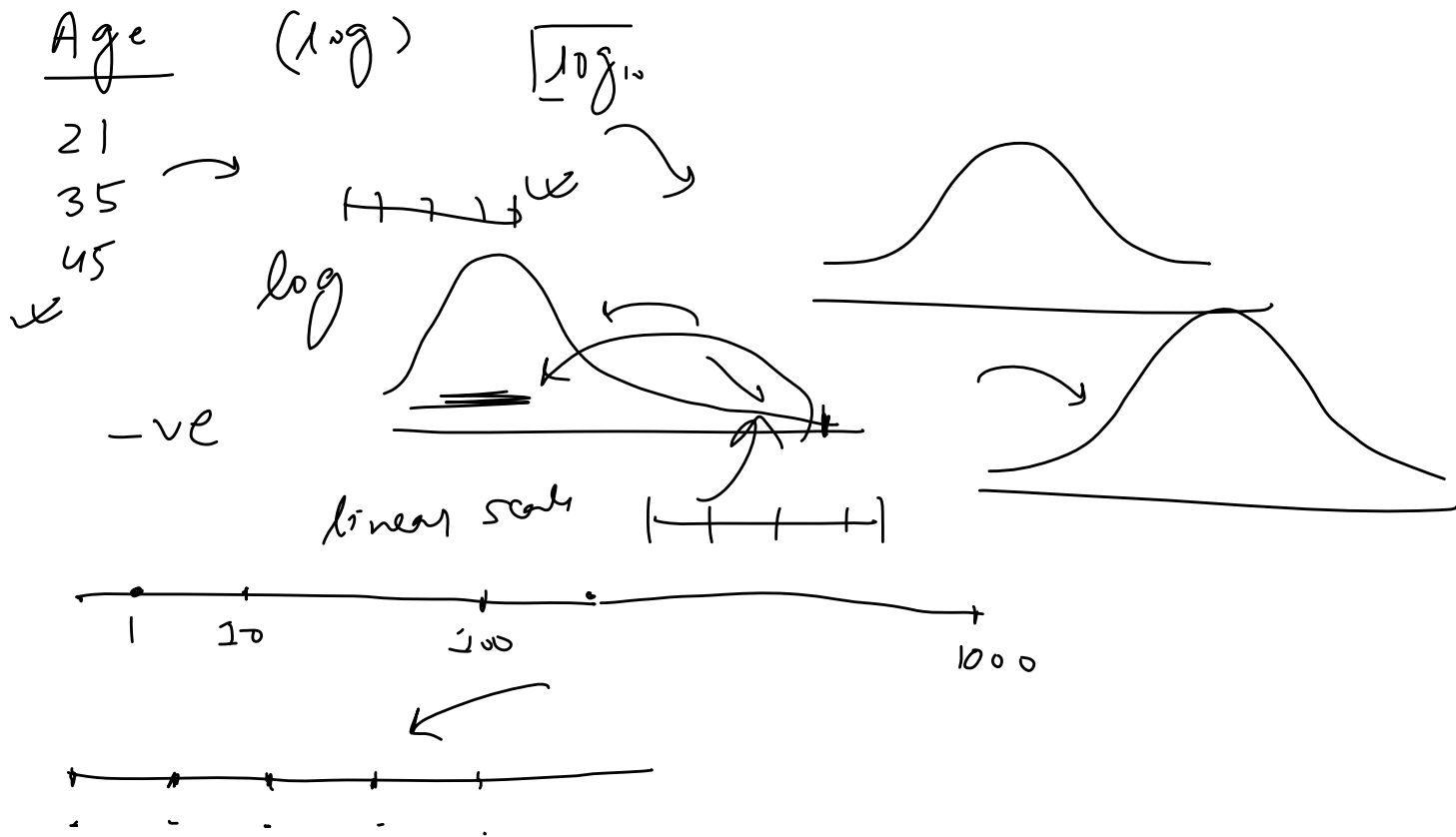


Normal Q-Q Plot



Log Transform

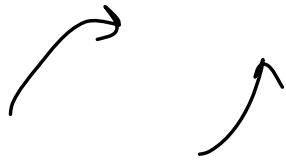
Friday, April 16, 2021 5:22 PM



Other Transforms

Friday, April 16, 2021 6:23 PM

Reciprocal $1/x$



Function
Transform

sq (x^2)



left skewed



sq + $\sqrt{x}$

sq + $\sqrt{x}$

Example

Friday, April 16, 2021 5:22 PM

Titanic



age | fare | survived

np.log

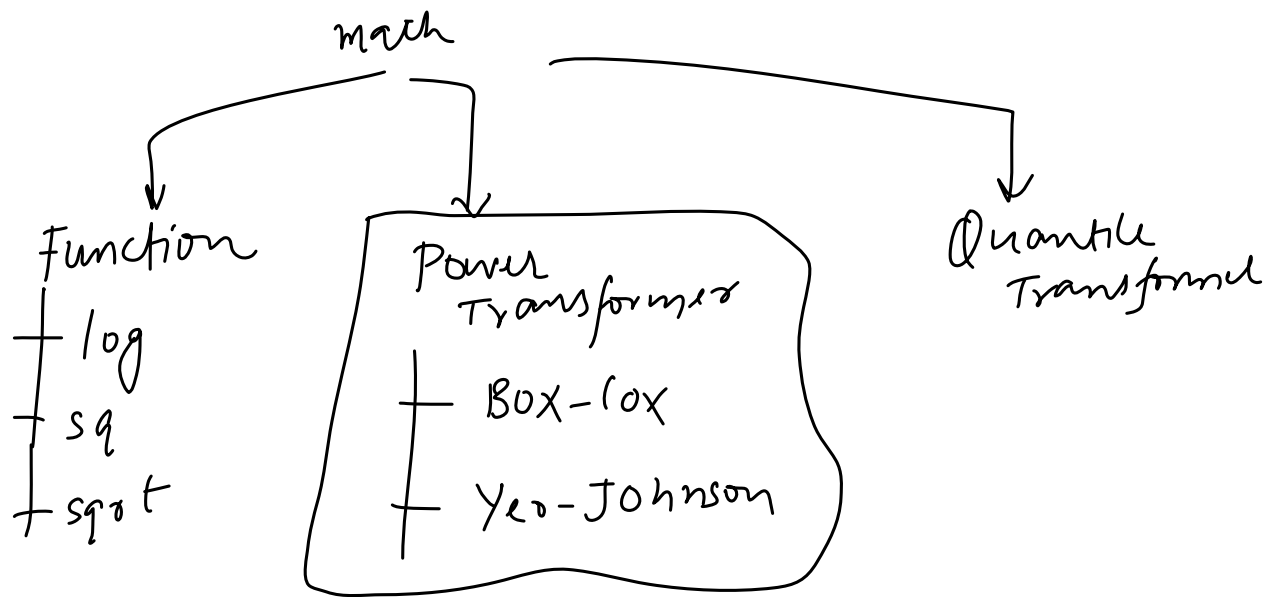
y_0

np.log1p

$(x+1)$

Power Transformer

Saturday, April 17, 2021 4:55 PM



Box Cox Transform

Saturday, April 17, 2021 4:56 PM

$$\underline{x_i^{(\lambda)}} = \begin{cases} \frac{x_i^{\lambda} - 1}{\lambda} & \text{if } \lambda \neq 0, \\ \ln(x_i) & \text{if } \lambda = 0, \end{cases}$$

↪ $\boxed{\eta > 0}$ $\textcircled{-ve}$

The exponent here is a variable called lambda (λ) that varies over the range of -5 to 5, and in the process of searching, we examine all values of λ . Finally, we choose the optimal value (resulting in the best approximation to a normal distribution) for your variable.

given dist.
↳ Normal dist

$x^2, x^3, x^{1.5}$

1.5, 1.1.

1) Max. likelihood
2) Bayesian

Yeo - Johnson Transform

Saturday, April 17, 2021 4:56 PM

$$x_i^{(\lambda)} = \begin{cases} [(x_i + 1)^\lambda - 1]/\lambda & \text{if } \lambda \neq 0, x_i \geq 0, \\ \ln(x_i) + 1 & \text{if } \lambda = 0, x_i \geq 0 \\ -[(-x_i + 1)^{2-\lambda} - 1]/(2 - \lambda) & \text{if } \lambda \neq 2, x_i < 0, \\ -\ln(-x_i + 1) & \text{if } \lambda = 2, x_i < 0 \end{cases}$$

This transformation is somewhat of an adjustment to the Box-Cox transformation, by which we can apply it to negative numbers.

Power Transformer

Example

Saturday, April 17, 2021 4:56 PM

concrete - mat
↳ Regr. → (LR) ↓

Cement

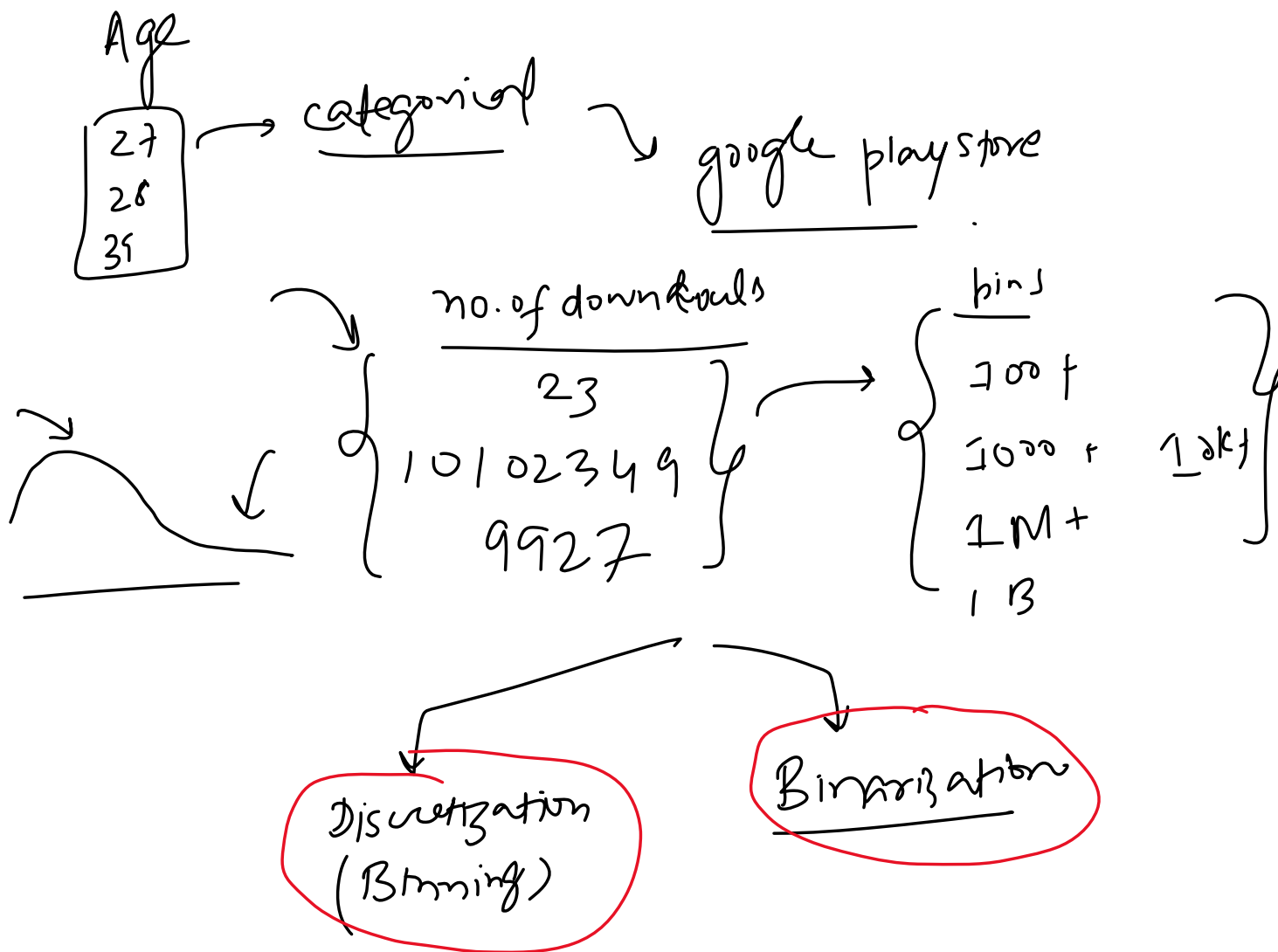
$$(540)^{0.177} \rightarrow$$

$$(30)^{(0.02)}$$

1. Encoding Numerical Features

Monday, April 19, 2021

4:18 PM



2. Discretization

Discretization is the process of transforming continuous variables into discrete variables by creating a set of contiguous intervals that span the range of the variable's values. Discretization is also called binning, where bin is an alternative name for interval.

Why use Discretization:

1. To handle Outliers
2. To improve the value spread

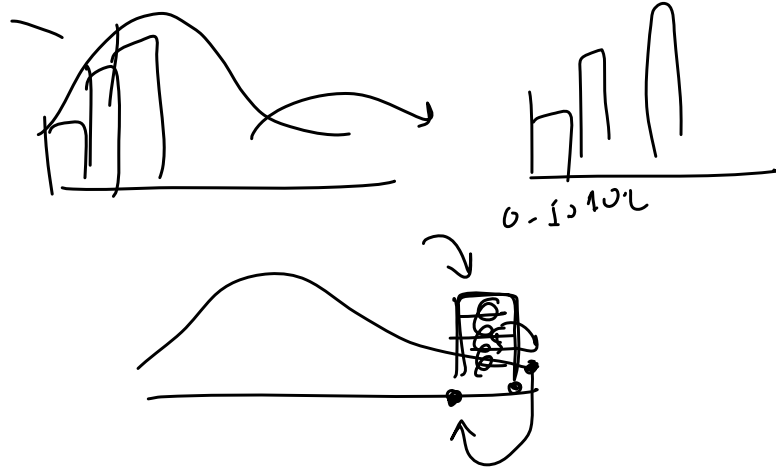
Age

23 42, 57 81 . . . 110

↓

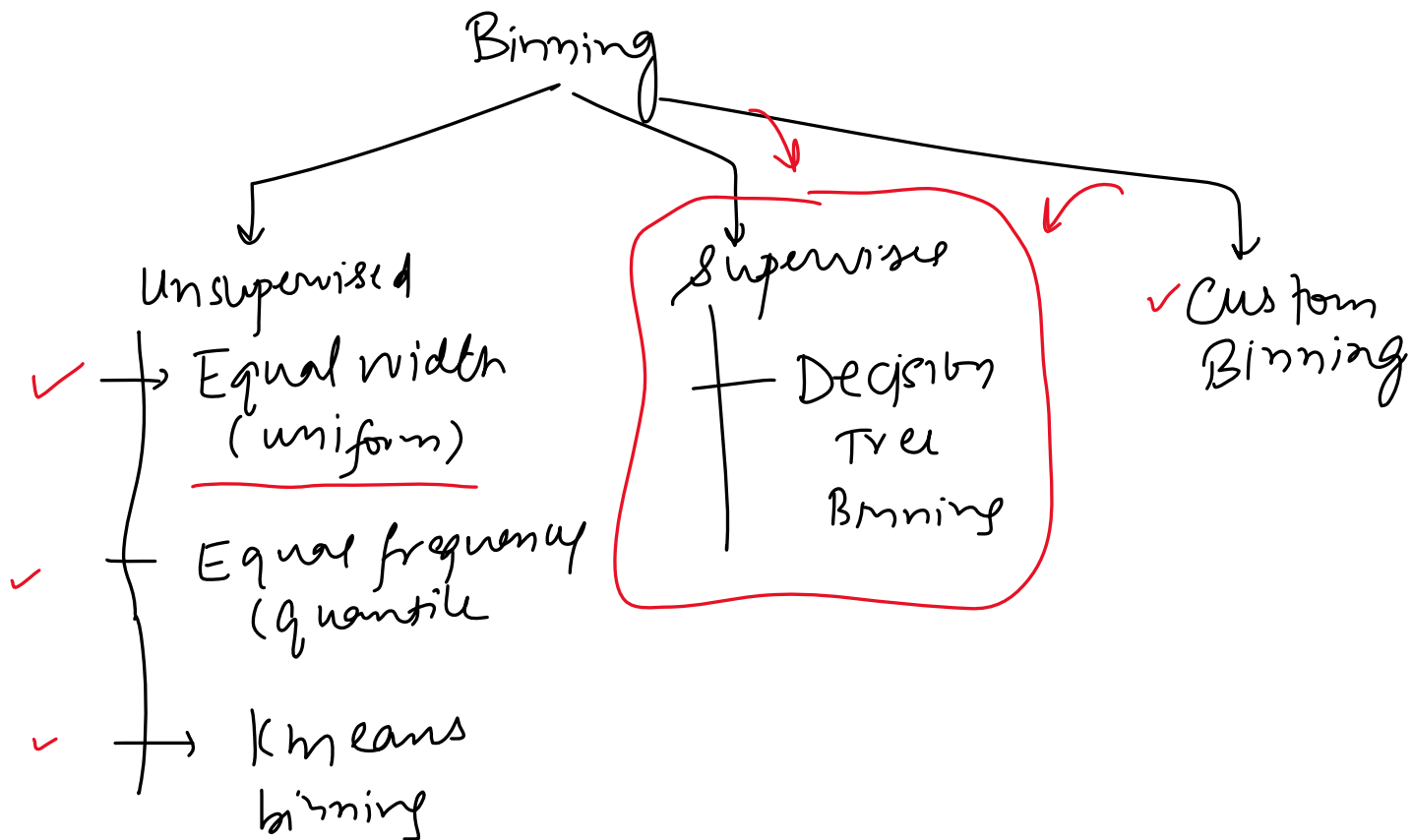
0-10, 10-20, 20-30 . . .

5 6 10



3. Types of Discretization

Monday, April 19, 2021 3:56 PM



4. Equal Width/Uniform Binning

Monday, April 19, 2021 3:56 PM

Age

- 1) Outliers
- 2) No change spread

(27), 32, 84, 56, ... max 100

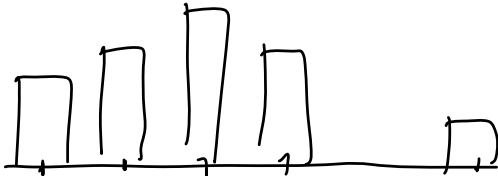
Bins = 10



min 0

$$\frac{\text{max} - \text{min}}{\text{bins}} = \frac{100 - 0}{10} = 10$$

$(0-10)$, $(10-20)$, $(20-30)$, ... $(90-100)$
 5 16 17 1 1 5

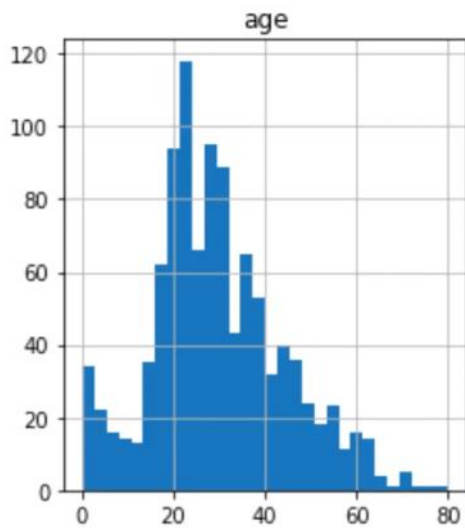


equal binning

	age	age_trf	age_labels
314	<u>43.0</u>	<u>5.0</u>	(<u>40.21</u> , <u>48.168</u>)
523	<u>44.0</u>	<u>5.0</u>	(<u>40.21</u> , <u>48.168</u>)
352	<u>15.0</u>	<u>1.0</u>	(<u>8.378</u> , <u>16.336</u>)
534	<u>30.0</u>	<u>3.0</u>	(<u>24.294</u> , <u>32.252</u>)
211	35.0	4.0	(32.252, 40.21]
530	2.0	0.0	(0.42, 8.378]
786	18.0	2.0	(16.336, 24.294]
827	1.0	0.0	(0.42, 8.378]
372	19.0	2.0	(16.336, 24.294]
518	36.0	4.0	(32.252, 40.21]

5. Equal Frequency/Quantile Binning

Monday, April 19, 2021 3:56 PM



Intervals = 10

Each interval contains 10% of total observations

Intervals: 0-16; 16-20; 20-22; 22-25; ...
50-74

- 1) Outliers
- 2) Value spread

0 - (?)
10th percent

0-16
10%

16-20
20%

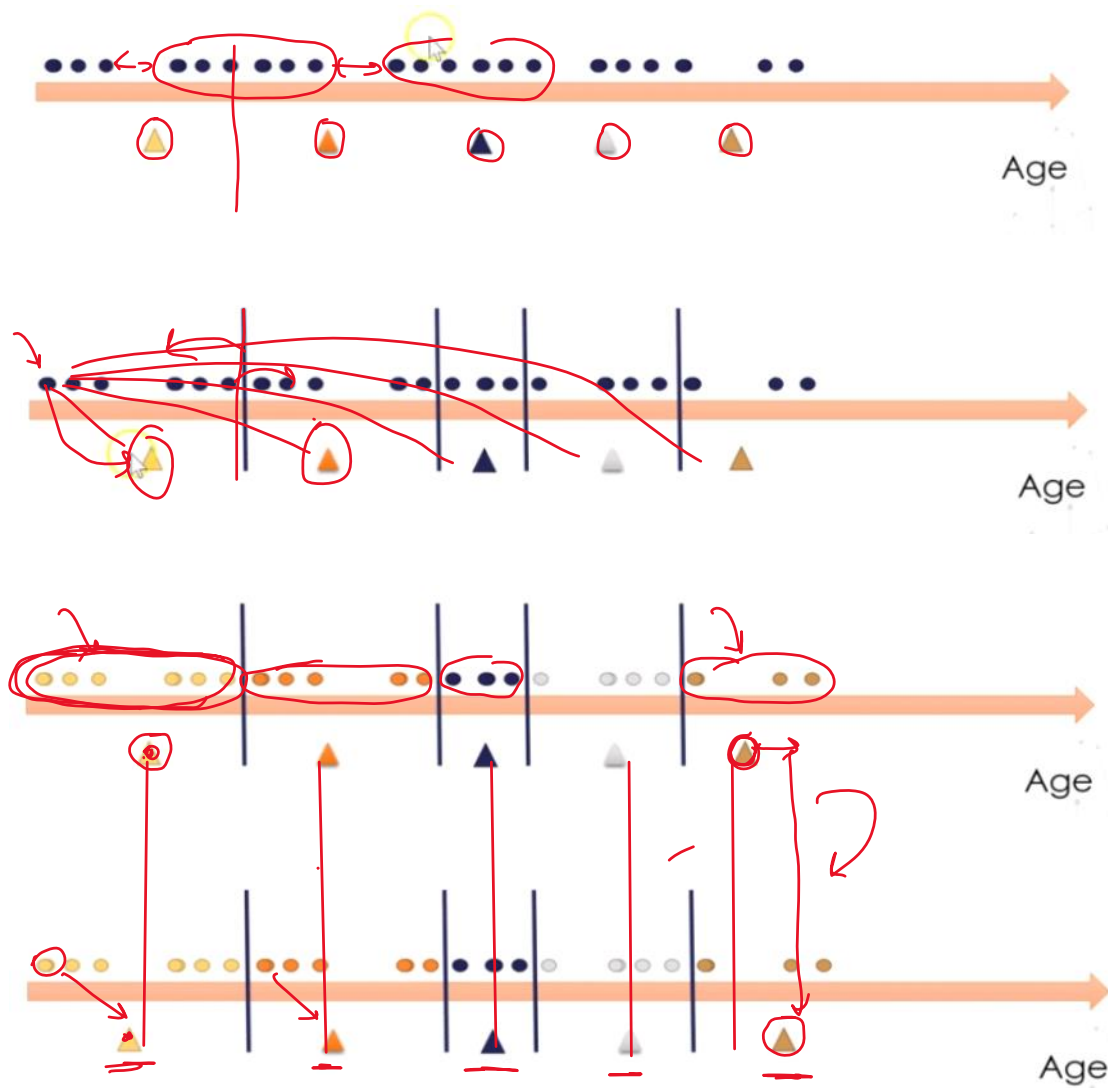
20-22
30%



6. KMeans Binning

Monday, April 19, 2021 3:57 PM

centroid
interval-5



K Means
↓
clustering
2D
40D

7. Encoding the discretized variable

Monday, April 19, 2021 3:58 PM

SKlearn

↳ KBinsDiscretizer()

bins = ?

strategy

+ uniform

+ quantile

+ kmeans

encoding

+ ordinal

+ onehotencoding

8. Example

Monday, April 19, 2021 3:58 PM

Titanic
↳

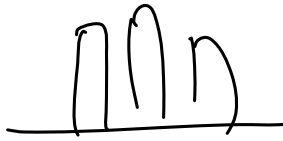
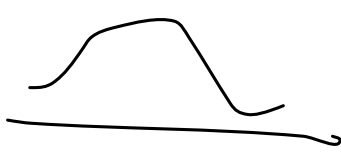
9. Custom/Domain Based Binning

Monday, April 19, 2021 3:58 PM

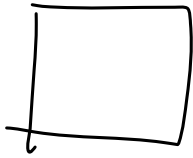
$\left\{ \begin{array}{l} [0-18] \rightarrow \text{Kids} \\ [18-60] \rightarrow \\ [60-80] \rightarrow \end{array} \right\}$ $\swarrow$ sklearn $\rightarrow$ Pandas

10. Binarization

Monday, April 19, 2021 4:17 PM



image



0, 1

0-255
color



0

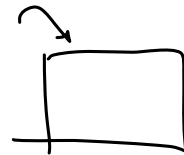
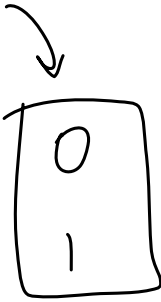
127.5
0 1

Annual income

62

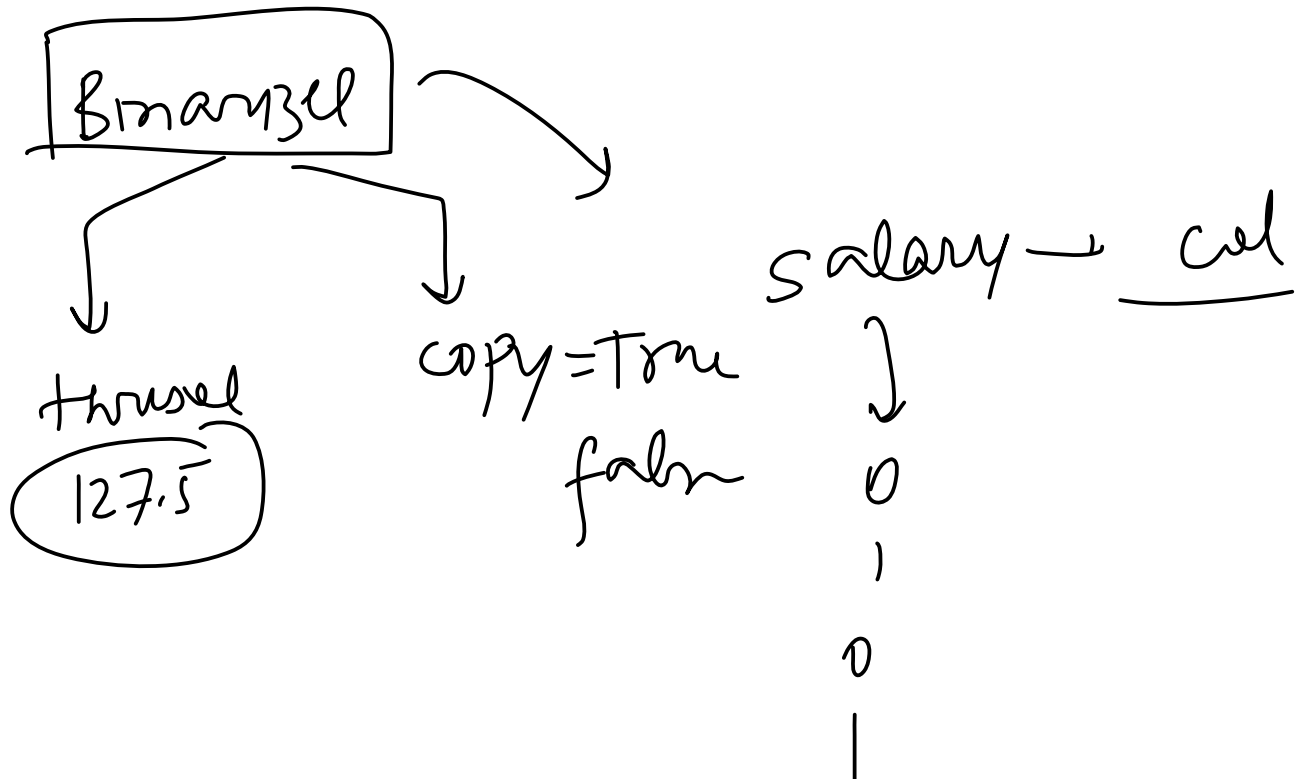
<

>



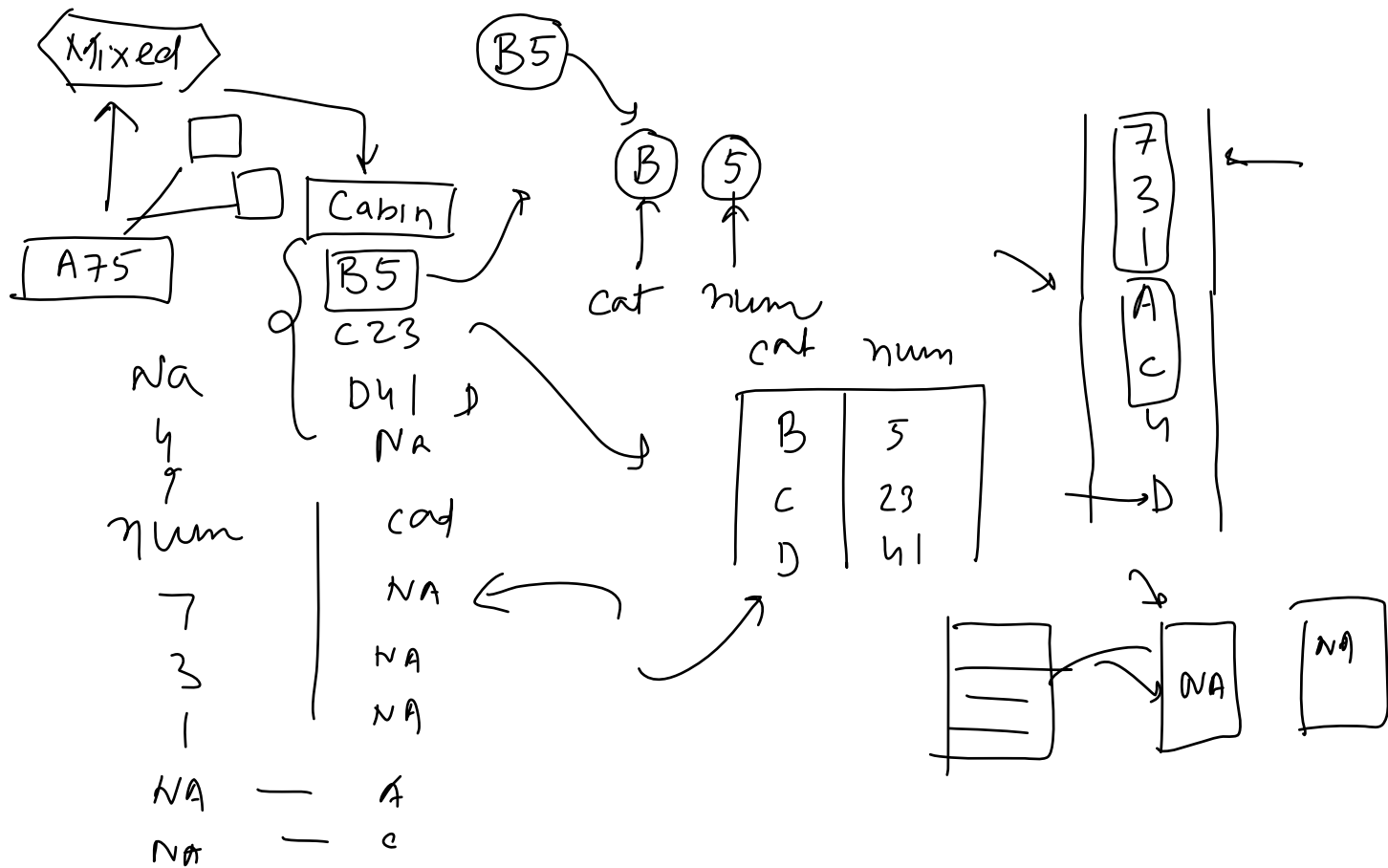
11. Example

Monday, April 19, 2021 4:17 PM



Mixed Data

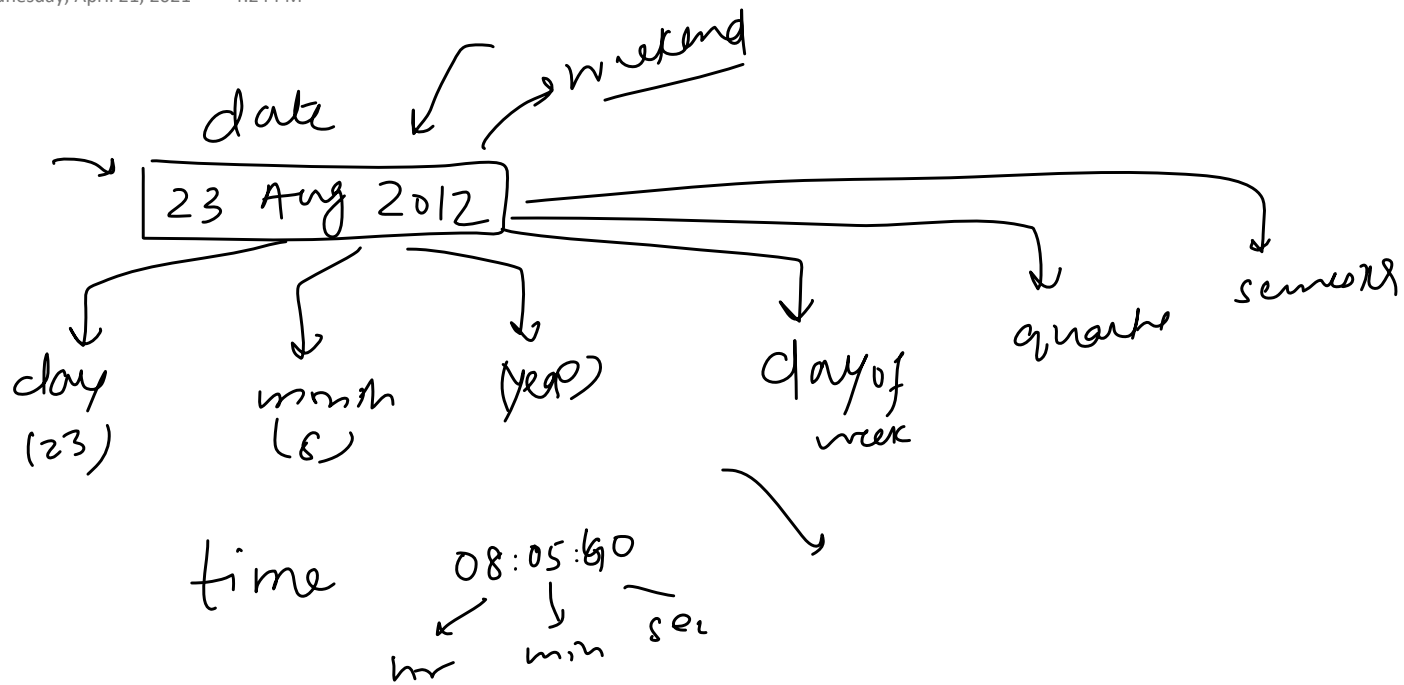
Tuesday, April 20, 2021 4:24 PM



Working with date and time

Wednesday, April 21, 2021

4:24 PM



Handling Missing Data

Thursday, April 22, 2021 1:03 PM

	1	2	3	4
1				
2				
3	-	-	-	-
4				
5				

(4)

CCA

Remove them

Missing Values

Impute

Simple Impute

univariate

Multivariate

KNN imputer

Iterative imputer
MICE

- ① → Remove values
- ② → Simple Imp
- ③ KNN imputer
- ④ Iterative
- ⑤ Missing indicator

num

- + mean median
- + Random
- + End of dis

cat

- + mode
- + Missing

Thursday, April 22, 2021 12:55 PM

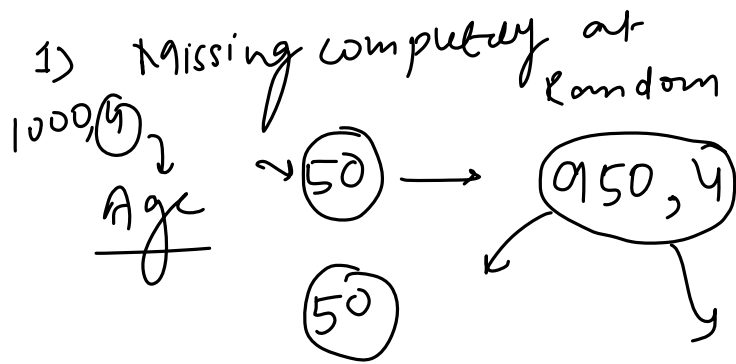
80w

اسی

all

Assumption For CCA

Thursday, April 22, 2021 12:58 PM



Advantage/Disadvantage

Thursday, April 22, 2021 12:59 PM

Advantage

1. Easy to implement as no data manipulation required
2. Preserves variable distribution (if data is MCAR, then the distribution of the variables of the reduced dataset should match the distribution in the original dataset)

Disadvantage

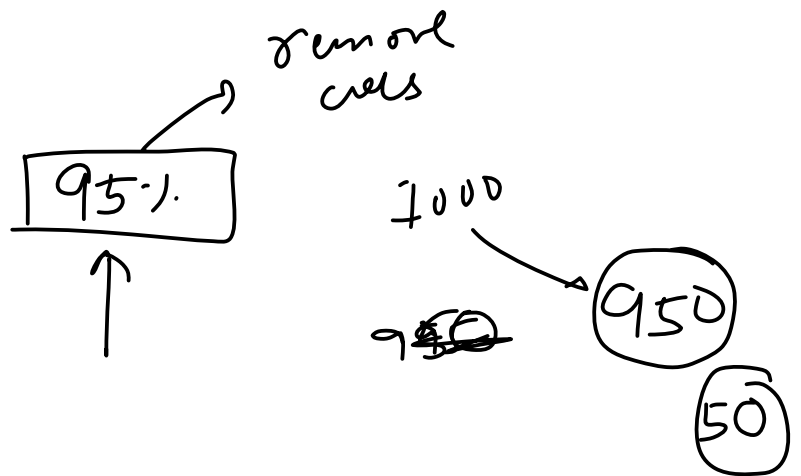
1. It can exclude a large fraction of the original dataset (if missing data is abundant)
2. Excluded observations could be informative for the analysis (if data is not missing at random)
3. When using our models in production, the model will not know how to handle missing data

When to use CCA?

Thursday, April 22, 2021 1:02 PM

1) MCAIR ✓

2) $\boxed{5\%} <$

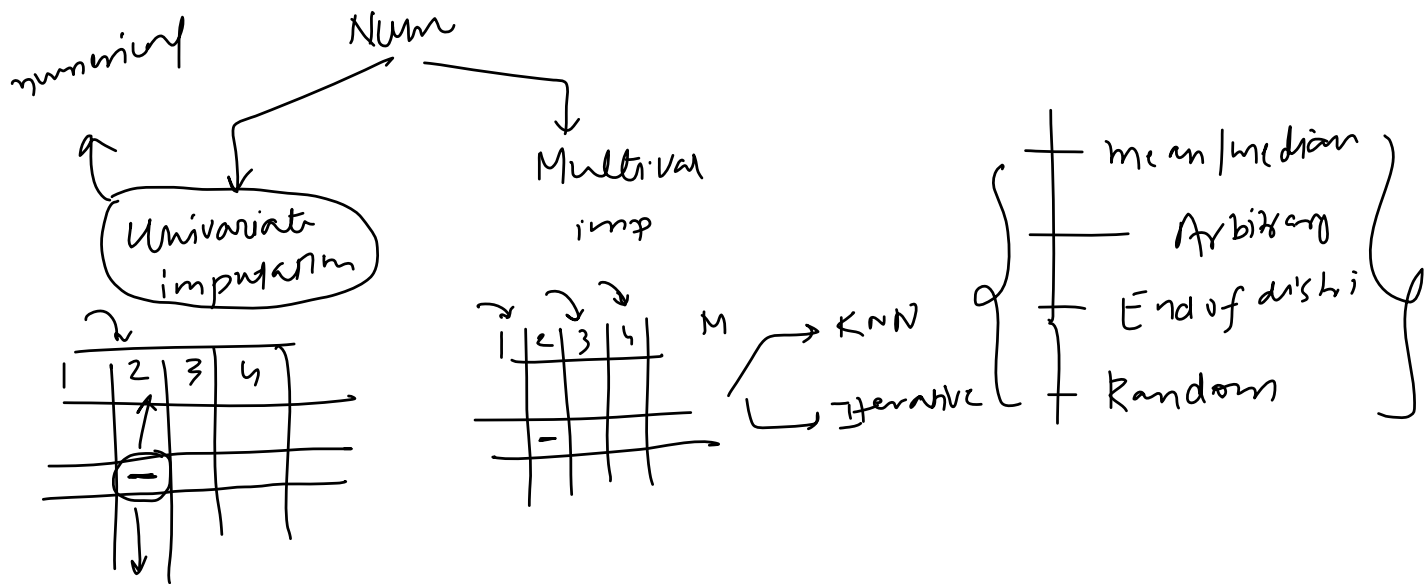


Example

Thursday, April 22, 2021 1:03 PM

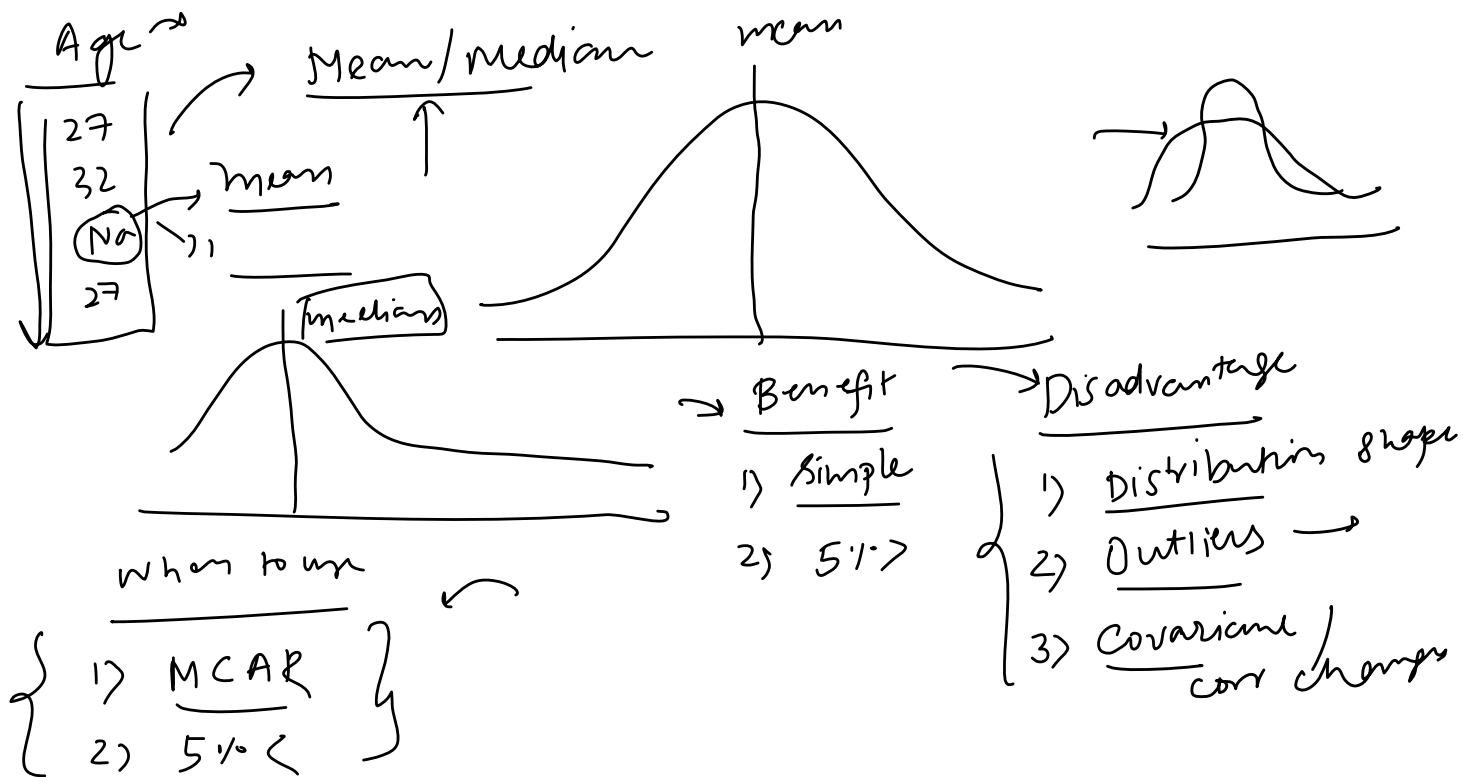
1. Handling Missing Numerical Data

Friday, April 23, 2021 11:28 AM



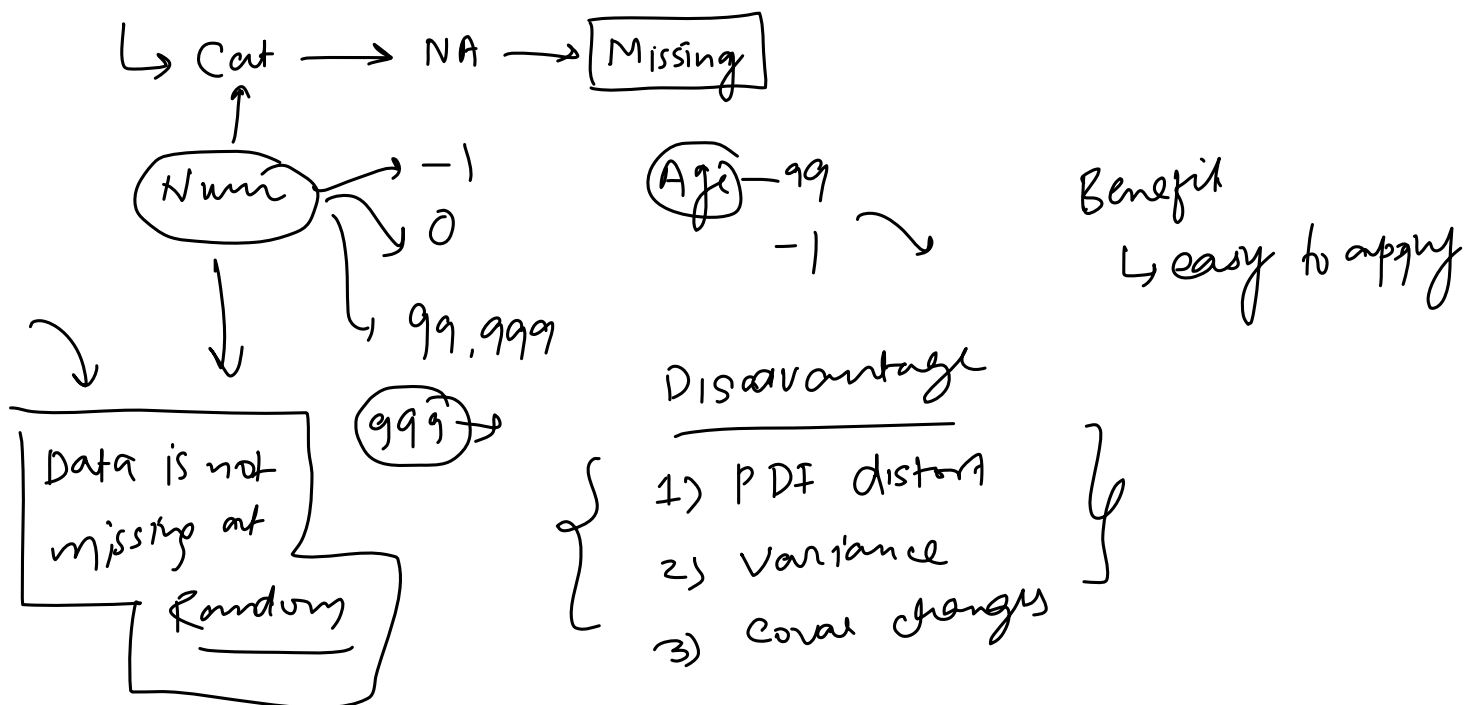
2. Mean/Median Imputation

Friday, April 23, 2021 11:31 AM



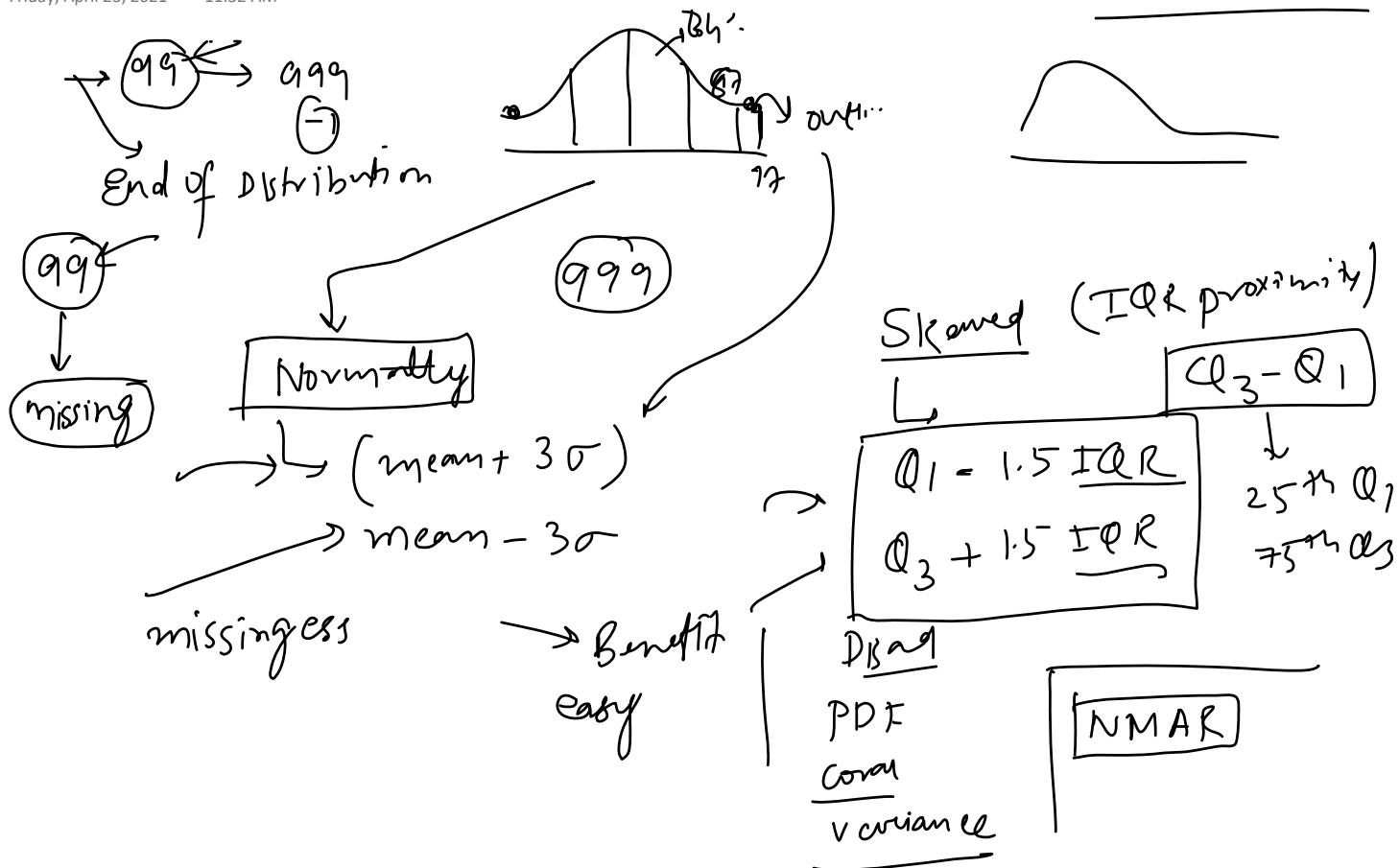
3. Arbitrary Value Imputation

Friday, April 23, 2021 11:32 AM



4. End of Distribution Imputation

Friday, April 23, 2021 11:32 AM



5. Random Sample Imputation

Friday, April 23, 2021

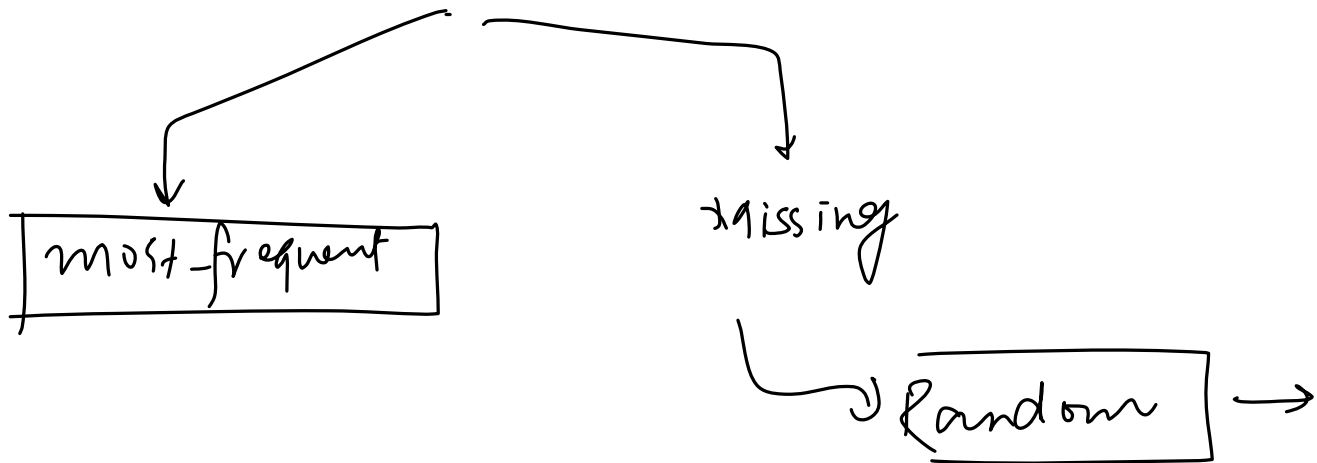
11:33 AM

6. Automatically select best imputation technique

Friday, April 23, 2021 11:32 AM

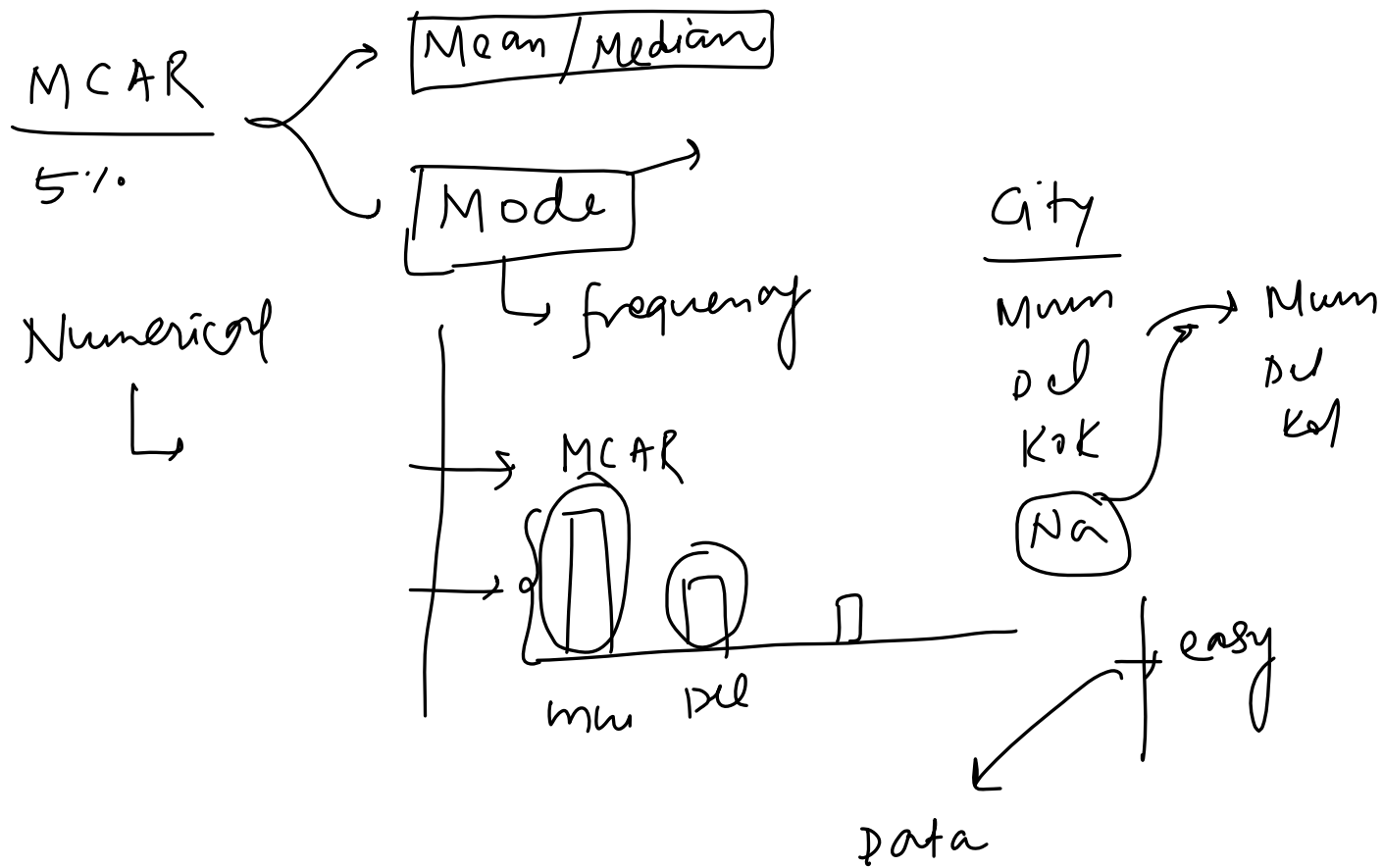
Handling Categorical Missing Data

Saturday, April 24, 2021 5:16 PM



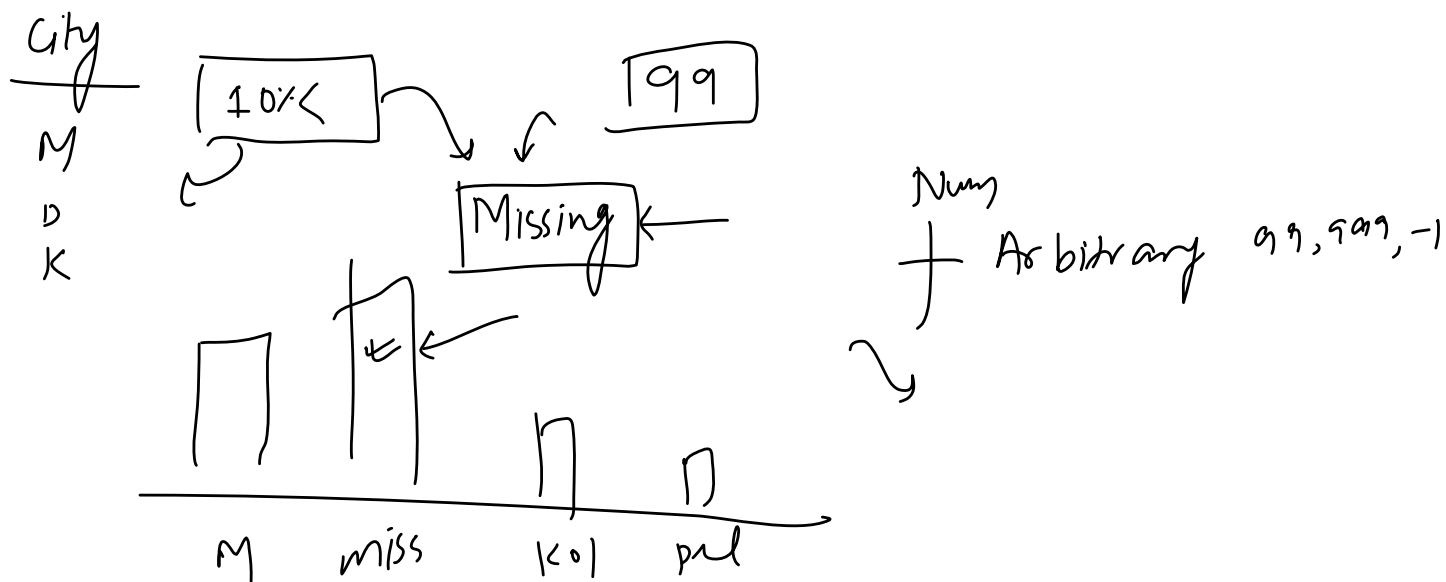
Most Frequent Value Imputation

Saturday, April 24, 2021 5:17 PM



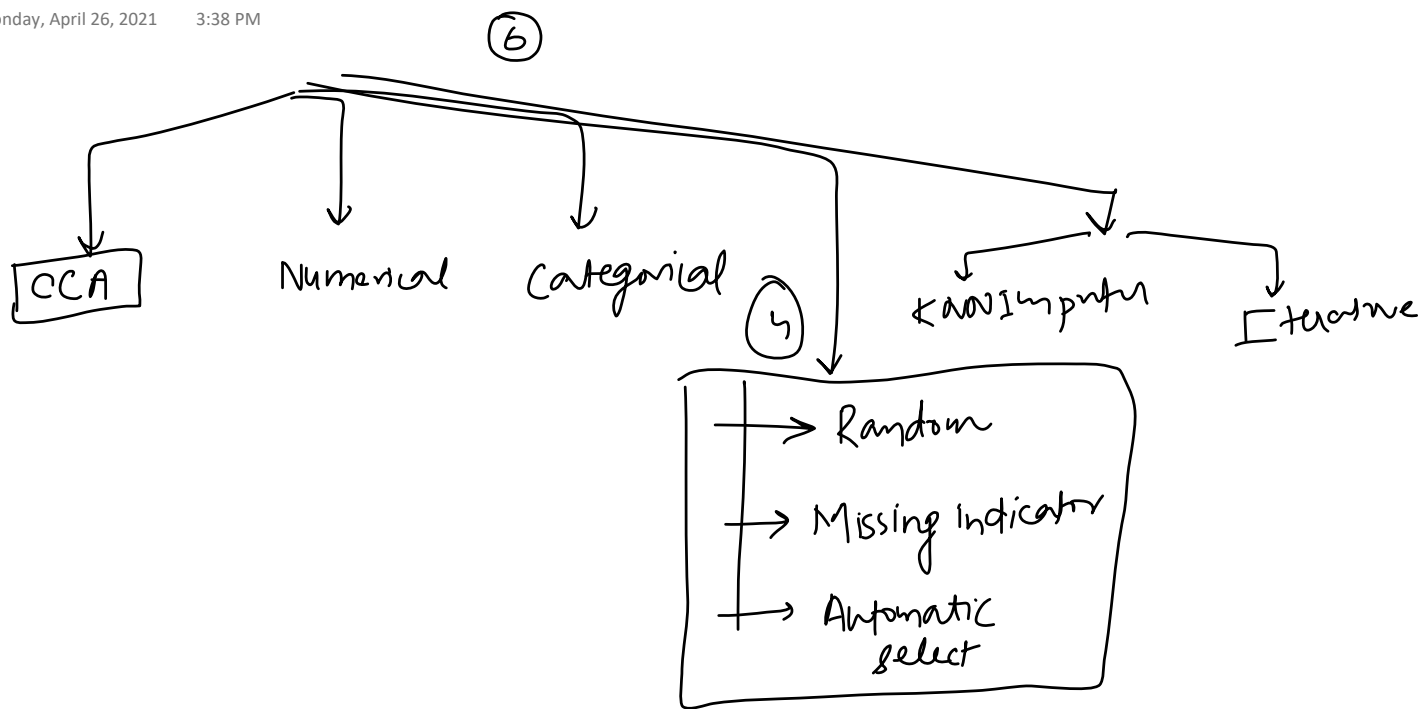
Missing Category Imputation

Saturday, April 24, 2021 5:17 PM



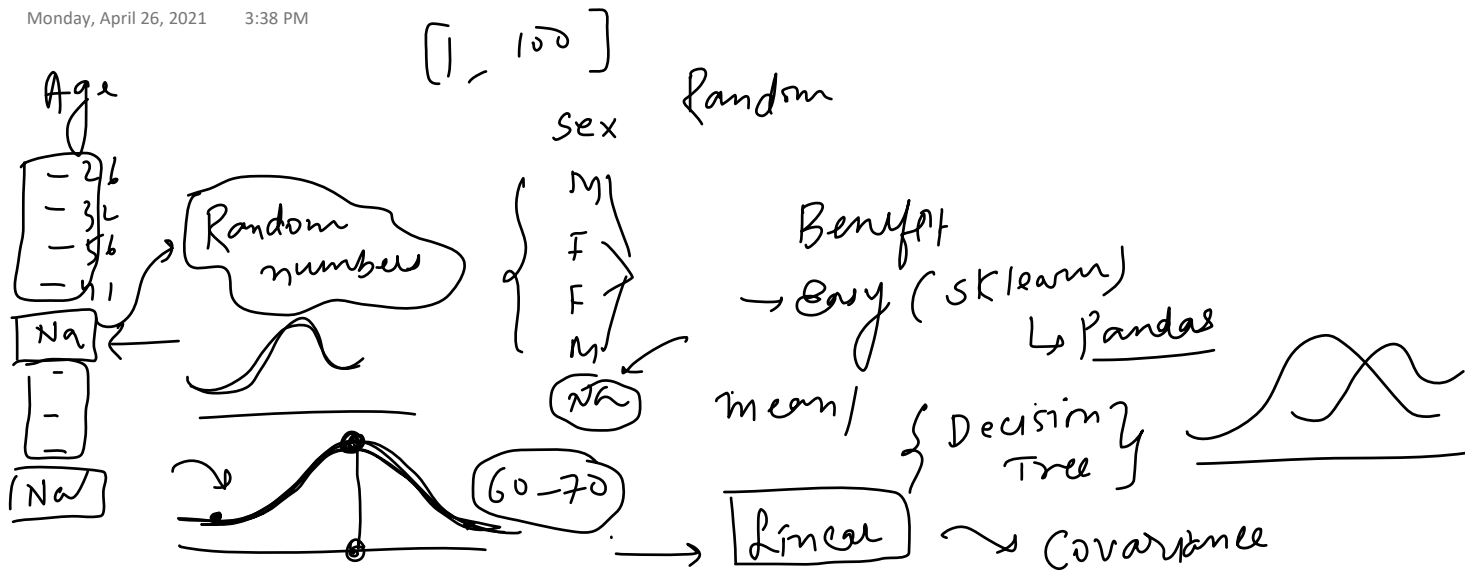
Topics to be covered

Monday, April 26, 2021 3:38 PM

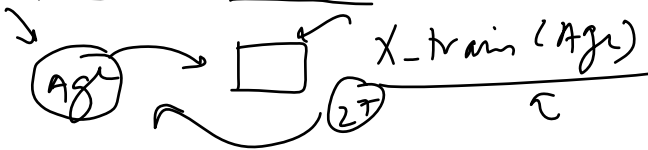


Random Imputation

Monday, April 26, 2021 3:38 PM



1. Preserves the variance of the variable (But Why?)
2. Memory heavy for deployment, as we need to store the original training set to extract values from and replace the NA in coming observations
3. Well suited for linear models as it does not distort the distribution, regardless of the % of NA



Missing Indicator

Monday, April 26, 2021 3:41 PM

Age	Surv	Age-Na
27	32	F
41	35	F
Na	41	T
62	32	F

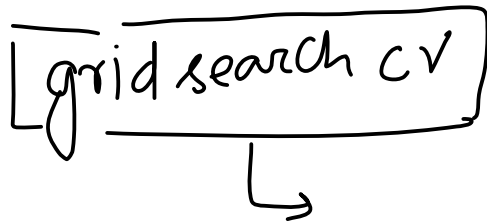
Annotations:

- Arrows point to the **Age** and **Surv** headers.
- An arrow points to the **Age-Na** header.
- A box labeled **KPD** has an arrow pointing to the **Age-Na** column.
- A box labeled **model** has an arrow pointing to the **Age-Na** column.

Automatically select value for Imputation

Monday, April 26, 2021 3:41 PM

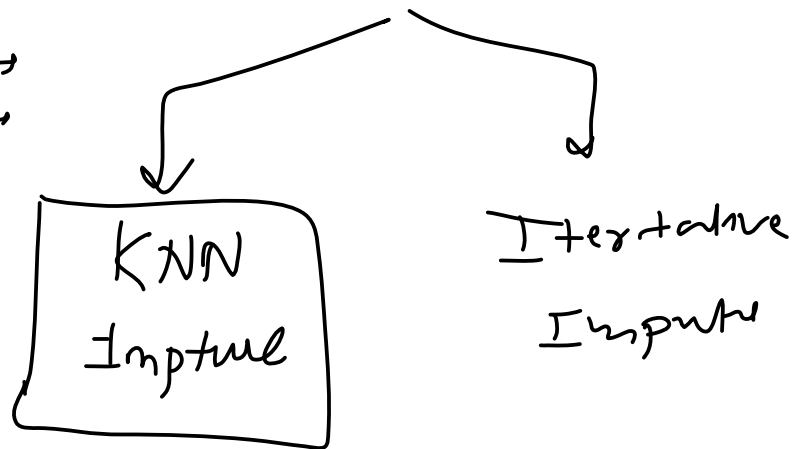
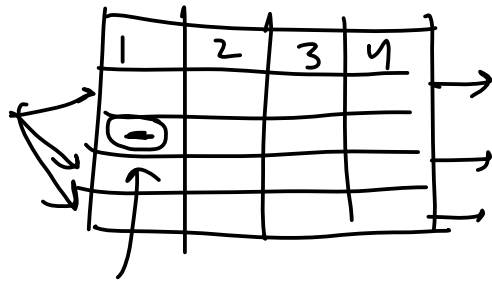
grid search cv

A handwritten diagram consisting of a rectangular box with the text "grid search cv" inside. A vertical line extends from the bottom center of the box, ending in a right-pointing arrowhead.

KNN Imputer

Tuesday, April 27, 2021 12:19 PM

mean



$K=2$

S NO	Feature 1	Feature 2	Feature 3	Feature 4
1	33	-----	67	21
2	31.5	45	68	12
3	23	51	71	18
4	40	-----	81	-----
5	35	60	79	-----

Distance from pt 1

$$d = \sqrt{\frac{3}{2}((68-33)^2 + (12-21)^2)}$$

$$= \sqrt{1.5(4+81)} = \sqrt{127.5} = 11.29$$

Distance from pt 3

$$d = \sqrt{\frac{3}{2}((51-45)^2 + (71-68)^2 + (18-12)^2)}$$

$$= \sqrt{36 + 9 + 36} = 9$$

Distance from pt 4

$$d = \sqrt{\frac{3}{1}((81-68)^2)}$$

$$d = \sqrt{3(9)} = \sqrt{27} = 5.19$$

Distance from pt 5

$$d = \sqrt{\frac{3}{2}((60-45)^2 + (79-68)^2)}$$

$$= \sqrt{1.5(25+121)} = \sqrt{219} = 14.79$$

Advantage & Disadvantage

- 1) More accurate ←
 - 2) More no. of calculations →
 - 3) Prediction
-

$K=2$

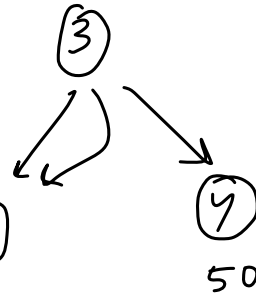
1 →	70				
2 →					
3 →	19				
4 →	50				

$$\frac{70 + 50}{2} \rightarrow$$

60

→ 10

$\frac{1}{dis}$



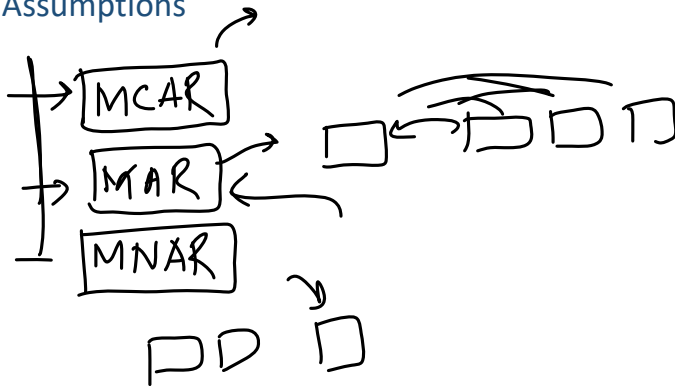
$$\frac{1}{10} \times \cancel{70} + \frac{1}{50} \times \cancel{50}$$

$$\frac{7 + 1}{2} = 4$$

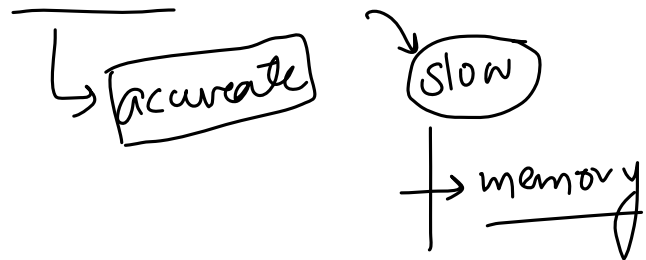
distance

MICE stands for **Multivariate Imputation by Chained Equations**

Assumptions



Advantages & Disadvantage



How it works?

Wednesday, April 28, 2021 11:41 AM

1. Actual Dataset

	R&D Spend	Administration	Marketing Spend	Profit
21	8.0	15.0	30.0	11.0
37	4.0	5.0	20.0	9.0
2	15.0	10.0	41.0	19.0
14	12.0	16.0	26.0	13.0
44	2.0	15.0	3.0	7.0

2. Removing the target column

	R&D Spend	Administration	Marketing Spend
21	8.0	15.0	30.0
37	4.0	5.0	20.0
2	15.0	10.0	41.0
14	12.0	16.0	26.0
44	2.0	15.0	3.0

3. Introduced some fake nan values

	R&D Spend	Administration	Marketing Spend
21	8.0	15.0	30.0
37	NaN	5.0	20.0
2	15.0	10.0	41.0
14	12.0	NaN	26.0
44	2.0	15.0	NaN

Step 1 - Fill all the NaN values with mean of respective cols

	R&D Spend	Administration	Marketing Spend
21	8.00	15.00	30.00
37	9.25	5.00	20.00
2	15.00	10.00	41.00
14	12.00	11.25	26.00
44	2.00	15.00	29.25

Step 2 - Remove all col1 missing values

Step 3 - Predict the missing values of col1 using other cols

	R&D Spend	Administration	Marketing Spend		Administration	Marketing Spend		R&D Spend	Administration	Marketing Spend
21	8.0	15.00	30.00	21	15.00	30.00	21	8.00	15.00	30.00
37	NaN	5.00	20.00	37			37	23.14	5.00	20.00
2	15.0	10.00	41.00	2	10.00	41.00	2	15.00	10.00	41.00
14	12.0	11.25	26.00	14	11.25	26.00	14	12.00	11.25	26.00
44	2.0	15.00	29.25	44	15.00	29.25	44	2.00	15.00	29.25

Step 4 - Remove all col2 missing values

Step 5 - Predict the missing values of col2 using other cols

	R&D Spend	Administration	Marketing Spend		R&D Spend	Marketing Spend		R&D Spend	Administration	Marketing Spend
21	8.00	15.0	30.00	21	8.00	30.00	21	8.00	15.00	30.00
37	23.14	5.0	20.00	37	23.14	20.00	37	23.14	5.00	20.00
2	15.00	10.0	41.00	2	15.00	41.00	2	15.00	10.00	41.00
14	12.00	NaN	26.00	14	2.00	29.25	14	12.00	11.06	26.00
44	2.00	15.0	29.25	44			44	2.00	15.00	29.25

Step 6 - Remove all col3 values

	R&D Spend	Administration	Marketing Spend		R&D Spend	Administration		R&D Spend	Administration	Marketing Spend
21	8.00	15.00	30.0	21	8.00	15.00	21	8.00	15.00	30.00
37	23.14	5.00	20.0	37	23.14	5.00	37	23.14	5.00	20.00
2	15.00	10.00	41.0	2	15.00	10.00	2	15.00	10.00	41.00
14	12.00	11.06	26.0	14	12.00	11.06	14	12.00	11.06	26.00
44	2.00	15.00	NaN	44			44	2.00	15.00	31.56

Iteration 0

Iteration 1

Difference

	R&D Spend	Administration	Marketing Spend		R&D Spend	Administration	Marketing Spend		R&D Spend	Administration	Marketing Spend
21	8.00	15.00	30.00	21	8.00	15.00	30.00	21	0.00	0.00	0.00
37	9.25	5.00	20.00	37		5.00	20.00	37	13.89	0.00	0.00
2	15.00	10.00	41.00	2	15.00	10.00	41.00	2	0.00	0.00	0.00
14	12.00	11.25	26.00	14	12.00		26.00	14	0.00	-0.19	0.00
44	2.00	15.00	29.25	44	2.00	15.00	31.56	44	0.00	0.00	2.31

Iteration 1

Iteration 2

Difference

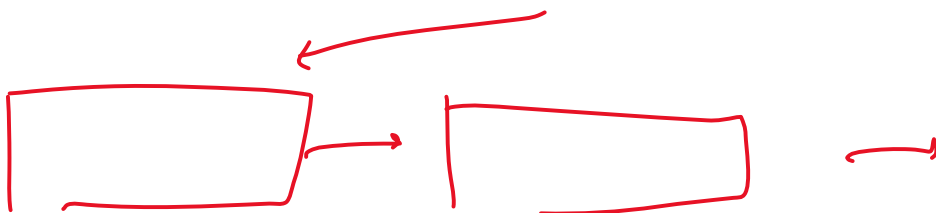
	R&D Spend	Administration	Marketing Spend		R&D Spend	Administration	Marketing Spend		R&D Spend	Administration	Marketing Spend
21	8.00	15.00	30.00	21	8.00	15.00	30.00	21	0.00	0.00	0.0
37	23.14	5.00	20.00	37	23.78	5.00	20.00	37	0.64	0.00	0.0
2	15.00	10.00	41.00	2	15.00	10.00	41.00	2	0.00	0.00	0.0
14	12.00	11.06	26.00	14	12.00	11.22	26.00	14	0.00	0.16	0.0
44	2.00	15.00	31.56	44	2.00	15.00	31.56	44	0.00	0.00	0.0

Iteration 2

Iteration 3

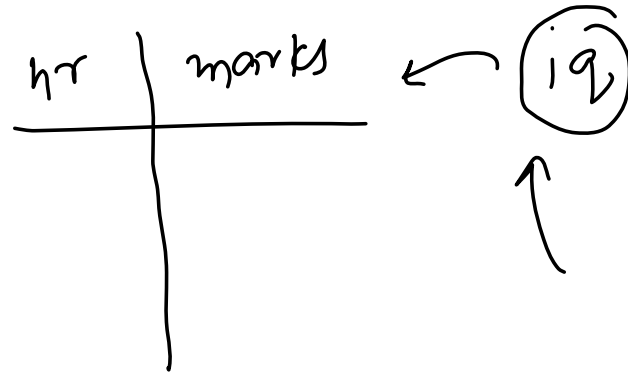
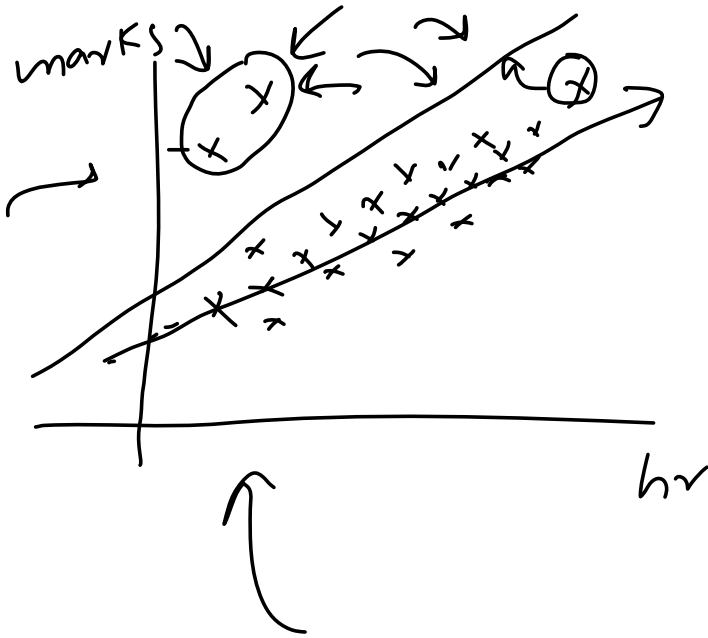
Difference

R&D Spend Administration Marketing Spend				R&D Spend Administration Marketing Spend				R&D Spend Administration Marketing Spend			
21	8.00	15.00	30.00	21	8.00	15.00	30.00	21	0.00	0.00	0.00
37	23.78	5.00	20.00	37	24.57	5.00	20.00	37	0.79	0.00	0.00
2	15.00	10.00	41.00	2	15.00	10.00	41.00	2	0.00	0.00	0.00
14	12.00	11.22	26.00	14	12.00	11.37	26.00	14	0.00	0.15	0.00
44	2.00	15.00	31.56	44	2.00	15.00	45.53	44	0.00	0.00	13.97



What are Outliers?

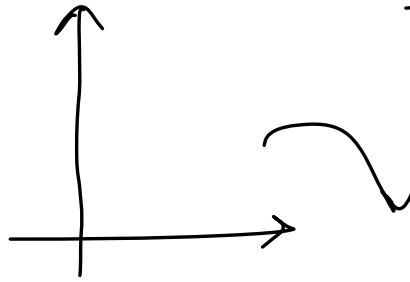
Thursday, April 29, 2021 3:25 PM



When is outlier dangerous?

Thursday, April 29, 2021 3:26 PM

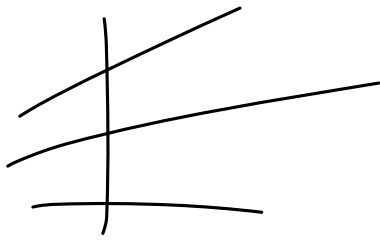
Age
| 300 |



Anomaly
detection → outliers

Effect of Outliers on ML algorithms

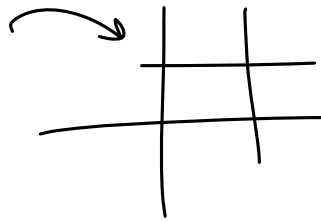
Thursday, April 29, 2021 3:27 PM



Linear regression
Logistic regression
weights

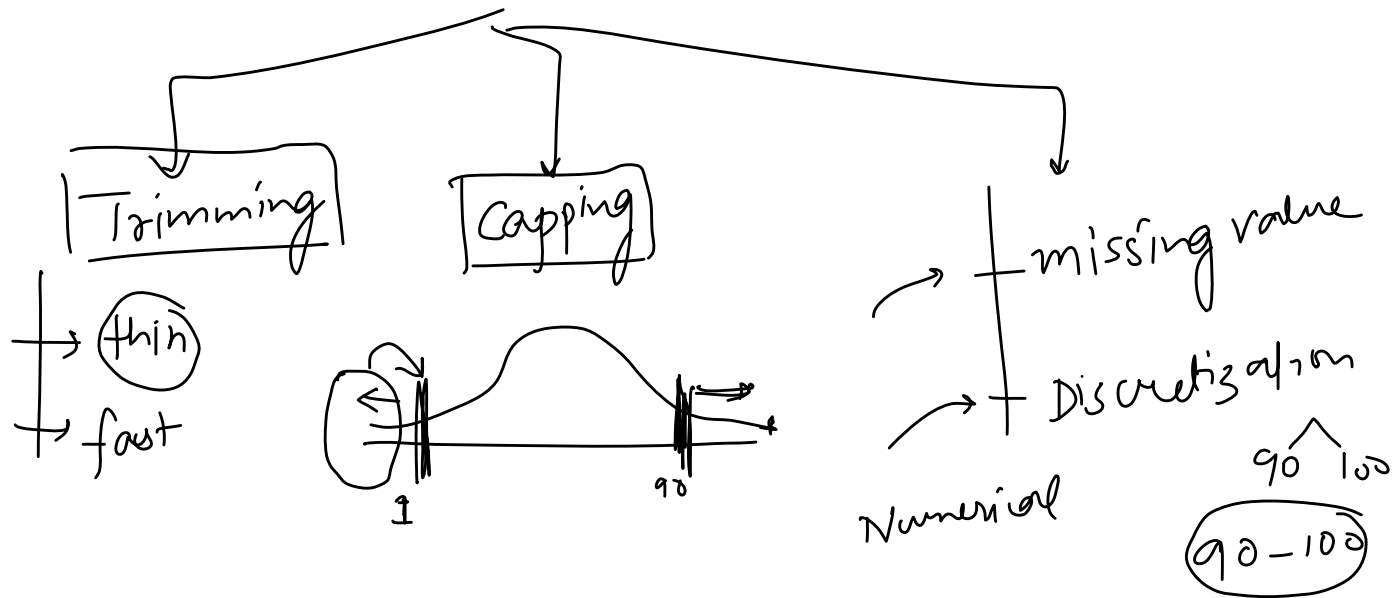
Adaboost ✓
Deep Learning ✓

Tree plans



How to treat Outliers?

Thursday, April 29, 2021 3:27 PM



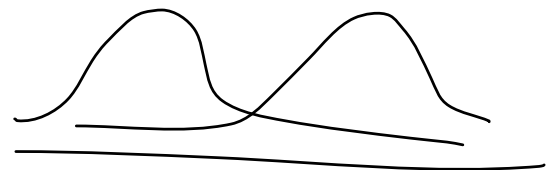
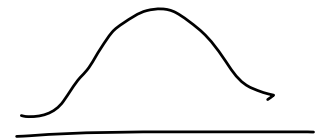
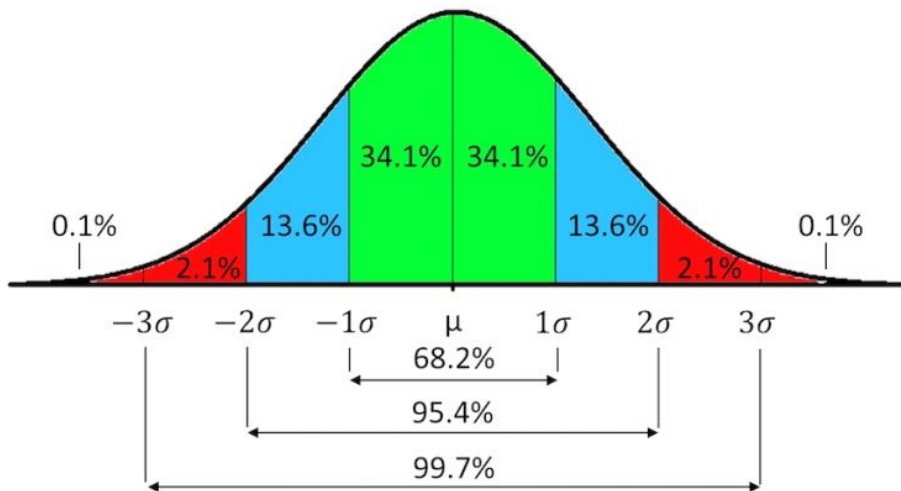
How to detect Outliers?

Thursday, April 29, 2021 3:27 PM

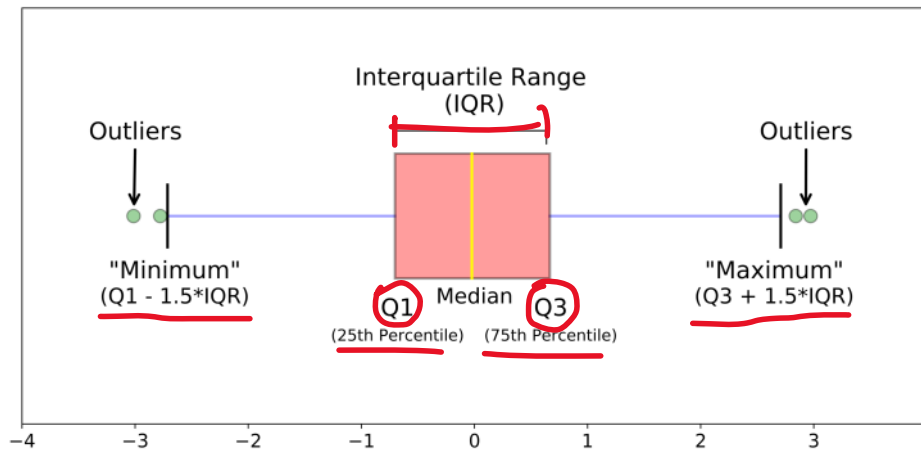
$$\begin{aligned} (\mu + 3\sigma) &> \\ (\mu - 3\sigma) &< \end{aligned}$$

No

1. Normal Distribution

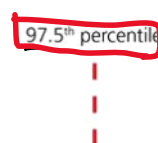


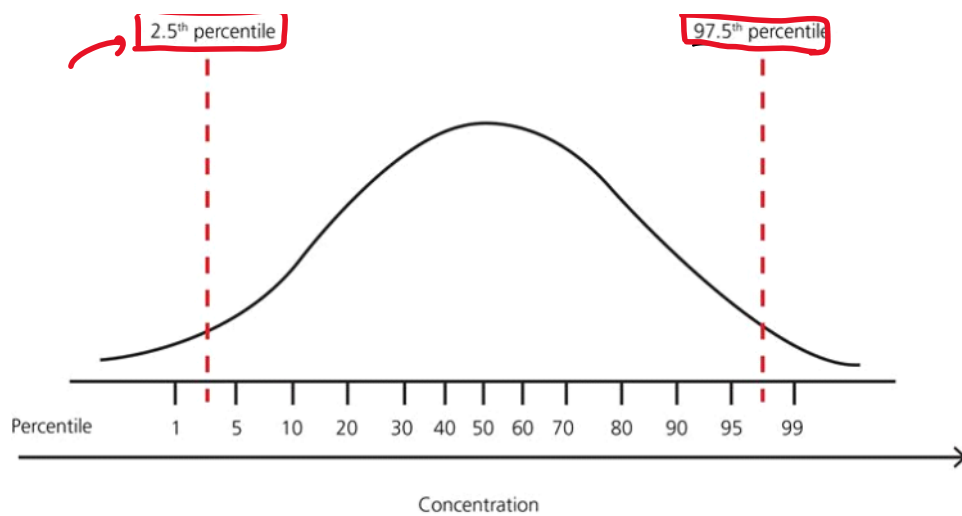
2. Skewed Distribution



$$\begin{aligned} \text{min} &= Q1 - 1.5 IQR \\ Q3 &+ 1.5 IQR \end{aligned}$$

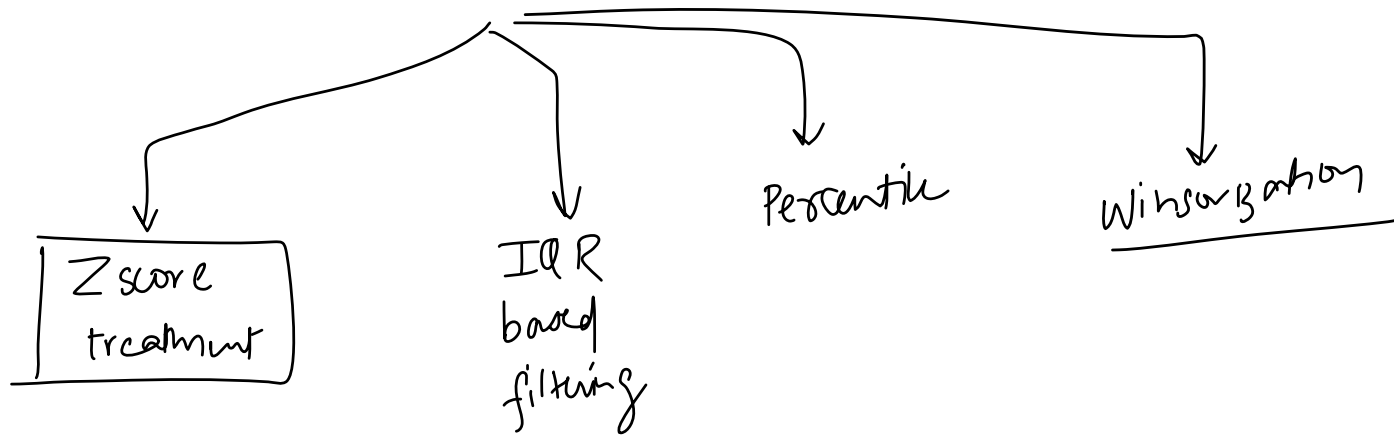
3. Other Distributions





Techniques for Outlier Detection and Removal

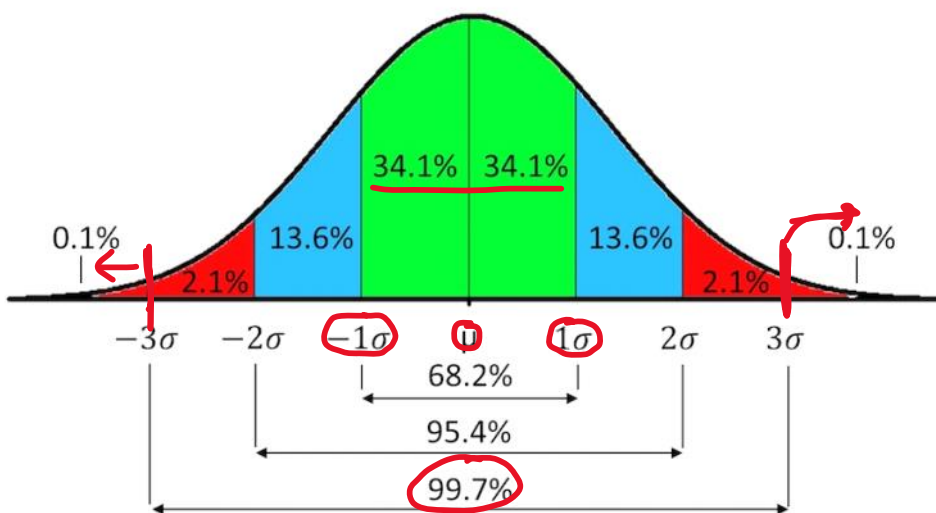
Thursday, April 29, 2021 3:35 PM



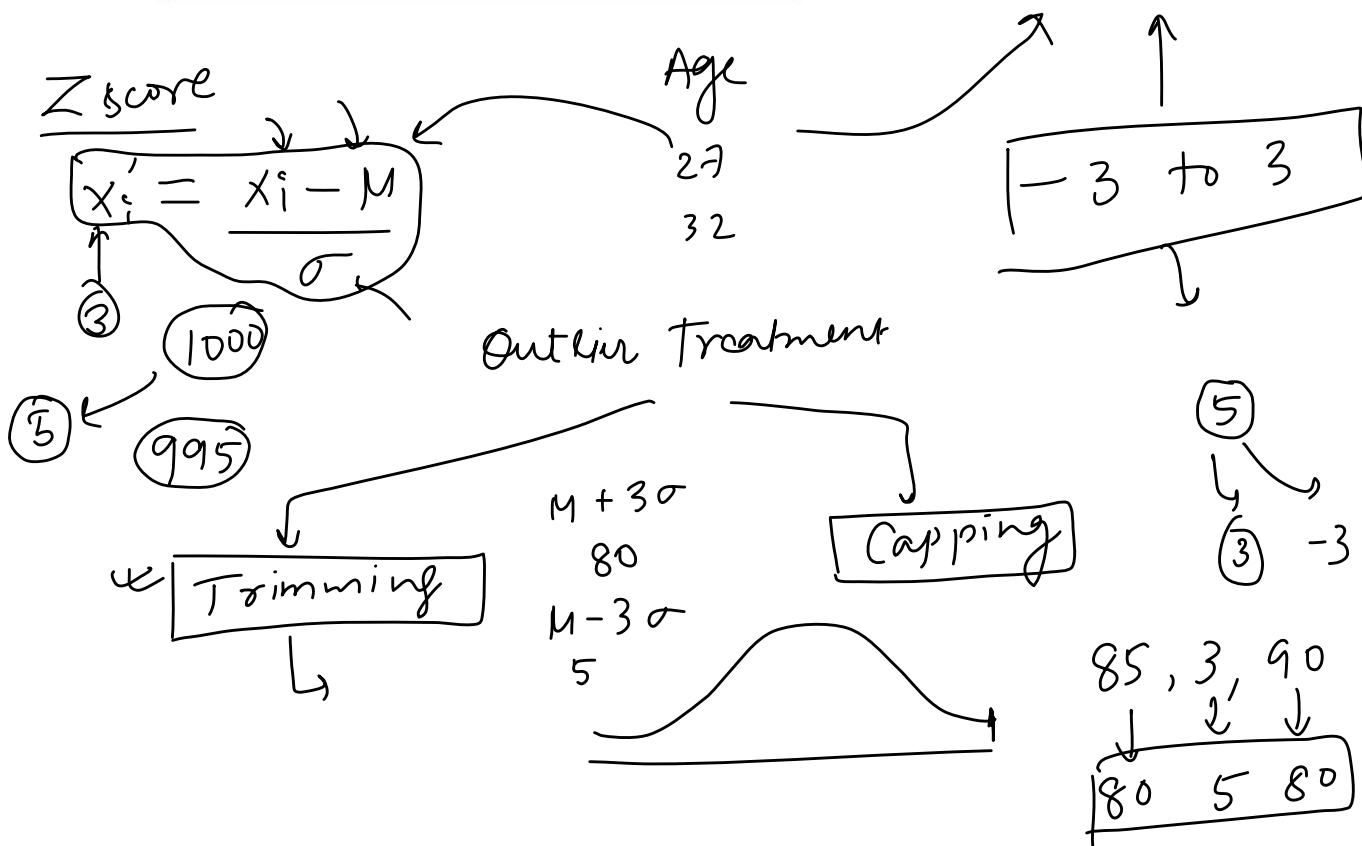
1

Outlier removal using Z score

Friday, April 30, 2021 1:33 PM



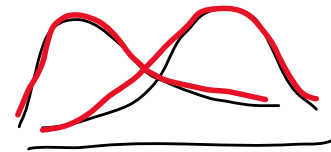
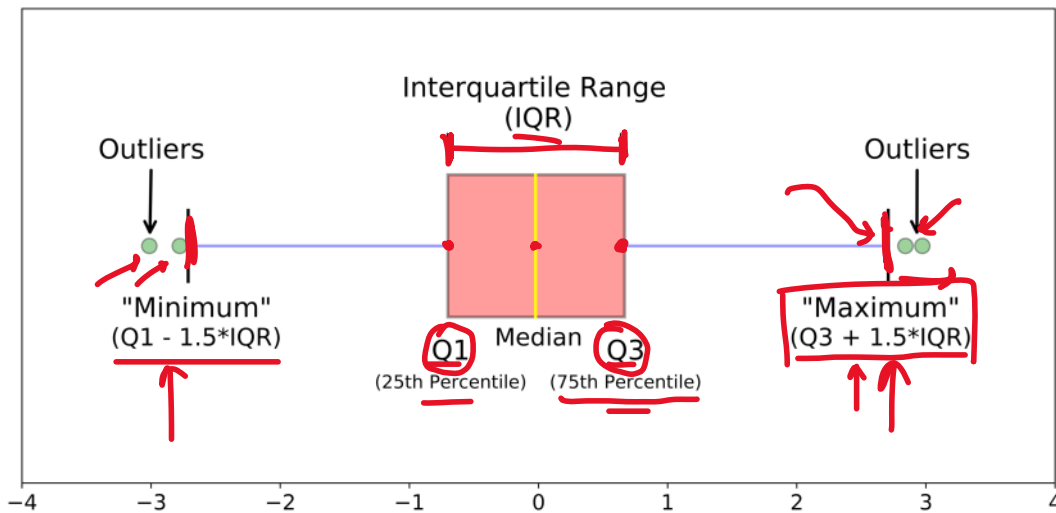
~~$\mu + 1\sigma$~~
 $\left\{ \begin{matrix} \mu + \sigma \\ \mu - \sigma \end{matrix} \right\}$ 68%
 $\left\{ \begin{matrix} \mu + 2\sigma \\ \mu - 2\sigma \end{matrix} \right\}$ 95%



Outlier Detection using IQR

Friday, April 30, 2021 3:48 PM

[Boxplot IQR]



percentiles

25
50 Median
75
100

Age

72
89
15
16

100% → 89 → max

0% → min 16

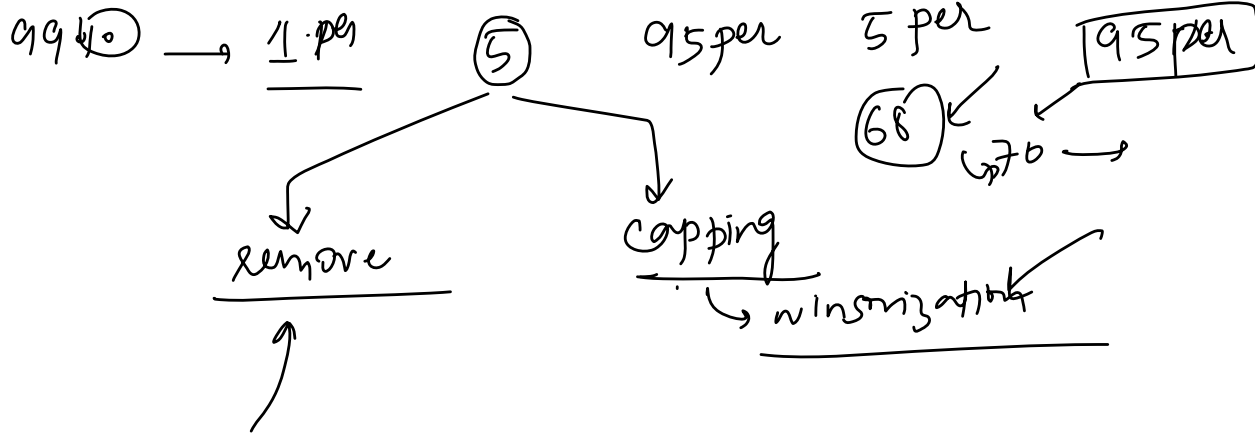
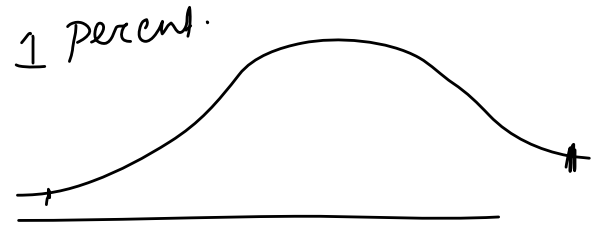
25th { ± IQR Proximity Rule }

Intro

Monday, May 3, 2021 11:49 AM

max - 95 → 100%
 min - 10 → 0%
 50% → median

Age
 27
 35
 72

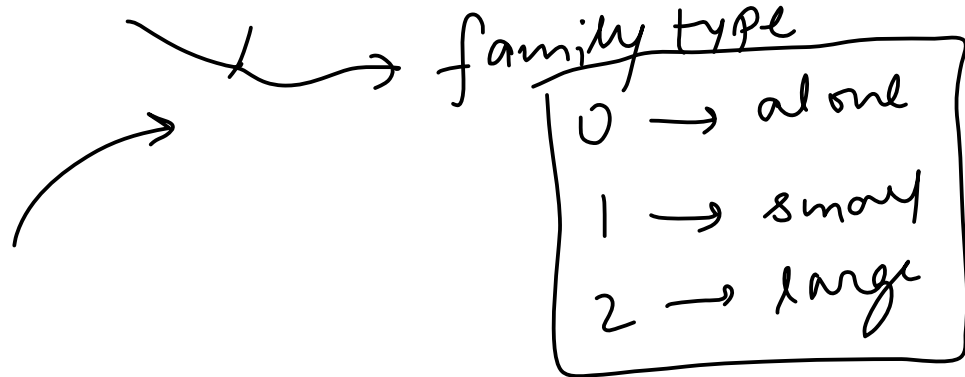


Feature Construction

Tuesday, May 4, 2021 11:13 AM

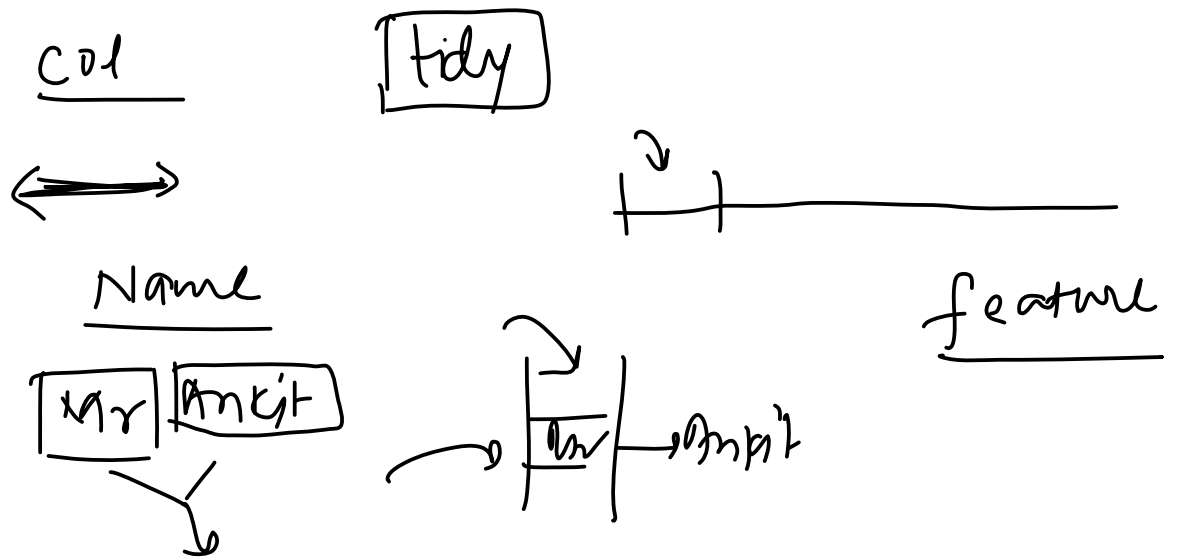
Titanic

SibSp | Parent



Feature Splitting

Tuesday, May 4, 2021 11:13 AM



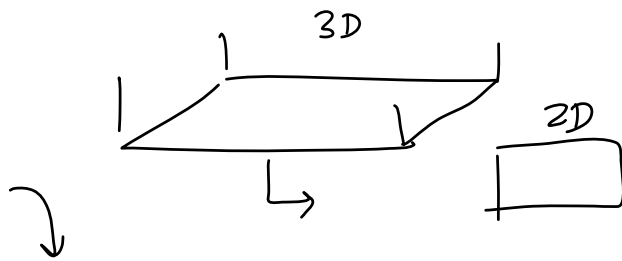
Curse of Dimensionality

Wednesday, May 5, 2021 1:02 PM

1. Introduction

Wednesday, May 5, 2021 1:02 PM

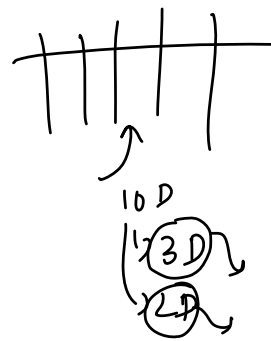
feature extraction $\rightarrow$ CoD $\downarrow$
 $\hookrightarrow$ PCA $\rightarrow$ fL



benefits

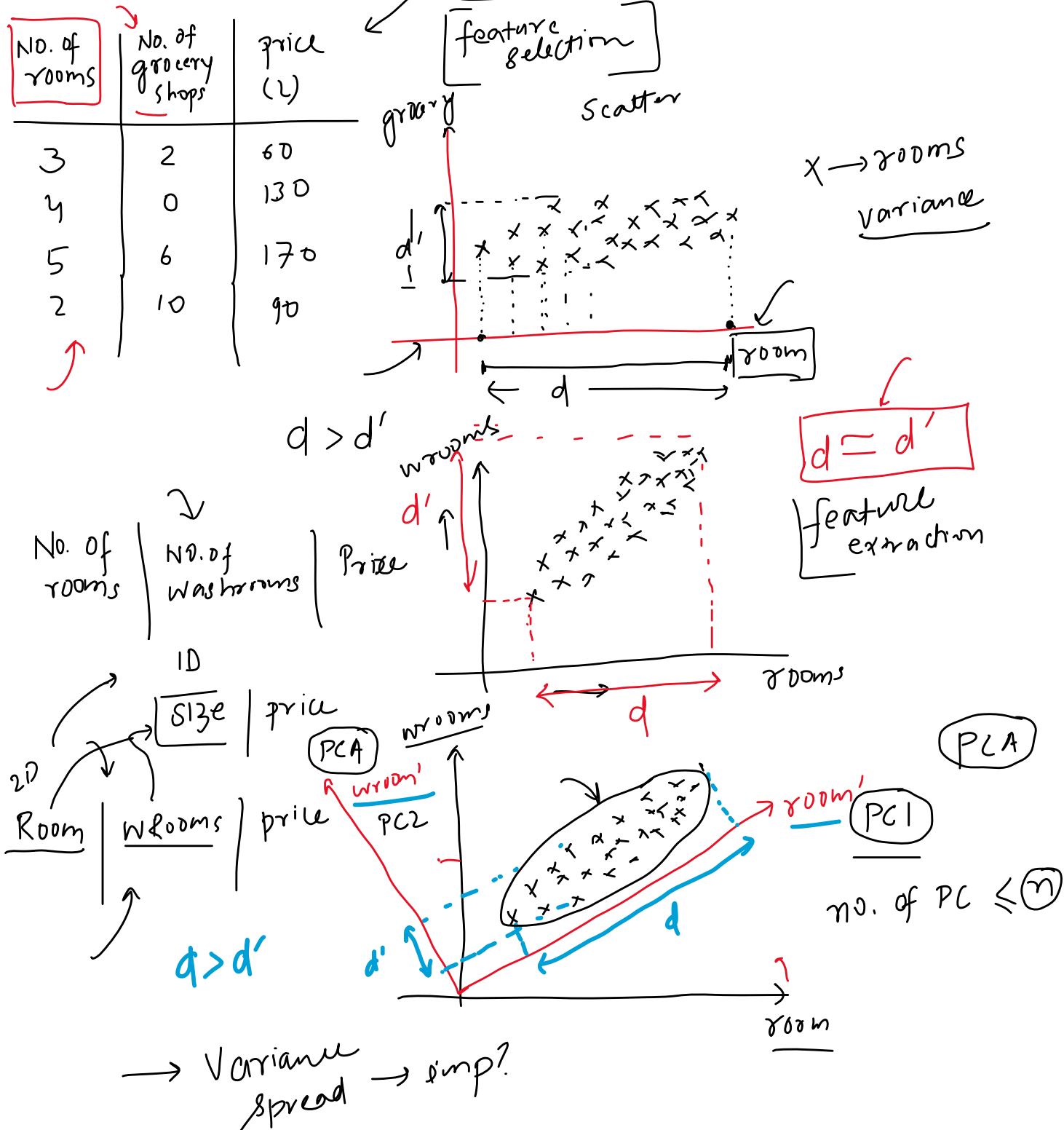
- 1) faster exe of algo
- 2) visualization

784 $\rightarrow$ 3D

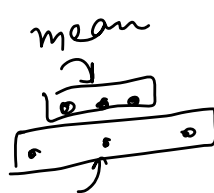


2. Geometric Intuition

Wednesday, May 5, 2021 1:02 PM



Wednesday, May 5, 2021 1:03 PM

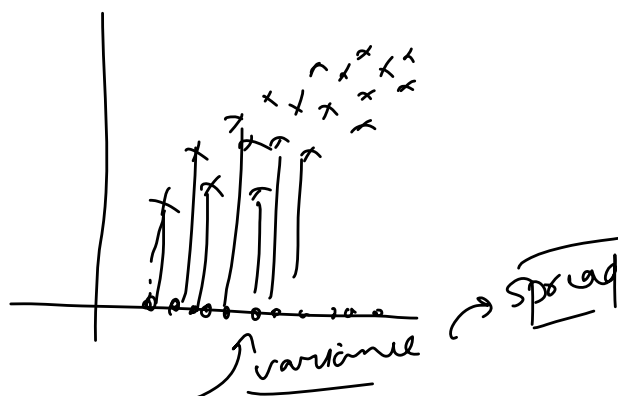


$$\text{variancia} = \frac{\sum_{i=1}^n (x_i - \bar{x})^2}{n}$$

variance

$$v = \frac{25 + 0 + 25}{3} = \frac{50}{3}$$

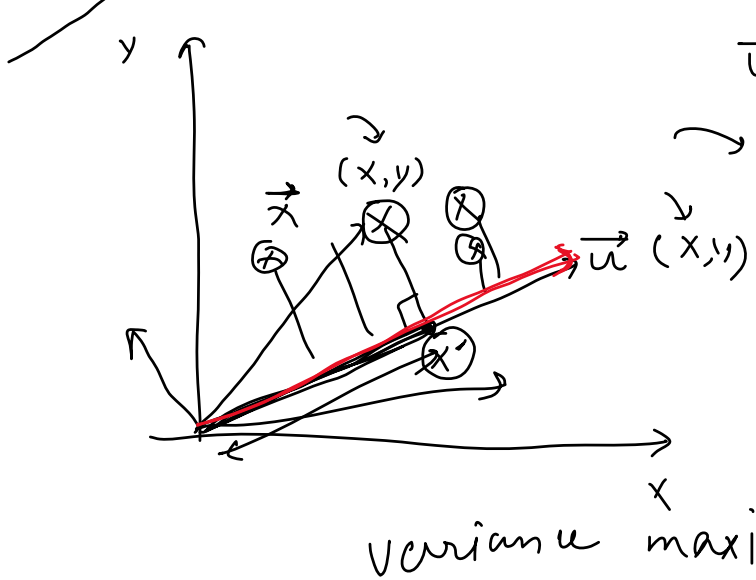
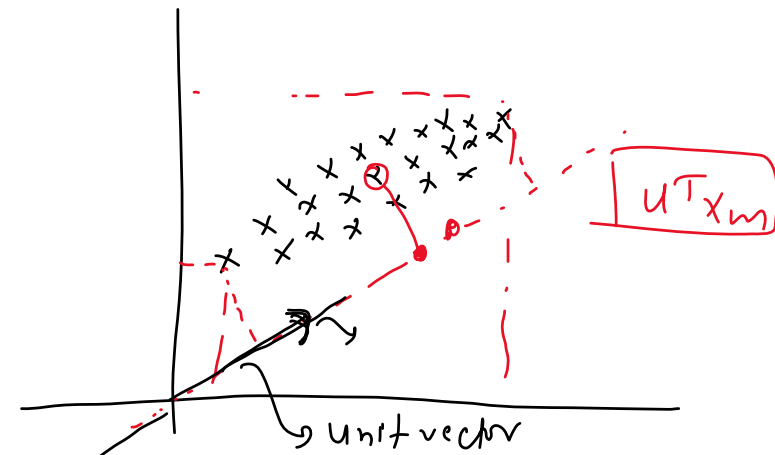
$$V = \frac{100 + 0 + 100}{3} = \sqrt{\frac{200}{3}} = 5d$$



4. Problem Formulation

Wednesday, May 5, 2021 1:03 PM

2D $\rightarrow$ 1D

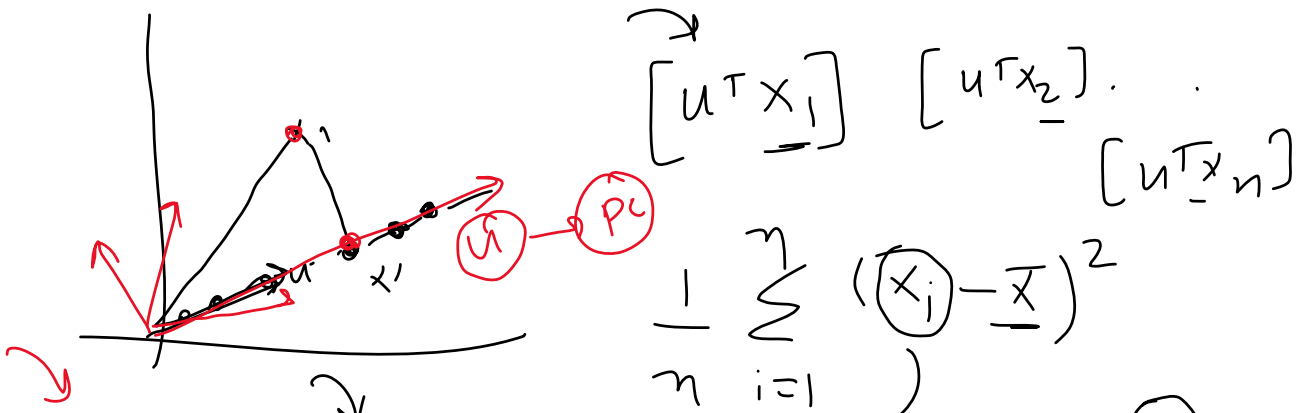


$$\frac{u \cdot \vec{x}}{|u|=1} = \vec{u} \cdot \vec{x} = \boxed{u^T x}$$

$$\begin{bmatrix} x & y \end{bmatrix} \begin{bmatrix} x \\ y \end{bmatrix}$$

$$\begin{bmatrix} x_1 & x_2 \end{bmatrix} \begin{bmatrix} y_1 \\ y_2 \end{bmatrix}$$

$$= \boxed{x_1 x_2 + y_1 y_2} - \text{scalar}$$



$$\begin{bmatrix} u^T x_1 \end{bmatrix} \begin{bmatrix} u^T x_2 \end{bmatrix} \dots \begin{bmatrix} u^T x_n \end{bmatrix}$$

$$\frac{1}{n} \sum_{i=1}^n (x_i - \bar{x})^2$$

$$\sum_{i=1}^n (u^T x_i - u^T \bar{x})^2$$

max
variance u

Optimization

Cov mat

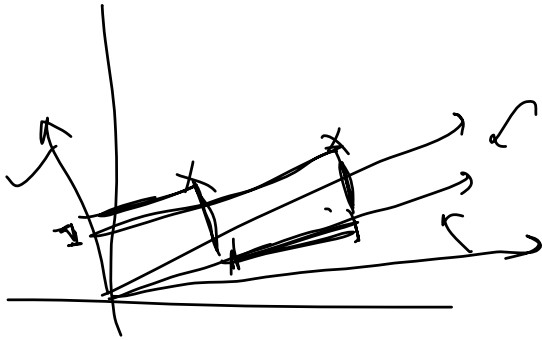
L'

η



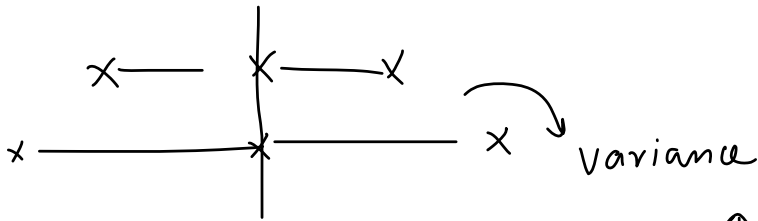
L

cov mat
eigen

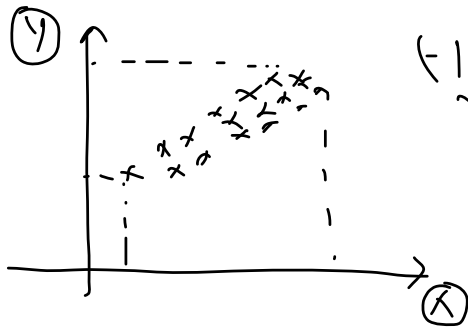


Covariance and Covariance Matrix

Thursday, May 6, 2021 7:02 AM

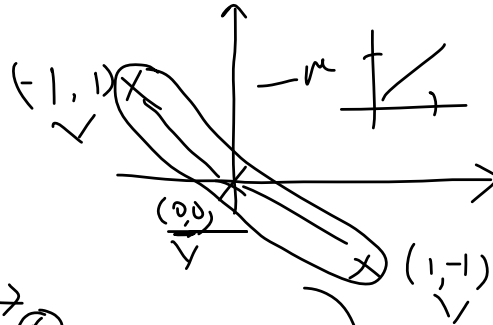


Covariance



Correlation

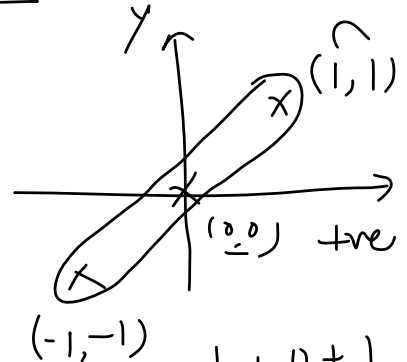
-1 to 1



$$\frac{1+0+1}{3}$$

$$\left(-\frac{2}{3}\right)$$

$$\frac{-1+0-1}{3}$$



$$\frac{1+0+1}{3}$$

$$=\frac{2}{3}$$

x_1	x_2
x_1	
	x_2

cov mat

$$\begin{bmatrix} & \\ & \end{bmatrix} 2 \times 2$$

(ve)

$$\begin{bmatrix} \text{cov}(x_1, x_1) & \text{cov}(x_2, x_1) \\ \text{cov}(x_1, x_2) & \text{cov}(x_2, x_2) \end{bmatrix}$$

cov mat

$$\text{cov}(x_1, x_1)$$

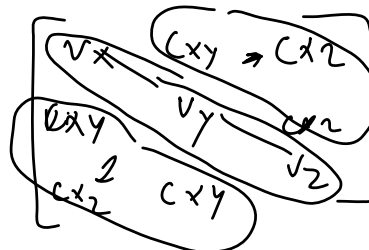
$$\downarrow$$

$$\text{var}(x_1)$$

$$\begin{bmatrix} \text{Var } x_1 & \text{cov}(x_2, x_1) \\ \text{cov}(x_1, x_2) & \text{Var } x_2 \end{bmatrix}$$

square symmetric

$$\text{cov}(a, b) = \text{cov}(b, a)$$



$$x | y | z$$

x	y	z
$\sqrt{x}$	$C_{x,y}$	$C_{x,z}$

cov matrix

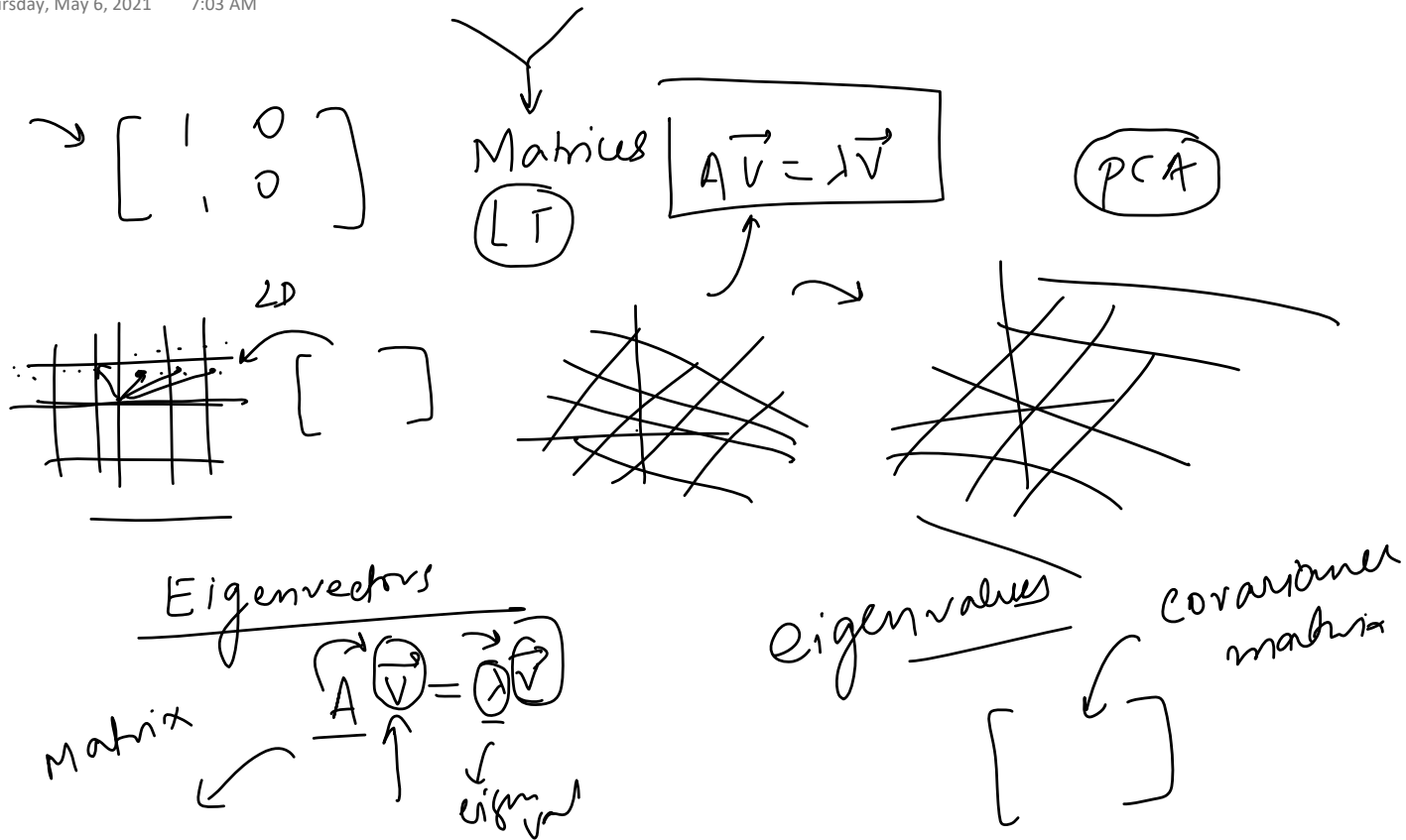
x	$\sqrt{\lambda_1}$	$C_{x,y}$	$C_{x,z}$
y	$C_{y,x}$	$\sqrt{\lambda_2}$	$C_{y,z}$
z	$C_{z,x}$	$C_{z,y}$	$\sqrt{\lambda_3}$
	x	y	z

→ cov matrix

eigen decomposition
of
covariance matrix

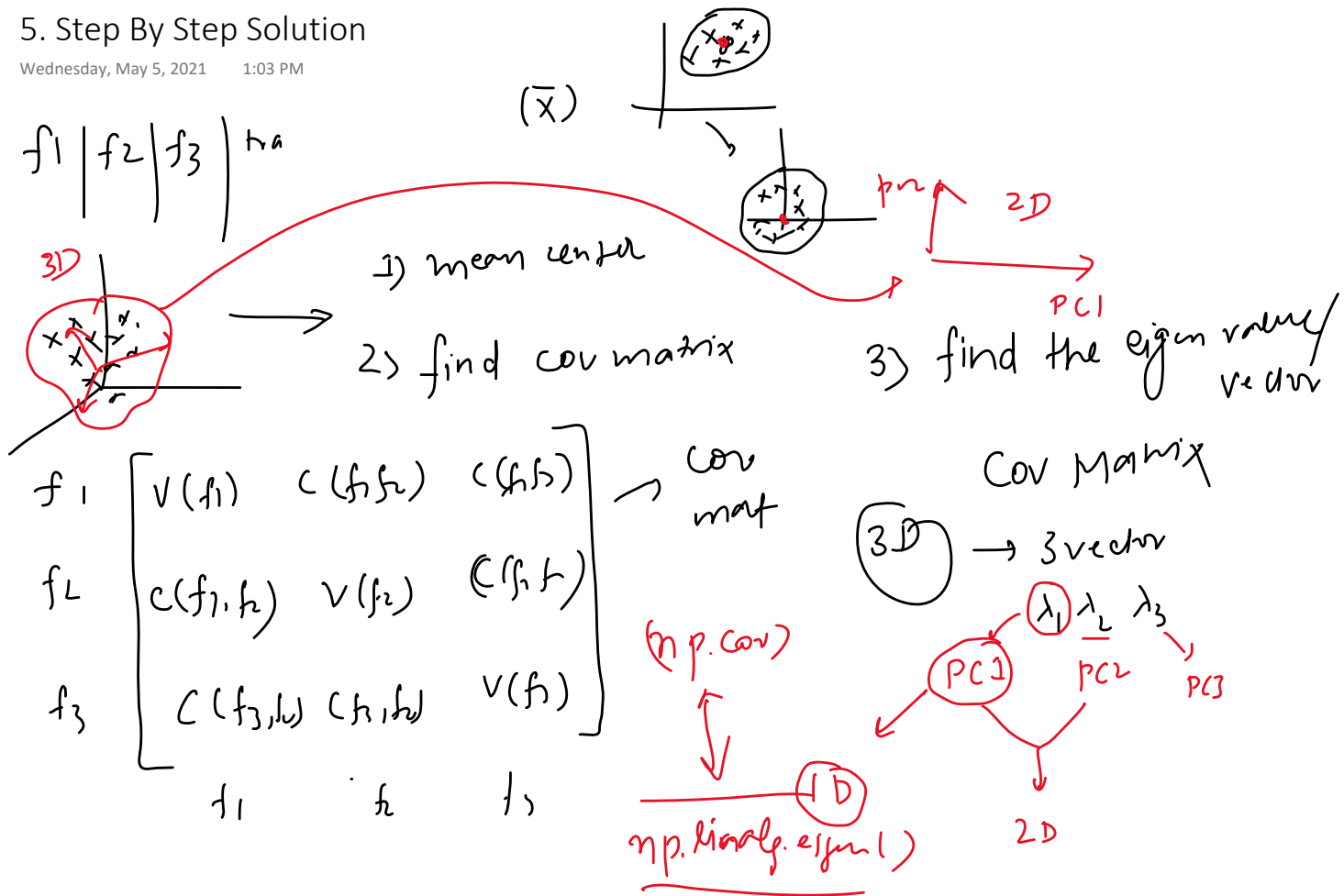
Linear Transformation, Eigen Vectors and Eigen Values

Thursday, May 6, 2021 7:03 AM



5. Step By Step Solution

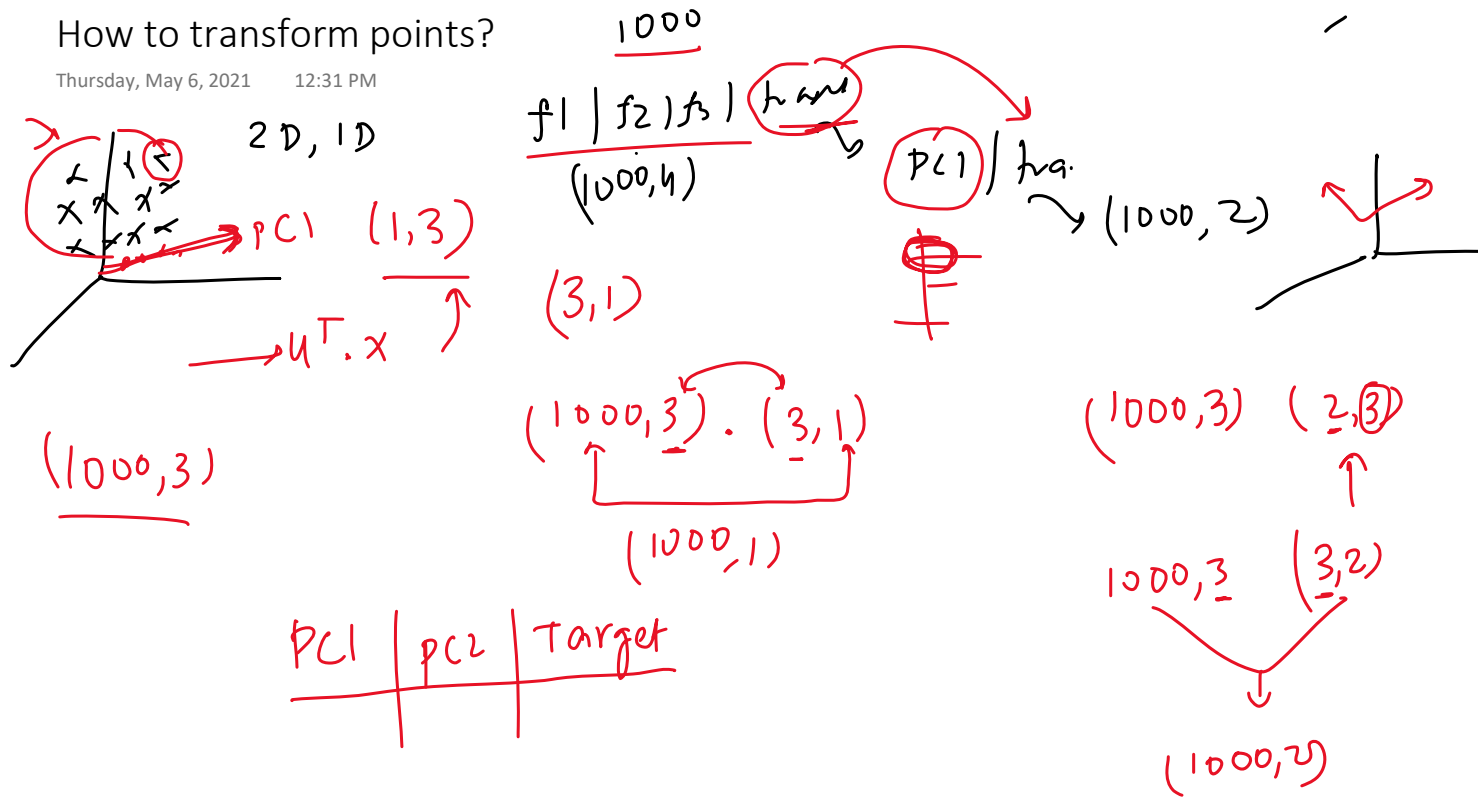
Wednesday, May 5, 2021 1:03 PM



How to transform points?

Thursday, May 6, 2021

12:31 PM

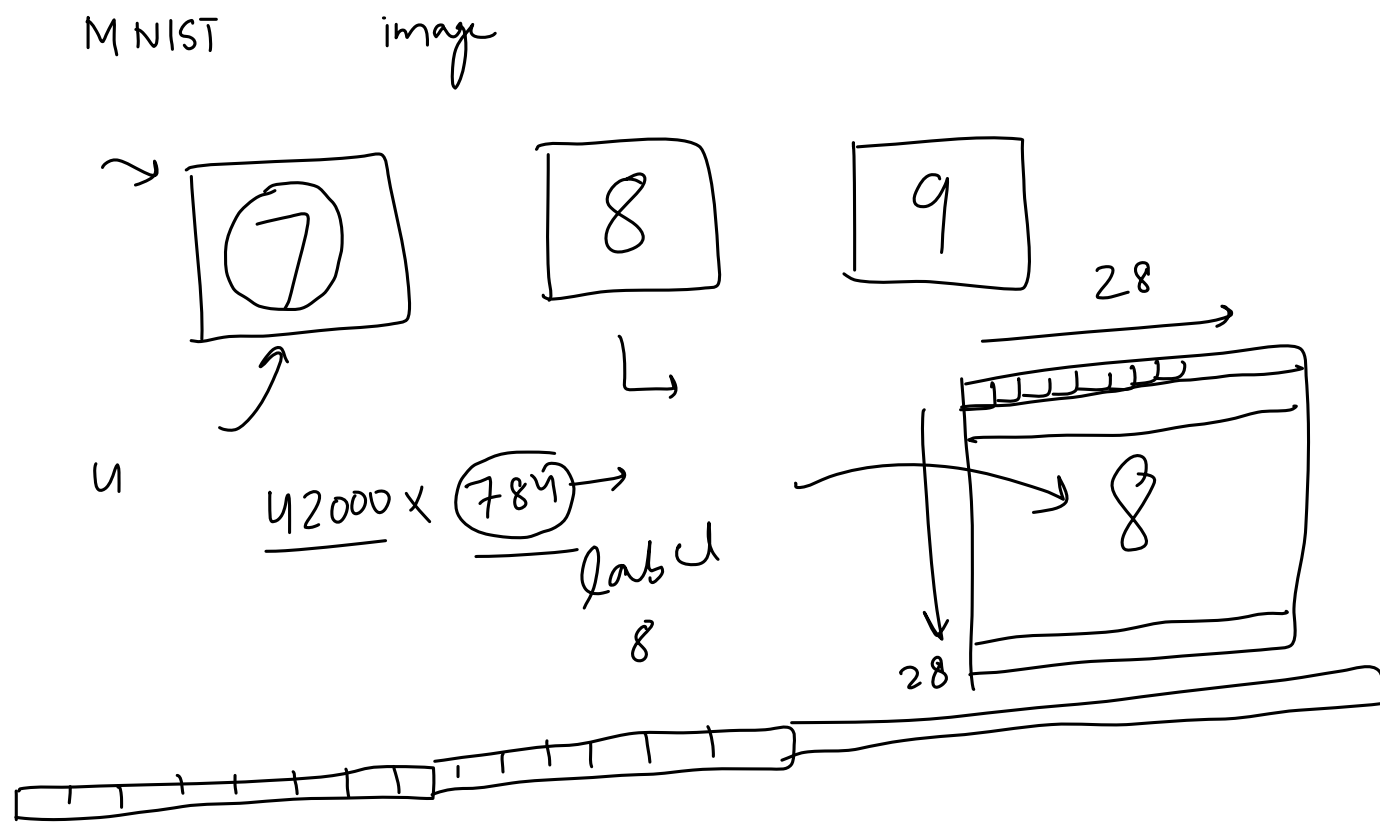


6. Coding the Steps

Wednesday, May 5, 2021 1:03 PM

7. Practical Example on MNIST Dataset

Wednesday, May 5, 2021 1:04 PM

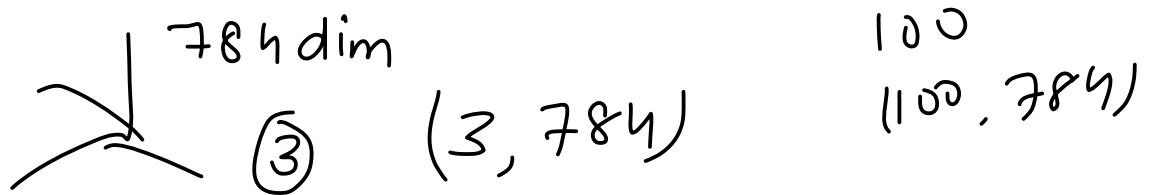
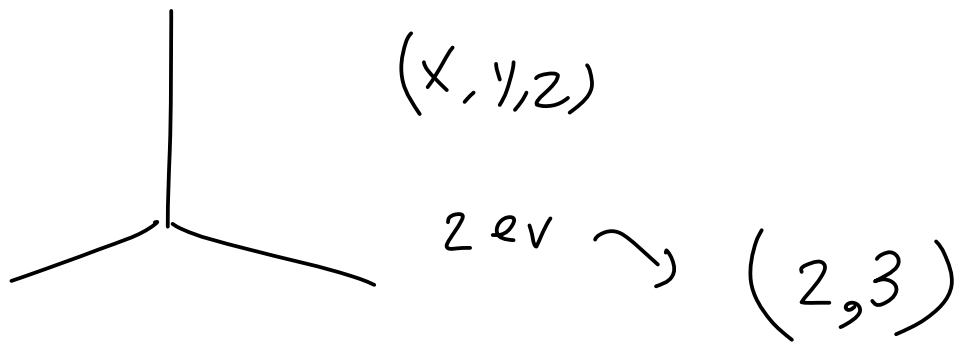


8. Visualization of MNIST Dataset

Wednesday, May 5, 2021 1:04 PM

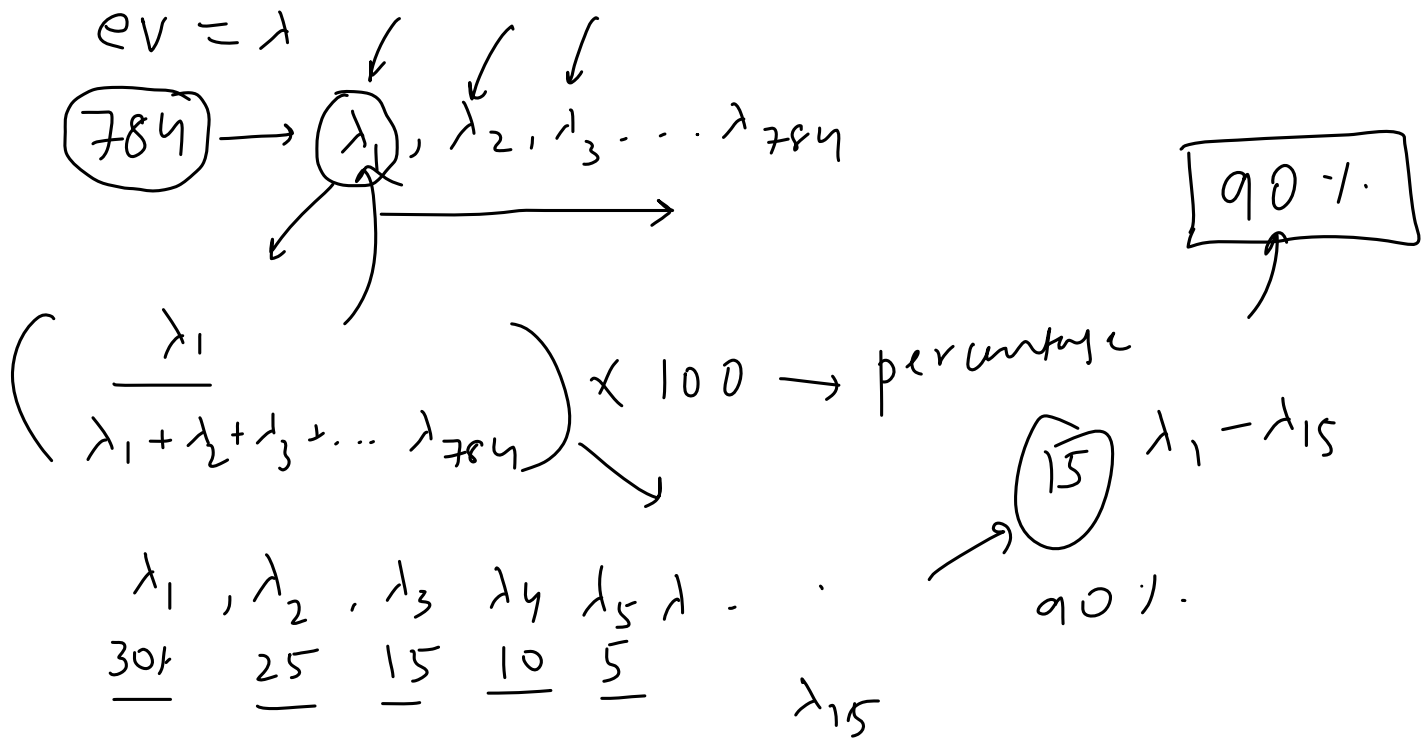
9. Explained Variance

Wednesday, May 5, 2021 1:04 PM



10. Finding optimum number of Principle Components

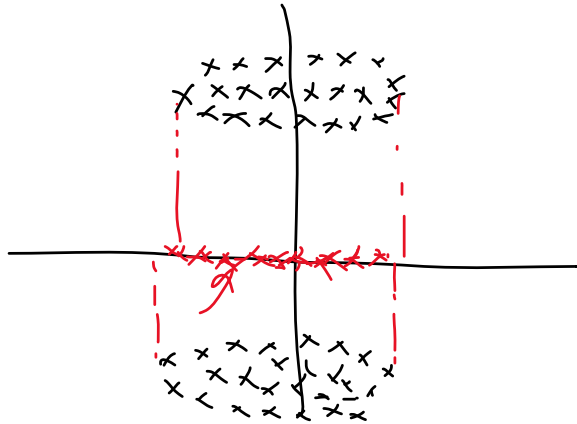
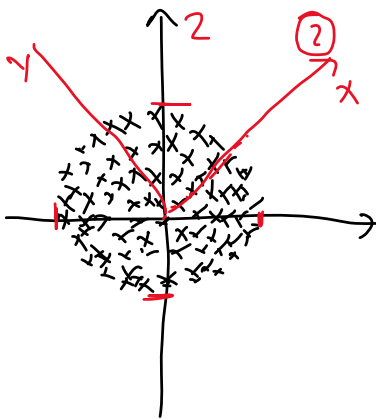
Wednesday, May 5, 2021 1:05 PM



11. When PCA does not work

Thursday, May 6, 2021 7:03 AM

$$y = x^2$$

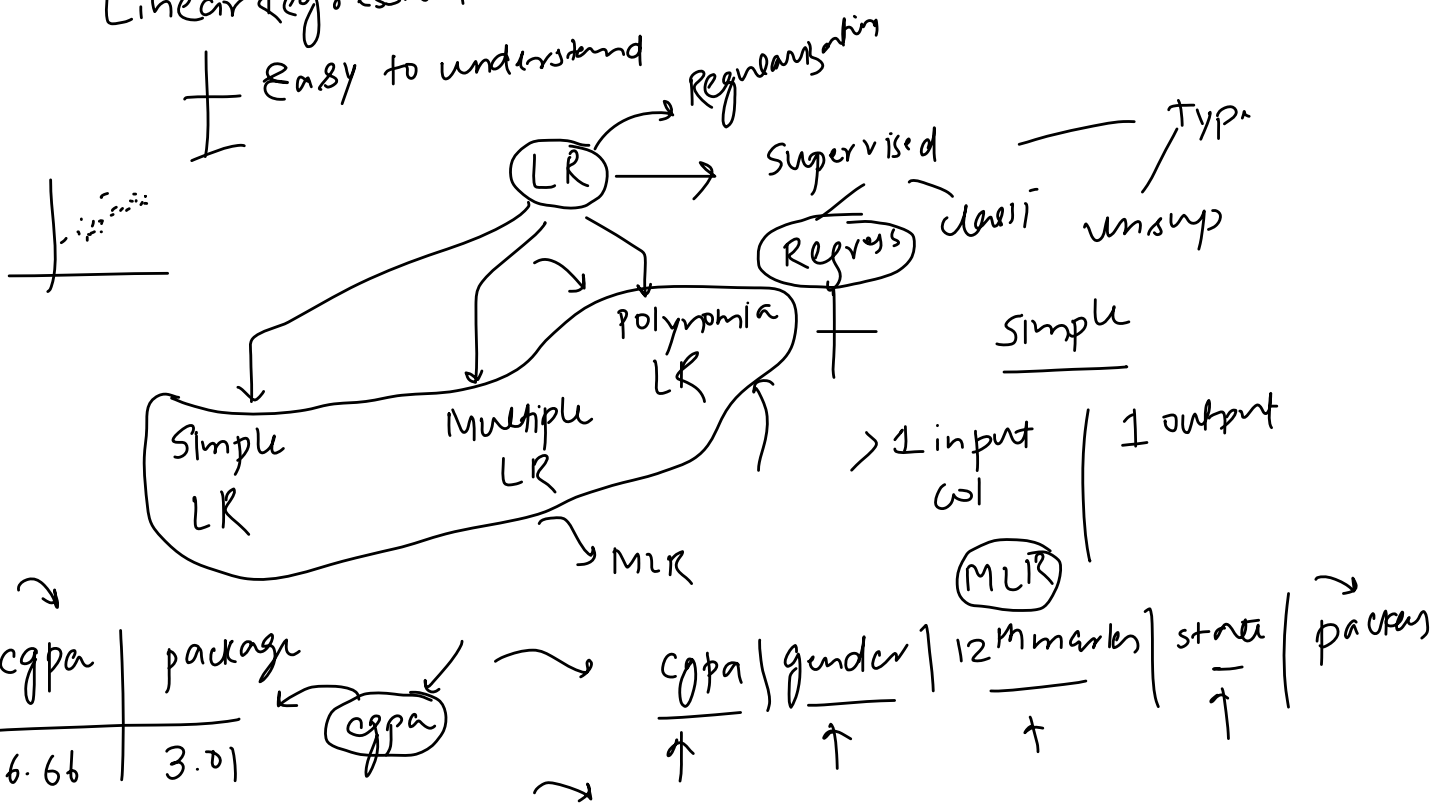


Introduction

Monday, May 10, 2021 3:53 PM

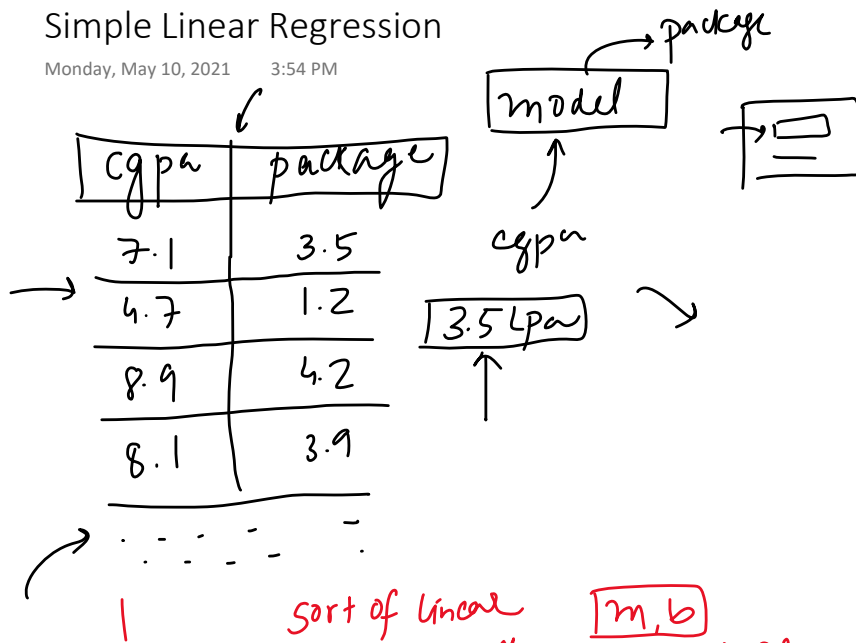
Linear Regression

± Easy to understand



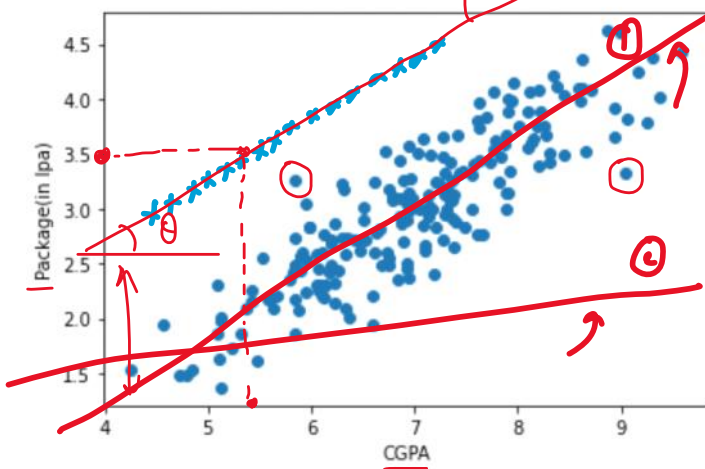
Simple Linear Regression

Monday, May 10, 2021 3:54 PM



sort of linear package m, b $y = m \otimes + b$ $\rightarrow$ cgpa

2008 students



Why?

Real word dataset

$m \rightarrow$ slope
 $b \rightarrow$ y intercept

Best fit line

LR

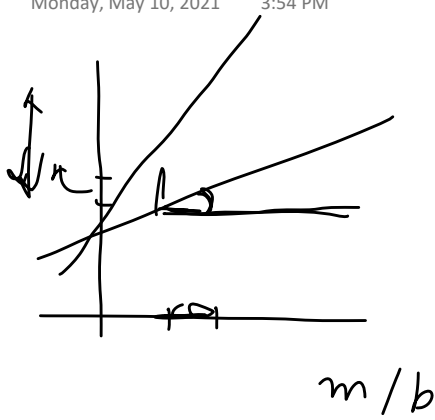
m, b

Code Example

Monday, May 10, 2021 3:54 PM

Intuition

Monday, May 10, 2021 3:54 PM

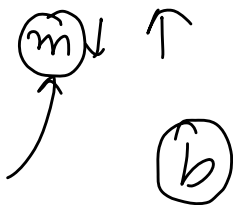


$$y = mx + b$$

$\text{package} = m \times \text{cgpa} + b$

$b = 0$

$m \rightarrow \text{weightage}$



$$\text{package} = m \times \text{cgpa} \quad \text{exp} \rightarrow 0$$

package

$$\text{package} = (m \times \text{exp}) + b = 0$$

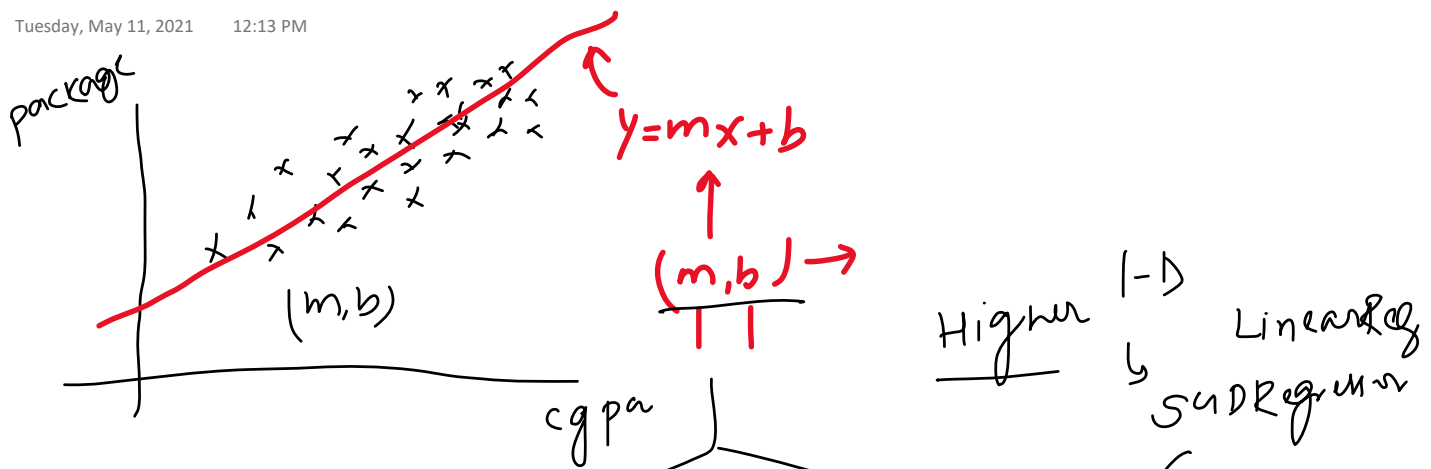
$p = m \times \text{exp}$

$p = 0$

$\text{exp} \mid \text{package offset}$

How to find m and b?

Tuesday, May 11, 2021 12:13 PM



Higher $\rightarrow$ Linear Regression

direct formula

Closed form solution

Non-closed form

OLS

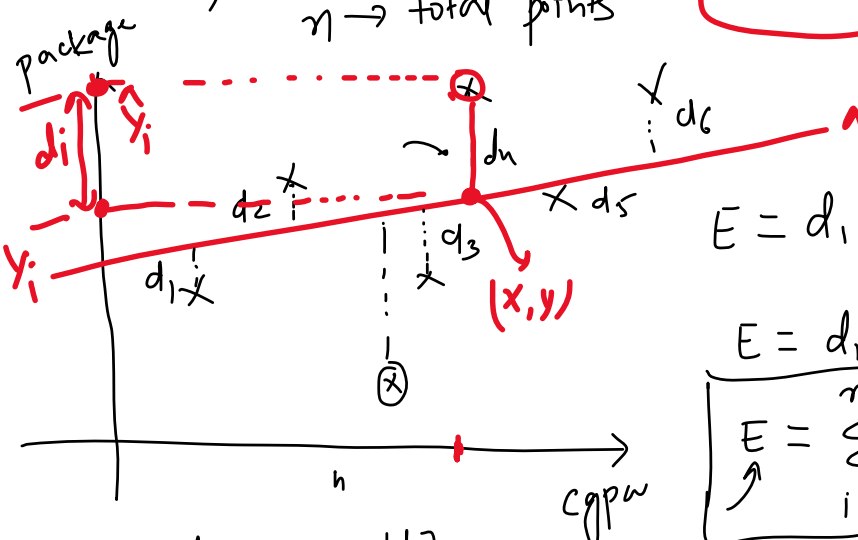
$y \rightarrow$ pack
 $x \rightarrow$ cgpa

Gradient descent

$$b = \bar{y} - m\bar{x}$$

$$m = \frac{\sum_{i=1}^n (x_i - \bar{x})(y_i - \bar{y})}{\sum_{i=1}^n (x_i - \bar{x})^2}$$

$\bar{x}$
 $\bar{y}$ $\rightarrow$ mean (m, b)
 $n \rightarrow$ total points



$$E = d_1 + d_2 + d_3 + \dots + d_n$$

$$E = d_1^2 + d_2^2 + d_3^2 + \dots + d_n^2$$

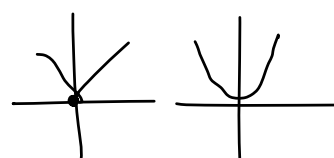
$$E = \sum_{i=1}^n d_i^2$$

Error function J

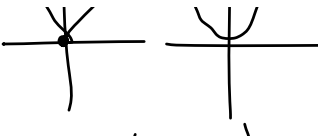
$$E = |d_1| + |d_2| + |d_3| + \dots$$

R1

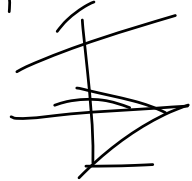
R2 $\rightarrow$



$$d_i = (y_i - \hat{y}_i)$$

$R^2 \rightarrow$  $d_i = (y_i - \hat{y}_i)$
 $\hookrightarrow (m, b)$

$$E = \sum_{i=1}^n (y_i - \hat{y}_i)^2$$

$\hat{y}_i = m x_i + b$ 

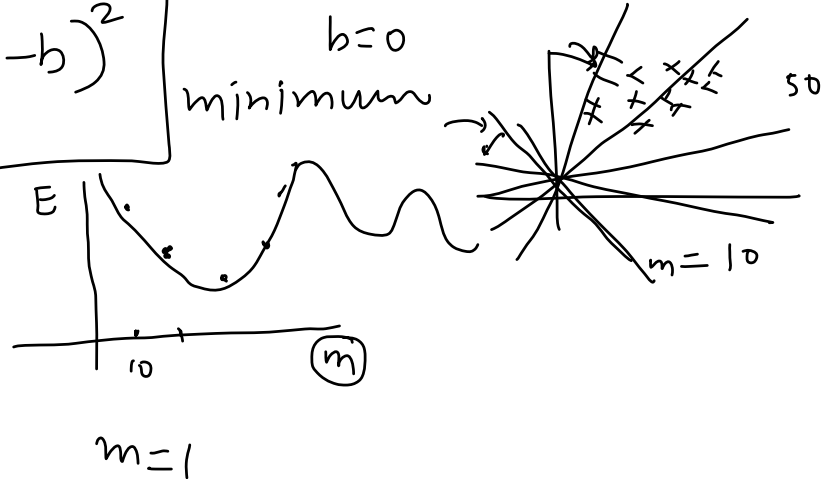
$E(m, b) = \sum_{i=1}^n (y_i - m x_i - b)^2$
 $\hookrightarrow m=0 \rightarrow \frac{E(m, b)}{y = f(x)}$

$$E(m, b) = \sum_{i=1}^n (y_i - m x_i - b)^2$$

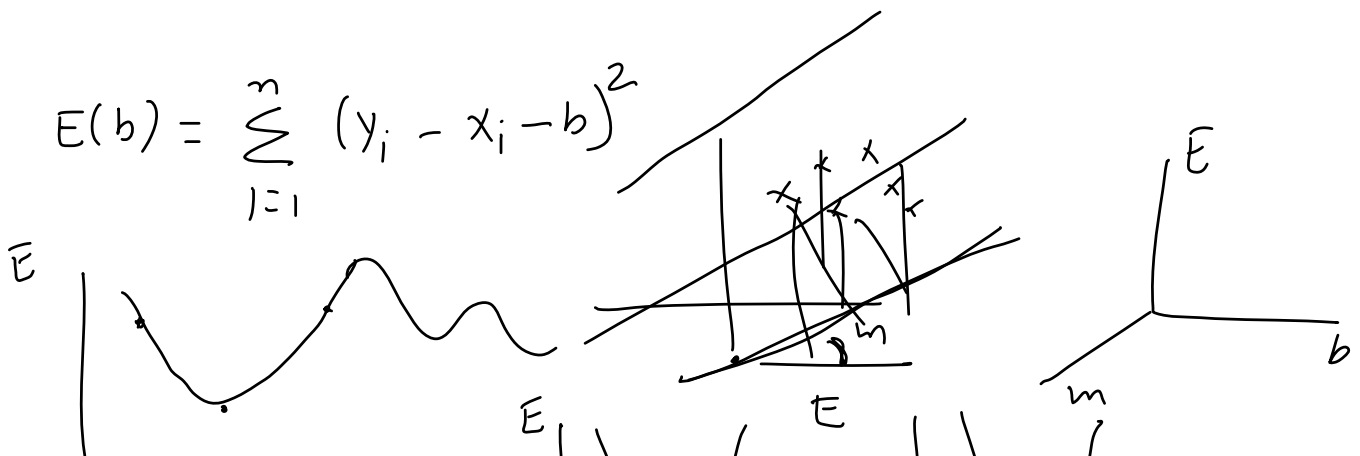
$y = f(x)$ $\hookrightarrow (m, b)$

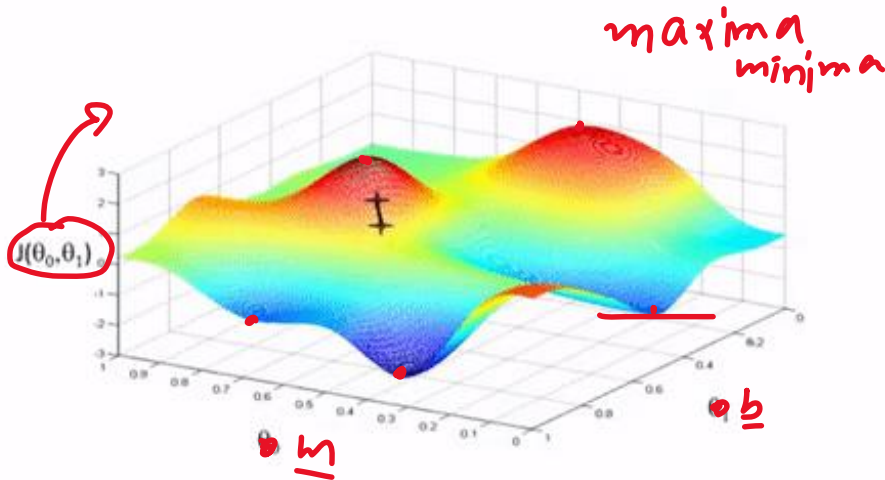
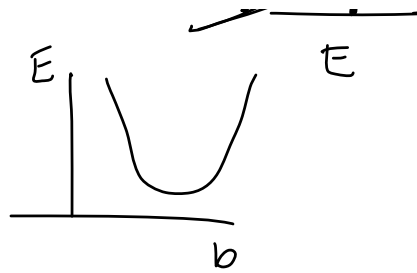
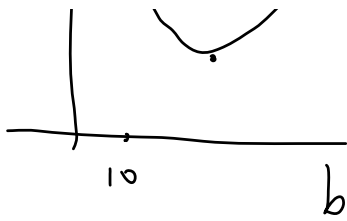
$b=0$

$$E(m) = \sum_{i=1}^n (y_i - m x_i)^2$$



$$E(b) = \sum_{i=1}^n (y_i - x_i - b)^2$$





maxima
minima

$\rightarrow E(x)$

$$\frac{dE}{dx} = 0$$

$$f(x, y)$$

$$\frac{\partial E}{\partial m} = 0, \frac{\partial E}{\partial b} = 0$$

Andrew Ng

$$\frac{\partial E}{\partial b} = \frac{\partial}{\partial b} \sum_{i=1}^n (y_i - mx_i - b)^2 = 0$$

$$= \sum_{i=1}^n \frac{\partial}{\partial b} (y_i - mx_i - b)^2 = 0$$

$$\Rightarrow \sum_{i=1}^n -2(y_i - mx_i - b) = 0$$

$$\Rightarrow \sum_{i=1}^n (y_i - mx_i - b) = 0$$

$$\underbrace{b + b + b + b + \dots + b}_{n \text{ times}} = nb$$

$$\frac{\sum_{i=1}^n y_i}{n} - \frac{\sum_{i=1}^n mx_i}{n} - \frac{\sum_{i=1}^n b}{n} = 0$$

$$\bar{y} - m\bar{x} - \frac{nb}{n} = 0$$

$$\bar{y} - m\bar{x} = b$$

$$\boxed{b = \bar{y} - m\bar{x}}$$

$$E = \sum (y_i - mx_i - \bar{y} + m\bar{x})^2$$

$$\frac{\partial E}{\partial m} = \sum \frac{\partial}{\partial m} (y_i - mx_i - \bar{y} + m\bar{x})^2 = 0$$

$$\Rightarrow \sum 2(y_i - mx_i - \bar{y} + m\bar{x})(-x_i + \bar{x}) = 0$$

$$= \sum -2 (y_i - mx_i - \bar{y} + m\bar{x}) (x_i - \bar{x}) = 0$$

$$= \sum \underbrace{(y_i - mx_i - \bar{y} + m\bar{x})}_{\leftarrow} (x_i - \bar{x}) = 0$$

$$= \sum \left[(y_i - \bar{y}) - m(x_i - \bar{x}) \right] (x_i - \bar{x}) = 0$$

$$= \sum \left[(y_i - \bar{y})(x_i - \bar{x}) - m(x_i - \bar{x})^2 \right] = 0$$

$$= \sum (y_i - \bar{y})(x_i - \bar{x}) - m \sum (x_i - \bar{x})^2$$

$$m = \frac{\sum_{i=1}^n (x_i - \bar{x})(y_i - \bar{y})}{\sum_{i=1}^n (x_i - \bar{x})^2}$$

Code from scratch

Tuesday, May 11, 2021 12:14 PM

Regression Metrics

Thursday, May 13, 2021 11:56 AM

1) MAE

2) MSE

3) RMSE

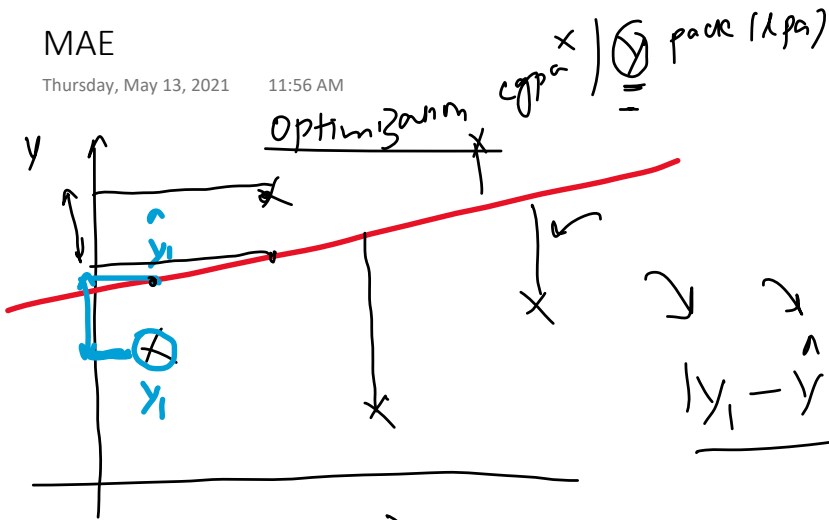
4) R^2 score

5) Adjusted R^2 score

MAE

Thursday, May 13, 2021

11:56 AM



$$|y_1 - \hat{y}_1| + |y_2 - \hat{y}_2| + \dots + |y_n - \hat{y}_n|$$

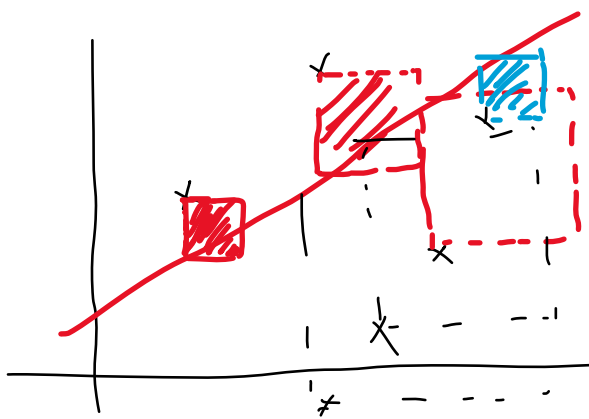
$$\text{mae} = \frac{\sum_{i=1}^n |y_i - \hat{y}_i|}{n}$$

- Advantage
- 1) same unit
 - 2) Robust outliers
- Disadvantage
- 1.5 lpa
-

MSE

Thursday, May 13, 2021 11:56 AM

mean squared error $\sum_{i=1}^n \frac{(y_i - \hat{y}_i)^2}{n}$



11.25

$$mae = \frac{\sum_{i=1}^n |y_i - \hat{y}_i|}{n}$$



mse

$$(y_i - \hat{y}_i)^2$$

Advantage

Disadvantage

→ function

Robust to outliers

mse

$$mse = \frac{\sum_{i=1}^n (y_i - \hat{y}_i)^2}{n}$$

→ function

$$y - lpa$$

$$mse - (lpa)^2$$

RMSE

Thursday, May 13, 2021 11:56 AM

$$RMSE = \sqrt{MSE}$$

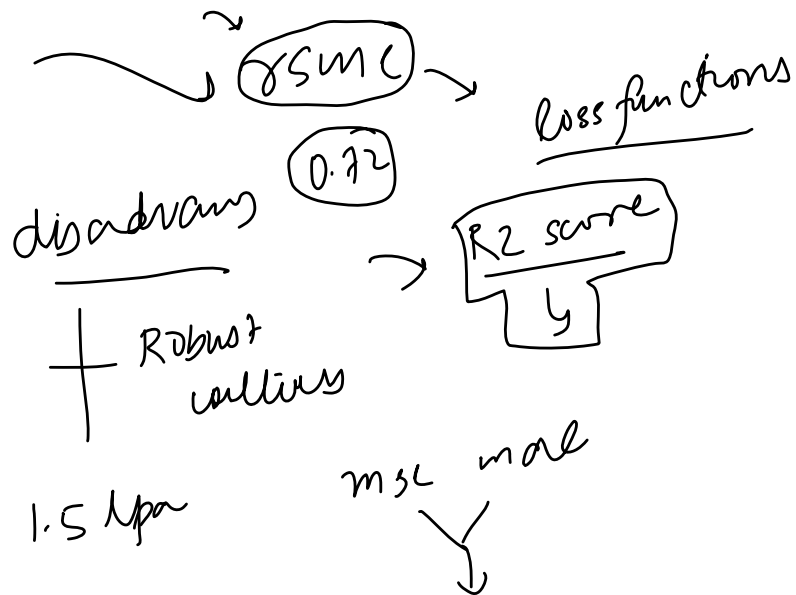
(MSE)

$$= \sqrt{\frac{\sum_{i=1}^n (y_i - \hat{y}_i)^2}{n}}$$

RMSE - lpa $\rightarrow$ benefit

$y - lpa$

1.5

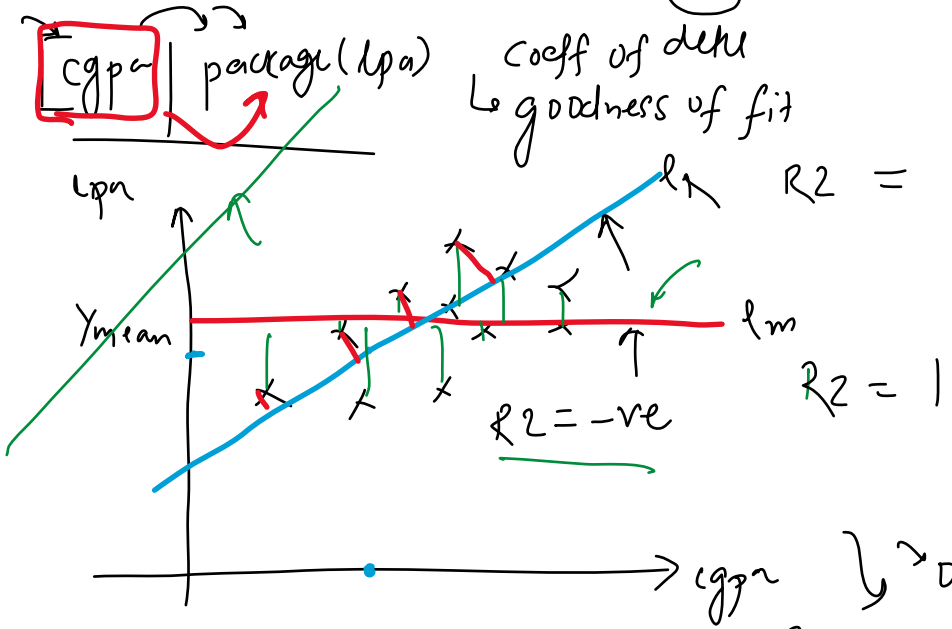


R2 Score

Thursday, May 13, 2021 11:56 AM

100 → mean
3.4

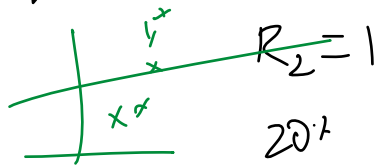
coeff of det
goodness of fit



$$R^2 = 1 - \frac{SS_R}{SS_M}$$

$$R^2 = 1 - \frac{\sum_{i=1}^n (y_i - \hat{y}_i)^2}{\sum_{i=1}^n (y_i - \bar{y})^2}$$

Sum of squared error
Rg



0 → R2 → 1

$$SS_R > SS_M$$

80% -

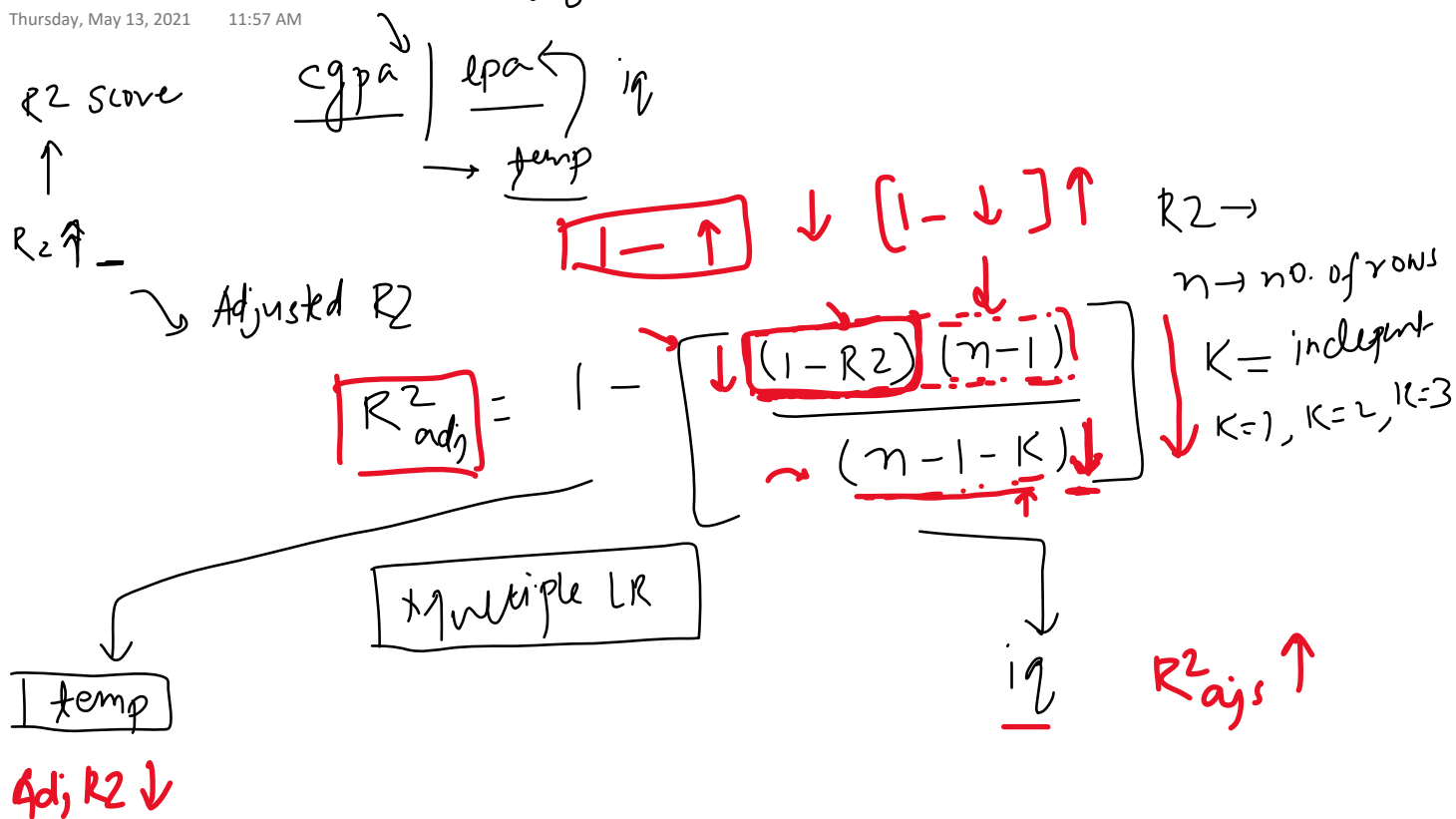
$R^2 \rightarrow 0.80$
cgpa | lpa
↑ 20%
cgpa explains 80% of variance in lpa

cgpa | iq | lpa
↓
lpa

Adjusted R2 score

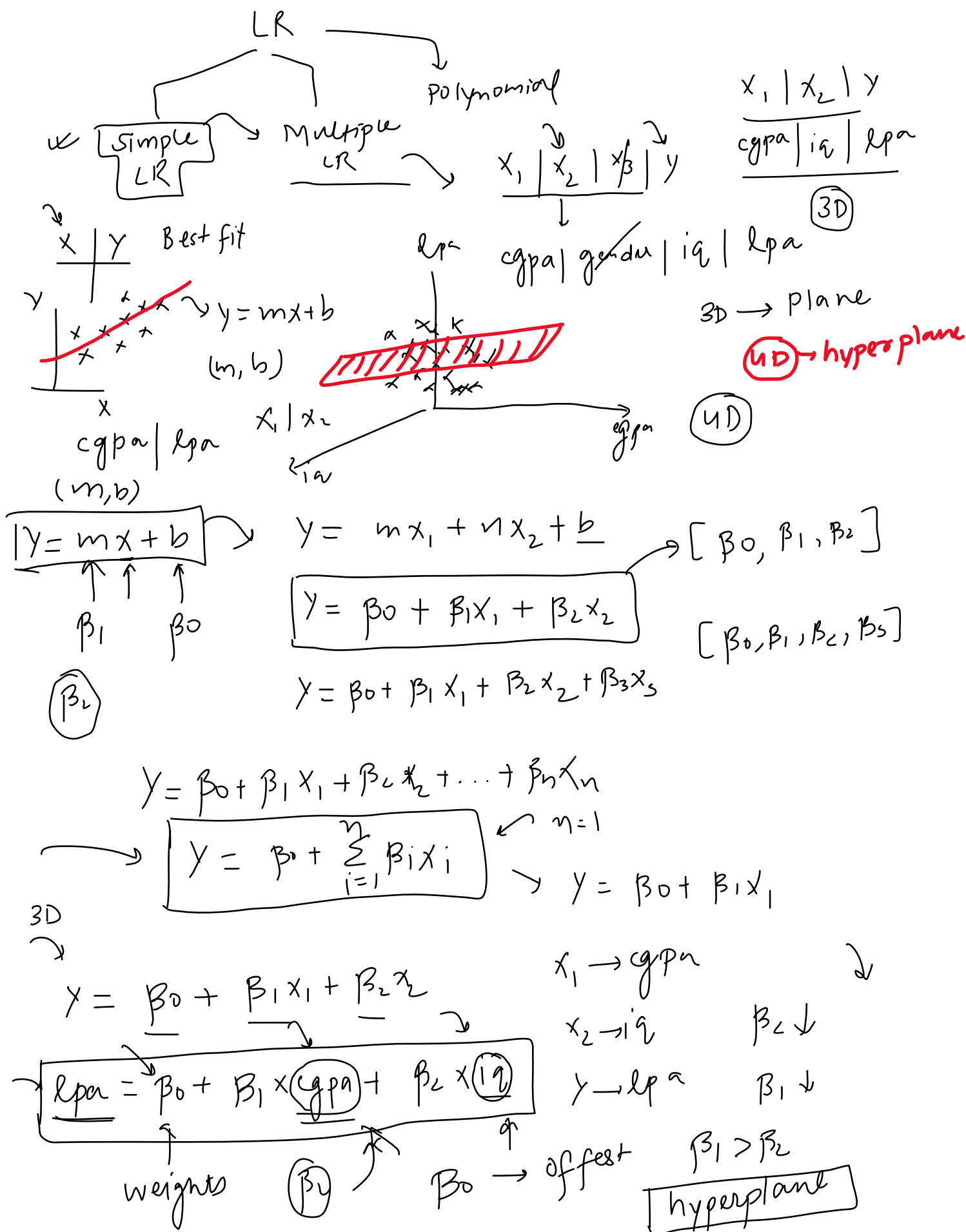
Thursday, May 13, 2021 11:57 AM

0.80 0.90



Multiple Linear Regression

Friday, May 14, 2021 4:31 PM

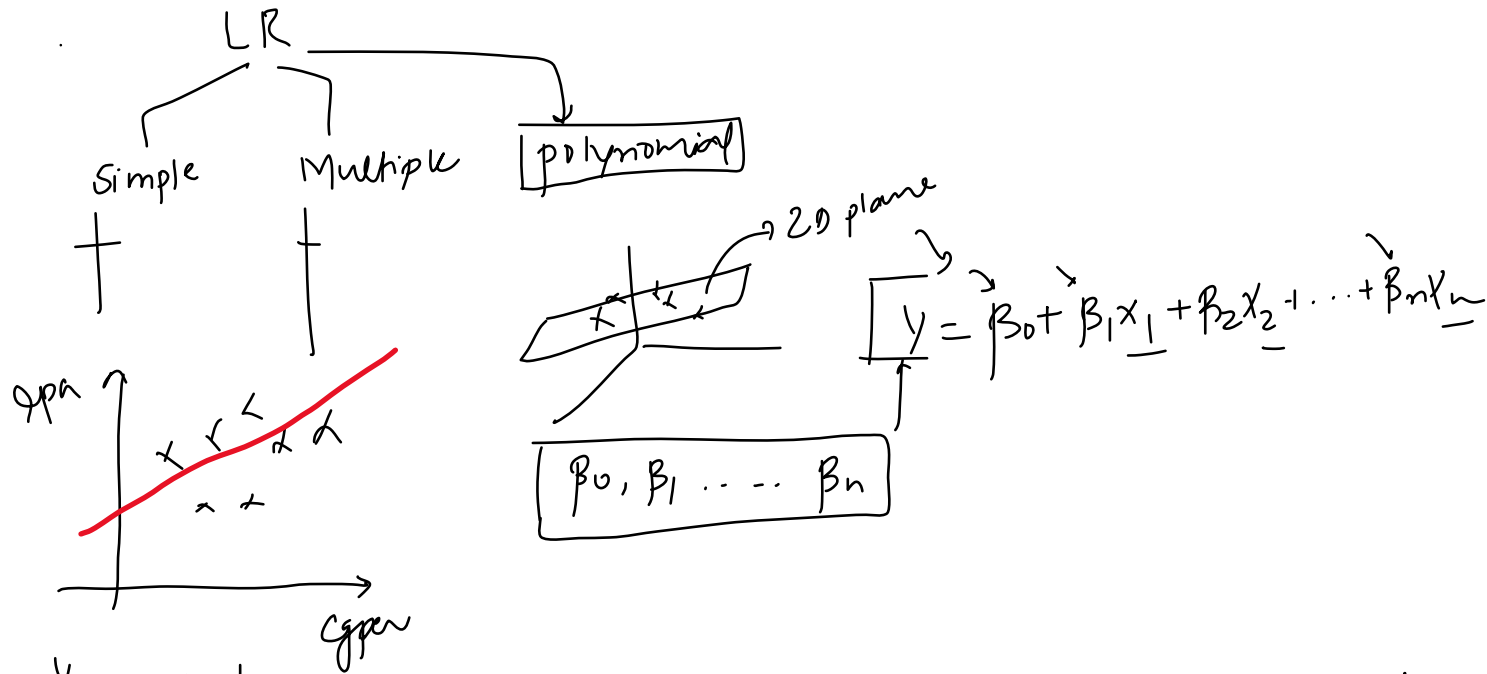


Code Example

Friday, May 14, 2021 4:31 PM

Mathematical Formulation

Saturday, May 15, 2021 4:02 PM



$$y = mx + b$$

$\hookrightarrow (m, b)$

$$y = mx + b$$

$$y = \beta_0 + \beta_1 x$$

$$\beta_0, \beta_1, \beta_2, \beta_3$$

cgpa | iq | gender | lpa

$x_1 \quad x_2 \quad x_3$

actual

100 students
(100, 4)

(V)

$\hat{y} - \hat{x}$

$$\hat{y} = \beta_0 + \beta_1 x_1 + \beta_2 x_2 + \beta_3 x_3$$

$$\hat{y}_1 = \beta_0 + \beta_1 x_{11}$$

$$\hat{Y} = \begin{bmatrix} \hat{y}_1 \\ \hat{y}_2 \\ \vdots \\ \hat{y}_{100} \end{bmatrix} = \begin{bmatrix} \beta_0 & \beta_1 x_{11} & \beta_2 x_{12} & \beta_3 x_{13} \\ \beta_0 & \beta_1 x_{21} & \beta_2 x_{22} & \beta_3 x_{23} \\ \vdots & \vdots & \vdots & \vdots \\ \beta_0 & \beta_1 x_{1001} & \beta_2 x_{1002} & \beta_3 x_{1003} \end{bmatrix}$$

1 row
1 col | 4 cols | 1

100 rows $\rightarrow$ n rows

4 cols $\rightarrow$ (3) col $\rightarrow$ m cols

$$A B = A^T B$$

$$\hat{Y} = \begin{bmatrix} \hat{y}_1 \\ \hat{y}_2 \\ \vdots \end{bmatrix} = \begin{bmatrix} \beta_0 & \beta_1 x_{11} & \beta_2 x_{12} & \beta_3 x_{13} & \dots & \beta_m x_{1m} \\ \beta_0 & \beta_1 x_{21} & \beta_2 x_{22} & \beta_3 x_{23} & \dots & \beta_m x_{2m} \\ \vdots & \vdots & \vdots & \vdots & \dots & \vdots \end{bmatrix}$$

$$\begin{bmatrix} \vdots \\ \hat{y}_n \end{bmatrix} = \begin{bmatrix} \vdots \\ \beta_0 \quad \beta_1 x_{n1} \quad \beta_2 x_{n2} \quad \beta_3 x_{n3} \quad \dots \quad \beta_m x_{nm} \end{bmatrix}$$

$$= \begin{bmatrix} \text{A} & \text{B} \\ \begin{bmatrix} 1 & x_{11} & x_{12} & x_{13} & \dots & x_{1m} \end{bmatrix} \\ \begin{bmatrix} 1 & x_{21} & x_{22} & x_{23} & \dots & x_{2m} \end{bmatrix} \\ \vdots \\ \begin{bmatrix} 1 & x_{n1} & x_{n2} & x_{n3} & \dots & x_{nm} \end{bmatrix} \end{bmatrix} \begin{bmatrix} \beta_0 \\ \beta_1 \\ \beta_2 \\ \vdots \\ \beta_m \end{bmatrix}$$

$$\boxed{\hat{y} = X \beta} \quad \text{--- ①}$$

$\begin{bmatrix} \beta_0 \\ \beta_1 \end{bmatrix} \begin{bmatrix} b \\ m \end{bmatrix} \begin{bmatrix} \beta_0 \\ \beta_1 \\ \beta_2 \end{bmatrix}$

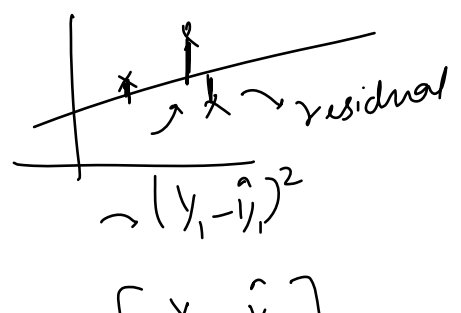
$\begin{matrix} \text{cgr} & | & \text{grad} & | & \text{ig} \\ x_1 & x_2 & x_3 \end{matrix}$

x_{11}	x_{12}	x_{13}
x_{21}	x_{22}	x_{23}
x_{31}	x_{32}	x_{33}

$\boxed{df} \rightarrow X, y$

$$y = \begin{bmatrix} y_1 \\ y_2 \\ \vdots \\ y_n \end{bmatrix} \quad e = y - \hat{y} = \begin{bmatrix} y_1 \\ y_2 \\ \vdots \\ y_n \end{bmatrix} - \begin{bmatrix} \hat{y}_1 \\ \hat{y}_2 \\ \vdots \\ \hat{y}_n \end{bmatrix}$$

$$\boxed{e} = \begin{bmatrix} y_1 - \hat{y}_1 \\ y_2 - \hat{y}_2 \\ \vdots \\ y_n - \hat{y}_n \end{bmatrix}$$



Single LR

single LN

$$E = \sum_{i=1}^n (y_i - \hat{y}_i)^2 \rightarrow E = \underline{e}^T \underline{e}$$

$$\begin{bmatrix} y_1 - \hat{y}_1 \\ y_2 - \hat{y}_2 \\ \vdots \\ y_n - \hat{y}_n \end{bmatrix} = \underline{e}$$

$$\begin{bmatrix} (y_1 - \hat{y}_1) & (y_2 - \hat{y}_2) & (y_3 - \hat{y}_3) & \dots & (y_n - \hat{y}_n) \end{bmatrix} \begin{bmatrix} (y_1 - \hat{y}_1) \\ (y_2 - \hat{y}_2) \\ (y_3 - \hat{y}_3) \\ \vdots \\ (y_n - \hat{y}_n) \end{bmatrix} \rightarrow \textcircled{1} \times 1$$

(1 x n) (n x 1)

$$(y_1 - \hat{y}_1)^2 + (y_2 - \hat{y}_2)^2 + (y_3 - \hat{y}_3)^2 + \dots + (y_n - \hat{y}_n)^2$$

$$\sum_{i=1}^n (y_i - \hat{y}_i)^2$$

$$(A + B)^T = A^T + B^T$$

$$(A - B)^T = A^T - B^T$$

$$E = e^T e = (y - \hat{y})^T (y - \hat{y})$$

$$= (y^T - \hat{y}^T) (y - \hat{y})$$

$$= \left[y^T - (X\beta)^T \right] (y - X\beta)$$

$$E = y^T y - \underline{y^T X \beta} - \underline{(X\beta)^T y} + (X\beta)^T X\beta$$

$$E = y^T y - 2 y^T X \beta + \beta^T X^T X \beta$$

max diff

$$\frac{dE}{d\beta} = \frac{d}{d\beta} \left[y^T y - \underline{2 y^T X \beta} + \beta^T X^T X \beta \right] = 0$$

$$\frac{dJ}{d\beta} = \frac{d}{d\beta} \left[\frac{y^T y}{2} - \underline{y^T X \beta} + \frac{1}{2} \beta^T X^T X \beta \right] = 0$$

$\frac{dy}{dA} = 2 X^T A^T$
 $y = \frac{A^T X A}{1}$

$$= -2 y^T X + 2 X^T X \beta^T = 0$$

$$= \cancel{X^T X} \beta^T = \cancel{y^T X} \frac{1}{[X^T X]}$$

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

$$(\beta^T)^T = [y^T X (X^T X)^{-1}]^T$$

$$\beta = [(X^T X)^{-1}]^T (y^T X)^T$$

$X^T X$

$$\beta = [(X^T X)^{-1}]^T X^T y$$

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

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

$$\begin{bmatrix} \beta_0 \\ \beta_1 \\ \vdots \\ \beta_n \end{bmatrix}$$

$(m+1) \times 1$

$$[(m+1) \times (m+1)] [(m+1) \times n] [n \times 1]$$

$X^T X$ $m \rightarrow \omega's$

$X = n \times (m+1)$
 $(m+1) \times n \quad n \times (m+1)$
 $\downarrow$
 $(m+1) \times (m+1)$

$X \rightarrow \begin{cases} X_{train} \\ \vdots \end{cases}$
 $y = y_{train}$
 $(X^T X)^{-1}$

$n \times (m+1)^T$
 $(m+1) \times n$

$$(AB)^T = B^T A^T \quad Y = A \quad X \beta = B \quad A^T A = A$$

$$\boxed{Y^T X \beta} = (X \beta)^T Y$$

$$A^T B = \underbrace{B^T A}_{\text{RHS}} - (1)$$

$$(A^T B)^T = \underbrace{B^T A}_{\text{RHS}} - (1)$$

$$\Rightarrow A^T B = (A+B)^T \rightarrow A^T B = C$$

$$\boxed{C = C^T}$$

$$A+B \quad C = \boxed{Y^T X \beta}$$

n students

$$Y = \begin{bmatrix} - \\ - \\ - \\ - \\ - \end{bmatrix}$$

$$(n \times 1) \rightarrow (1 \times n)$$

$$(Y^T X \beta)^T = Y^T X \beta$$

$$\downarrow$$

$$\underbrace{(1 \times n) (n \times (m+1))}_{1 \times (m+1)} \underbrace{[(m+1) \times 1]}_{(m+1) \times 1}$$

$$1 \times 1 = [7]^T$$

$$[7]$$

(X)
n rows
m cols

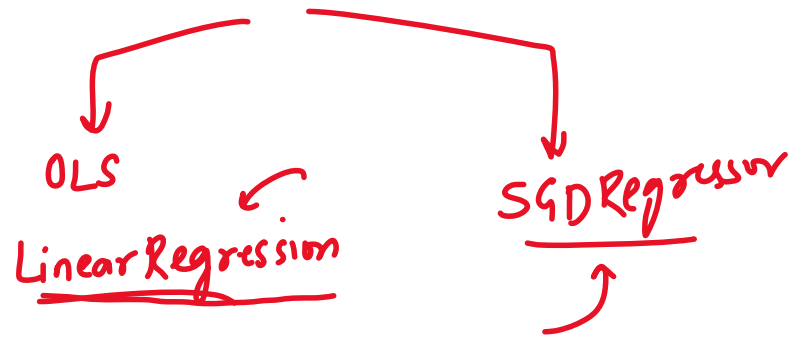
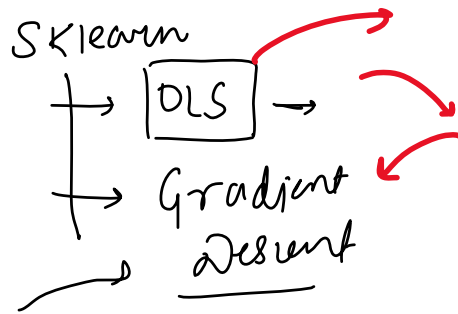
$$n \times (m+1)$$

$$\begin{bmatrix} \beta_0 \\ \beta_1 \\ \vdots \\ \beta_m \end{bmatrix}_{m+1}$$

$$(m+1) \times 1$$

Why Gradient Descent

Saturday, May 15, 2021 5:35 PM



Code From Scratch

Monday, May 17, 2021 12:15 PM

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

(100,3) (100)
 ↓
 100 (100,4)
 ↓

$\underline{X} \rightarrow \text{matrix}$

$\underline{X} \rightarrow \underline{X}_{\text{train}}$

$\underline{y} \rightarrow \underline{y}_{\text{train}}$

cgpa	iq	gender	lpa
1	-	-	-
1	-	-	-

diabetes $\rightarrow$ sklearn $\rightarrow R^2$
 my class $\rightarrow R^2 \rightarrow$ same

$\underline{X}_{\text{test}}$ β_0 $\beta_1 \rightarrow \beta_0$

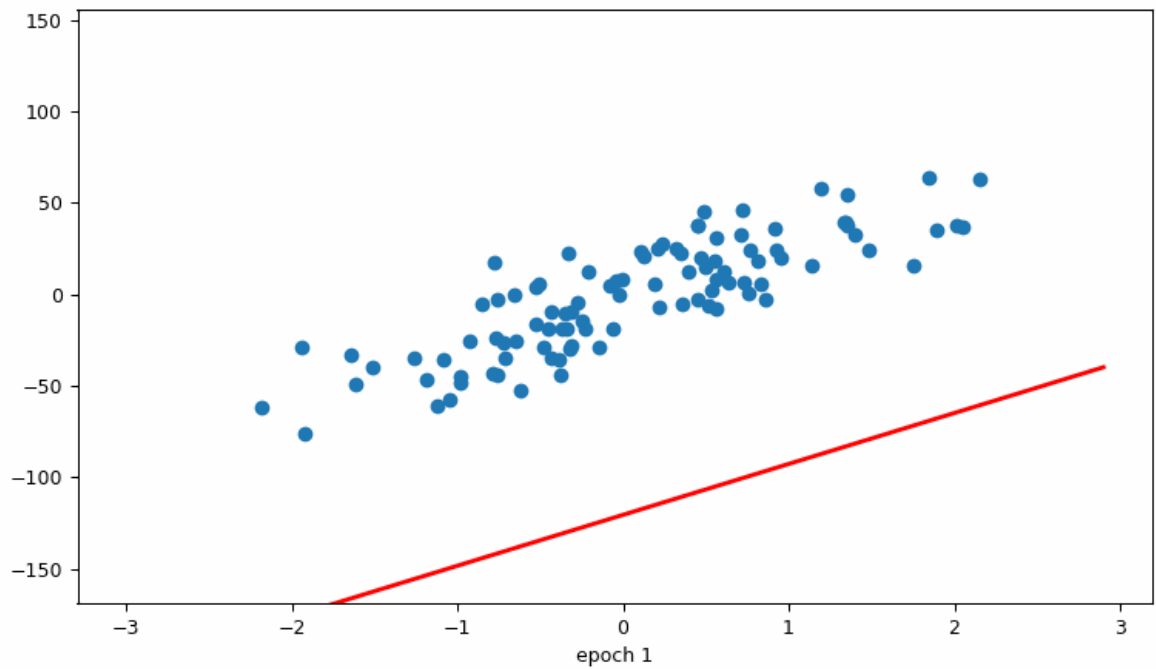
3 $\rightarrow$

$$y = \beta_0 + \beta_1 x_1 + \beta_2 x_2 + \beta_3 x_3$$

$$= \beta_0 + \text{np.dot}(\beta, \underline{x}_{\text{test}})$$

$\underline{X}_{\text{test}}$ (coeff (100,1)
 (89,10)

(89,10) + β_0
 89 $\rightarrow$ ① $\rightarrow$ (89,1)
 y pred



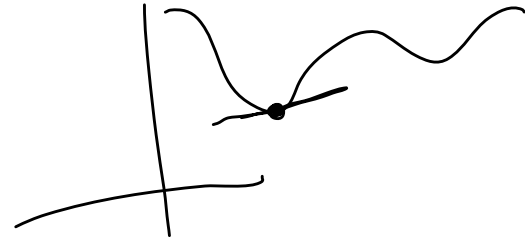
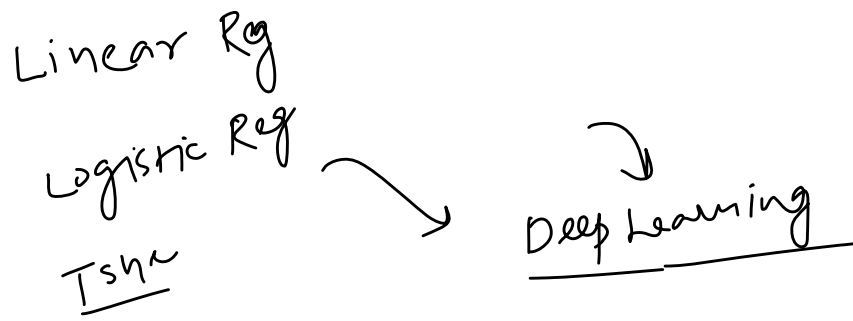
$$\sum (y_i - mx_i - b)^2$$

$$\sum -2 (y_i - mx_i - b) x_i$$

$$-2$$

What is Gradient Descent?

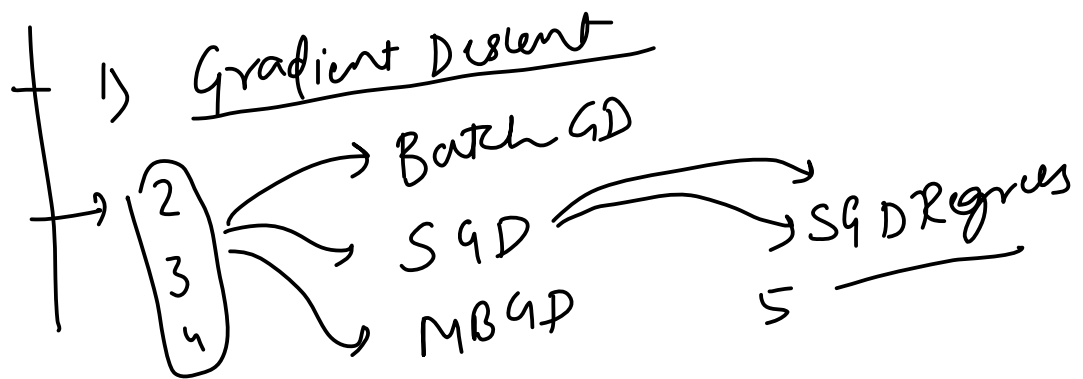
Thursday, May 20, 2021 1:46 PM



The Plan

Thursday, May 20, 2021 1:46 PM

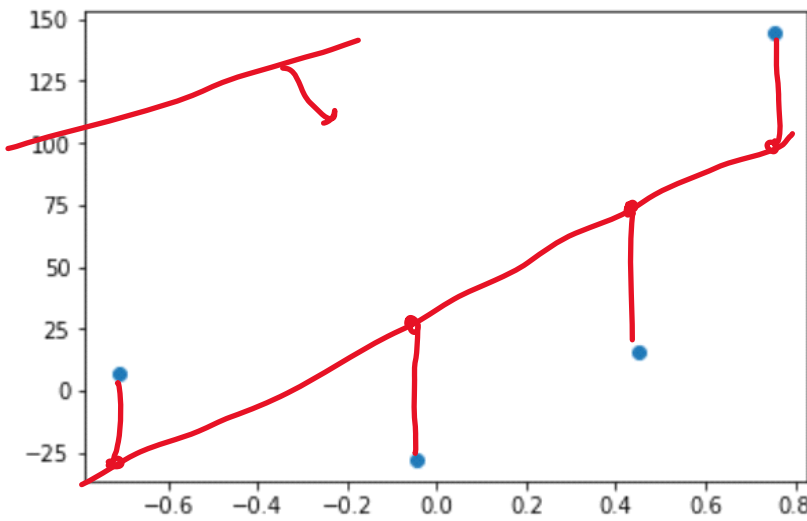
5 videos



Intuition

Thursday, May 20, 2021 1:46 PM

2 cols
4 rows



cgpa | lpa

OLS →

$$\hat{y}_i = mx_i + b$$

$$m = 78.35$$

$$L = \sum_{i=1}^n (y_i - \hat{y}_i)^2$$

$$L = \sum_{i=1}^n (y_i - mx_i - b)^2$$

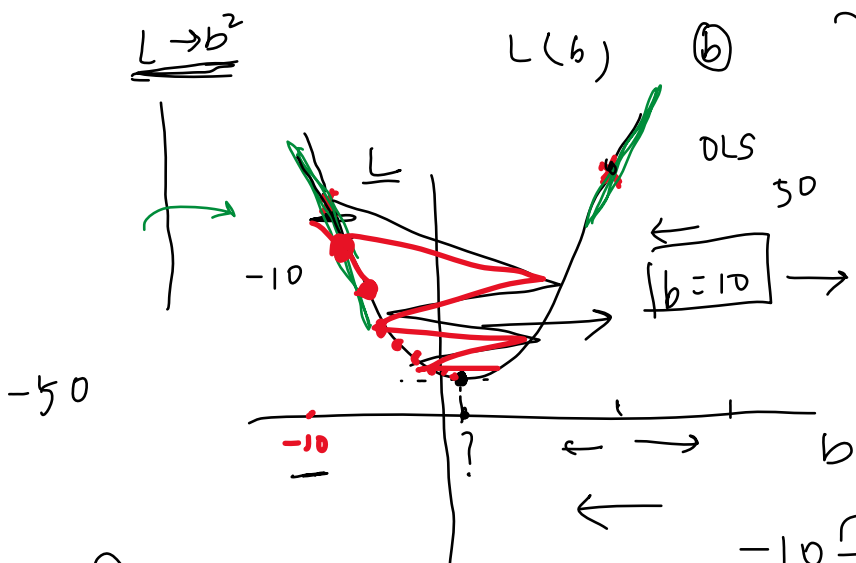
$$L = \sum_{i=1}^n (y_i - 78.35x_i - b)^2$$

Step 1 - select a random b_{min}

b increment
b decrement

L_{min}

$$b = -10 \quad (x = 5)$$



$$b = -10 \rightarrow \text{Slope} = -ve$$

$$b_{new} = b_{old} - \text{slope}$$

1

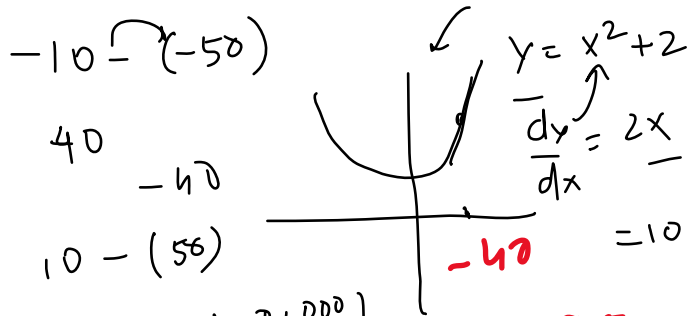
$$b_{new} = b_{old} - \eta \text{slope}$$

$$b_{new} = -9.5 - (0.01 \times -40) = -9.5 + 0.4 = -9.1$$

When to stop

$$|b_{old} - b_{new}| \leq 0.0001$$

Iteration 1000 100



$$b_{new} = -10 + (0.01 \times 50) = -10 + 0.5 = -9.5$$

$$b_{new} - b_{old} = 0.0001$$

$$\Rightarrow 0.0001$$

2) Iteration $\rightarrow$ 1000, 100,

↑ ↑ ↑

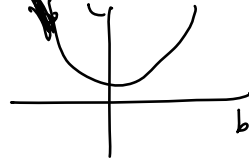
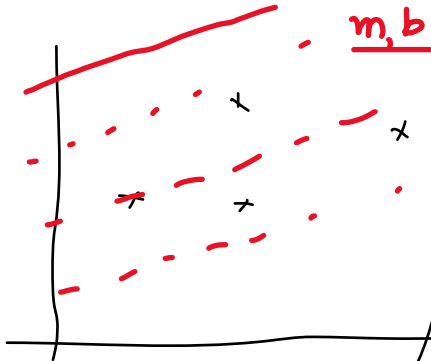
epochs

$$\boxed{b_{\text{new}} - b_{\text{old}} \leq 0}$$

Mathematical Formulation

Thursday, May 20, 2021 1:47 PM

$m = 78.35$



Step $\rightarrow$ Start with a random

for i in epochs;

$$b_{new} = b_{old} - \eta \times \text{slope} \quad (b=0)$$

$i = 0$ η, γ

$$\eta = 0.01$$

$$b = 0$$

$$L = \sum_{i=1}^n (y_i - \hat{y}_i)^2$$

$$\frac{dL}{db} = \frac{d}{db} \left(\sum_{i=1}^n (y_i - \hat{y}_i)^2 \right) = 2 \sum_{i=1}^n (y_i - mx_i - b) (-1)$$

$$\frac{d}{db} \sum_{i=1}^n (y_i - mx_i - b)^2 \quad \text{slope} = -2 \sum_{i=1}^n (y_i - mx_i - b)$$

$$= -2 \sum_{i=1}^n (y_i - 78.35 \times x_i - 0)$$

epoch

$$b_{new} = b_{old} - \eta \times \text{slope}(b=0)$$

steps

Example

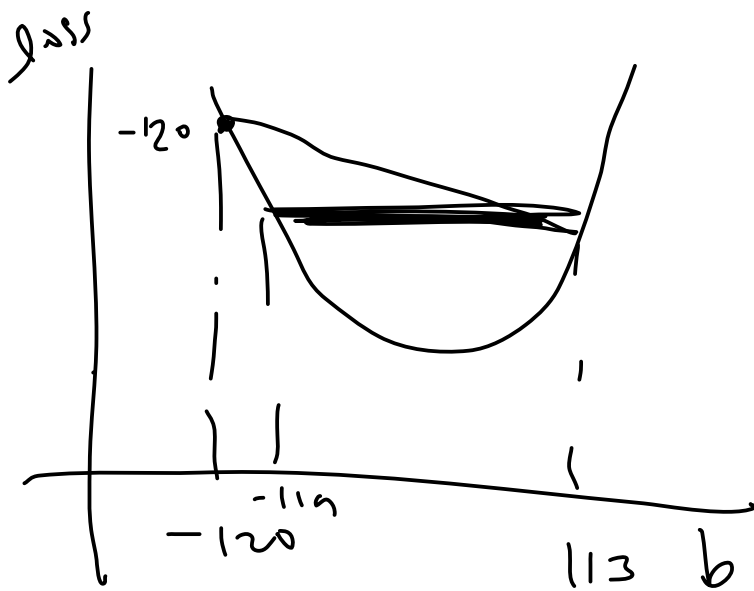
Thursday, May 20, 2021 1:47 PM

Code from Scratch

Thursday, May 20, 2021 1:48 PM

Visualization 1

Thursday, May 20, 2021 1:52 PM

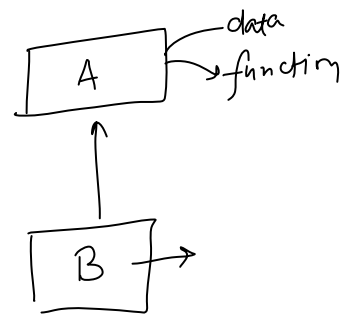


- Inheritance ✓
- Super ✓
- Abstraction ✓
- Static method ✓
- MCQs

→ OOP

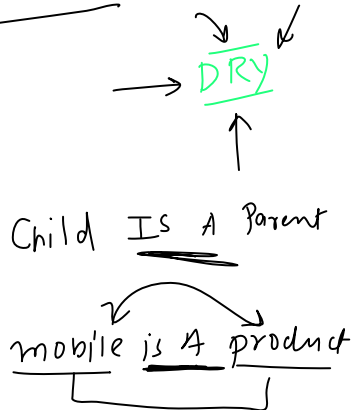
DS Algo

Society



Biology

database



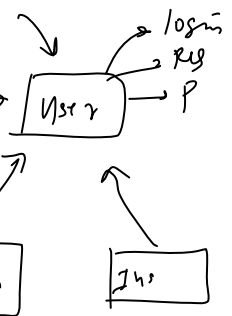
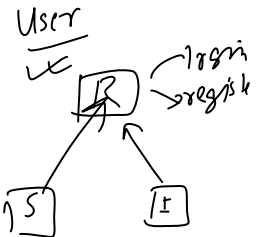
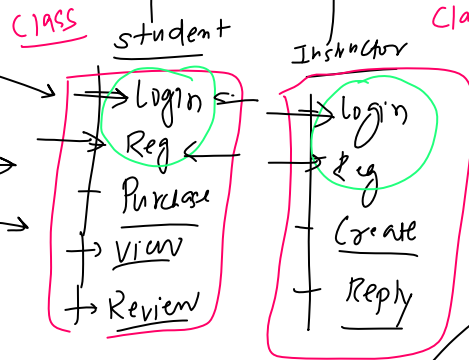
- can be used
- data (property)
 - function
 - constructor

Student is a user

can't be used

→ private members (directly)

Udemy

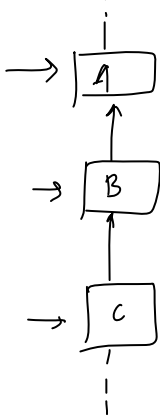


Types

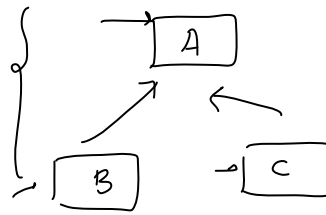
Single



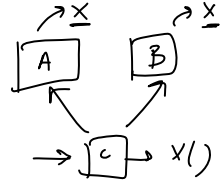
Multi-level



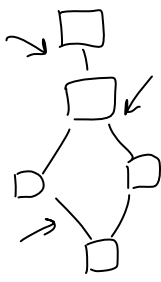
Hierachical

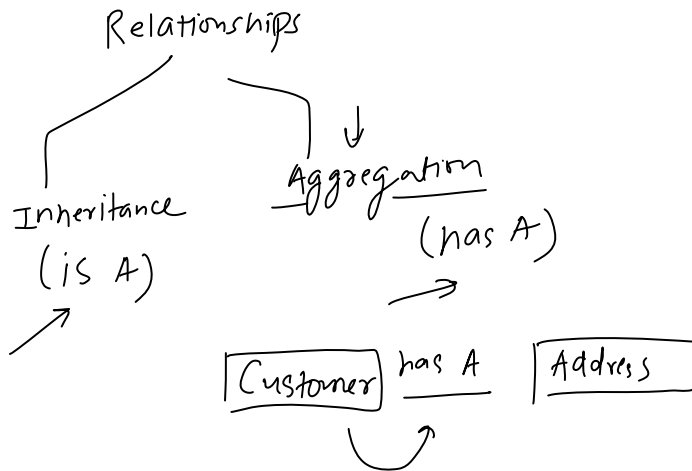
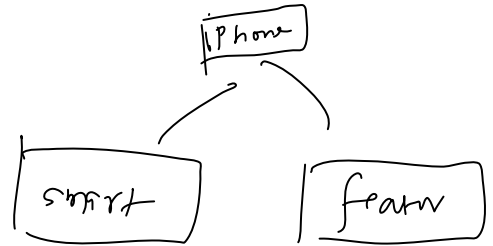
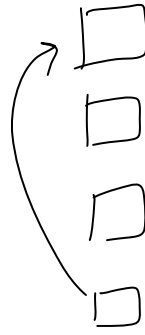
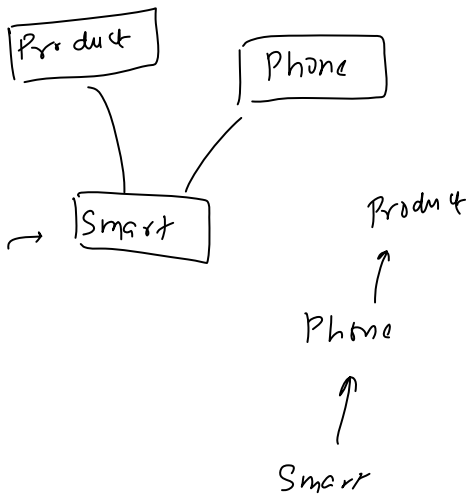


Multiple ✓

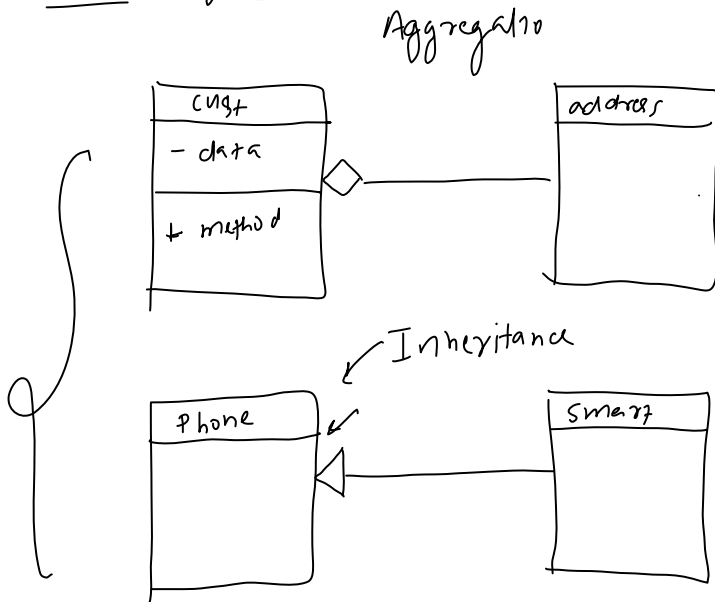


Hybrid

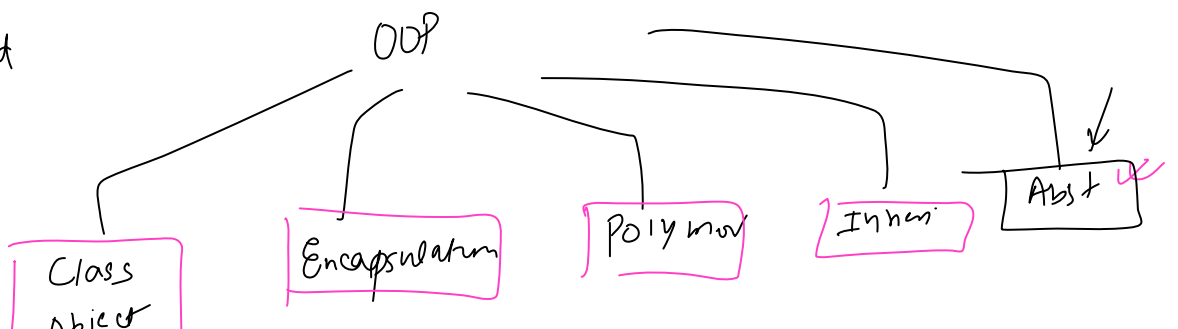




Class diagram



Construct
static method
inst vs static
super



Class
Object

Encapsulation

Python

(20) → 9 → DSA Algo → (8) classes

{ 17 18 19 20 }
SQL

Class 9 → Time Complexity
Class 10 → Dynamic Arrays (List)
Class 11 → LL
Class 12 → Stacks + Queue
Class 13 → Hashing - searching and sorting
Class 14 → Tree + Graph
Class 15 → Dynamic + Greedy
Class 16 → Exam + MCQ + qws

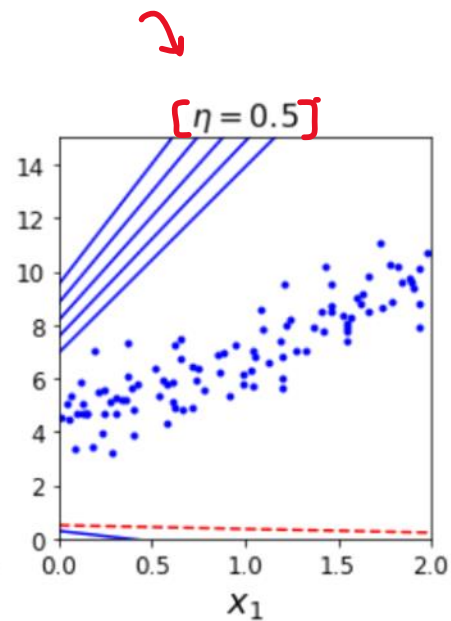
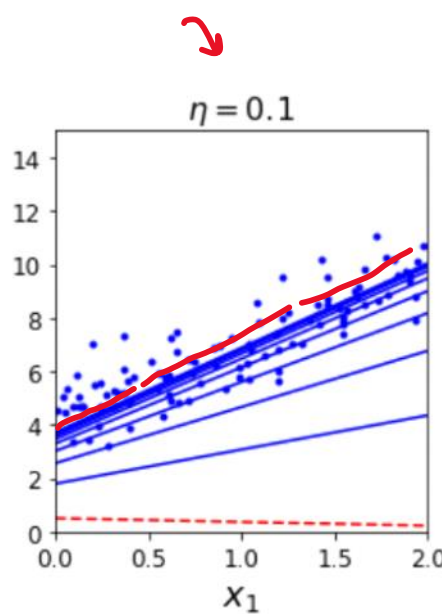
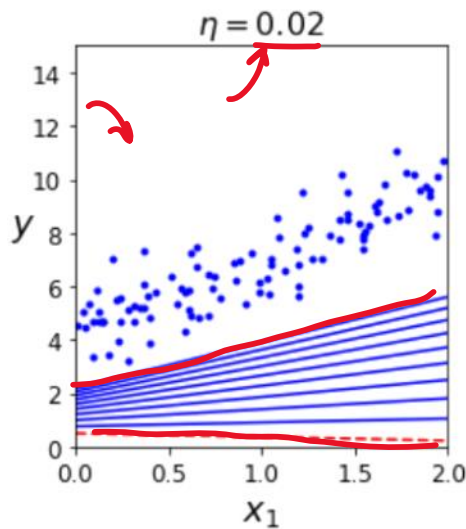
Python → Ex apl
Python → file
Python → exam

Few Discussions

Thursday, May 20, 2021 4:47 PM

1. Effect of Learning rate
2. The universality of Gradient Descent

epoch = 10



$b = 0$

$b = b_{old} - \eta \text{ (slope)}$

$\frac{dL}{db} = \left[\sum (y_i - \hat{y}_i)^2 \right] \text{ (LR)}$

LR $\rightarrow$ function

Adding m into the mix

Thursday, May 20, 2021 1:48 PM

Steps

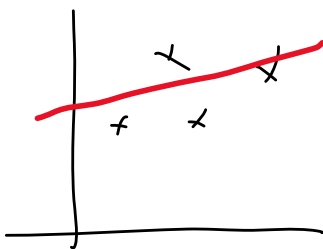
1) init random vals for m and b
 $m = 1$ and $b = 0$

2) epochs = 100, $\eta = 0.01$

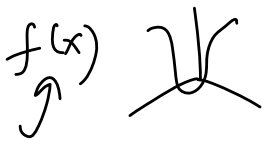
for i in epochs:

$$b = b - \eta \text{ slope}$$

$$m = m - \eta \text{ slope}$$



$$z = f(x, y)$$



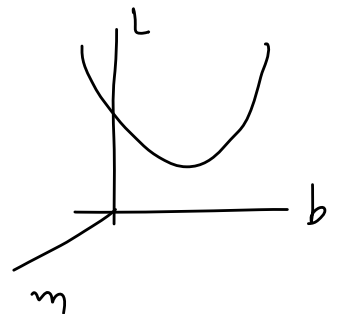
$$b = 0$$

$$L(m, b)$$

$$b_slope = \frac{\partial L}{\partial b}$$

$$m_slope = \frac{\partial L}{\partial m}$$

$$\sum (y_i - mx_i - b)^2$$



2D graph (m, b)

$$\begin{aligned} \frac{\partial L}{\partial b} &= -2 \sum (y_i - mx_i - b) \\ &= -2 \sum (y_i - mx_i - b) \\ &= \text{slope}_b \text{ at } b = 0 \end{aligned}$$

$$\begin{aligned} \frac{\partial L}{\partial m} &= -2 \sum (y_i - mx_i - b) x_i \\ &= -2 \sum (y_i - mx_i - b) x_i \\ \text{slope}_m \text{ at } m = 1 \end{aligned}$$

Code

Thursday, May 20, 2021 1:48 PM

Visualization 2

Thursday, May 20, 2021 1:52 PM

Effect of Learning Data

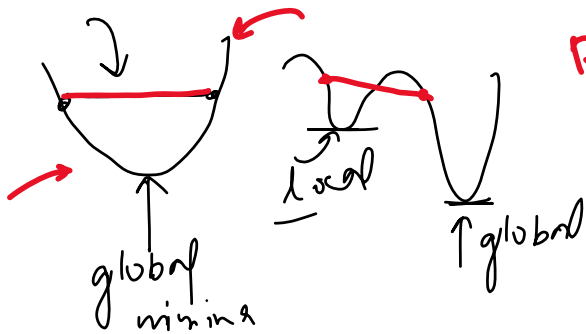
Thursday, May 20, 2021 1:53 PM

Effect of Loss Function

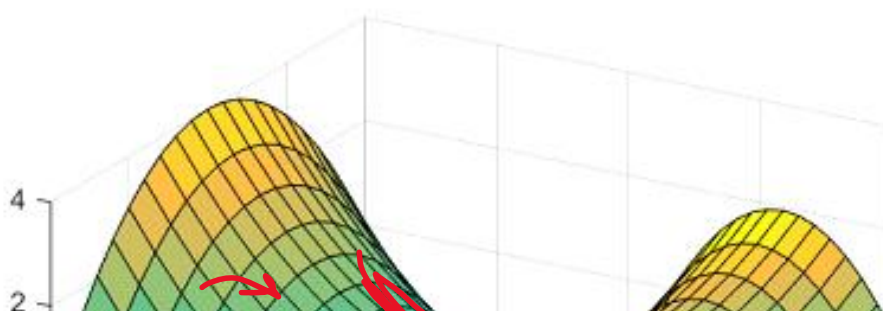
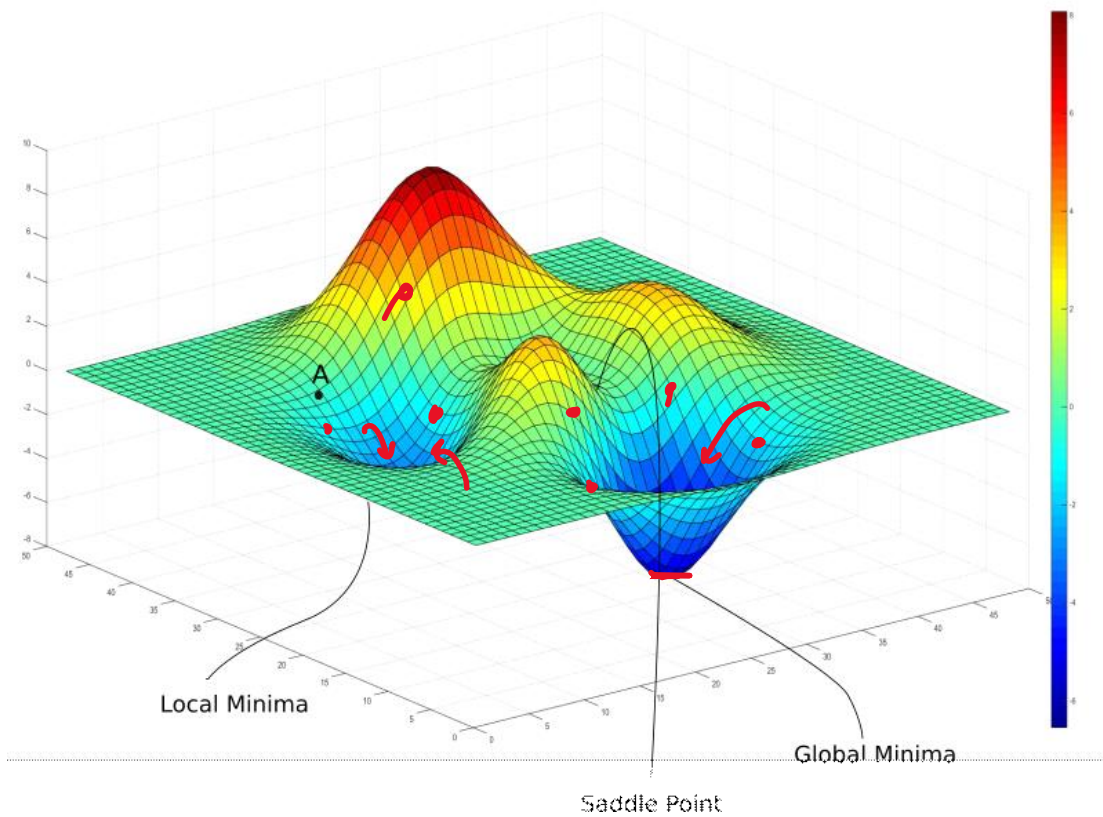
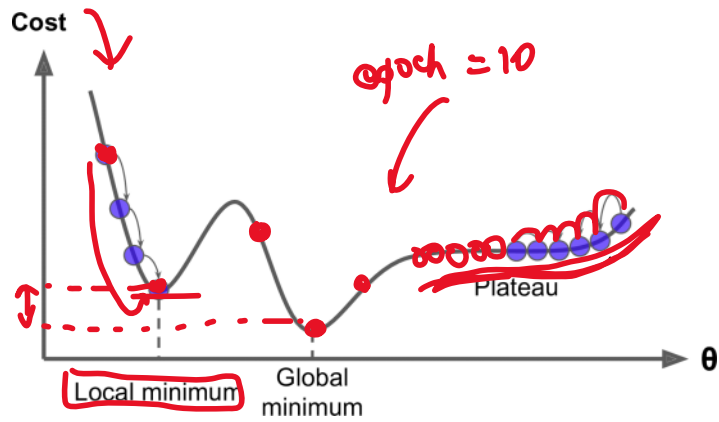
Thursday, May 20, 2021 1:53 PM

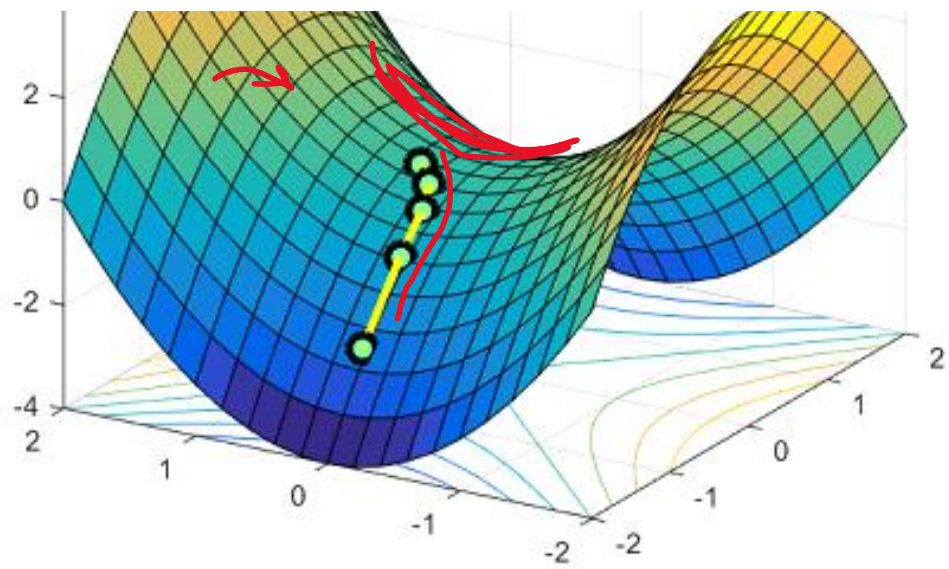
$$L = \sum (y_i - \hat{y}_i)^2$$

convex function



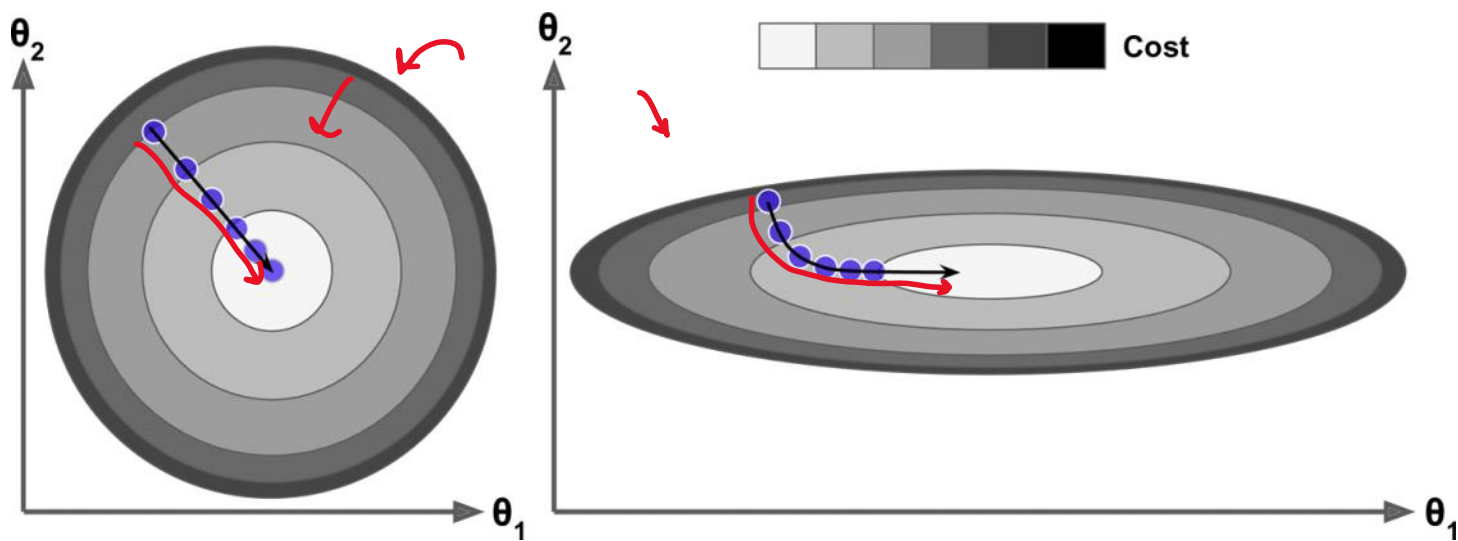
10 cols
deep
learn





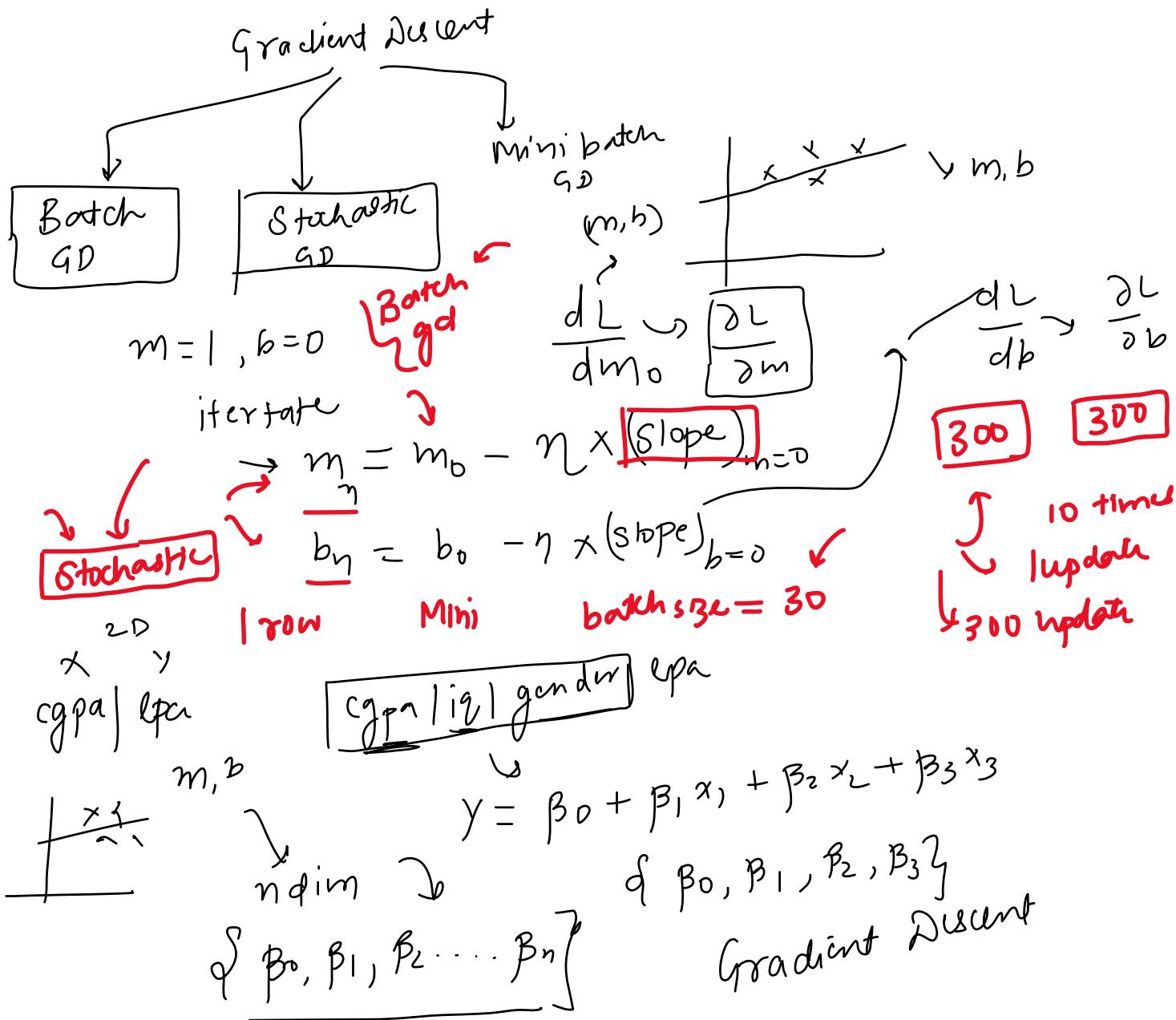
Effect of Data

Thursday, May 20, 2021 1:53 PM



Types of Gradient Descent

Saturday, May 22, 2021 1:30 PM



x cgp / iq / lpa

(2,3)

η -dim - dataset 3-cols
 $\hat{y}_i$
 $y = \beta_0 + \beta_1 x_1 + \beta_2 x_2$
 (lpa) (cgp) (iq)

1) Random values

$$\beta_0 = 0, \beta_1, \beta_2 = 1$$

2) epoch = 100, lr = 0.1

$$\begin{cases} \beta_0 = \beta_0 - \eta \text{slope} \\ \beta_1 = \beta_1 - \eta \text{slope} \\ \beta_2 = \beta_2 - \eta \text{slope} \end{cases}$$

MSE

$$L = \frac{1}{2} \sum_{i=1}^n (y_i - \hat{y}_i)^2$$

{ row=2, cols=2+1 }

$$\hat{y}_i = \beta_0 +$$

$$= \frac{1}{2} \left[(y_1 - \hat{y}_1)^2 + (y_2 - \hat{y}_2)^2 \right]$$

$$L = \frac{1}{2} \left[(y_1 - \beta_0 - \beta_1 x_{11} - \beta_2 x_{12})^2 + (y_2 - \beta_0 - \beta_1 x_{21} - \beta_2 x_{22})^2 \right]$$

	x_1	x_2	y
	8.1	9.3	3.2
	7.5	9.5	3.5

$$\hat{y} = \beta_0 + \beta_1 x_1 + \beta_2 x_2$$

$$\frac{\partial L}{\partial \beta_0} = \frac{1}{2} \left[2(y_1 - \hat{y}_1)(-1) + 2(y_2 - \hat{y}_2)(-1) \right]$$

$$\hat{y}_1 = \beta_0 + \beta_1 x_{11} + \beta_2 x_{12}$$

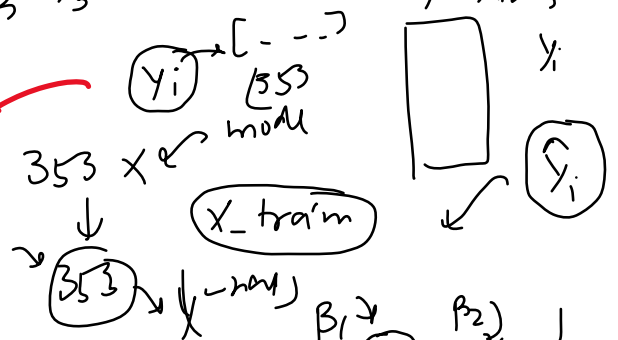
$$\hat{y}_2 = \beta_0 + \beta_1 x_{21} + \beta_2 x_{22}$$

$$\frac{\partial L}{\partial \beta_0} = - \left[(y_1 - \hat{y}_1) + (y_2 - \hat{y}_2) \right]$$

$$= - \frac{2}{n} \left[(y_1 - \hat{y}_1) + (y_2 - \hat{y}_2) + (y_3 - \hat{y}_3) + \dots + (y_n - \hat{y}_n) \right]$$

$$= \frac{-2}{n} \sum_{i=1}^n (y_i - \hat{y}_i) x_{i1}$$

$$= \frac{-2}{n} \sum_{i=1}^n (y_i - \hat{y}_i) x_{i1} = \frac{\partial L}{\partial \beta_0}$$



$$L = \frac{1}{2} \sum_{i=1}^n (y_i - \hat{y}_i)^2$$

$$L = \frac{1}{2} [(y_1 - \hat{y}_1)^2 + (y_2 - \hat{y}_2)^2]$$

$$L = \frac{1}{2} [(y_1 - \beta_0 - \beta_1 x_{11} - \beta_2 x_{12})^2 + (y_2 - \beta_0 - \beta_1 x_{21} - \beta_2 x_{22})^2]$$

$$\frac{\partial L}{\partial \beta_1} = \frac{1}{2} [2(y_1 - \hat{y}_1)(-x_{11}) + 2(y_2 - \hat{y}_2)(-x_{21})]$$

	x_{11}	x_{21}	x_{31}	x_{n1}
--	----------	----------	----------	----------

$$\frac{\partial L}{\partial \beta_1} = \frac{-2}{n} [(y_1 - \hat{y}_1)(x_{11}) + (y_2 - \hat{y}_2)(x_{21}) + (y_3 - \hat{y}_3)(x_{31}) + \dots + (y_n - \hat{y}_n)(x_{n1})]$$

$$\frac{\partial L}{\partial \beta_1} = \frac{-2}{n} \sum_{i=1}^n (y_i - \hat{y}_i) x_{i1}$$

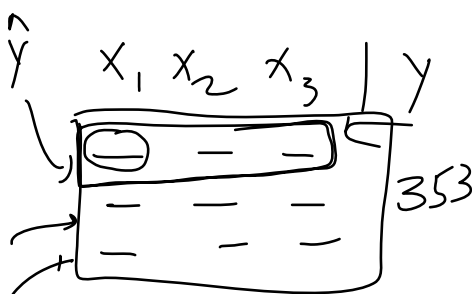
x_{ij} → 1 col data
 β_1 → values of 1 col.

$$\frac{\partial L}{\partial \beta_2} = \frac{-2}{n} \sum_{i=1}^n (y_i - \hat{y}_i) x_{i2}$$

m cols
 $\beta_0 - \beta_m$

$$\left\{ \frac{\partial L}{\partial \beta_m} \right\} = \frac{-2}{n} \sum_{i=1}^n (y_i - \hat{y}_i) x_{im}$$

code



$$\hat{y}_1 = \beta_0 + \beta_1 x_{11} + \beta_2 x_{12} + \beta_3 x_{13}$$

$$\hat{y}_2 = \beta_0 + \beta_1 x_{21} + \beta_2 x_{22} + \beta_3 x_{23}$$

$$\beta_0 + \text{np.dot}(X_{\text{train}}, \text{coef}) \hat{y} = \beta_0 + [x_{11} \ x_{12} \ x_{13}] \begin{bmatrix} \beta_1 \\ \beta_2 \\ \beta_3 \end{bmatrix}$$

$(353, 10)$ $(10, 1)$
 $\beta_0 + (353, 1) \rightarrow (353, 1)$
 $\hat{y} = \text{np.dot}(\text{coef}, X_{\text{train}}) + \beta_0$
 y_{pred}

	x_1	x_2	y	$\hat{y}$
1	1	2	5	6
2	3	4	7	8

$y - \hat{y} = [-1 \ -1]$

$$\frac{\partial L}{\partial \beta_1} = -\frac{2}{n} \sum_{i=1}^n (y_i - \hat{y}_i) x_{i1} = -\frac{2}{n} \begin{bmatrix} -1 & -1 \end{bmatrix} \begin{bmatrix} 1 \\ 3 \end{bmatrix}$$

$$\frac{\partial L}{\partial \beta_1} = [y - \hat{y}] \begin{bmatrix} 2 \\ n \end{bmatrix}$$

$[-4 \ -4] \times \frac{-2}{n} = 8$
 $(1, 2) \times (2, 2) = 8$
 $[y - \hat{y}] \begin{bmatrix} 1 & 2 \\ 3 & 4 \end{bmatrix} \times \frac{-2}{n}$

$$\frac{\partial L}{\partial \beta_1} \dots \frac{\partial L}{\partial \beta_0} = \left[(y_i - \hat{y}_i) X_{\text{train}} \right] \times \frac{-2}{n}$$

y_{train}

$$y_i = 353$$

$(353, 1)$ $(353, 10)$ $(1, 353)$

$(1, 353)$ $(353, 10)$
 $(1, 10) \times \begin{bmatrix} -2 \\ n \end{bmatrix} = (1, 10)$
 coef_del

The Problem with Batch GD

Tuesday, May 25, 2021 6:43 AM

$$\eta = 100000 \rightarrow 10^5$$

$$\text{col} \rightarrow 5 \rightarrow \text{6 coeffs } 10^2$$

$$\text{epoch} \rightarrow 50 \rightarrow 10^3 \quad 6$$

①

$$\begin{array}{r} 50 \\ 1 \rightarrow 1000 \\ 6 \rightarrow 6000 \end{array}$$

2) Hardware

total $\rightarrow 10^{10}$

slow on
big data

Deep Learning
CNN $\rightarrow$ images
RNN $\rightarrow$ text

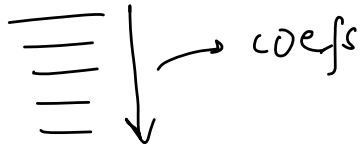
Stochastic GD

Tuesday, May 25, 2021 6:44 AM

5, 10

100 epochs

Batch



stochastic

row $\rightarrow$ update

n rows

1 epoch

n updates n rows

1 epoch

1 update

$\rightarrow$ random

steady state

faster

faster convergence

stochastic

Code

Tuesday, May 25, 2021 6:44 AM

$$\frac{\partial L}{\partial \beta_0} = \frac{-2}{m} \sum_{i=1}^m (y_i - \hat{y}_i)$$

$i = id x$

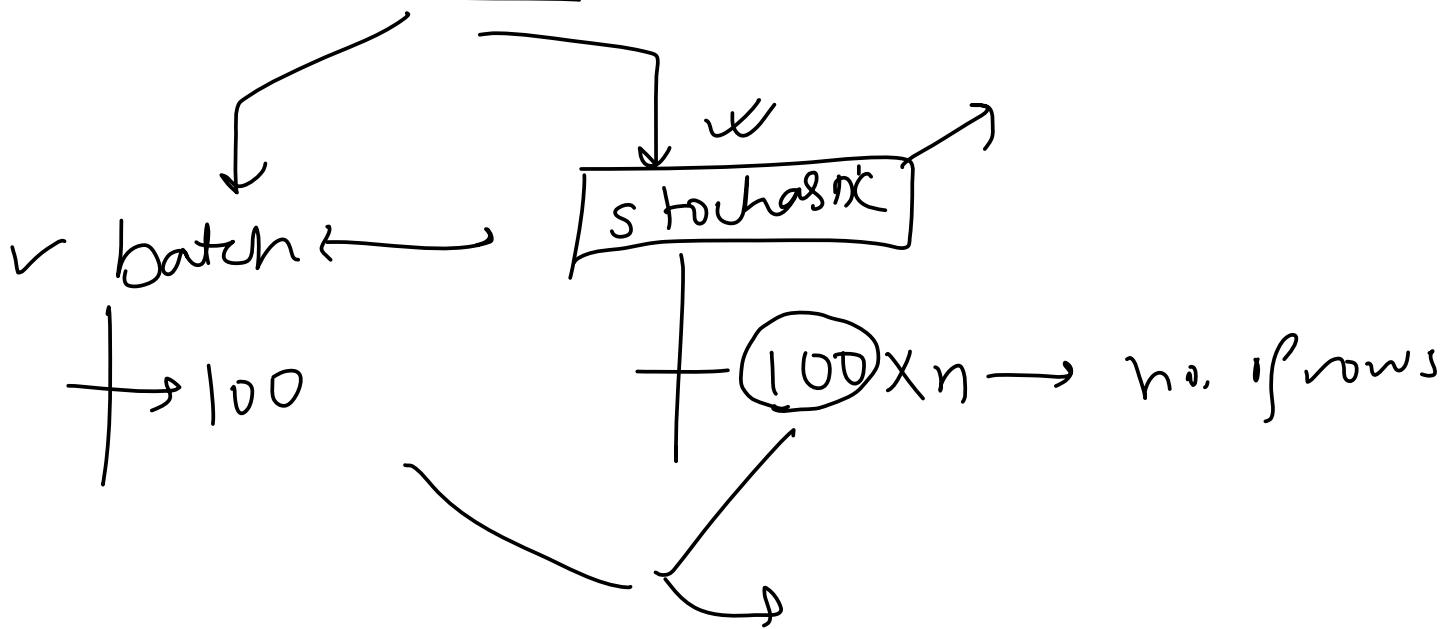
$$= -2 (y_i - \hat{y}_i)$$

Time Comparison

Tuesday, May 25, 2021 12:39 PM

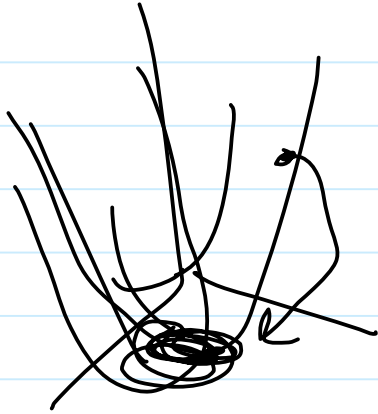
no. of epoch $\rightarrow$ fixed

$$E = 100$$



Visualizations

Tuesday, May 25, 2021 6:44 AM



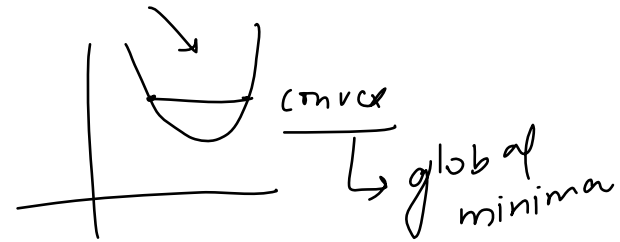
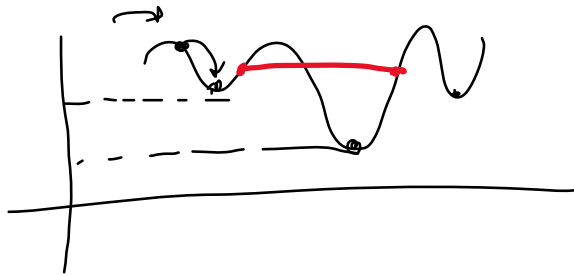
When to use Stochastic GD

Tuesday, May 25, 2021 6:44 AM

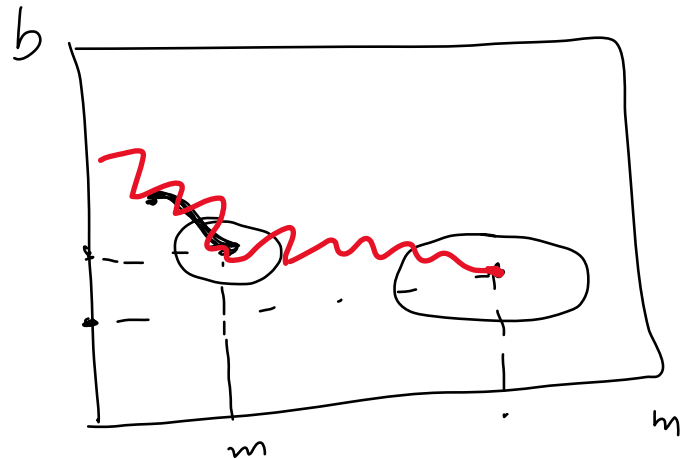
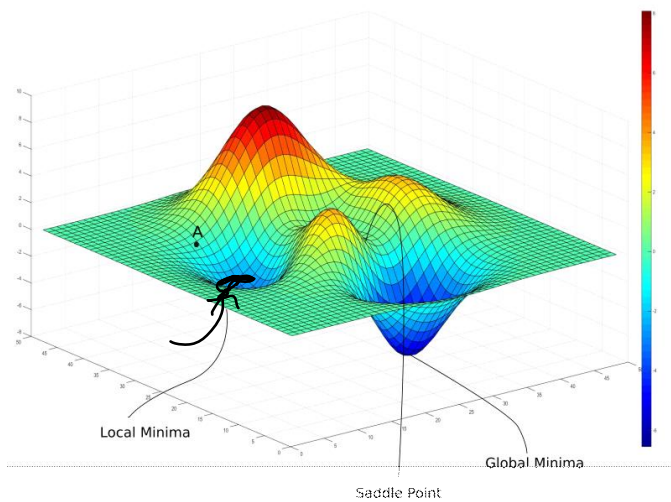
1) Big data $\rightarrow$ SGD

2) Non convex function

non convex



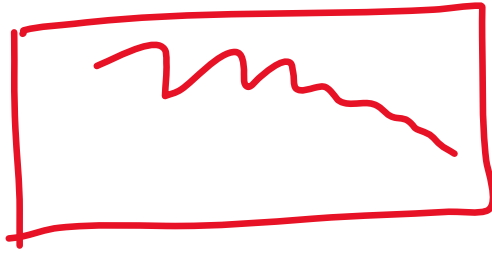
learning



Learning Schedules

Tuesday, May 25, 2021 7:00 AM

(PL)



$\eta = 100$

epoch = 1

→

$lr = 0.1$

$lr = 0.03$

```
t0, t1 = 5.50
def learning_rate(t):
    return t0 / (t + t1)

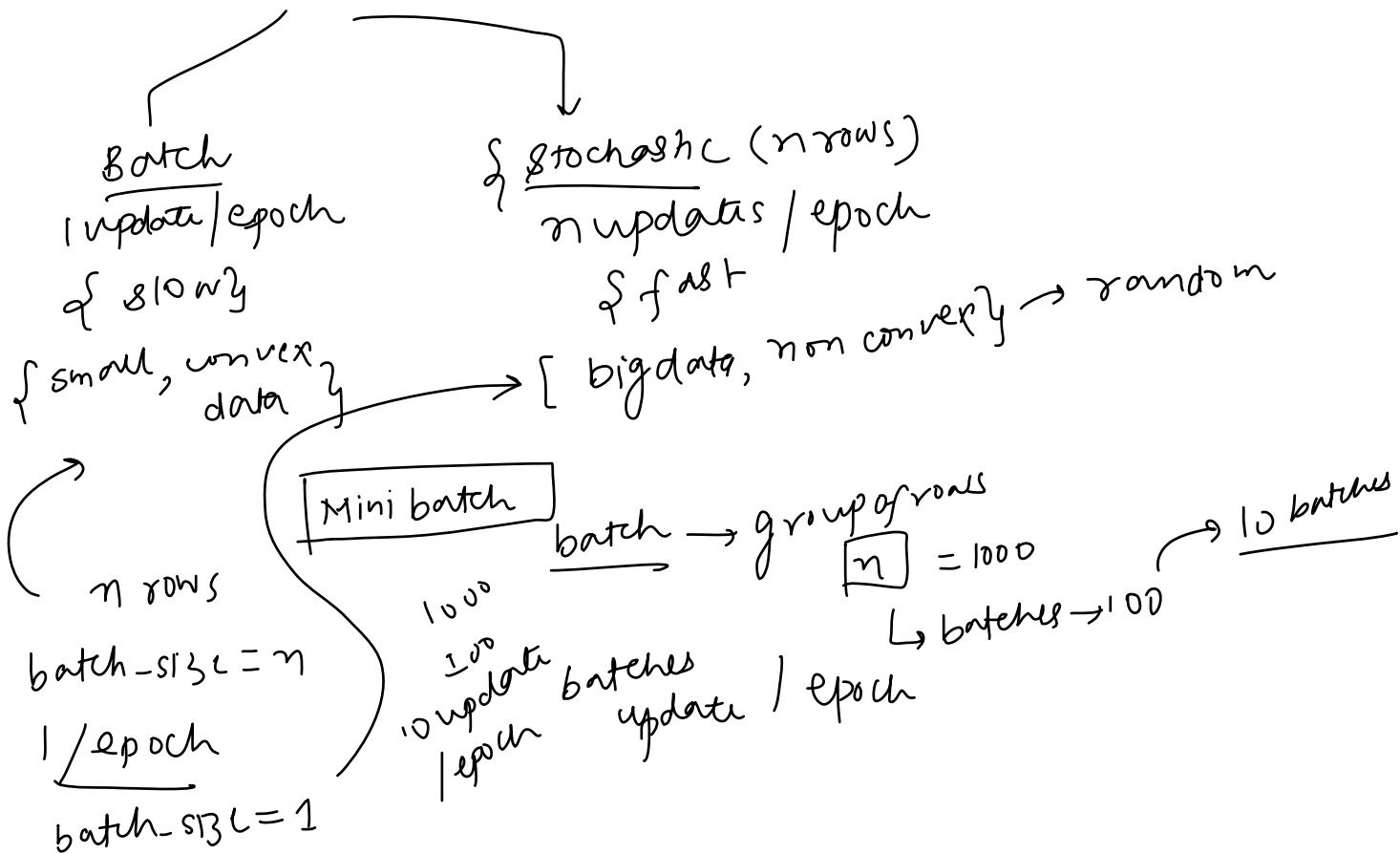
for i in range(epochs):
    for j in range(X.shape[0]):
        lr = learning_rate(i * X.shape[0] + j)
```

Sklearn Implementation

Tuesday, May 25, 2021 2:16 PM

Mini-Batch Gradient Descent

Wednesday, May 26, 2021 4:47 PM



Code

Wednesday, May 26, 2021 4:47 PM

Visualization

Wednesday, May 26, 2021 4:47 PM

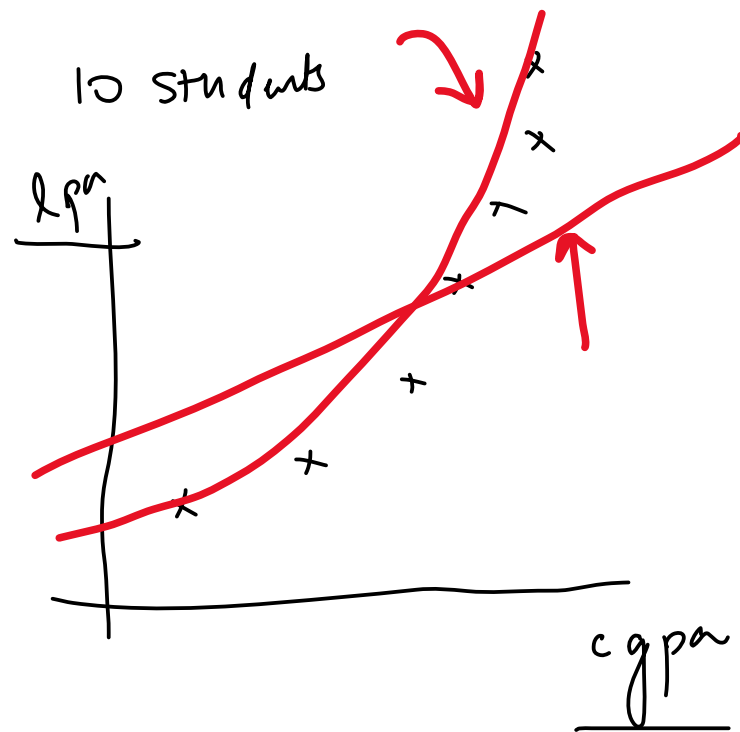
Sklearn Implmentation

Wednesday, May 26, 2021 4:47 PM

Introduction

Friday, May 28, 2021 1:47 PM

cgpa | lpa

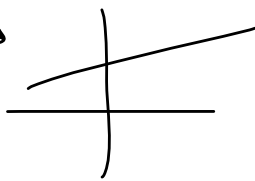


Intuition

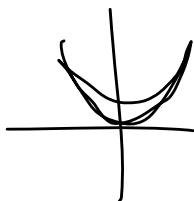
Friday, May 28, 2021 1:47 PM

polynomials

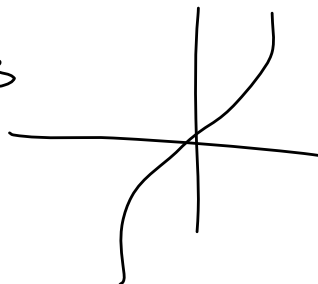
$$y = mx + b$$



$$y = x^2$$



$$y = x^3$$



$$x^3 + x^2$$

$$y = \textcircled{a}x^4 + \textcircled{b}x^3 + \textcircled{c}x^2 + dx + 2$$

↑ ↑ ↑ ↑
degree

$$\begin{matrix} 1 \\ x \end{matrix} \bigg| y \rightarrow y = mx + b$$

② ↑
3

transform polynomial → degree = 2

$$\begin{matrix} 3 \\ x^0 \end{matrix} \bigg| \begin{matrix} x^1 \\ 2 \\ 3 \end{matrix} \bigg| \begin{matrix} x^2 \\ 4 \\ 9 \end{matrix}$$

$$Y = \beta_0 + \beta_1 x^0 + \beta_2 x^2$$

Code

Friday, May 28, 2021

1:48 PM

Multiple Polynomial Regression

Friday, May 28, 2021 1:48 PM

Why is it linear

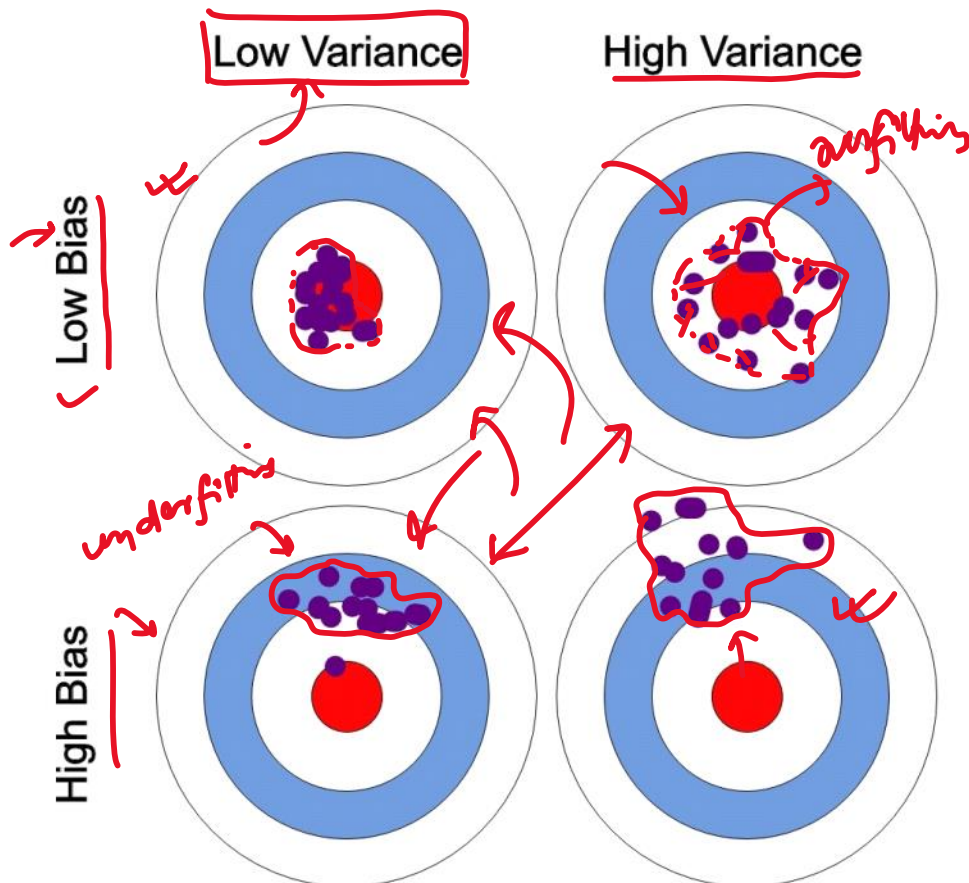
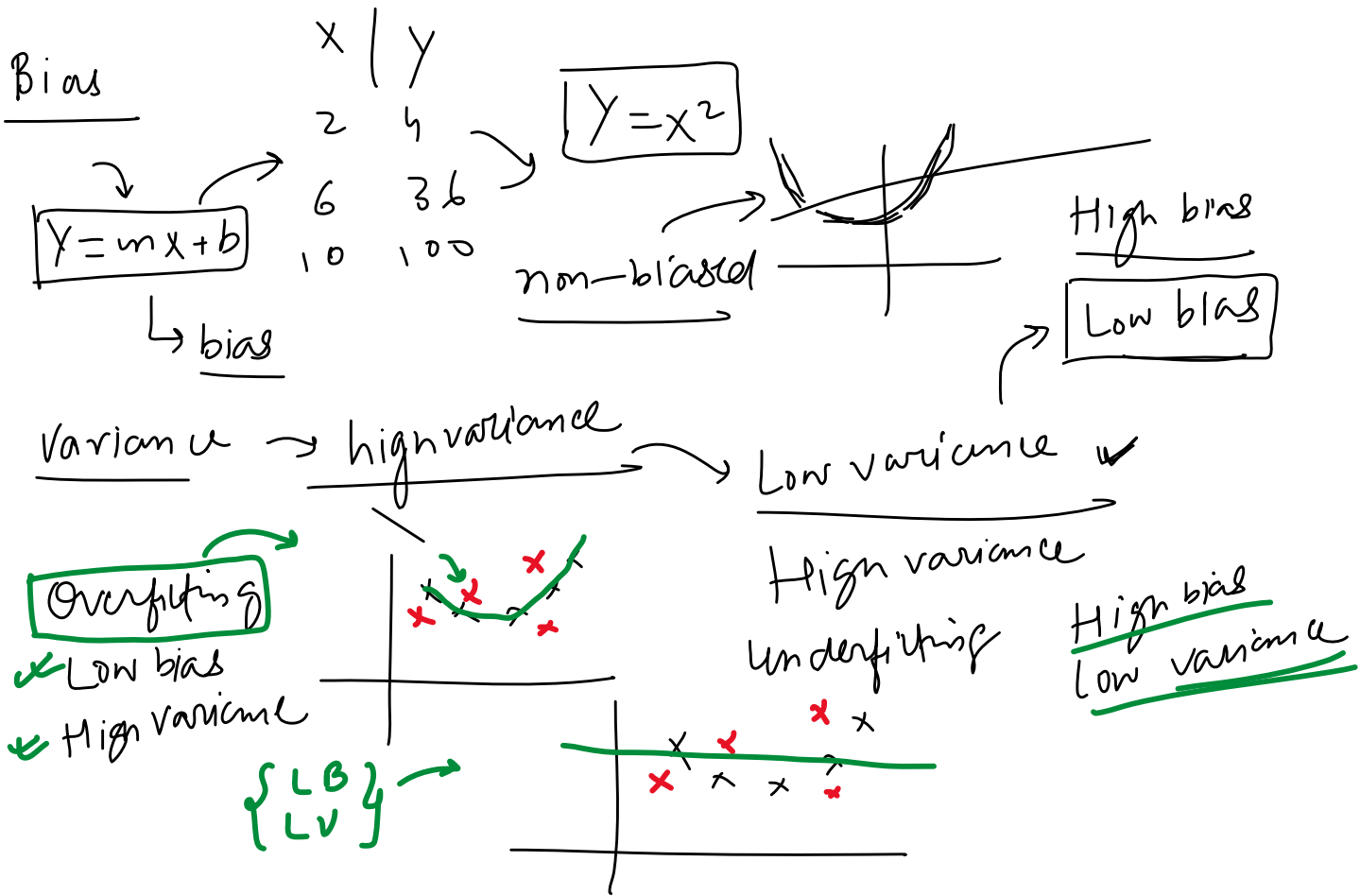
Friday, May 28, 2021 1:48 PM

When to use Polynomial Regression

Friday, May 28, 2021 1:48 PM

Bias And Variance

Tuesday, June 1, 2021 12:20 PM



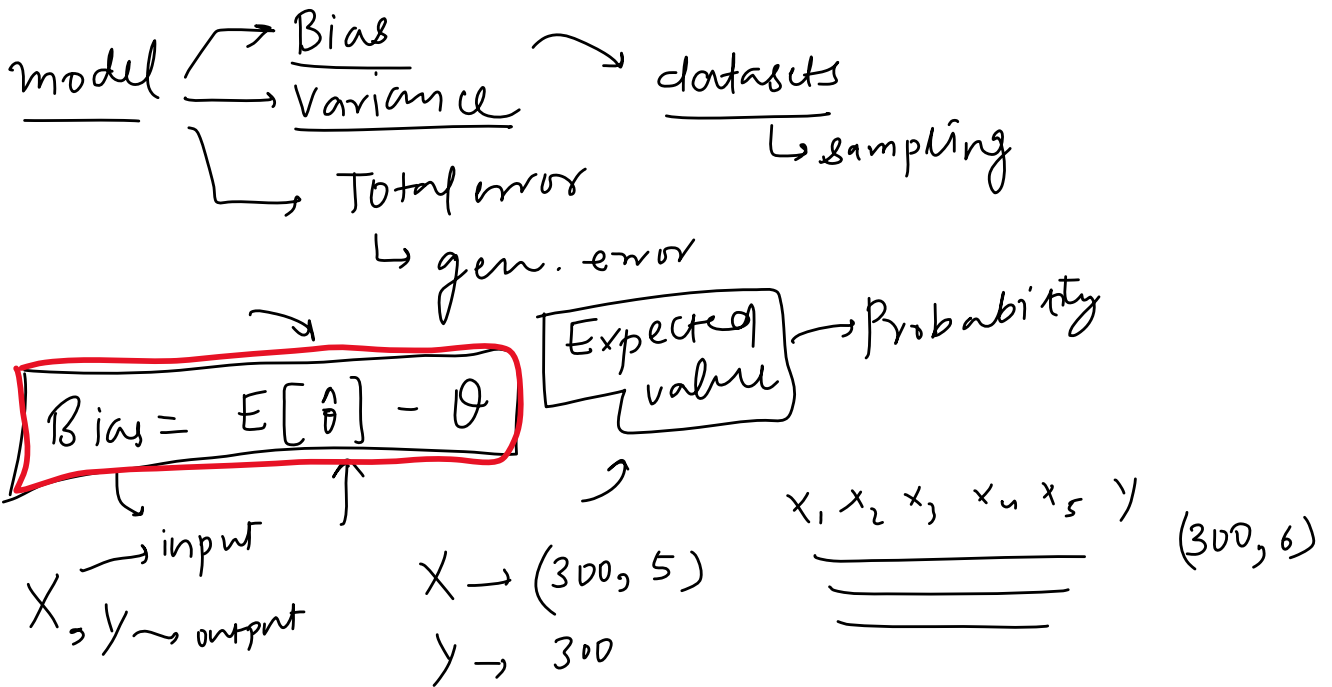
Practical example

under

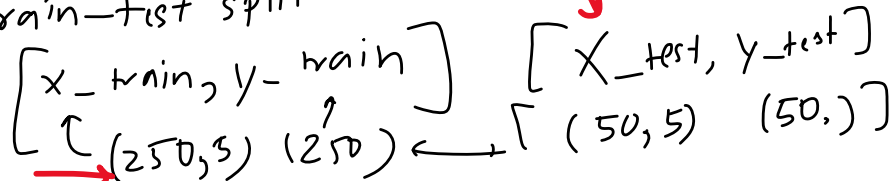
overfitting

How to calculate Bias and Variance

Tuesday, June 1, 2021 1:44 PM



Step 1 - train-test split



1
50 times

$x_{\text{train}}, y_{\text{train}}$

Sampling with replacement

50 datasets $x_{\text{train}}, y_{\text{train}}$

Step 2 $\rightarrow$ model $\rightarrow$ degree $\rightarrow$ 1

50 times

Training your model on different datasets

$m_1 - m_{50}$

Step 3 50 models

$x_{\text{test}}, y_{\text{test}}$ (50 rows)

$m_1 \rightarrow x_{\text{test}} \rightarrow y_{\text{pred1}}$

$m_2 \rightarrow x_{\text{test}} \rightarrow y_{\text{pred2}}$

50 values $\rightarrow [y_{\text{pred1}}]$
 $\rightarrow [y_{\text{pred2}}]$
 $\rightarrow [y_{\text{pred4}}]$

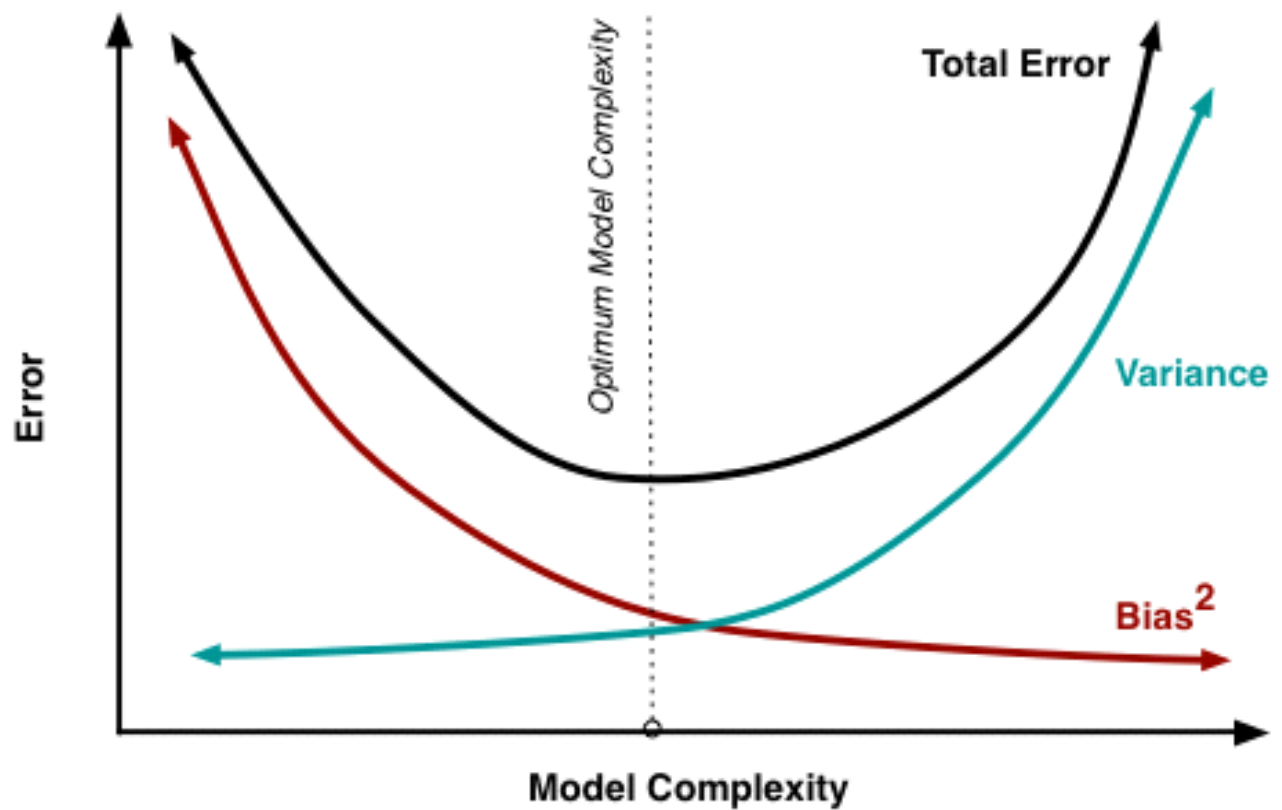
(50, 50)
cal with mean

Code

Monday, May 31, 2021 9:26 AM

Bias Variance Decomposition for Squared Error

Monday, May 31, 2021 9:26 AM

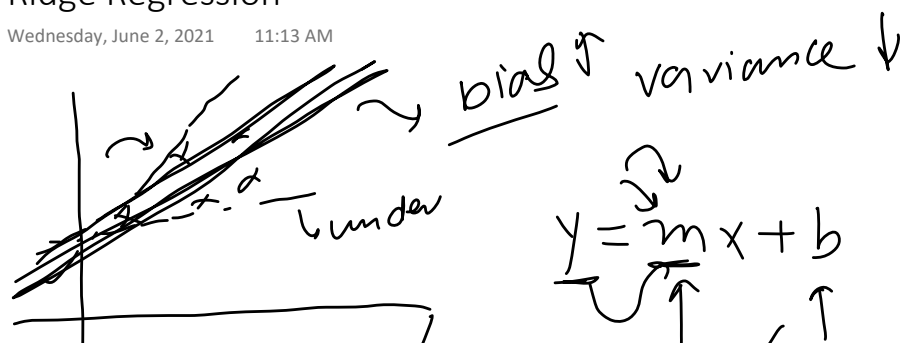


Relationship between Bias and Variance

Monday, May 31, 2021 9:26 AM

Ridge Regression

Wednesday, June 2, 2021 11:13 AM



$$L = \sum_{i=1}^n (y_i - \hat{y}_i)^2 + \lambda m^2$$

$$L = \sum_{i=1}^n (y_i - mx_i - b)^2 + \lambda m^2$$

$$b = \bar{y} - m\bar{x}$$

$\bar{y} \rightarrow y\text{-mean}$
 $\bar{x} \rightarrow x\text{-mean}$
 $m \rightarrow \text{slope}$

$$\frac{\partial L}{\partial b} = 0$$

$$\frac{\partial L}{\partial m} = 0$$

$$L = \sum_{i=1}^n (y_i - mx_i - \bar{y} + m\bar{x})^2 + \lambda m^2$$

$$\frac{\partial L}{\partial m} = 2 \sum_{i=1}^n (y_i - mx_i - \bar{y} + m\bar{x}) (-x_i + \bar{x}) + 2\lambda m = 0$$

$$= -2 \sum_{i=1}^n (y_i - \bar{y} - mx_i + m\bar{x}) (x_i - \bar{x}) + 2\lambda m = 0$$

$$= \lambda m - \sum_{i=1}^n [(y_i - \bar{y}) - m(x_i - \bar{x})] (x_i - \bar{x}) = 0$$

$$= \lambda m - \sum_{i=1}^n (y_i - \bar{y})(x_i - \bar{x}) - m (x_i - \bar{x})^2 = 0$$

$$= \lambda m - \sum_{i=1}^n (y_i - \bar{y})(x_i - \bar{x}) + m \sum_{i=1}^n (x_i - \bar{x})^2 = 0$$

$$= \lambda m + m \sum_{i=1}^n (x_i - \bar{x})^2 = \sum_{i=1}^n (y_i - \bar{y})(x_i - \bar{x})$$

$$m = \frac{\sum_{i=1}^n (y_i - \bar{y})(x_i - \bar{x})}{\sum_{i=1}^n (x_i - \bar{x})^2 + \lambda}$$

λ → hyperparam.
alpha

$$\lambda = 0 = \lambda = 10, 100$$

$m \rightarrow$

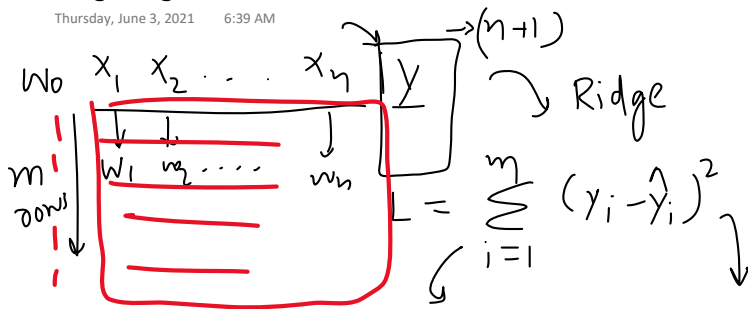
$$b = \bar{y} - m\bar{x}$$

Ridge Regression for 2D data

Thursday, June 3, 2021 6:38 AM

Code

Thursday, June 3, 2021 6:39 AM



$$L = (XW - Y)^T (XW - Y)$$

m rows

$$Y = \begin{bmatrix} y_1 \\ y_2 \\ \vdots \\ y_m \end{bmatrix} \quad W = \begin{bmatrix} w_0 \\ w_1 \\ \vdots \\ w_n \end{bmatrix} \quad X = \begin{bmatrix} 1 & x_{11} & x_{12} & \dots & x_{1n} \\ 1 & x_{21} & x_{22} & \dots & x_{2n} \\ \vdots & \vdots & \vdots & \ddots & \vdots \\ 1 & x_{m1} & x_{m2} & \dots & x_{mn} \end{bmatrix}$$

$(n \times 1)$

Normal LR $\rightarrow$ Ridge

$$L = (XW - Y)^T (XW - Y) + \lambda \|W\|^2$$

$\lambda w_0^2 + \lambda w_1^2 + \lambda w_2^2 + \dots + \lambda w_n^2$

$$[w_0 \ w_1 \ w_2 \ \dots \ w_n] \begin{bmatrix} w_0 \\ w_1 \\ \vdots \\ w_n \end{bmatrix} \quad W^T W \rightarrow \lambda (w_0^2 + w_1^2 + w_2^2 + \dots + w_n^2)$$

1 $(a-b)^T = a^T - b^T$

2 $L = (XW - Y)^T (XW - Y) + \lambda W^T W$

$\frac{dL}{dW} \quad (ab)^T = b^T a^T$

$$y^T W X \rightarrow W X^T y^T$$

3 $L = [(XW)^T - (Y)^T] (XW - Y) + \lambda W^T W$

4 $= (W^T X^T - Y^T) (XW - Y) + \lambda W^T W$

5 $= W^T X^T X W - \underbrace{W^T X^T Y - Y^T X W}_{=0} + Y^T Y + \lambda W^T W$

$$L = \underbrace{W^T X^T X W}_{\text{red}} - 2 \underbrace{W^T X^T Y}_{\text{red}} + \underbrace{Y^T Y}_{\text{red}} + \lambda W^T W$$

$$\frac{dL}{dW} = 2 X^T X W - 2 X^T Y + 0 + 2 \lambda W = 0$$

$$X^T X W + \lambda W = X^T Y$$

3 (4x4)

$$X^T X W + \lambda W = X^T y$$

$$(X^T X + \lambda \mathbf{I}) W = X^T y$$

3 (4x4)
(n x 1, n x 1)

$$\begin{bmatrix} 0 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & 1 \end{bmatrix}$$

$$W = (X^T X + \lambda \mathbf{I})^{-1} X^T y$$

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

$$\begin{bmatrix} 1 \\ 1 \end{bmatrix}$$

Code

Thursday, June 3, 2021 6:39 AM

Vector form loss

$$L = \sum_{i=1}^n (y_i - \hat{y}_i)^2 \rightarrow L = (XW - y)^T (XW - y) + \lambda \|W\|^2$$

$$L = (XW - y)^T (XW - y) + \lambda W^T W$$

$X = \begin{bmatrix} 1 & x_{11} & x_{12} & x_{13} & \dots & x_{1n} \\ 1 & x_{21} & x_{22} & x_{23} & \dots & x_{2n} \\ \vdots & \vdots & \vdots & \vdots & \ddots & \vdots \\ 1 & x_{m1} & x_{m2} & x_{m3} & \dots & x_{mn} \end{bmatrix}$
 $\begin{matrix} \text{rows} \\ \text{columns} \end{matrix}$
 $y = \begin{bmatrix} y_1 \\ y_2 \\ \vdots \\ y_m \end{bmatrix}$
 $W = \begin{bmatrix} w_0 \\ w_1 \\ w_2 \\ \vdots \\ w_n \end{bmatrix}$

$w_0, w_1, \dots, w_n$ (parameters)

$$w_0 = w_0 - \eta \frac{\partial L}{\partial w_0} \quad w_1 = w_1 - \eta \frac{\partial L}{\partial w_1} \quad \dots \quad w_n = w_n - \eta \frac{\partial L}{\partial w_n}$$

$$w_{\text{new}} = w_{\text{old}} - \eta \left[\frac{\partial L}{\partial w} \right] \rightarrow \text{gradient} \begin{bmatrix} \frac{\partial L}{\partial w_0} \\ \frac{\partial L}{\partial w_1} \\ \vdots \\ \frac{\partial L}{\partial w_n} \end{bmatrix}$$

$$L = \frac{1}{2} (XW - y)^T (XW - y) + \frac{1}{2} \lambda W^T W$$

$$= \frac{1}{2} (W^T X^T - y^T) (XW - y) + \frac{1}{2} \lambda W^T W$$

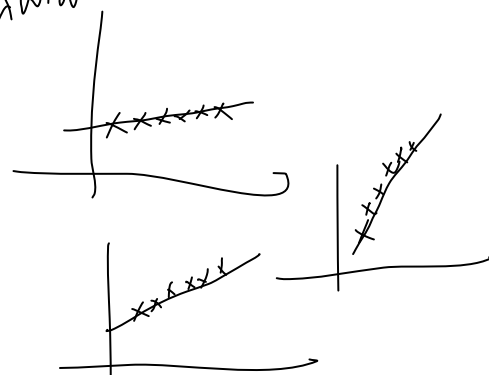
$$= \frac{1}{2} \left[W^T X^T X W - \underbrace{W^T X^T y}_{\text{red}} - \underbrace{y^T X W}_{\text{red}} + y^T y \right] + \frac{1}{2} \lambda W^T W$$

$$= \frac{1}{2} \left[W^T X^T X W - \underbrace{2 y^T X W}_{\text{red}} + y^T y \right] + \frac{1}{2} \lambda W^T W$$

$$\frac{dL}{dW} = \frac{1}{2} \left[X^T X W - \underbrace{X^T y}_{\text{red}} \right] + \frac{1}{2} \lambda W$$

$$= \underbrace{X^T X W - X^T y + \lambda W}_{\text{red}} = \frac{dL}{dW} \left(\frac{\partial L}{\partial W} \right)$$

$$W = \begin{bmatrix} w_0 & w_1 & \dots & w_n \\ 0 & 1 & \dots & 1 \end{bmatrix} \text{ starting}$$



$$W = \begin{bmatrix} w_0 & w_1 & \dots & w_n \\ 0 & 1 & \dots & 1 \end{bmatrix} \quad \text{starting}$$

$$\underbrace{W}_{\text{Epochs}} = \underbrace{W}_{\text{Epochs}} - \underbrace{\eta}_{\text{epoch times}} \frac{dL}{dw}$$

$W \rightarrow$ final answer
 epoch times

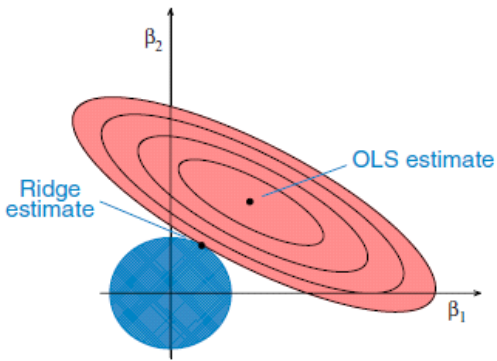
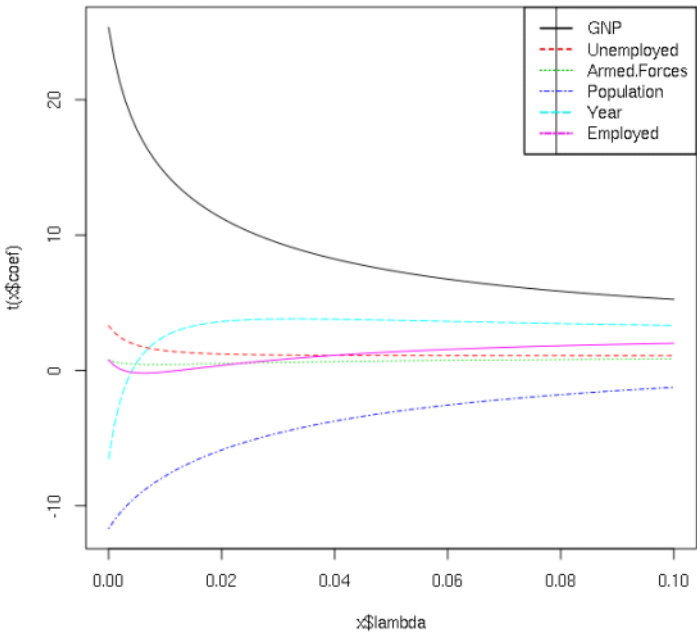
$$\hookrightarrow W = W - \eta \frac{dL}{dw}$$

$$\frac{dL}{dw} = \underline{X^T X W} - \underline{X^T y} + \underline{\lambda W}$$

Notes

Friday, June 4, 2021 4:47 PM

Why is it called ridge



5 Key Understandings

Saturday, June 5, 2021

4:20 PM

$$\sum_{i=1}^n (y_i - \hat{y}_i)^2 + \boxed{\lambda \|w\|^2}$$

↑

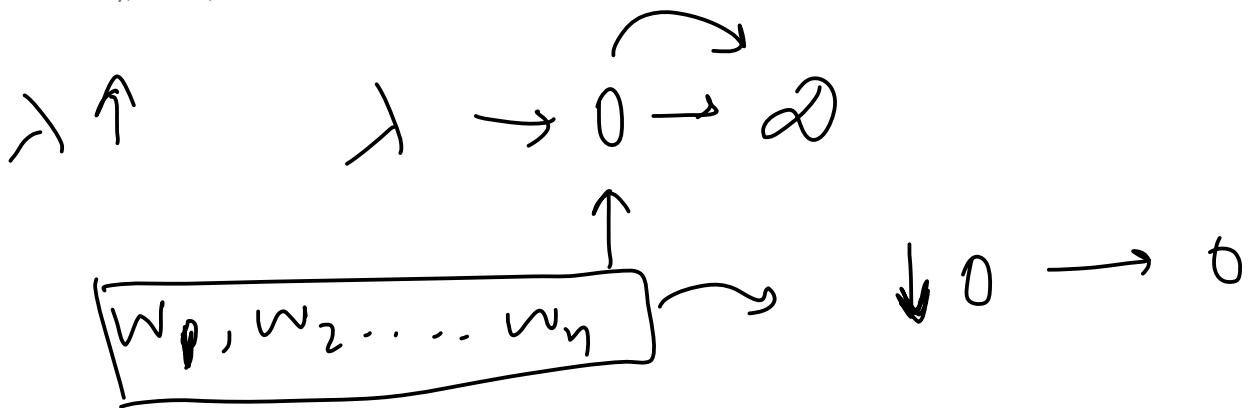
↓ $\lambda (w_1^2 + w_2^2 + \dots + w_n^2)^2$

Shrinkage
coef

↓
Overfitting ↓

1. How the coefficients get affected?

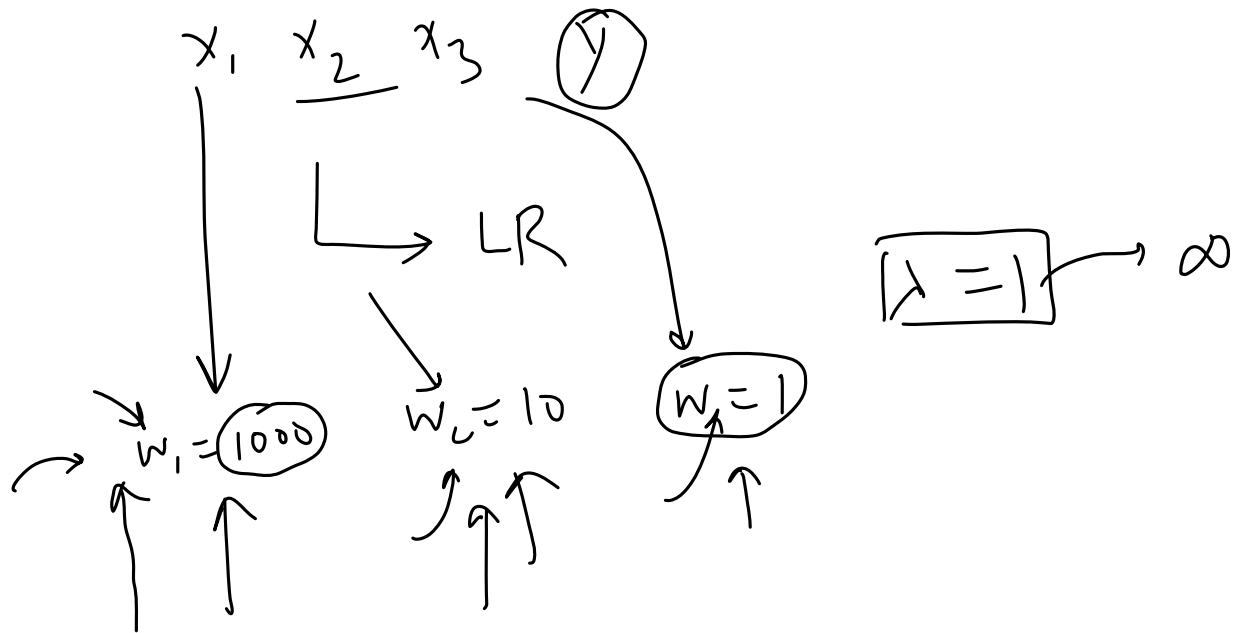
Saturday, June 5, 2021 4:20 PM



2. Higher Values are impacted more

Saturday, June 5, 2021 4:21 PM

Never reaches 0



3. Bias Variance Tradeoff

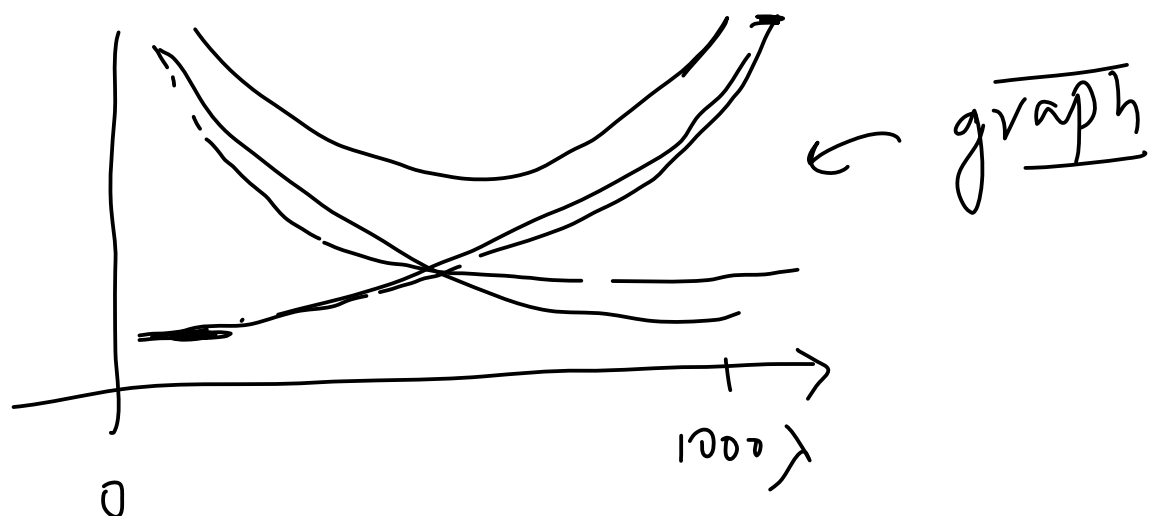
Saturday, June 5, 2021 4:21 PM

Bias Variance

λ $\downarrow \uparrow$

Bias $\downarrow$ overfit Variance $\uparrow$

Bias $\uparrow$ underfitting Variance $\downarrow$



4. Impact on the Loss Function

Saturday, June 5, 2021 4:21 PM

$$\lambda \rightarrow \mathcal{L} = \sum_{i=1}^n (y_i - \hat{y}_i)^2 + \lambda ||w||^2$$

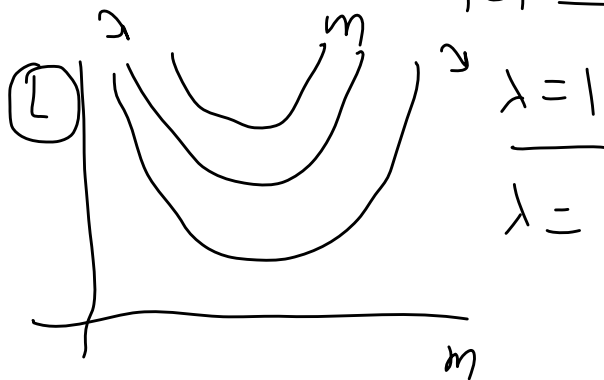
$x, y \rightarrow m, b$ $\rightarrow b$ is constant

$\downarrow$ $\underline{b = \bar{b}}$

m

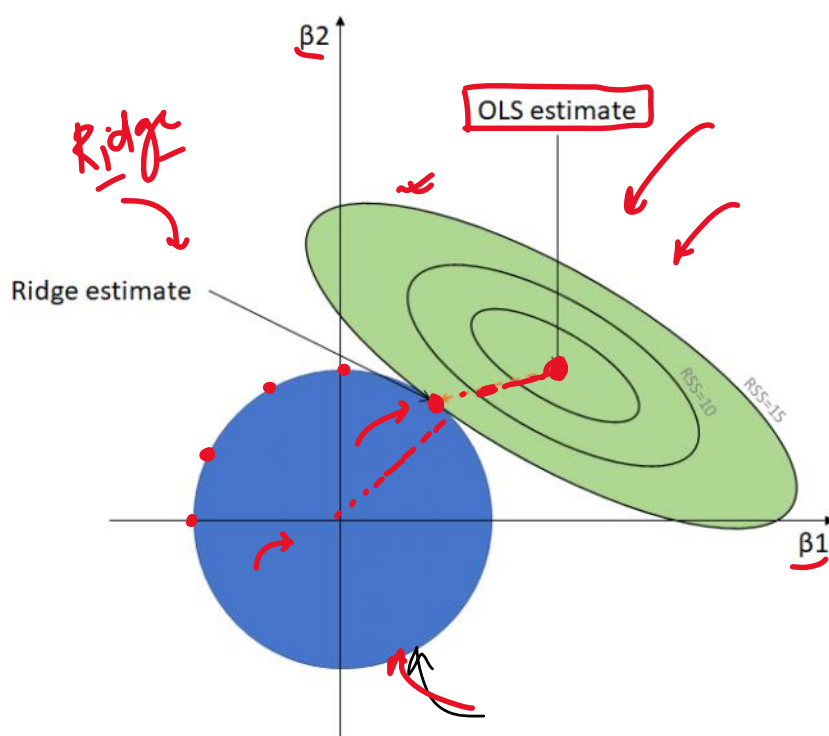
$$b = -2.29$$

$$\mathcal{L} = \sum_{i=1}^n (\bar{y}_i - \underline{m} x_i)^2 + \lambda \underline{m}^2$$



5. Why called Ridge

Saturday, June 5, 2021 4:22 PM



Hard constraint
Ridge constraint

2 coef, β_1 β_2 β_0
 $L = \text{MSE} + \lambda \|\mathbf{w}\|^2$

Contour

$(y_i - \hat{y}_i)^2$
 $\sum_{i=1}^n \frac{(y_i - (\beta_0 + \beta_1 x_{i1} + \beta_2 x_{i2}))^2}{\text{MSE}}$
 $\lambda (\beta_1^2 + \beta_2^2)$

Practical Tip

Monday, June 7, 2021 1:20 PM

Use ridge when there are more than 2 input cols

Ridge

$x_1, x_2, x_3, \dots$
→
coefs

x_1, y

$>= 2$

Lasso Regression

Thursday, June 10, 2021 6:42 AM

L1 Regularization | overfitting → L2 reg

$$L = \sum_{i=1}^n (y_i - \hat{y}_i)^2 + \lambda \|w\|^2$$

↳ $\lambda (w_1^2 + w_2^2 + \dots + w_n^2)$

$$\lambda \geq 0$$

→ 0 $w_1 \rightarrow w_n \rightarrow$ coeff

overfitting | under

alpha

↳ Lasso

L1 norm

$$[|w_1| + |w_2| + |w_3| + \dots + |w_n|]$$

$$L = \text{MSE} + \lambda \|w\|$$

→ code implement (sklearn)

→ Key point →

→ difference Ridge Vs Lasso

underfitting

$\lambda = 0 \rightarrow$ Linear

$x_1, x_2, \dots, x_n$ overfitting degree

ridge
lasso → $\lambda \rightarrow \uparrow$

dim <

decrease

feature selection

lasso

$\lambda \uparrow$

overfitting ↓

Bias ↑

Variance ↑

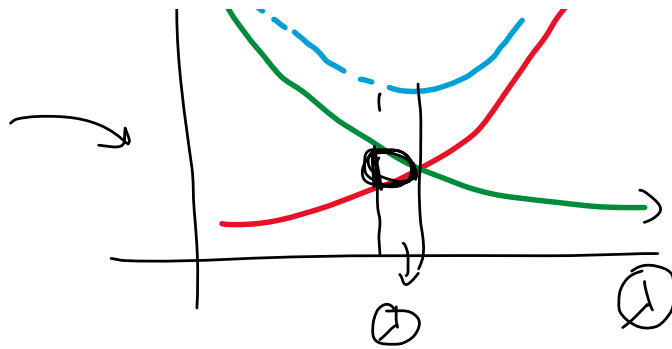
Bias ↓

values

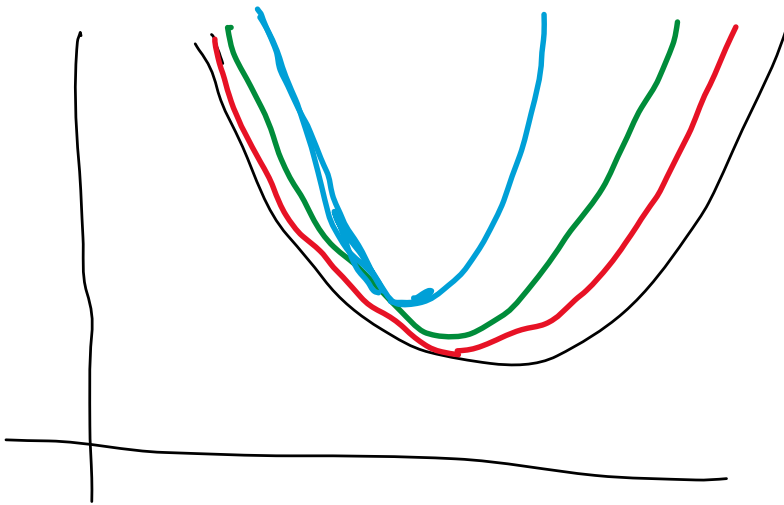
Variance ↓



Red → bias
green - variance
blue - total error



σ
blue - total error



alpha	age	sex	bmi	bp	s1	s2	s3	s4	s5	s6
0.0000	-9.160885	-205.462260	516.684624	340.627341	-895.543609	561.214533	153.884786	126.734316	861.121400	52.419828
0.0001	-9.118336	-205.337133	516.880570	340.556792	-883.415291	551.553259	148.578680	125.355917	856.480254	52.467627
0.0010	-8.763583	-204.321125	518.371729	339.975385	-787.690766	475.274718	106.786540	114.632063	819.739542	52.872100
0.0100	-6.401088	-198.669767	522.048548	336.348363	-383.709187	152.663678	-66.060583	75.611090	659.869402	55.828128
0.1000	6.642753	-172.242166	485.523872	314.682122	-72.939323	-80.590053	-174.466515	83.616653	484.363285	73.584154
1.0000	42.242217	-57.305508	282.170831	198.061386	14.363544	-22.551274	-136.930053	102.023193	260.104308	98.552274
10.0000	21.174004	1.659796	63.659772	48.493240	18.421492	12.875448	-38.915435	38.842464	61.612405	35.505355
100.0000	2.858979	0.629452	7.540604	5.849997	2.710879	2.142134	-4.834047	5.108223	7.448466	4.576129
1000.0000	0.295726	0.069290	0.769004	0.597829	0.282900	0.225936	-0.495607	0.527031	0.761497	0.471029
10000.0000	0.029674	0.006995	0.077054	0.059915	0.028412	0.022715	-0.049686	0.052870	0.076321	0.047241

Lasso
sparsity

$\lambda \uparrow \quad w \rightarrow 0$

single $x|y \rightarrow$

alpha	age	sex	bmi	bp	s1	s2	s3	s4	s5	s6
0.0000	-9.160885	-205.462260	516.684624	340.627341	-895.543596	561.214523	153.884780	126.734314	861.121395	52.419828
0.0001	-9.071288	-205.337332	516.780313	340.539730	-888.652320	555.952271	150.585260	125.453044	858.639860	52.379002
0.0010	-8.264924	-204.213177	517.641106	339.751339	-826.653342	508.609613	120.899583	113.924518	836.314382	52.011583
0.0100	-1.361404	-192.944226	526.348511	332.649058	-430.205495	191.277876	-44.048113	68.990747	688.384976	47.939528
0.1000	0.000000	-113.976046	526.737112	292.635423	-82.691928	-0.000000	-152.691332	0.000000	551.077200	7.169852
1.0000	0.000000	0.000000	363.882636	27.278420	0.000000	0.000000	-0.000000	0.000000	336.135971	0.000000
10.0000	0.000000	0.000000	0.000000	0.000000	0.000000	0.000000	-0.000000	0.000000	0.000000	0.000000
100.0000	0.000000	0.000000	0.000000	0.000000	0.000000	0.000000	-0.000000	0.000000	0.000000	0.000000
1000.0000	0.000000	0.000000	0.000000	0.000000	0.000000	0.000000	-0.000000	0.000000	0.000000	0.000000
10000.0000	0.000000	0.000000	0.000000	0.000000	0.000000	0.000000	-0.000000	0.000000	0.000000	0.000000

feature
selection

$$m = \frac{\sum_{i=1}^n (y_i - \bar{y})(x_i - \bar{x})}{\sum_{i=1}^n (x_i - \bar{x})^2 + \lambda}$$

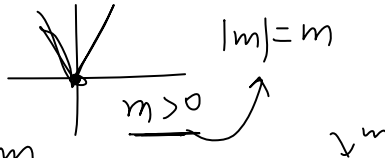
simple
 $x|y$

$$y = mx + b$$

$$b = \bar{y} - m\bar{x}$$

$\bar{y} \rightarrow \text{mean}(y)$
 $\bar{x} \rightarrow \text{mean}(x)$

$$m = \frac{\sum_{i=1}^n (y_i - \bar{y})(x_i - \bar{x})}{\sum_{i=1}^n (x_i - \bar{x})^2}$$



$$b = \bar{y} - m\bar{x}$$

$m = ?$

$$L = \sum_{i=1}^n (y_i - \hat{y}_i)^2 + \lambda |m|$$

$$\frac{d}{dm} \sum_{i=1}^n (y_i - m x_i - \bar{y} + m \bar{x})^2 + 2\lambda |m|$$

$$\frac{dL}{dm} = \sum_{i=1}^n (y_i - m x_i - \bar{y} + m \bar{x})^2 + 2\lambda |m| = 2 \sum (y_i - m x_i - \bar{y} + m \bar{x})(-x_i + \bar{x}) + 2\lambda = 0$$

$$-2 \sum [(y_i - \bar{y}) - m(x_i - \bar{x})](x_i - \bar{x}) + 2\lambda = 0$$

$$m \sum (x_i - \bar{x})^2 = \sum (y_i - \bar{y})(x_i - \bar{x}) - \lambda$$

$$-2 \geq [(y_i - \bar{y}) (x_i - \bar{x}) - m(x_i - \bar{x})^2] + \lambda = 0$$

$$- \sum [(y_i - \bar{y}) (x_i - \bar{x}) - m(x_i - \bar{x})^2] + \lambda = 0$$

$$- \sum (y_i - \bar{y}) (x_i - \bar{x}) + m \sum (x_i - \bar{x})^2 + \lambda = 0$$

$$m = \frac{\sum (y_i - \bar{y}) (x_i - \bar{x}) - \lambda}{\sum (x_i - \bar{x})^2}$$

Lasso coeff sparsity

for $m > 0$

$$m = \frac{\sum (y_i - \bar{y}) (x_i - \bar{x}) - \lambda}{\sum (x_i - \bar{x})^2}$$

for $m \leq 0$

$$m = \frac{\sum (y_i - \bar{y}) (x_i - \bar{x})}{\sum (x_i - \bar{x})^2}$$

for $m \leq 0$

$$m = \frac{\sum (y_i - \bar{y}) (x_i - \bar{x}) + \lambda}{\sum (x_i - \bar{x})^2}$$

$m = \frac{YX - \lambda}{X^2}$

$\lambda = 100$
 $X^2 = 50$
 $m = \frac{100 - \lambda}{50}$

$\lambda = 0$ $m = 2$
 $\lambda = 10$ $m = \frac{9}{5}$
 $\lambda = 50$ $m = 1$

$\lambda = 100$ $m = 0$
 $\lambda > 100$ $m = -1$

$$m = \frac{YX + \lambda}{X^2} = \frac{100 + \lambda}{50}$$

$$= \frac{100 + 150}{50} =$$

$$m = 5$$

$m < 0$

$$m = \frac{\sum (y_i - \bar{y}) (x_i - \bar{x}) + \lambda}{\sum (x_i - \bar{x})^2}$$

$$m = \frac{-100 - \lambda}{50}$$

$$= \frac{-100 - 150}{50} = -5$$

$$m = \frac{-100 + \lambda}{50}$$

$\lambda = 0$ $m = -2$
 $\lambda = 50$ $m = -1$

$\lambda = 150$, $m = -5$

$$m = 1$$

$\lambda = 100$ $m = 0 \rightarrow 1$
 $\lambda = 150$

Lasso $\lambda \rightarrow$ numerical

$m = \frac{\sum (y_i - \bar{y}) (x_i - \bar{x})}{\sum (x_i - \bar{x})^2 + \lambda}$

Ridge $\lambda < 0$
 $= 0$

$\lambda = 10000000$

Denominator

Ridge

$$\lambda(w_1^2 + w_2^2 + \dots + w_n^2)$$

$$L = \underbrace{\sum_{i=1}^n (y_i - \hat{y}_i)^2}_{\text{mse}} + \underbrace{\lambda \|w\|_2^2}_{\text{overfitting}}$$

$$\lambda \uparrow \quad w \downarrow 0$$

100 cols

Lasso

$$\lambda(|w_1| + |w_2| + \dots + |w_n|)$$

$$L = \underbrace{\sum_{i=1}^n (y_i - \hat{y}_i)^2}_{\text{mse}} + \underbrace{\lambda \|w\|_1}_{L_1}$$

feature selection

$$\lambda \uparrow \quad w \rightarrow 0$$

EN Reg

$$L = \sum (y_i - \hat{y}_i)^2 + a \|w\|^2 + b \|w\|$$

$$\begin{aligned} \lambda &= a + b \\ \text{ratio} &= \frac{a}{a+b} \\ l_1 &= \frac{a}{\lambda} \\ a &= l_1 \times \lambda \\ b &= \lambda - a \end{aligned}$$

$$\lambda = 1 \quad \text{ratio} = 0.5$$

$$a = 0.5 \quad b = 0.5$$

↑

$$\text{ratio} > 0.9$$

$$\lambda, \text{ratio}$$

$$l_1, \lambda$$

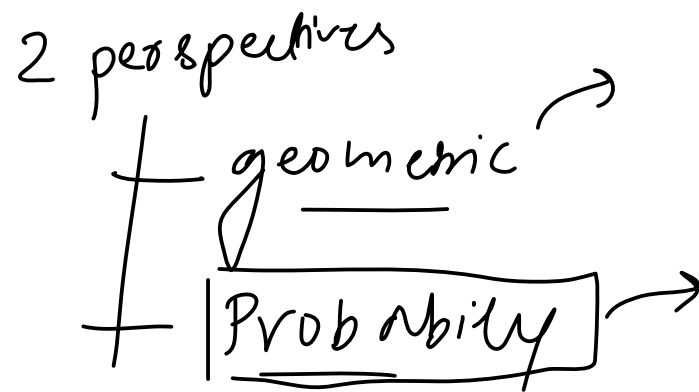
90% rid 10% lasso

Input cols → multicollinearity (ElasticNet)

$$\underline{x_1} \mid \underline{x_2}$$

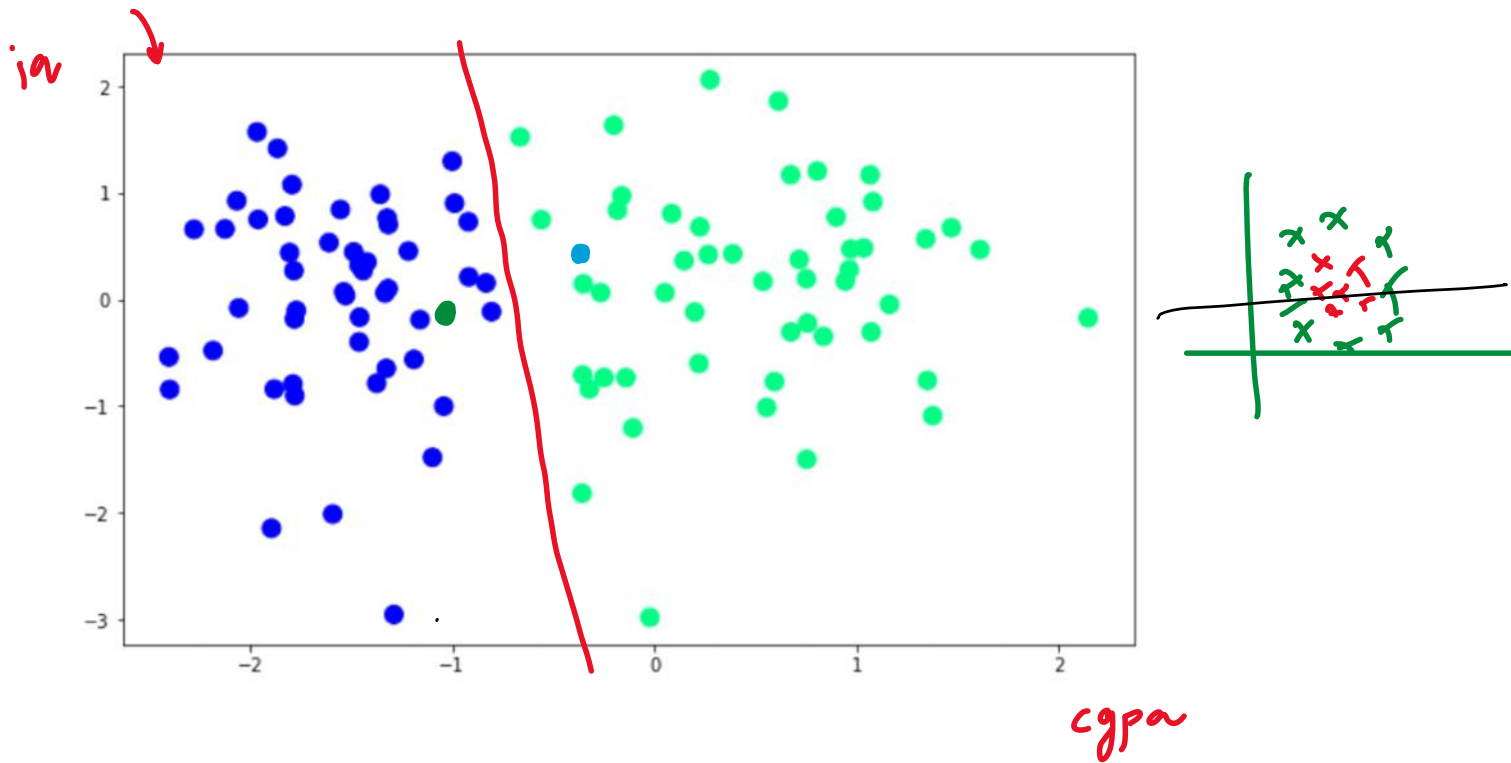
Introduction

Tuesday, June 15, 2021 12:07 PM



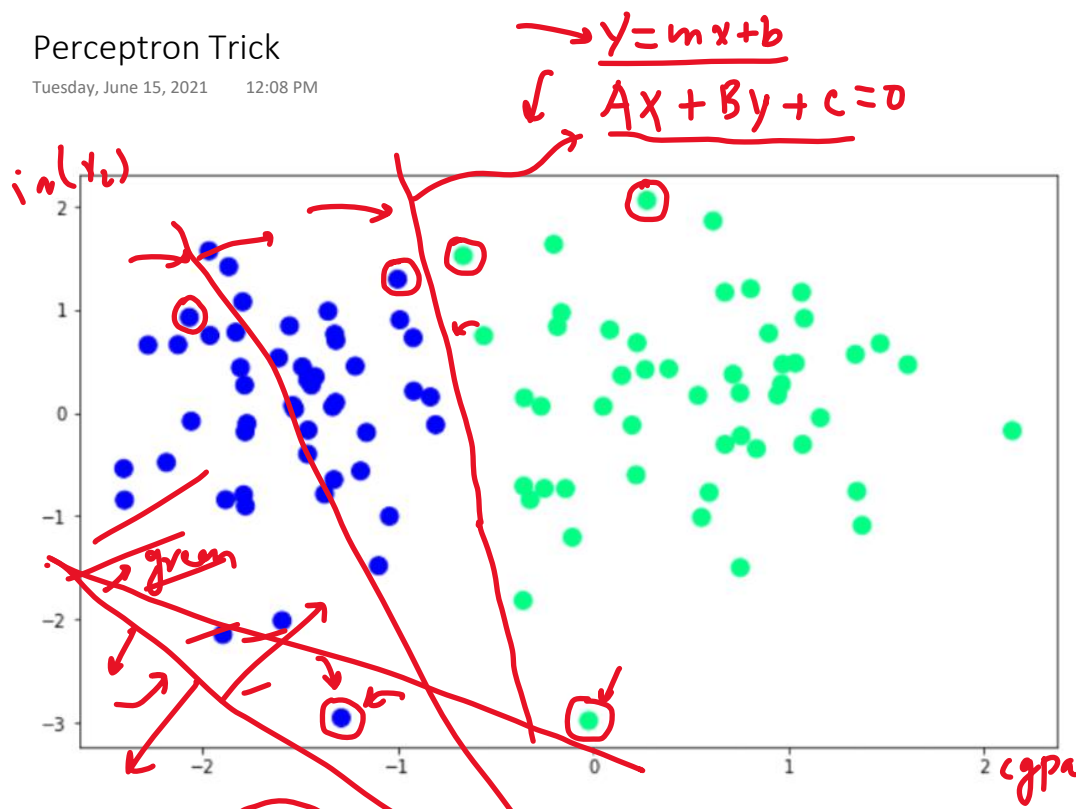
Requirement

Tuesday, June 15, 2021 12:07 PM



Perceptron Trick

Tuesday, June 15, 2021 12:08 PM



$$Ax_1 + Bx_2 + C = 0$$

$$\downarrow \quad \downarrow$$

$$Ax_1 + Bx_2 + Cx_3 + d = 0$$

$$\boxed{A, B, C}$$

$$\underline{A} = 1, \underline{B} = 1$$

$$\underline{C} = 0$$

loop $\rightarrow$ 1000

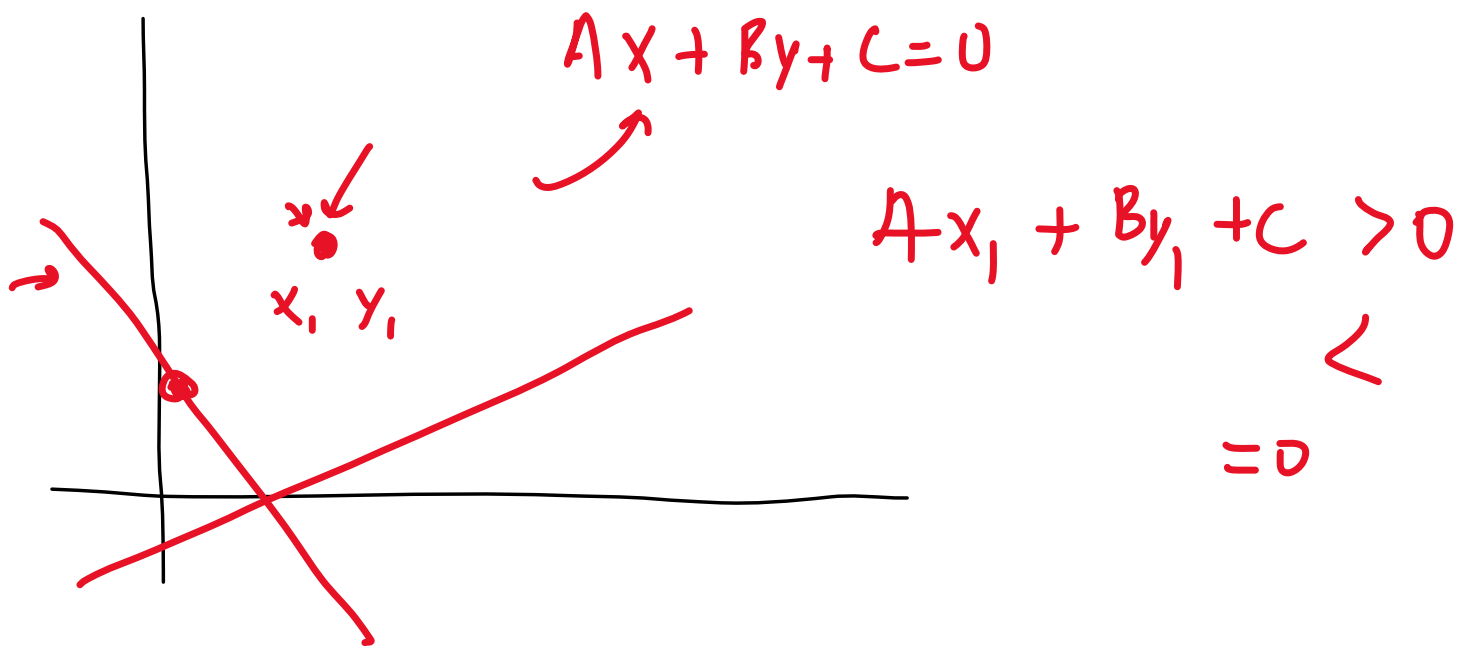
$\rightarrow$ random student

$\rightarrow$ 1000

while converge

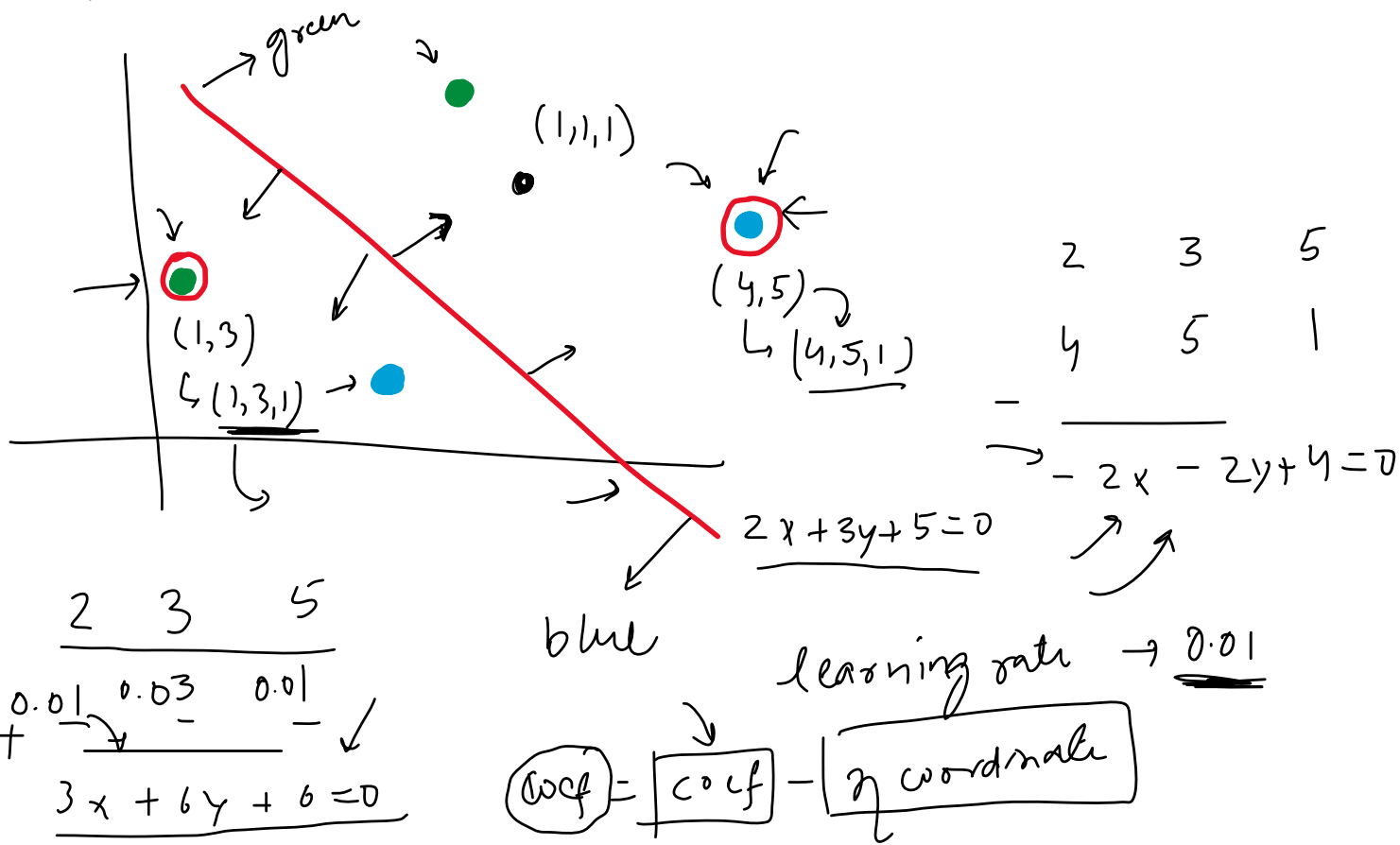
How to label regions?

Tuesday, June 15, 2021 1:12 PM



Transformations

Tuesday, June 15, 2021 1:31 PM



Algorithm

Tuesday, June 15, 2021 2:31 PM

x_0	(x_1)	(x_2)	y
	cgpa	ia	placed
1	7.5	61	1
1	8.9	109	1
1	7.0	81	0

$$Ax + By + C = 0$$

$$w_0 + w_1 x_1 + w_2 x_2 = 0$$

$$w_0 = C, w_1 = A, w_2 = B$$

$$w_0 x_0 + w_1 x_1 + w_2 x_2 = 0$$

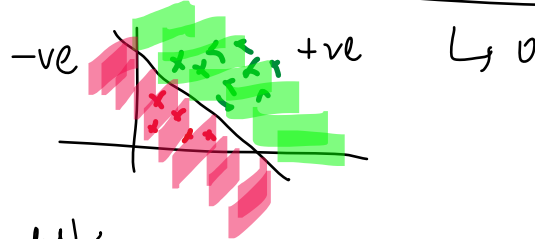
$$\sum_{i=0}^2 w_i x_i = 0$$

$$[w_0, w_1, w_2]$$

$$w_0 x_0 + w_1 x_1 + w_2 x_2 = 0 \quad [w_0, w_1, w_2] \begin{bmatrix} x_0 \\ x_1 \\ x_2 \end{bmatrix}$$

$$= \geq 0 \rightarrow 1$$

$$< 0 \rightarrow 0$$



epoch $\rightarrow$ 1000, $\eta = 0.01$

for i in range (epochs):

randomly select a student

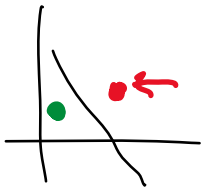
if $x_i \in N$ and $\sum_{i=0}^2 w_i x_i \geq 0$

$$w_{\text{new}} = w_{\text{old}} - \eta x_i$$

if $x_i \in P$ and $\sum_{i=0}^2 w_i x_i < 0$

$$w_{\text{new}} = w_{\text{old}} + \eta x_i$$

$\begin{cases} x_i \in N & w_i x_i < 0 \\ x_i \in P & w_i x_i \geq 0 \end{cases}$



$$[w_0, w_1, w_2]$$

Simplified Algo

Tuesday, June 15, 2021 2:44 PM

$$\text{if } x_i \in N \text{ and } \sum w_i x_i \geq 0$$

$$w_n = w_0 - \eta x_i$$

$$\text{if } x_i \in P \text{ and } \sum w_i x_i < 0$$

$$w_n = w_0 + \eta x_i$$

for i in 1000
random student

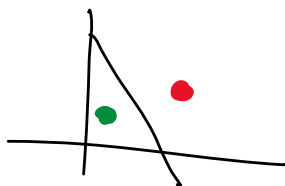
$$w_n = w_0 + \eta (y_i - \hat{y}_i) x_i$$

$$w_n = w_0$$

$$w_n = w_0 + \eta x_i$$

$$w_n = w_0 - \eta x_i$$

x_i	$\hat{y}_i$	$y_i - \hat{y}_i$
1	1	0
0	0	0
1	0	1
0	1	-1



for i in range(epochs):
select a random student (i)
 $w_n = w_0 + \eta (y_i - \hat{y}_i) x_i$

A, B, C

$$A x + B y + C = 0$$

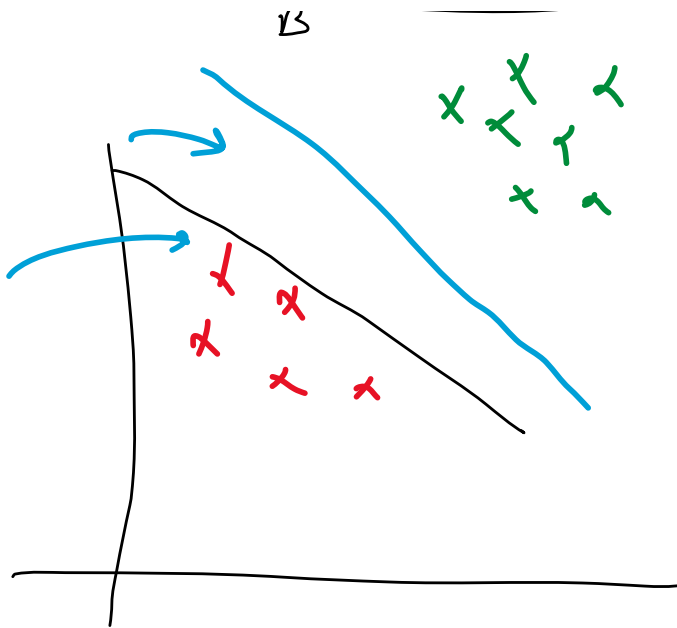
$$y = m x + b$$

$$m = -\frac{A}{B}$$

$$C = -\frac{C}{B}$$

Def [-] []

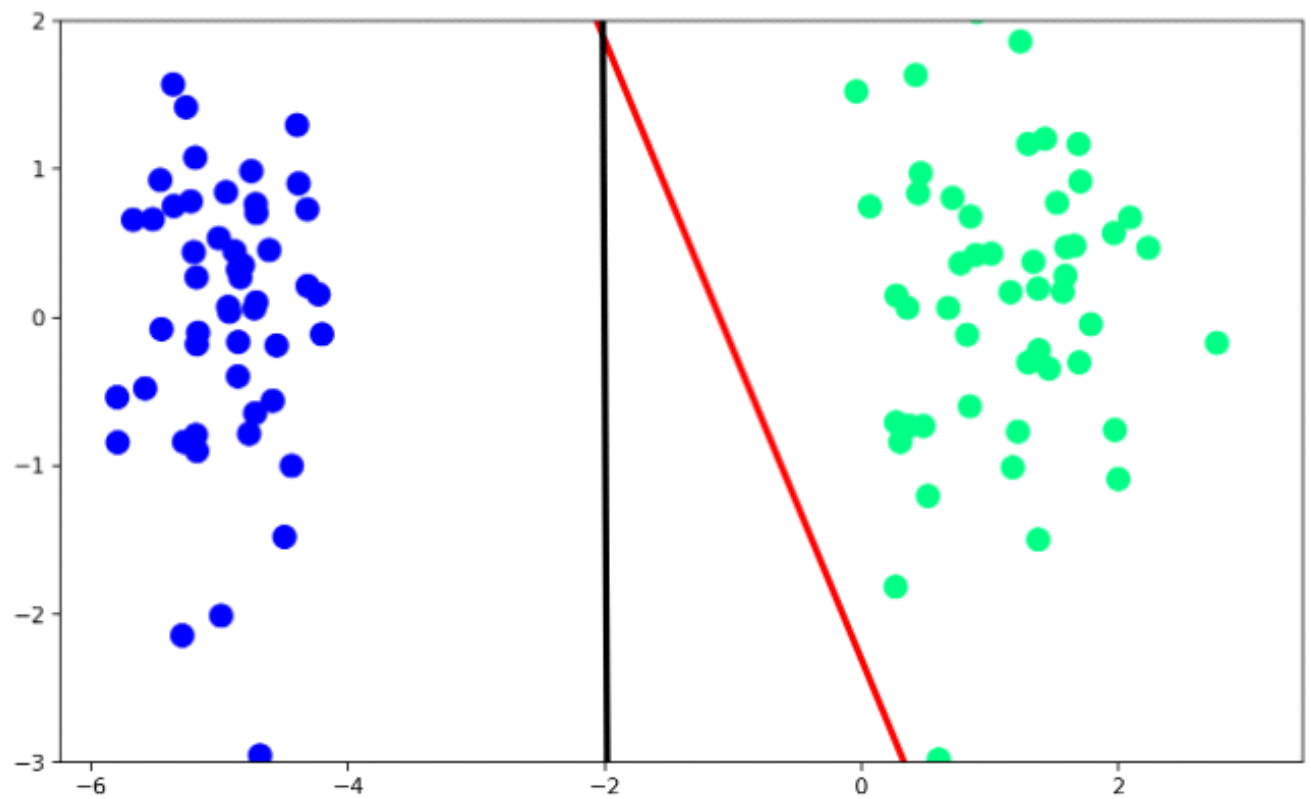
✓ ✗ ✗



training error
test

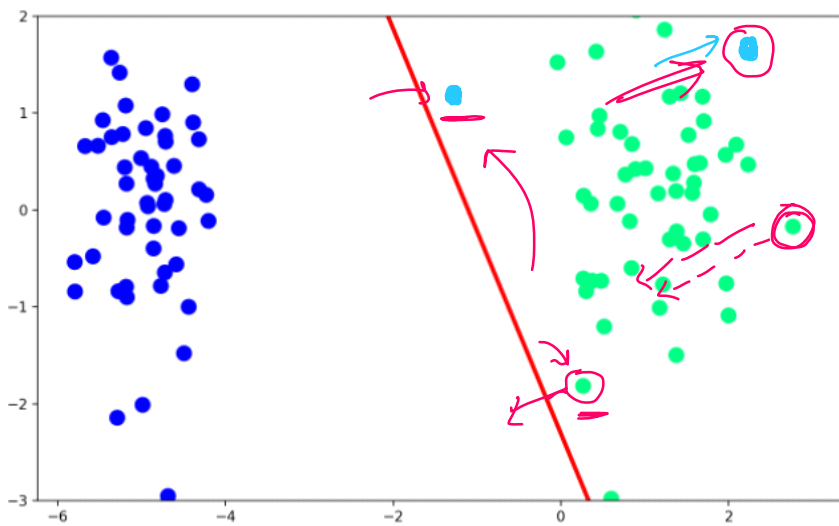
Problem with Perceptron

Thursday, June 17, 2021 12:18 PM



Possible Solution?

Thursday, June 17, 2021 12:18 PM



$$w_{\eta} = w_0 + (y_i - \hat{y}_i) x_i$$

→ misclassified line - pull

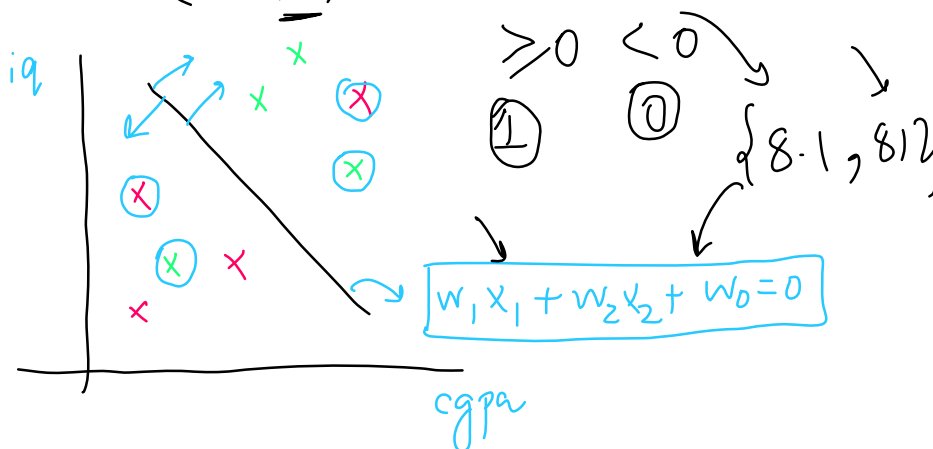
→ correctly line push

push pull

$$w_{\eta} = w_0 + \underbrace{\eta (y_i - \hat{y}_i)}_0 x_i$$

$(y_i - \hat{y}_i) \neq 0$ model predict
 $\sum w_i x_i = [0, 1]$
 $\rightarrow (y_i - \hat{y}_i)$

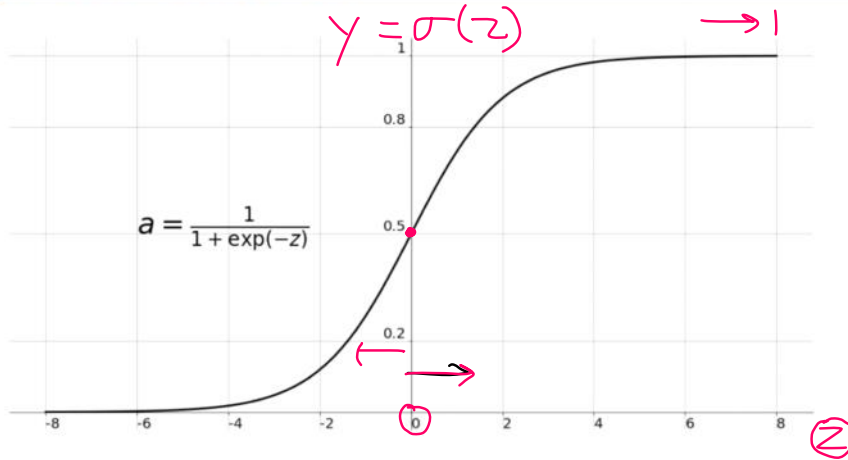
$$w_1 x_{8.1} + w_2 x_{8.1} + w_0 =$$



y_i	$\hat{y}_i$	$y_i - \hat{y}_i$
1	1	0
0	0	0
1	0	1
0	1	-1

cgpa	iq	(x_i) placed
9	91	0
8.8	78	1
8.1	102	1
7.9	98	1

Sigmoid Function



$$\sigma(z) = \frac{1}{1 + e^{-z}}$$

$$-\infty < z < \infty$$

$$0 < y < 1$$

$\{7.5, 81\} \rightarrow w_1 w_2 w_0$

$$y_1 = w_1 \times 7.5 + w_2 \times 81 + w_0 =$$



$$\begin{cases} z \geq 0 \rightarrow 1 \\ z < 0 \rightarrow 0 \end{cases}$$

$\sigma > 0.5$

$$z = \sum w_i x_i = \sigma(z)$$

$\sigma(z) < 0.5$

$\{7.5, 81\} \rightarrow z \rightarrow \sigma(z) = 0.5$

if 1 $p(N) = 1 - p(y)$
else 0

$0.5 < p(y) =$

$p(y) = 0.7$

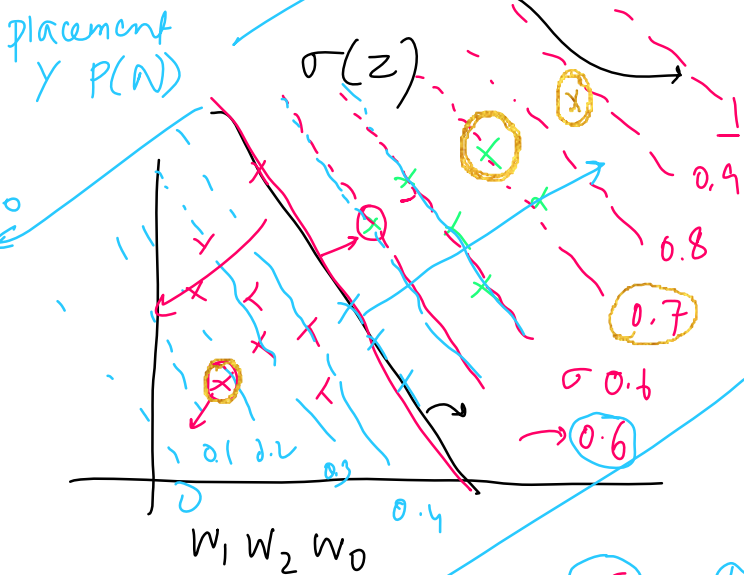
$\sum w_i x_i = 0 \quad \sigma = 0.5 \quad p(y) = 0.6$

$\sum w_i x_i > 0 \quad \sigma > 0.5$

$\sum w_i x_i \quad \sigma > 0.6$

$\sigma() = p()$ $p(y) p(N)$

$p(0.5) \rightarrow p(y) = 0.5$



Impact of Sigmoid

Thursday, June 17, 2021 3:27 PM

$$w_{\eta} = w_0 + \eta (y_i - \hat{y}_i) x_i$$

$$\hat{y}_i = \sigma(z)$$

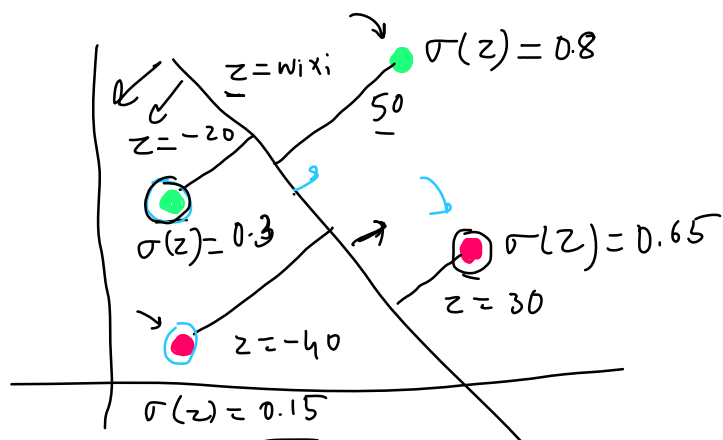
$$\text{where } z = \sum w_i x_i$$

$$w_{\eta} = w_0 + \eta \times 0.2 \times x_i$$

$$w_{\eta} = w_0 - \eta \times 0.65 \times x_i$$

$$w_{\eta} = w_0 + \eta \times 0.7 \times x_i$$

$$w_{\eta} = w_0 - \eta \times 0.15 \times x_i$$



y_i	$\hat{y}_i$	$y_i - \hat{y}_i$
1	0.8	0.2
0	0.65	-0.65
1	0.3	0.7
0	0.15	-0.15

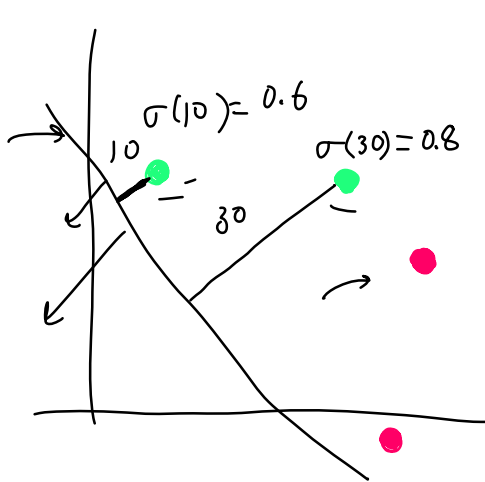
$$w_{\eta} = w_0 + \eta (y_i - \hat{y}_i) x_i$$

$$\hat{y}_i = \sigma(z)$$

$$\text{where } z = \sum w_i x_i$$

$$w_{\eta} = w_0 + \boxed{\eta \times 0.4 \times x_i} = x_1$$

$$w_{\eta} = w_0 + \boxed{\eta \times 0.2 \times x_i} = x_2$$



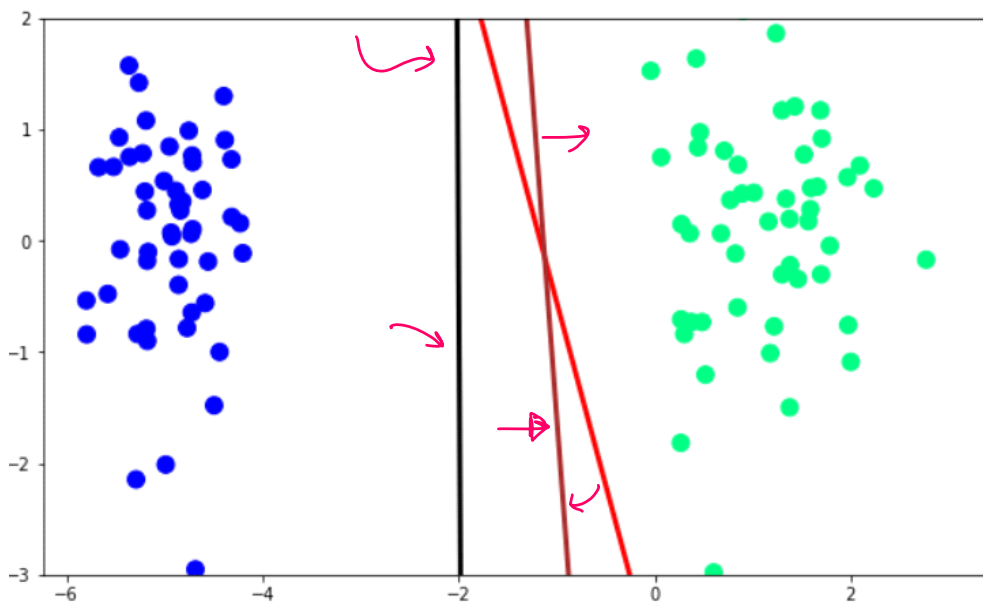
y_i	$\hat{y}_i$	$y_i - \hat{y}_i$
1	0.6	0.4
1	0.8	0.2
0		
0		

$x_1 > x_2$

code implement

The Problem

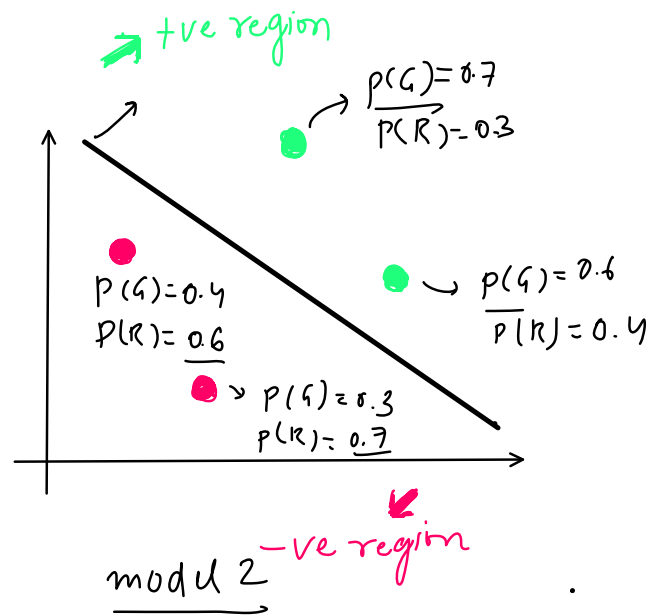
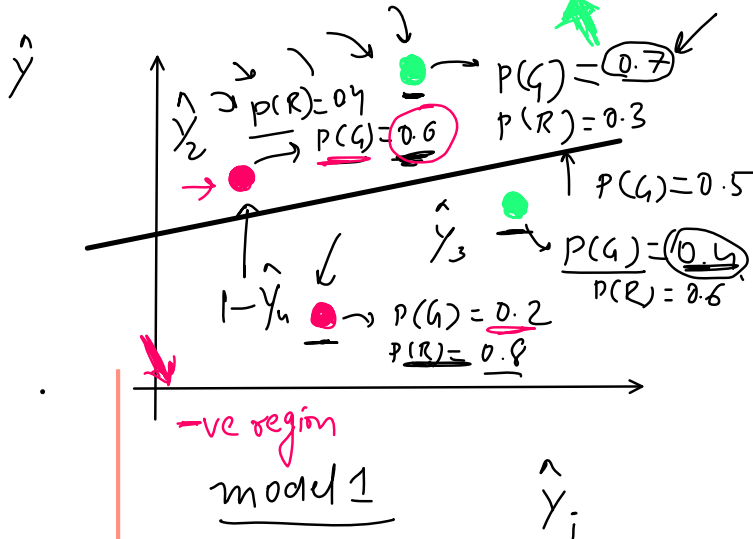
Friday, June 18, 2021 2:20 PM



red → step
brown → sigmoid
black → sklearn
→ for i in 1000
 random
Loss function
error function

Maximum Likelihood and Cross-Entropy

Thursday, June 17, 2021 12:19 PM



$$\hat{y} = \sigma(z)$$

$$z = \sum w_i x_i$$

$$\text{model 2} \rightarrow 0.7 \times 0.6 \times 0.6 \times 0.7 = 0.176$$

$$\text{model 1} \rightarrow 0.7 \times 0.4 \times 0.4 \times 0.8 = 0.089$$

$$\log(ab) = \log a + \log b$$

$$\log(\max) = -\log(0.7) - \log(0.4) - \log(0.4) - \log(0.8)$$

$$0-1 = -ve$$

cross entropy minimize

No

$$-\log(\hat{y}_1) - \log(\hat{y}_2) - \log(\hat{y}_3) - \log(\hat{y}_4)$$

$$(1 - \hat{y})$$

$$y_i = 1$$

$$-y_i \log(\hat{y}_i) - \frac{(1 - y_i) \log(1 - \hat{y}_i)}{1}$$

$$-y_i \log(\hat{y}_i) = -\log(\hat{y}_i)$$

$$= -\log(0.7)$$

$$y_2 = 0$$

$$y_3 = 1$$

$$x_2 = 0 \quad x_3 = 1$$

$$= -\log(0.7)$$

$$-x_2 \log(\hat{y}_2) - (1-x_2) \log(1-\hat{y}_2)$$

$$-x_3 \log(\hat{y}_3)$$

$$-\log(\hat{y}_3) = -\log(0.9)$$

$$x_4 = 0$$

$$-(1-x_4) \log(1-\hat{y}_4)$$

$$-\log(1-\hat{y}_4)$$

$$-\log(0.8)$$

$$L = \sum_{i=1}^n -x_i \log(\hat{y}_i) - (1-x_i) \log(1-\hat{y}_i)$$

MSE

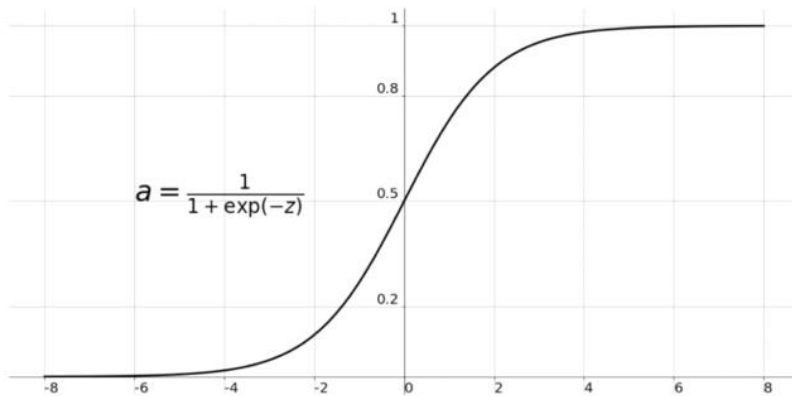
closed form
gradient descent

$$L = -\frac{1}{n} \sum_{i=1}^n x_i \log(\hat{y}_i) + (1-x_i) \log(1-\hat{y}_i)$$

min
 w_1, w_2, w_0

log-loss error
binary cross entropy

Sigmoid Function



$$\sigma(x) = \frac{1}{1 + e^{-x}}$$

$$\sigma'(x) = \frac{d}{dx} \left(\frac{1}{1 + e^{-x}} \right)$$

$$\frac{d}{dx} \left(\frac{1}{x} \right) = \frac{d}{dx} (x)^{-1}$$

$$= -x^{-2} = -\frac{1}{x^2}$$

$$\frac{d}{dx} \left[\frac{1}{1 + e^{-x}} \right] = \frac{d}{dx} \left[(1 + e^{-x})^{-1} \right] = -\frac{1}{(1 + e^{-x})^2} \frac{d}{dx} (1 + e^{-x})$$

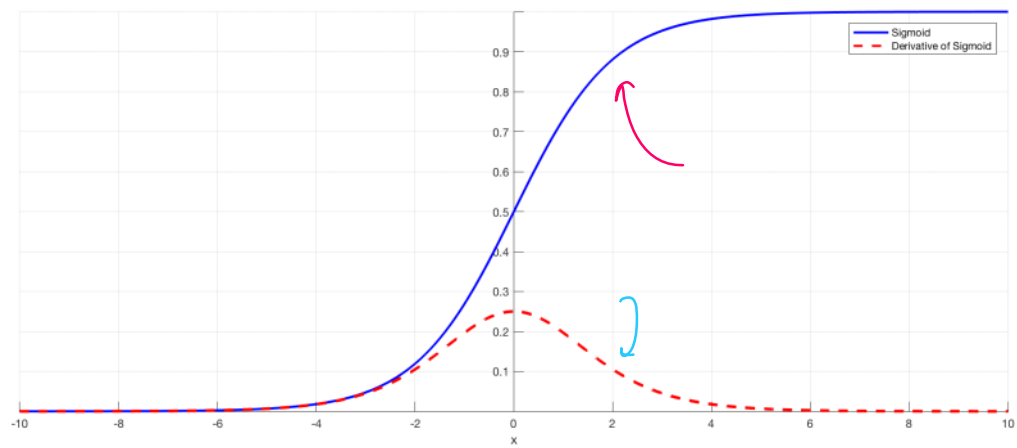
$e^{-x} = e^{-x}$ $-x = -1$

$$= -\frac{1}{(1 + e^{-x})^2} \frac{d}{dx} (e^{-x}) = -\frac{e^{-x}}{(1 + e^{-x})^2} \frac{d}{dx} (-x) = \frac{e^{-x}}{(1 + e^{-x})^2}$$

$$\frac{1 \cdot e^{-x}}{(1 + e^{-x})(1 + e^{-x})} = \frac{1}{1 + e^{-x}} \cdot \frac{e^{-x}}{1 + e^{-x}} = \sigma(x) \left[\frac{e^{-x}}{1 + e^{-x}} \right]$$

$$= \sigma(x) \left[\frac{1 + e^{-x} - 1}{1 + e^{-x}} \right] = \sigma(x) \left[\frac{1 + e^{-x}}{1 + e^{-x}} - \frac{1}{1 + e^{-x}} \right]$$

$$\sigma(x) [1 - \sigma(x)] \Rightarrow \sigma'(x) = \sigma(x) [1 - \sigma(x)]$$



Gradient Descent

Monday, June 21, 2021 11:50 AM

$\rightarrow$ classification $\{x_1, x_2\} \rightarrow y$
 $w_1 x_1 + w_2 x_2 + w_0 = 0$
 $x_0 \quad x_1 \quad x_2 \quad y = \{0, 1\}$
 $\hat{y}_i = \sigma(z) = \frac{1}{1 + e^{-z}}$
 $z = \sum_{j=0}^n w_j x_j$
 $\underline{w_1 x_1 + w_2 x_2 + w_0 x_0}$
 $\rightarrow \sum_{j=0}^2 w_j x_j = z$
 $L = -\frac{1}{m} \sum_{i=1}^m y_i \log(\hat{y}_i) + (1 - y_i) \log(1 - \hat{y}_i)$
 $L(\underline{w_0}, \underline{w_1}, \underline{w_2}) =$
 $\arg\min_{(w_0, w_1, w_2)}$
 $w_1, w_2, w_0 = 1$
 $\rightarrow$ gradient descent
 $\rightarrow$ for i in epochs 0.01
 $w_{new} = w_{old} - \eta \frac{\partial L}{\partial w_{old}}$

$w_0 = w_0 - \eta \frac{\partial L}{\partial w_0}, w_1 = w_1 - \eta \frac{\partial L}{\partial w_1}, w_2 = w_2 - \eta \frac{\partial L}{\partial w_2}$
 $n \text{ cols} \rightarrow n+1 \text{ derivatives}$

$w_n = w_n - \eta \frac{\partial L}{\partial w_n}$
 $j = 0, 1, 2, \dots, n$

$L = -\frac{1}{m} \sum_{i=1}^m y_i \log(\hat{y}_i) + (1 - y_i) \log(1 - \hat{y}_i)$

$\frac{\partial L}{\partial w_j} = -\frac{1}{m} \sum_{i=1}^m \left[\frac{\partial L}{\partial w_j} \left(y_i \log(\hat{y}_i) + (1 - y_i) \log(1 - \hat{y}_i) \right) \right]$
 $\hat{y}_i = \sigma(z) \quad z = \sum_{j=0}^n w_j x_j$

$\frac{\partial L}{\partial w_j} y_i \log(\hat{y}_i) = \frac{\partial L}{\partial w_j} y_i \log(\sigma(z))$
 $\rightarrow \frac{\partial L}{\partial w_j} = y_i - \hat{y}_i$
 $\rightarrow \frac{\partial L}{\partial w_j} = \sum_{i=1}^m (y_i - \hat{y}_i) x_j$

$$\frac{\partial L}{\partial w_j} = y_i \frac{\partial L}{\partial w_j} \sigma\left(\sum_{j=0}^n w_j x_j\right)$$

$$\frac{\partial L}{\partial w_j} \rightarrow y_i \log\left(\sigma\left(\sum_{j=0}^n w_j x_j\right)\right) = \frac{y_i}{\hat{y}_i} \hat{y}_i (1 - \hat{y}_i) \frac{\partial L}{\partial w_j} \sum_{j=0}^n w_j x_j$$

$$y_i (1 - \hat{y}_i) \sum_{j=0}^n \frac{\partial L}{\partial w_j} w_j x_j = \boxed{y_i (1 - \hat{y}_i) \sum_{j=0}^n x_j}$$

$$\frac{\partial L}{\partial w_j} (1 - y_i) \log(1 - \hat{y}_i) \Rightarrow (1 - y_i) \frac{\partial L}{\partial w_j} \log(1 - \hat{y}_i) \quad \hat{y}_i = \sigma\left(\sum_{j=0}^n w_j x_j\right)$$

$$\frac{(1 - y_i)}{(1 - \hat{y}_i)} \frac{\partial L}{\partial w_j} (1 - \hat{y}_i) \Rightarrow - \frac{(1 - y_i)}{(1 - \hat{y}_i)} \frac{\partial L}{\partial w_j} \hat{y}_i \Rightarrow - \frac{(1 - y_i)}{(1 - \hat{y}_i)} \frac{\partial L}{\partial w_j} \sigma(z)$$

$$\Rightarrow - \frac{(1 - y_i)}{(1 - \hat{y}_i)} \sigma(z) (1 - \sigma(z)) \frac{\partial L}{\partial w_j} \sigma(z) = - \frac{(1 - y_i)}{(1 - \hat{y}_i)} \hat{y}_i (1 - \hat{y}_i) \frac{\partial L}{\partial w_j} \sigma(z)$$

$$\Rightarrow - \hat{y}_i (1 - y_i) \frac{\partial L}{\partial w_j} \sum_{j=0}^n w_j x_j \Rightarrow - \hat{y}_i (1 - y_i) \sum_{j=0}^n \frac{\partial L}{\partial w_j} (w_j x_j)$$

$$\Rightarrow \boxed{- \hat{y}_i (1 - y_i) \sum_{j=0}^n x_j}$$

$$\frac{\partial L}{\partial w_j} = - \frac{1}{n} \sum_{i=1}^m \left[y_i (1 - \hat{y}_i) \sum_{j=0}^n x_j - \hat{y}_i (1 - y_i) \sum_{j=0}^n x_j \right]$$

$$= \frac{1}{n} \sum_{i=1}^m \left[y_i (1 - \hat{y}_i) - \hat{y}_i (1 - y_i) \right] \sum_{j=0}^n x_j$$

$$\frac{\partial L}{\partial w_j} = -\frac{1}{m} \sum_{i=1}^m \left[y_i (1 - \hat{y}_i) - \hat{y}_i (1 - y_i) \right] \sum_{j=0}^n x_j$$

$$= -\frac{1}{m} \sum_{j=1}^m \left[y_i - \cancel{y_i \hat{y}_i} - \hat{y}_i + \cancel{y_i \hat{y}_i} \right] \sum_{j=0}^n x_j$$

$$\frac{\partial L}{\partial w_j} = -\frac{1}{m} \sum_{i=1}^m (y_i - \hat{y}_i) \sum_{j=0}^n x_j$$

$$\frac{\partial L}{\partial w_0} = -\frac{1}{2} [1 + 0] [1 + 1]$$

$$= -\frac{1}{2} [1] [2] = -1$$

$$= -\frac{1}{2} [1 \ 0] \begin{bmatrix} 1 \\ 1 \end{bmatrix}$$

$$= -\frac{1}{2} [1] = -\frac{1}{2}$$

x_1	x_2	y_i	$\hat{y}_i$
1	2	1	0
3	4	0	0

no. of rows = $m = 2$

no. of cols = $n = 2$

n input

Rows = m Columns = n

	1	2	3	...	n
1	x_{11}	x_{12}	x_{13}	...	x_{1n}
2	x_{21}	x_{22}	x_{23}	...	x_{2n}
...	...	...	...	...	...
m	x_{m1}	x_{m2}	x_{m3}	...	x_{mn}

$\hat{y}_1 \rightarrow \hat{y}_1$
 $\hat{y}_2 \rightarrow \hat{y}_2$
 $\hat{y}_m \rightarrow \hat{y}_m$

Weights: $w_1, w_2, w_3, \dots, w_n, w_0$ (n+1)

$$\sigma(w_1 x_{11} + w_2 x_{12} + w_3 x_{13} + \dots + w_n x_{1n} + w_0) = \hat{y}_1$$

$$\sigma(w_1 x_{21} + w_2 x_{22} + w_3 x_{23} + \dots + w_n x_{2n} + w_0) = \hat{y}_2$$

$$\hat{\underline{y}} = \begin{bmatrix} \hat{y}_1 \\ \hat{y}_2 \\ \vdots \\ \hat{y}_m \end{bmatrix} = \begin{bmatrix} \sigma(w_0 + w_1 x_{11} + w_2 x_{12} + \dots + w_n x_{1n}) \\ \sigma(w_0 + w_1 x_{21} + w_2 x_{22} + \dots + w_n x_{2n}) \\ \vdots \\ \sigma(w_0 + w_1 x_{m1} + w_2 x_{m2} + \dots + w_n x_{mn}) \end{bmatrix}$$

$$\hat{\underline{y}} = \sigma \left(\begin{bmatrix} w_0 + w_1 x_{11} + w_2 x_{12} + \dots + w_n x_{1n} \\ w_0 + w_1 x_{21} + w_2 x_{22} + \dots + w_n x_{2n} \\ \vdots \\ w_0 + w_1 x_{m1} + w_2 x_{m2} + \dots + w_n x_{mn} \end{bmatrix} \right)$$

$$\hat{\underline{y}} = \sigma \left(\begin{bmatrix} 1 & x_{11} & x_{12} & \dots & x_{1n} \\ 1 & x_{21} & x_{22} & \dots & x_{2n} \\ \vdots & \vdots & \vdots & \ddots & \vdots \\ 1 & x_{m1} & x_{m2} & \dots & x_{mn} \end{bmatrix} \begin{bmatrix} w_0 \\ w_1 \\ \vdots \\ w_n \end{bmatrix} \right)$$

Labels: w_0 (bias), A (input matrix), X (input vector), w (weight vector).

$$\hat{y} = \sigma \left(\begin{bmatrix} 1 & x_{21} & x_{22} & \dots & x_{2n} \\ \vdots & \vdots & \vdots & \ddots & \vdots \\ 1 & x_{m1} & x_{m2} & \dots & x_{mn} \end{bmatrix} \begin{bmatrix} w_1 \\ \vdots \\ w_n \end{bmatrix} \right)$$

$$\hat{y} = \sigma(XW)$$

$$L = -\frac{1}{m} \sum_{i=1}^m y_i \log(\hat{y}_i) + (1-y_i) \log(1-\hat{y}_i)$$

$$L = -\frac{1}{m} \left[\sum_{i=1}^m y_i \log(\hat{y}_i) + \sum_{i=1}^m (1-y_i) \log(1-\hat{y}_i) \right]$$

$(1-y) \log$
 $\sigma(1-XW)$

$$\sum_{i=1}^m y_i \log(\hat{y}_i) = y_1 \log \hat{y}_1 + y_2 \log \hat{y}_2 + y_3 \log \hat{y}_3 + \dots + y_m \log \hat{y}_m$$

$$\begin{bmatrix} y_1 & y_2 & y_3 & \dots & y_m \end{bmatrix} \begin{bmatrix} \log \hat{y}_1 \\ \log \hat{y}_2 \\ \vdots \\ \log \hat{y}_m \end{bmatrix}$$

$$\underline{\begin{bmatrix} y_1 & y_2 & y_3 & \dots & y_m \end{bmatrix}} \log \left(\begin{bmatrix} \hat{y}_1 \\ \hat{y}_2 \\ \vdots \\ \hat{y}_m \end{bmatrix} \right)$$

$$y \log \hat{y} = y \log(\sigma(XW))$$

$$L = -\frac{1}{m} \left[y \log \hat{y} + (1-y) \log(1-\hat{y}) \right] \quad \text{min}$$

where $\hat{y} = \sigma(XW)$ $\uparrow$ GD $[W]$ find

where $\hat{y} = \sigma(XW)$ $L(GD) [W]$

Loss function in Matrix form

$$L = -\frac{1}{m} \left[y \log(\sigma(WX)) + (1-y) \log(1 - \sigma(WX)) \right]$$

minimum

for i in epochs:

gradient descent

$$W = W - \eta \frac{\Delta L}{\Delta W}$$

$$W = [$$

$$]$$

$$W_0 \rightarrow W_n$$

$$(\eta+1)$$

$$\left\{ \frac{\Delta L}{\Delta W} \right\}$$

$$\frac{\Delta L}{\Delta W} = \left[\frac{\partial L}{\partial W_0}, \frac{\partial L}{\partial W_1}, \frac{\partial L}{\partial W_2}, \dots, \frac{\partial L}{\partial W_n} \right]$$

$$\frac{\Delta L}{\Delta W}$$

$$L = -\frac{1}{m} \left[y \log \hat{y} + (1-y) \log(1-\hat{y}) \right]$$

$$\frac{dL}{dW} =$$

$$\frac{d}{dW} y \log \hat{y} \Rightarrow y \frac{d}{dW} \log \hat{y} \Rightarrow \frac{y}{\hat{y}} \frac{d}{dL} (\hat{y})$$

$$\Rightarrow \frac{y}{\hat{y}} \frac{d}{dL} \sigma(WX) \Rightarrow \frac{y}{\hat{y}} \sigma(WX) [1 - \sigma(WX)] \frac{d(WX)}{dW}$$

$$= \frac{y}{\hat{y}} \hat{y} (1 - \hat{y}) X = \boxed{y(1 - \hat{y}) X} \quad \hat{y} = \sigma(WX)$$

$$= \frac{1}{\hat{y}} \dots \dots \dots \underbrace{\dots \dots \dots}_{y - \hat{y}} \dots$$

$$\frac{d}{dw} (1-y) \log(1-y) \Rightarrow (1-y) \frac{d}{dw} \log(1-\hat{y}) \Rightarrow \frac{(1-y)}{(1-\hat{y})} \frac{d}{dw} [1-\hat{y}]$$

$$= - \frac{(1-y)}{(1-\hat{y})} \frac{d}{dw} \sigma(wx) \Rightarrow - \frac{(1-y)}{(1-\hat{y})} \left[\sigma(wx) \left[1 - \sigma(wx) \right] \frac{d}{dw} (wx) \right]$$

$$\Rightarrow - \frac{(1-y)}{(1-\hat{y})} \hat{y} (1-\hat{y}) x = \boxed{- \hat{y} (1-y) x}$$

$$\frac{dL}{dw} = - \frac{1}{m} \left[y(1-\hat{y})x - \hat{y}(1-y)x \right]$$

$$= - \frac{1}{m} \left[y(1-\hat{y}) - \hat{y}(1-y) \right] x$$

$$= - \frac{1}{m} \left[y - \cancel{y\hat{y}} - \hat{y} + \cancel{y\hat{y}} \right] x$$

$$\boxed{\frac{\Delta L}{\Delta w} = - \frac{1}{m} (y - \hat{y}) x}$$

gd

2 ↙

$$\boxed{w = \underline{w} + \eta \frac{1}{m} (y - \hat{y}) x}$$

$$\underline{w} = \underline{w} + \eta \frac{1}{m} (\underline{y} - \hat{\underline{y}}) \underline{X}^T$$

$$\underline{w}^{(n+1),1} = \begin{bmatrix} w_0 \\ w_1 \\ \vdots \\ w_n \end{bmatrix}$$

$$\underline{X} = \begin{bmatrix} 1 & x_{11} & x_{12} & \dots & x_{1n} \\ 1 & x_{21} & x_{22} & \dots & x_{2n} \\ \vdots & \vdots & \vdots & \ddots & \vdots \\ 1 & x_{m1} & x_{m2} & \dots & x_{mn} \end{bmatrix}$$

$m, n+1$

$$\underline{y} = \begin{bmatrix} y_1 \\ y_2 \\ \vdots \\ y_m \end{bmatrix}$$

$m, 1$

$$\hat{\underline{y}} = \begin{bmatrix} \hat{y}_1 \\ \hat{y}_2 \\ \vdots \\ \hat{y}_m \end{bmatrix}$$

$m, 1$

$$\underline{w}^{(n+1),1} = \underline{w}^{(n),1} + \left[\frac{\eta}{m} \right] (\underline{y} - \hat{\underline{y}}) \underline{X}^T$$

$(1, m) \quad m, (n+1) \rightarrow (1, n+1)$

$\downarrow$
 $(n+1), 1$

Accuracy

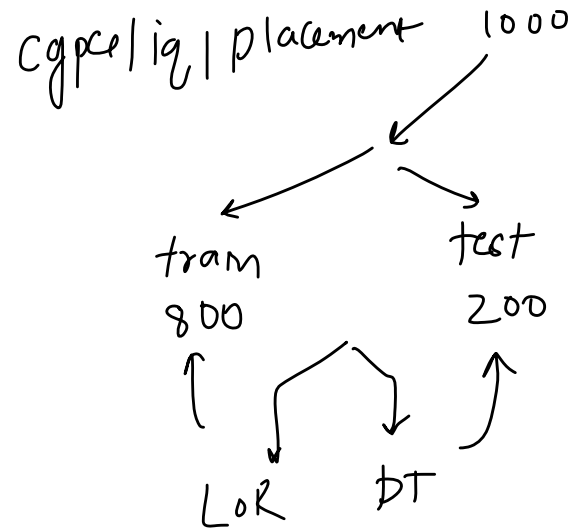
Tuesday, June 22, 2021 1:12 PM

binary ✓

Actual Label	Logistic Regression Prediction	Decision Tree Prediction
1	✓ 1	✓ 1
0	x 1	x 1
0	✓ 0	✓ 0
0	✓ 0	✓ 0
1	✓ 1	✓ 1
1	✓ 1	✓ 1
0	x 1	✓ 0
0	✓ 0	✓ 0
0	✓ 0	✓ 0
1	✓ 1	✓ 1

Accuracy = $\frac{\text{no. of } \checkmark}{\text{total prediction}}$

$\rightarrow \frac{9}{10} = 90\%$



$\frac{8}{10} = 0.8 \rightarrow 80\%$

Accuracy of multi-classification problem

Wednesday, June 23, 2021 8:44 AM



Actual Label	Logistic Regression Prediction	Decision Tree Prediction
0	✓ 0	0
0	✓ 0	0
0	✓ 0	0
2	✓ 2	2
0	✓ 0	0
2	✓ 2	2
0	✓ 0	0
2	✓ 2	2
1	✓ 1	1
1	✓ 1	1

iris
setosa, virginica / versicolour
0 1 2

$$\text{accuracy} = \frac{\# \text{ correct prediction}}{\# \text{ total}}$$

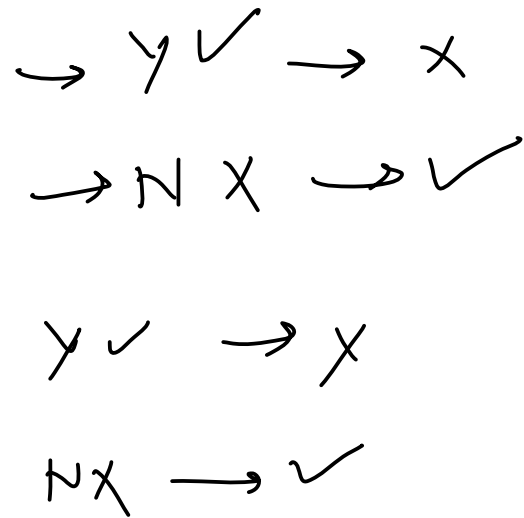
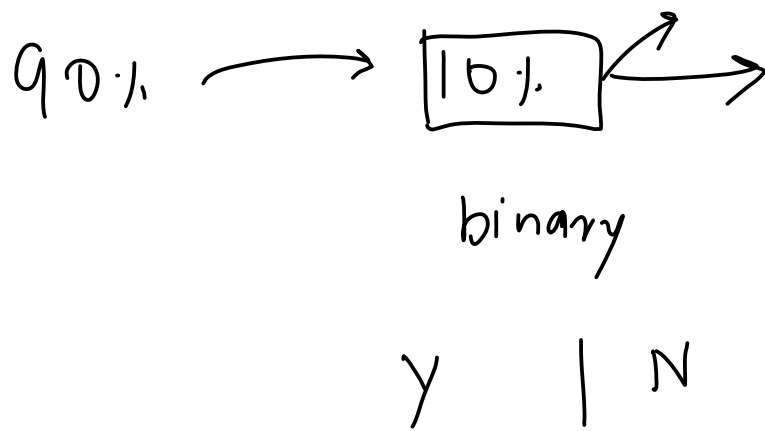
$$= \frac{10}{10} = 1 = 100\%$$

How much accuracy is good?

Wednesday, June 23, 2021 11:14 AM

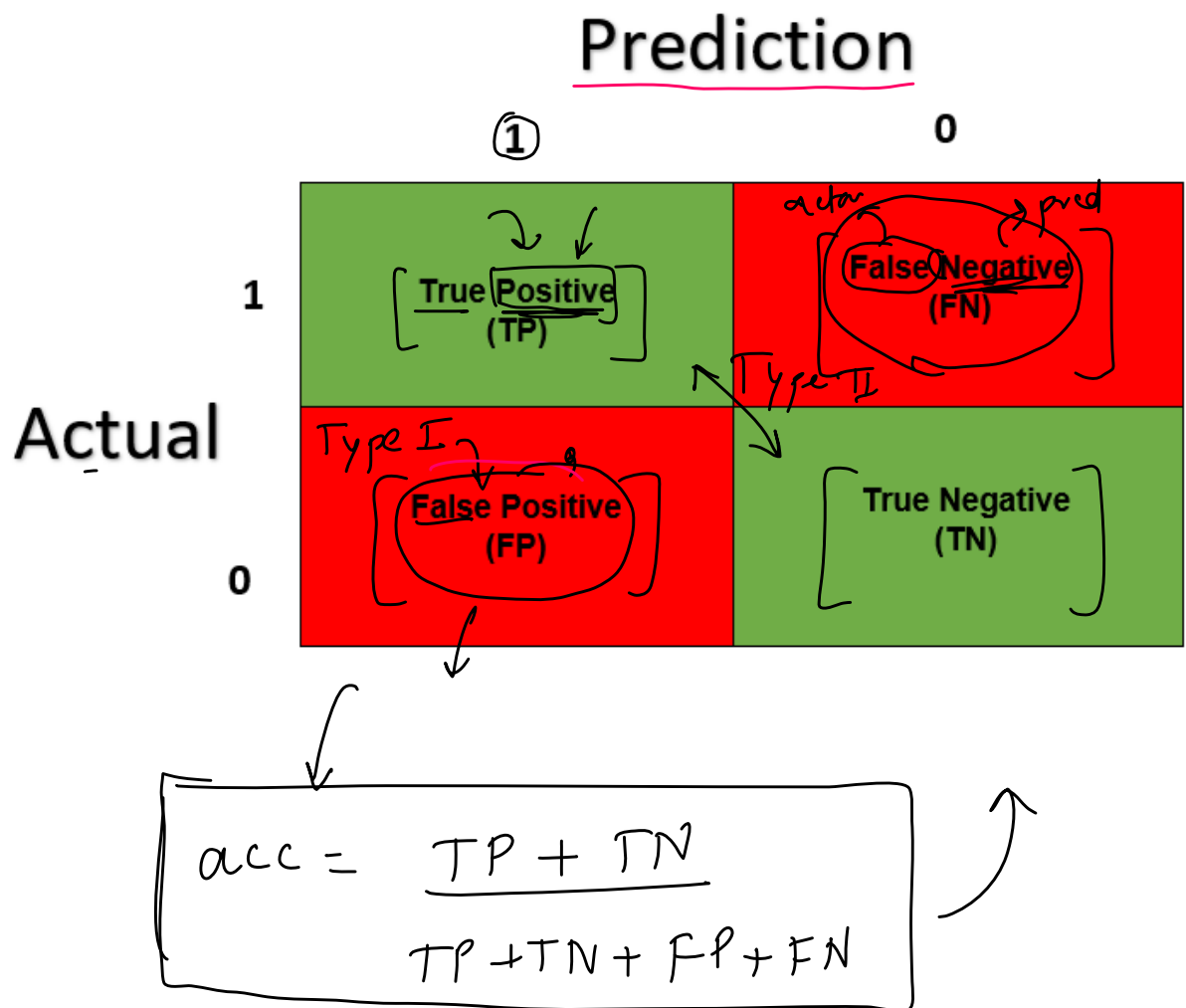
The Problem with Accuracy

Wednesday, June 23, 2021 11:14 AM



Confusion Matrix

Tuesday, June 22, 2021 1:12 PM



Extended

↳ { echo
or
not }

duplicate

↳ echo

$1 + 2 + 3 = 6$ → outcome

x	y
2	y

Type 1 Error

Wednesday, June 23, 2021 8:45 AM

Type 2 Error

Wednesday, June 23, 2021 8:45 AM

Confusion Matrix for Multi-classification Problem

Wednesday, June 23, 2021 8:45 AM

0, 1, 2 →

10x10

		Predicted		
		0	1	2
Actual	0	7	0	5
	1	2	21	6
	2	9	1	13

3x3

binary
2x2

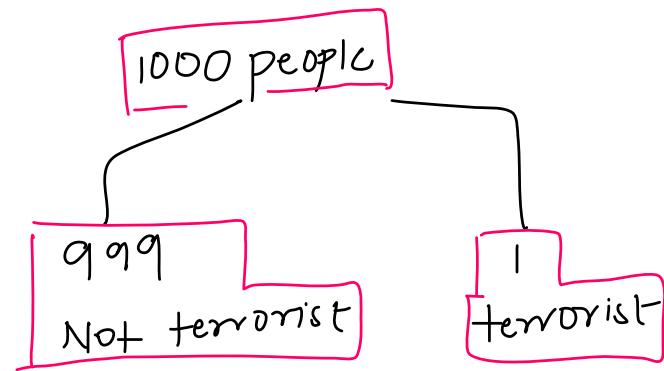
$$\frac{\text{True pred}}{\text{total}} = \text{accuracy}$$

3x3

When accuracy is misleading?

Wednesday, June 23, 2021 8:45 AM

Imbalanced Dataset



model $\rightarrow$ No one is terrorist

	Predicted	
	1	0
Actual	1 $\rightarrow$ 0 ✓	$\rightarrow$ 1
	$\rightarrow$ 0	$\rightarrow$ 999

$$\text{Accuracy} = \frac{999}{999 + 1} = 99.9\%$$

Precision

Wednesday, June 23, 2021 8:46 AM

{ spam: 1, not spam: 0 }

Actual

Predicted (A)

	Sent to Spam	Not sent to spam
Spam	100	170 FN
Not Spam	30 FP	700

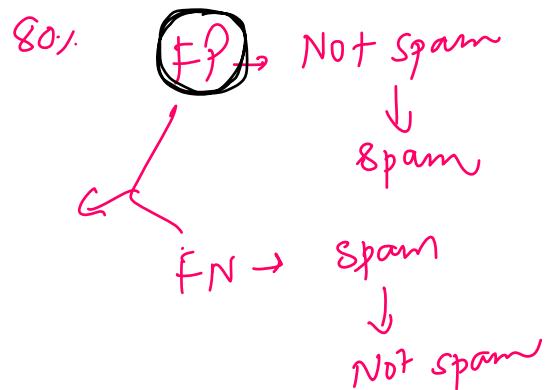
Predicted (B)

	Sent to Spam	Not sent to spam
Spam	100	190
Not Spam	10	700

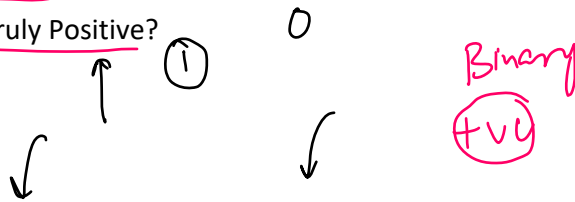
$$P_A = \frac{100}{100 + 30}$$

$$P_B = \frac{100}{100 + 10}$$

$$P_A < P_B$$



What proportion of predicted Positives is truly Positive?



	Sent to Spam	Not sent to spam
Spam	True Positive	False Negative
Not Spam	False Positive	True Negative

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

Recall

Wednesday, June 23, 2021 8:46 AM

has cancer: 1 , no cancer: 0

Predicted

Actual

	Detected Cancer	Not Detected
Has Cancer	1000	200 (FN)
No Cancer	800 FP	8000

	Detected Cancer	Not Detected
Has Cancer	1000	500 ← B
No Cancer	500	8000

✓ A 90%

B 90%

$$Recall = \frac{1000}{1200}$$

↑
↓
R_A > R_B

$$R_B = \frac{1000}{1500} \leftarrow$$

(FP)

What proportion of actual Positives is correctly classified?

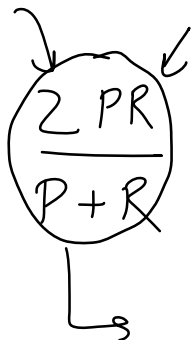
	Detected Cancer	Not detected Cancer
Has Cancer	<u>True Positive</u>	<u>False Negative</u>
No Cancer	False Positive	True Negative

$$Recall = TP / (TP + FN) \leftarrow$$

F1 Score

Wednesday, June 23, 2021 8:46 AM

F1 score =



$$P = 0$$

$$R = 100$$

$$F1 score = 50$$

$$F1 = 0$$

$$F1 score = \frac{P+R}{2}$$

$$\frac{2 \times 80 \times 80}{160} = 80$$

④

$$P = R = 80$$

$$\frac{80 + 80}{2} = 80 \quad F1 = 80$$

③

⑦

$$P = 60 \quad R = 100$$

$$\frac{100 + 60}{2} = 80$$

$$\frac{2 \times 60 \times 100}{160} = 75$$

Multi-class Precision and Recall

Wednesday, June 23, 2021 8:46 AM

Positive
① ✓
Yes
No

Binary
② class

Multi
③

Dog Cat Rabbit

$$\frac{2 \times 0.86 + 0.66}{0.86 + 0.66}$$

R_D, R_C, R_R

$$F1_D = \frac{2 P_D R_D}{P_D + R_D}$$

$$F1_C = \frac{2 P_C R_C}{P_C + R_C}$$

$$F1_R = \frac{2 P_R R_R}{P_R + R_R}$$

Actual \ Predicted				Total	Recall
	Dog	Cat	Rabbit		
Dog	25	5	10	40	0.62
Cat	0	30	4	34	0.88
Rabbit	4	10	20	34	0.58
Total	29	45	34		
Precision	0.86	0.66	0.58		

$$R_D = \frac{25}{40} = 0.62, R_C = \frac{30}{34} = 0.88, R_R = \frac{20}{34} = 0.58$$

+ macro recall
+ weighted recall

$F1_D, F1_C, F1_R$ + macro f1
+ weighted f1

Multi-class F1 Score

Thursday, June 24, 2021 11:14 AM

Softmax Regression

Monday, June 28, 2021 3:35 PM

cgpa | iq | place?

7.1

71

$\left\{ \begin{array}{l} 0 \\ 1 \\ 2 \end{array} \right\}$

8.5

85

9.5

95

$\left\{ \begin{array}{l} \text{Yes} \\ \text{No} \end{array} \right\} \begin{array}{l} 1 \\ 0 \end{array}$
Optout 2

→ classifier

{ 6.5, 65 }

→ Yes / No / Optout

Binary classification

$\left\{ \begin{array}{l} \text{Softmax} \\ \text{Multinomial} \\ \text{LR} \end{array} \right\} \rightarrow \text{Deep Learning}$

Softmax → general
↑
for 2 class

$$0 < \sigma < 1$$

Softmax function

$K = \text{no. of class } \{3\}$

Yes → 1

no → 2

opt → 3

(Yes)

$$\sigma(z)_i = \frac{e^{z_i}}{\sum_{j=1}^K e^{z_j}}$$

$$\sigma(z)_3 = \frac{e^{z_3}}{e^{z_1} + e^{z_2} + e^{z_3}} \rightarrow 3$$

$$\sigma(z)_1 = \frac{e^{z_1}}{e^{z_1} + e^{z_2} + e^{z_3}} \quad \left\{ \begin{array}{l} \text{Yes} \\ \text{No} \end{array} \right\} \quad \sigma(z)_2 = \frac{e^{z_2}}{e^{z_1} + e^{z_2} + e^{z_3}}$$

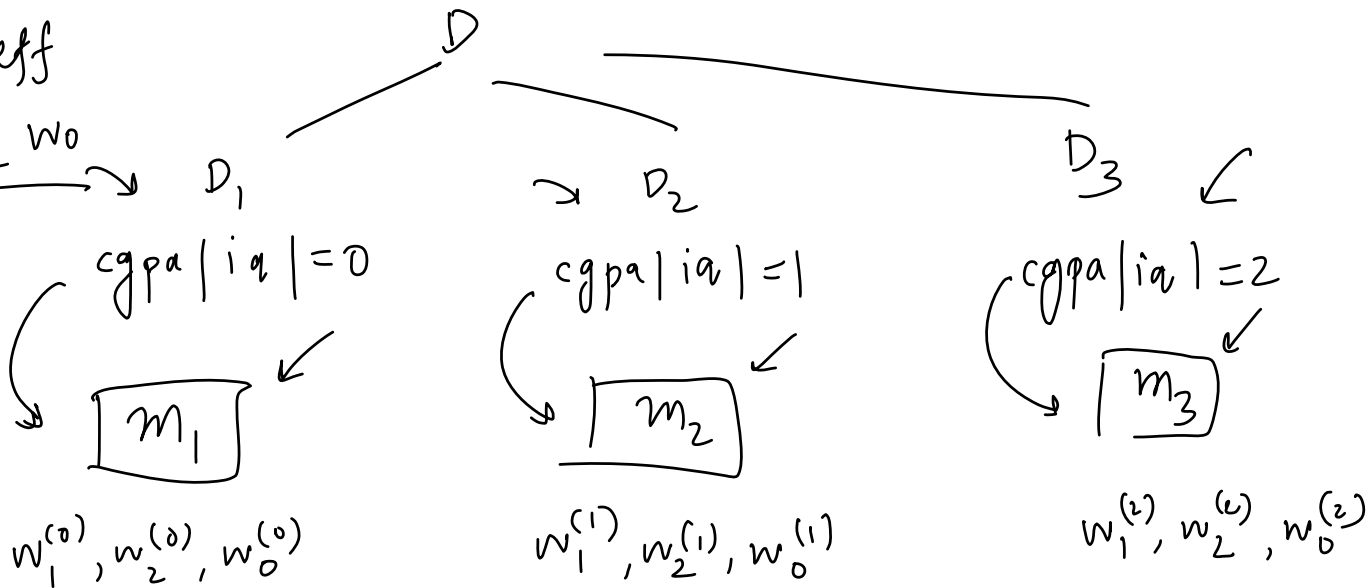
$$\{y_i^{(0)}, y_i^{(1)}, y_i^{(2)}\}$$

cgpa | iq | place?

-	-	$\frac{0}{1/2}$	transform	OHE	5
-	-	1			
-	-	2			
-	-				
cgpa	iq	= 0		= 1	= 2
✓	✓	1		0	0
		0		1	0
		0		0	1

3 coeff

w_1, w_2, w_0



{ loss function
gradient descent }

Loss Function

Monday, June 28, 2021 3:47 PM

$$L = -\frac{1}{m} \sum_{i=1}^m y_i \log(\hat{y}_i) + (1-y_i) \log(1-\hat{y}_i)$$

$$L = -\frac{1}{m} \sum_{i=1}^m \sum_{k=1}^K y_k^{(i)} \log(\hat{y}_k^{(i)})$$

x_1	x_2	y	$y_{k=1}$	$y_{k=2}$	$y_{k=3}$
$\rightarrow x_{11}$	x_{12}	1	1	0	0
$\rightarrow x_{21}$	x_{22}	2	0	1	0
$\rightarrow x_{31}$	x_{32}	3	0	0	1

$$\begin{aligned} & y_1^{(1)} \log(\hat{y}_1^{(1)}) + y_2^{(1)} \log(\hat{y}_2^{(1)}) + \\ & y_3^{(1)} \log(\hat{y}_3^{(1)}) + \\ & y_1^{(2)} \log(\hat{y}_1^{(2)}) + y_2^{(2)} \log(\hat{y}_2^{(2)}) + \\ & y_3^{(2)} \log(\hat{y}_3^{(2)}) + \\ & y_1^{(3)} \log(\hat{y}_1^{(3)}) + y_2^{(3)} \log(\hat{y}_2^{(3)}) + \\ & y_3^{(3)} \log(\hat{y}_3^{(3)}) \end{aligned}$$

$$L = y_1^{(1)} \log(\hat{y}_1^{(1)}) + y_2^{(2)} \log(\hat{y}_2^{(2)}) + y_3^{(3)} \log(\hat{y}_3^{(3)})$$

$$\hat{y}_1^{(1)}, \hat{y}_2^{(2)}, \hat{y}_3^{(3)}$$

softmax

$$\hat{y}_1^{(1)} = \sigma(\underline{w}_1^{(1)} x_{11} + \underline{w}_2^{(1)} x_{12} + \underline{w}_0^{(1)})$$

$$y_2^{(2)} = \sigma(\underline{w}_1^{(2)} x_{21} + \underline{w}_2^{(2)} x_{22} + \underline{w}_0^{(2)})$$

$$y_3^{(3)} = \sigma(\underline{w}_1^{(3)} x_{31} + \underline{w}_2^{(3)} x_{32} + \underline{w}_0^{(3)})$$

$$\begin{bmatrix} w_1^{(1)} & w_2^{(1)} & w_0^{(1)} \\ w_1^{(2)} & w_2^{(2)} & w_0^{(2)} \\ w_1^{(3)} & w_2^{(3)} & w_0^{(3)} \end{bmatrix}$$

(L)

$$\frac{\partial L}{\partial w_1^{(1)}}, \frac{\partial L}{\partial w_2^{(1)}}, \frac{\partial L}{\partial w_0^{(1)}} \dots qdu$$

gradient

q-values init=1

$$\begin{bmatrix} \text{---} \\ \text{---} \\ \text{---} \end{bmatrix} =$$

loop $\Rightarrow$ 1000 epochs

$$w_1^{(1)} = w_1^{(1)} - \eta \frac{\partial L}{\partial w_1^{(1)}}$$

$$w_2^{(1)} = w_2^{(1)} - \eta \frac{\partial L}{\partial w_2^{(1)}}$$

⋮

Prediction

Monday, June 28, 2021 3:35 PM

$$S_x = \{7, 70\} \Rightarrow \overline{Y, N, Opt}$$

m_1 Yes

$$\underline{w}_1^{(1)} \quad \underline{w}_2^{(1)} \quad \underline{w}_0^{(1)}$$

$$\underline{z}_1 = 7 \times \underline{w}_1^{(1)} + 70 \times \underline{w}_2^{(1)} + \underline{w}_0^{(1)}$$

m_2 No

$$\underline{w}_1^{(2)} \quad \underline{w}_2^{(2)} \quad \underline{w}_0^{(2)}$$

$$\underline{z}_2 = 7 \times \underline{w}_1^{(2)} + 70 \times \underline{w}_2^{(2)} + \underline{w}_0^{(2)}$$

m_3 Opt out

$$\underline{w}_1^{(3)} \quad \underline{w}_2^{(3)} \quad \underline{w}_0^{(3)}$$

$$\underline{z}_3 = 7 \times \underline{w}_1^{(3)} + 70 \times \underline{w}_2^{(3)} + \underline{w}_0^{(3)}$$

$$\sigma(y) = \frac{e^{z_1}}{e^{z_1} + e^{z_2} + e^{z_3}}$$

$$= \underline{0.40}$$

$$\sigma(N) = \frac{e^{z_2}}{e^{z_1} + e^{z_2} + e^{z_3}}$$

$$\underline{0.35}$$

$$e(o) = \frac{e^{z_3}}{e^{z_1} + e^{z_2} + e^{z_3}}$$

$$\underline{0.25}$$

Sigmoid Vs Softmax

Monday, June 28, 2021 3:47 PM

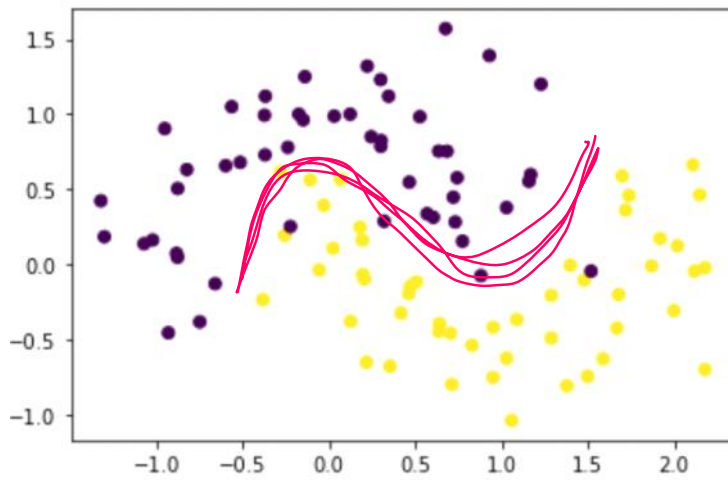
Code Sample

Monday, June 28, 2021 3:36 PM

Polynomial Logistic Regression

Monday, June 28, 2021 6:43 PM

x_1 , x_2 , $y \in \{0, 1\}$



Linear polynomial degree = 2

$x^0 \ x^1 \ x^2 \ x^3$

$\begin{bmatrix} x_1 & x_2 \end{bmatrix} \rightarrow 2 \text{ cols}$

$\rightarrow 6 \text{ cols}$

$\begin{matrix} x^0 & x^1 & x^2 & x^0 & x^1 & x^2 & y \\ \uparrow & \downarrow & \downarrow & | & | & | & \\ w_1 & w_2 & w_3 & & & & w_8 \end{matrix}$

w_0

Wisdom of the Crowd

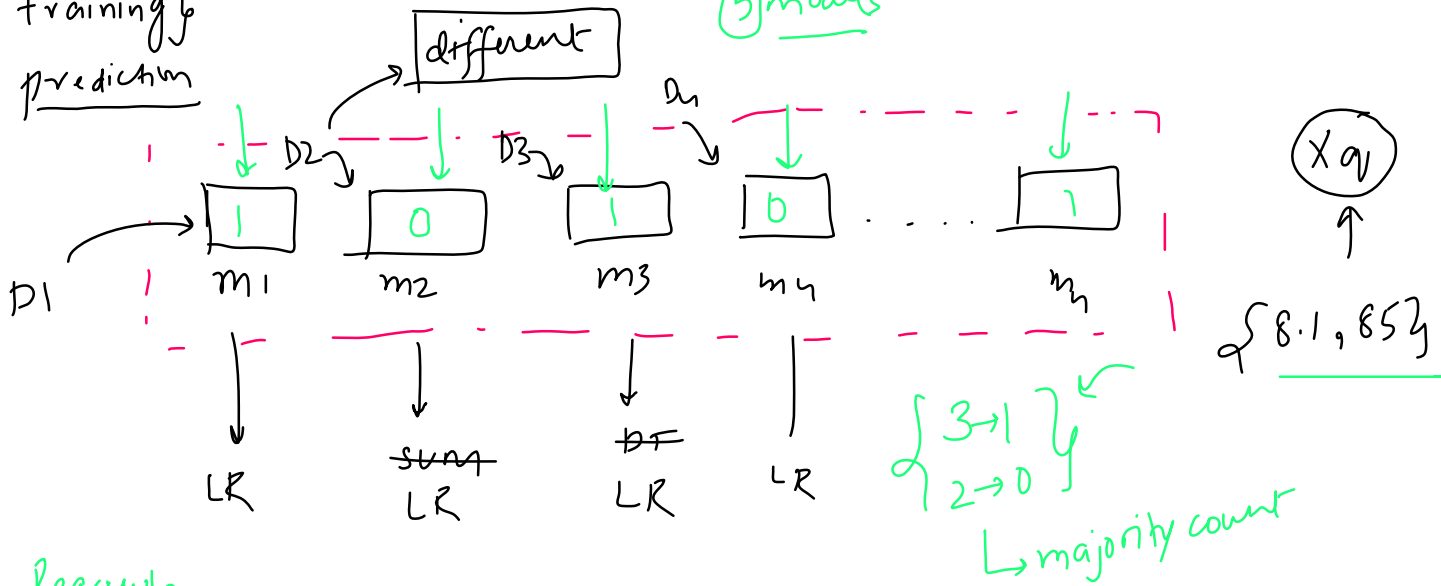
Monday, July 12, 2021 9:56 AM

Core Idea

Monday, July 12, 2021 9:45 AM

cgpa | iq | placement
lpa
5 models

training
prediction

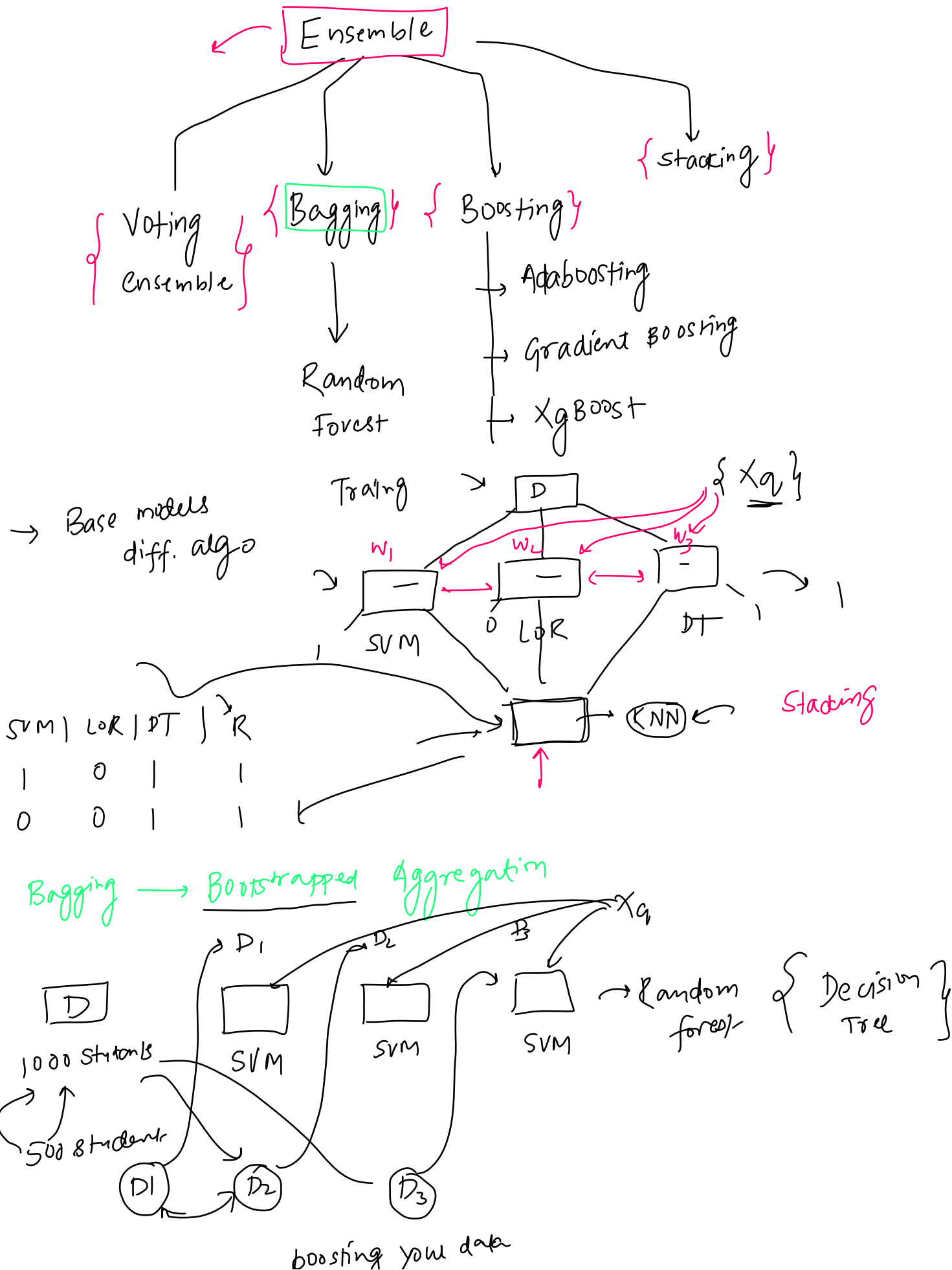


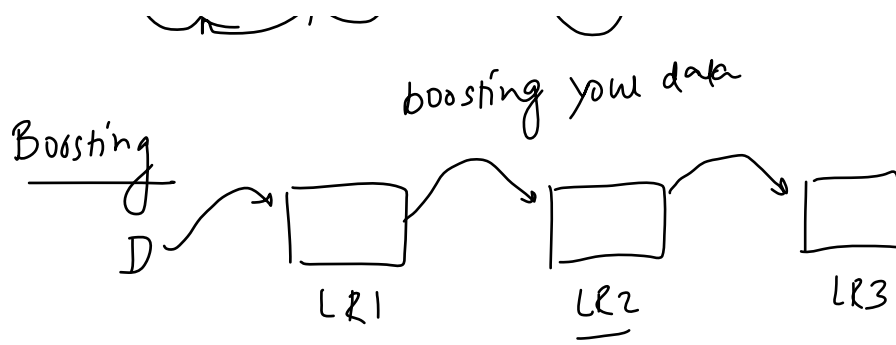
Regression

$$\begin{array}{c} \boxed{1} \quad \boxed{1} \quad \boxed{1} \quad \boxed{1} \\ \downarrow \quad \downarrow \quad \downarrow \quad \downarrow \\ 8.1 \quad + \quad 3.2 \quad + \quad 4.1 \quad + \quad 7.5 \\ \hline 4 \end{array} = \boxed{6.6} \text{ Lpa}$$

Type of Ensemble Learning

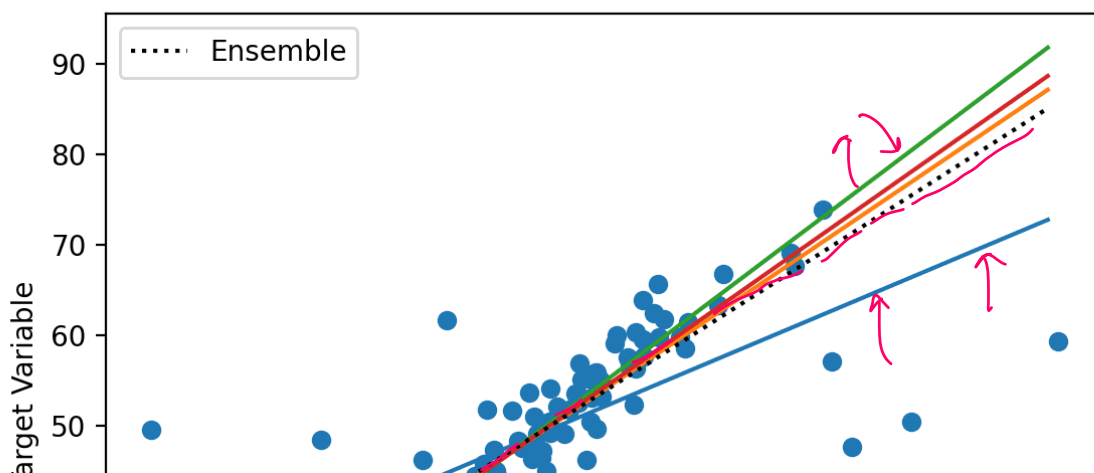
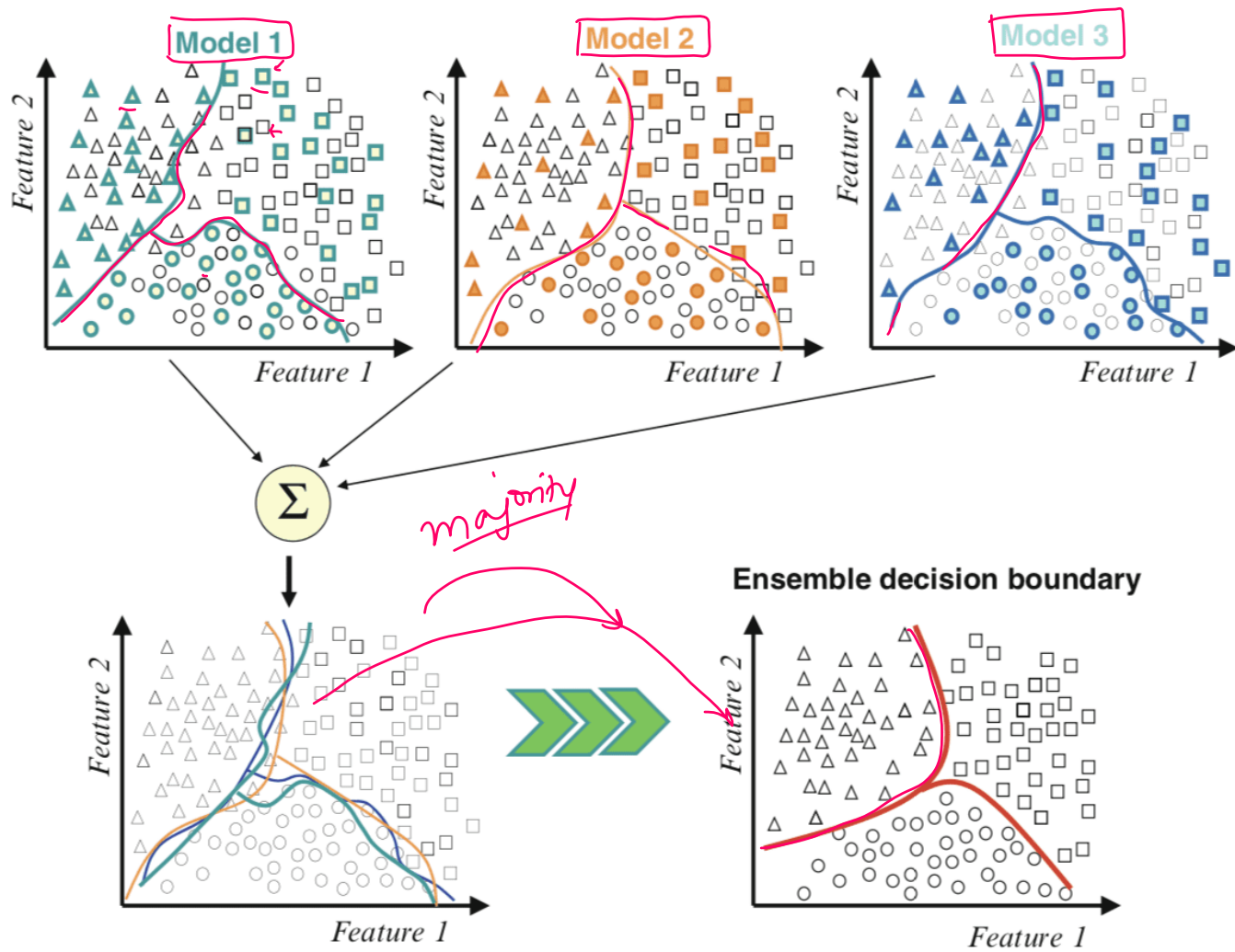
Monday, July 12, 2021 9:45 AM

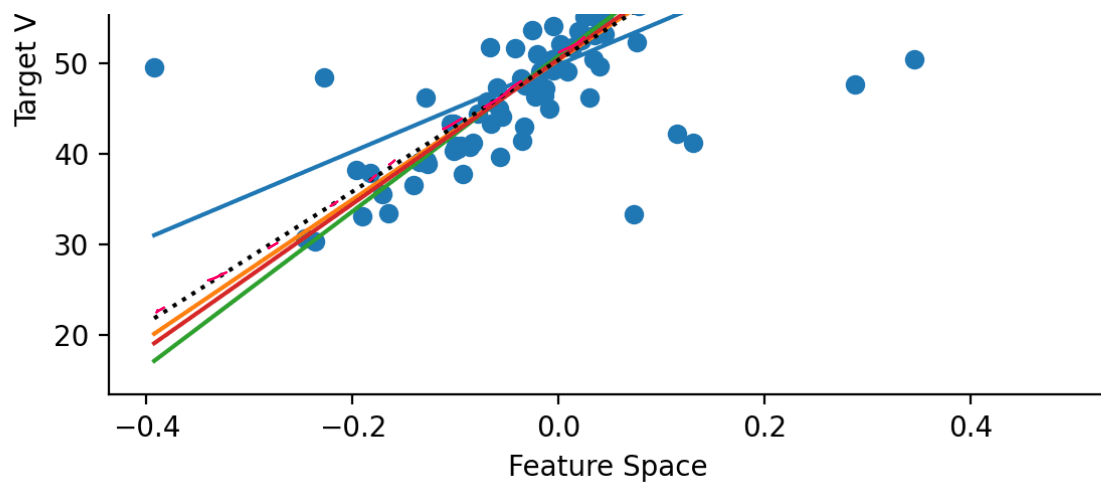




Why it works?

Monday, July 12, 2021 9:45 AM





Benefits

Monday, July 12, 2021 9:45 AM

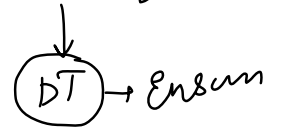
1) Improvement in performance

2) Bias Variance



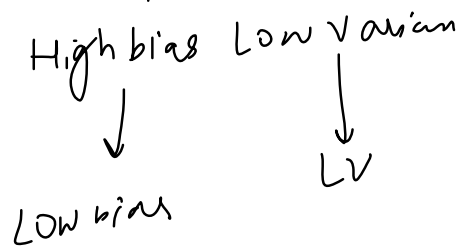
Low Bias + Low Variance

Low Bias High Variance



Low Bias → LV

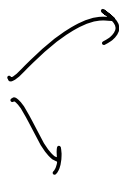
3) Robustness

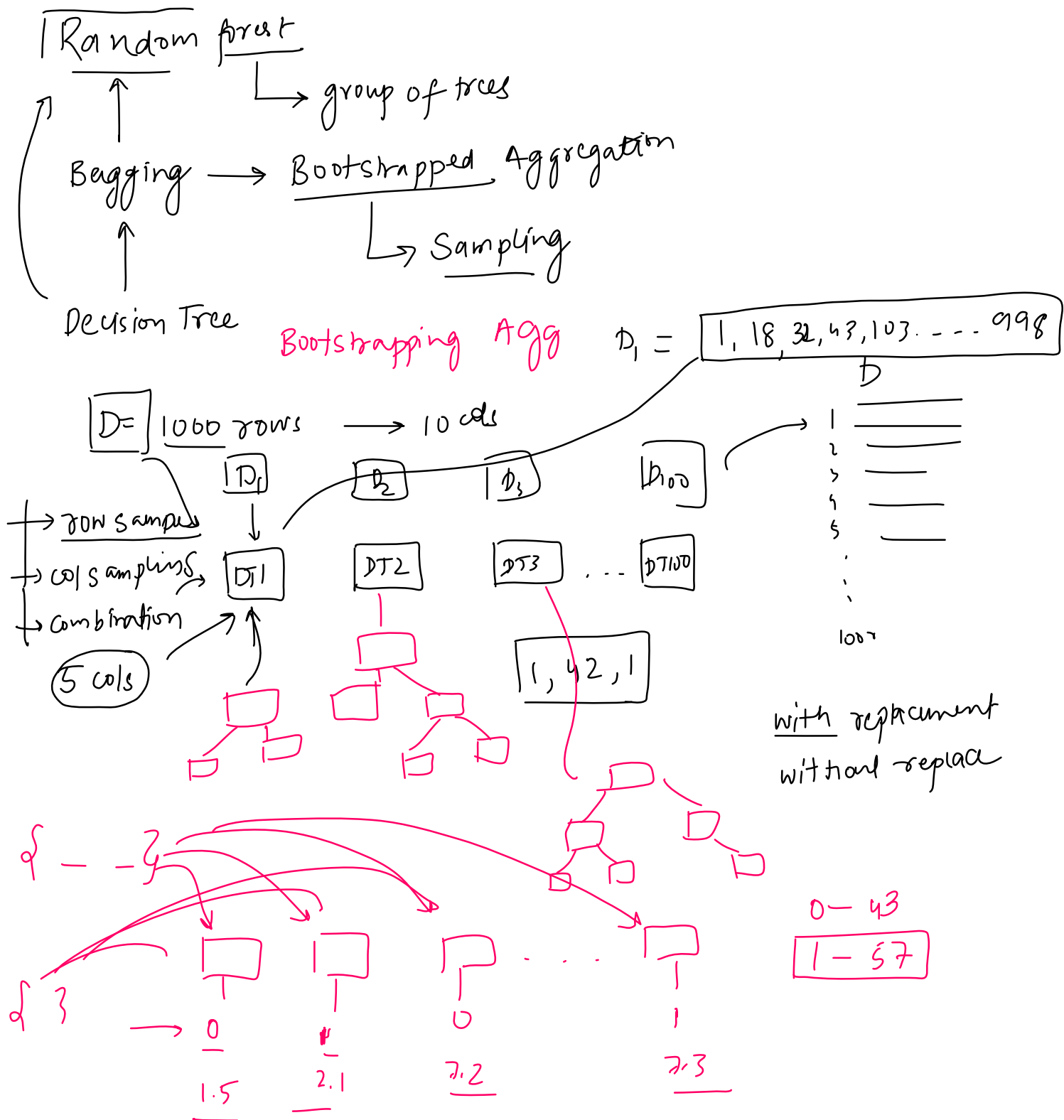


When to use?

Monday, July 12, 2021

9:46 AM

Always 



Demo

Monday, July 19, 2021

5:05 PM

Bias Variance and Random Forest

Tuesday, July 20, 2021 10:52 AM

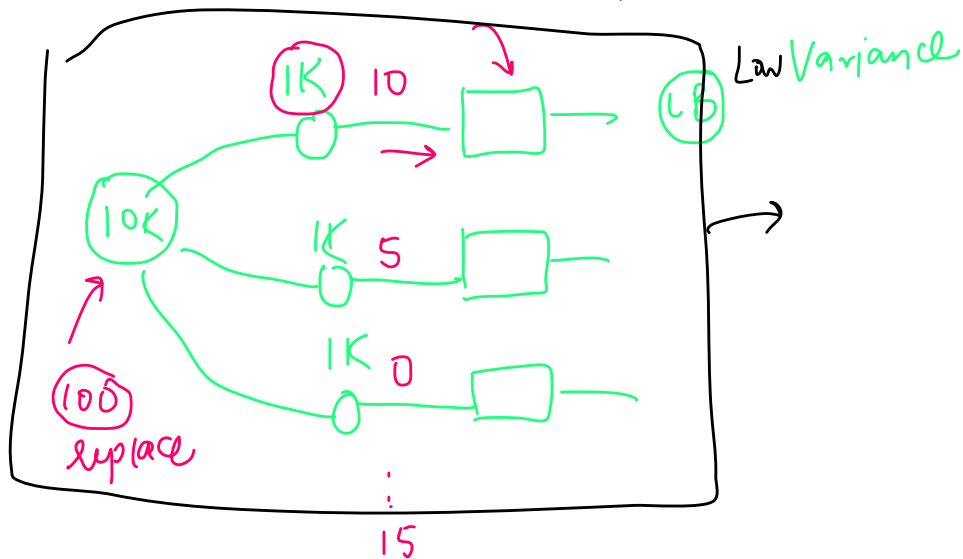
→ LB LV $B \propto \frac{1}{V}$

LB HV
→ Fully grown DT
→ SVM
→ KNN

HB LV
→ LR

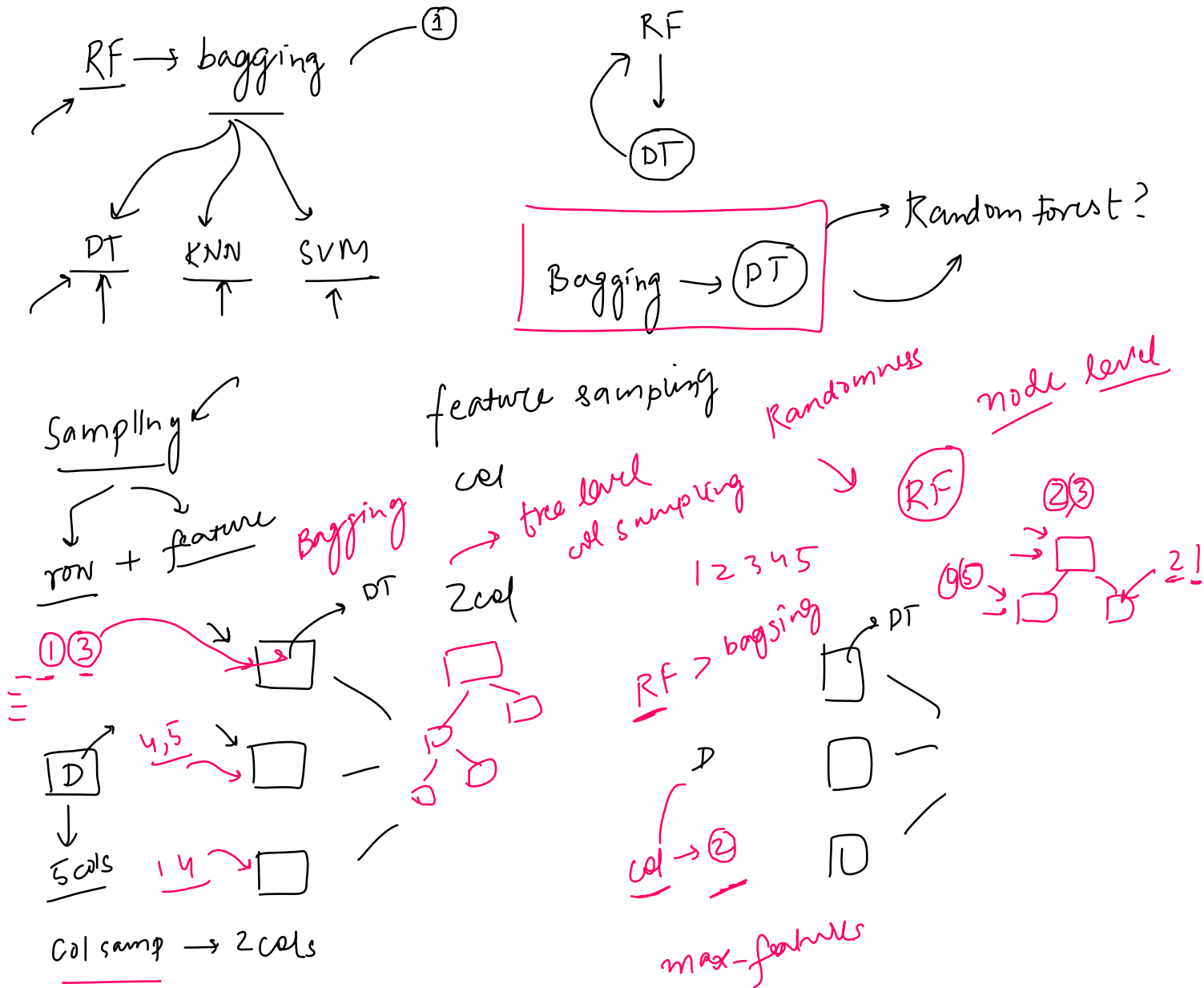
Random Forest

→ LB HV → LB LV



Bagging Vs Random Forest

Monday, July 19, 2021 4:53 PM



Hyperparameters

Monday, July 19, 2021 4:56 PM

Code Demo

Monday, July 19, 2021

5:05 PM

Monday, July 19, 2021 5:14 PM

Monday, July 19, 2021 5:14 PM



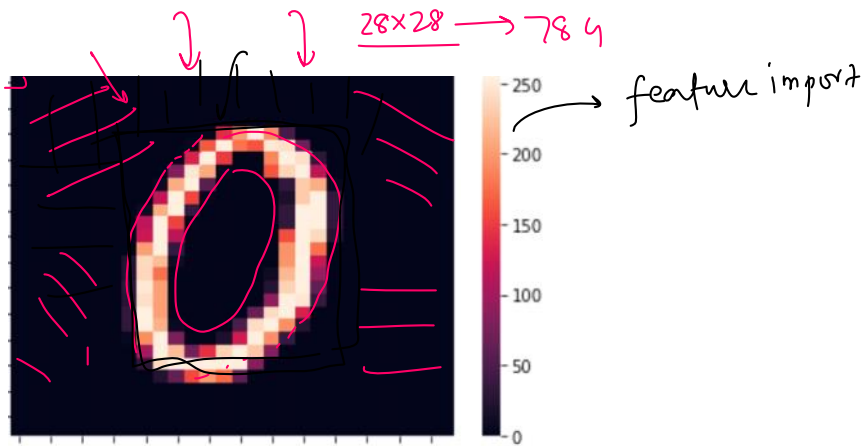
Feature Importance

Monday, July 19, 2021 4:56 PM

MNIST

label	pixel0	pixel1	pixel2	pixel3	pixel4	pixel5	pixel6	pixel7	pixel8	...	pixel774	pixel775	pixel776	pixel777	pixel778
0	1	0	0	0	0	0	0	0	0	...	0	0	0	0	0
1	0	0	0	0	0	0	0	0	0	...	0	0	0	0	0
2	1	0	0	0	0	0	0	0	0	...	0	0	0	0	0
3	4	0	0	0	0	0	0	0	0	...	0	0	0	0	0
4	0	0	0	0	0	0	0	0	0	...	0	0	0	0	0

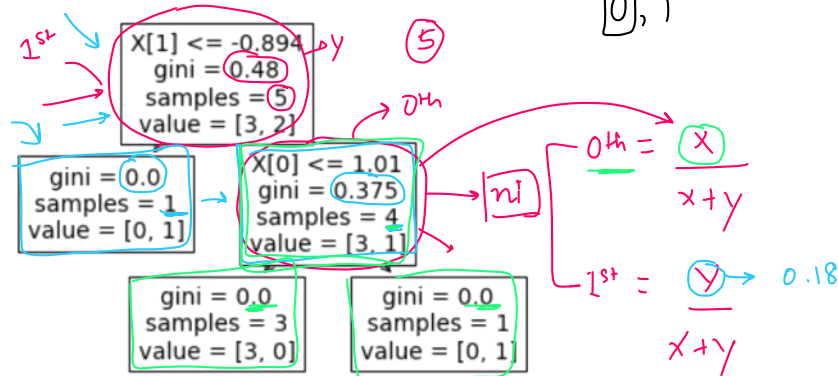
42000 images



Formula

$$ni = \frac{N-t}{N} \left[\text{impurity} - \left(\frac{N-t-L}{N-t} \times \text{right_impurity} \right) - \left(\frac{N-t-L}{N-t} \times \text{left_impurity} \right) \right]$$

$$fi_k = \frac{\sum_{j \in \text{node split on feature } k} ni}{\sum_{j \in \text{all nodes}} ni}$$



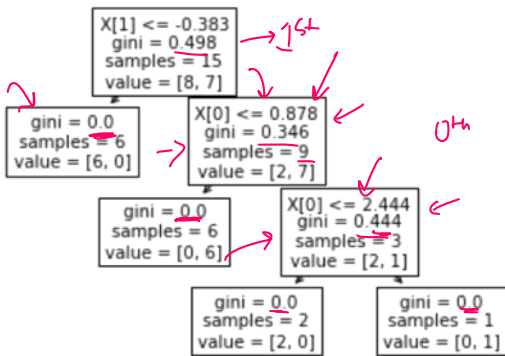
$$= \frac{5}{5} \left[0.48 - \frac{4}{5} \times 0.375 - \frac{1}{5} \times 0 \right] = Y \Rightarrow \left[0.48 - (0.8 \times 0.375) \right] \Rightarrow 0.48 - 0.30 = 0.18$$

$$\frac{4}{5} \left[0.375 \right] = X = 0.8 \times 0.375 = 0.30 = X$$

$$0^{th} = \frac{0.3}{0.3 + 0.18} = 0.625$$

$$0^{th} = \frac{0.3}{0.3 + 0.18} = 0.625$$

$$1^{st} = \frac{0.18}{0.3 + 0.18} = 0.375$$



0th Node

$$\frac{15}{15} \left[0.49 - \frac{9}{15} \times 0.37 \right]$$

$$= 0.290$$

$$\frac{9}{15} \left[0.34 - \frac{3}{15} \times 0.44 \right] = 0.118 \leftarrow$$

$$\frac{3}{15} [0.44] \leftarrow = 0.088$$

$$f_i[0] = \frac{0.11 + 0.08}{0.29 + 0.11 + 0.08} = \frac{0.19}{0.48} = 0.39$$

$$f_i[1] = \frac{0.29}{0.48} = 0.60$$

Extra Trees

Monday, July 19, 2021

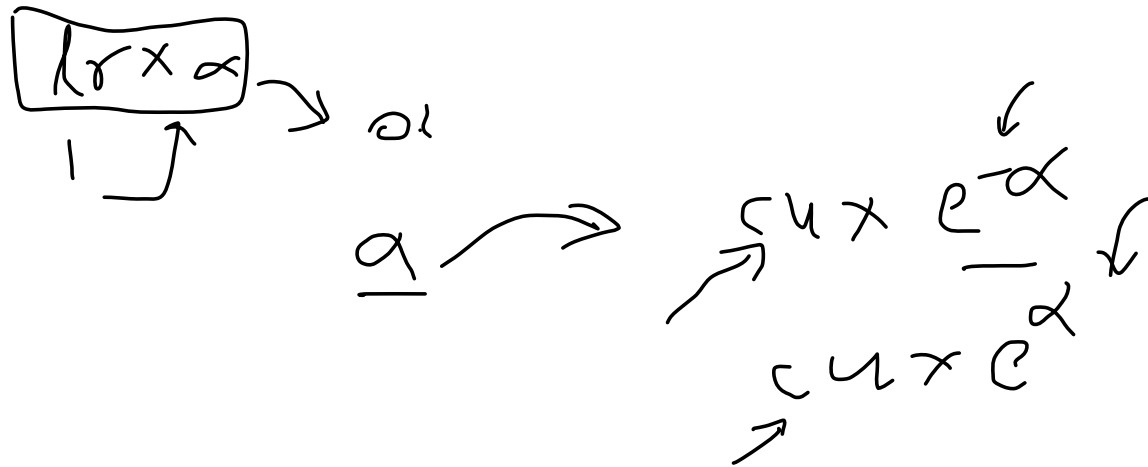
4:57 PM

Before Starting

Monday, August 9, 2021 10:42 AM

1) Weak classifiers ($> 50\%$)

$$\alpha = \frac{1}{2} \log \left(\frac{1 - e}{e} \right)$$



The Big Idea

Monday, August 9, 2021 10:43 AM

Step by Step Breakdown

Monday, August 9, 2021 10:43 AM

Points to remember

Monday, August 9, 2021 10:43 AM

Code Walkthrough

Monday, August 9, 2021 10:43 AM

Example and Hyperparameters

Monday, August 9, 2021 10:43 AM

Algorithm

Monday, September 20, 2021

8:23 AM

$$\rightarrow \{(x_i, y_i)\}_{i=1}^n \quad \eta=3 \quad x_i, y_i$$

Input: training set $\{(x_i, y_i)\}_{i=1}^n$ a differentiable loss function $L(y, F(x))$, number of iterations M .

→ 1. Initialize $f_0(x) = \arg \min_{\gamma} \sum_{i=1}^N L(y_i, \gamma)$.

→ 2. For $m = 1$ to M :

→ (a) For $i = 1, 2, \dots, N$ compute

$i/m \rightarrow$
row no

$$\rightarrow r_{im} = - \left[\frac{\partial L(y_i, f(x_i))}{\partial f(x_i)} \right]_{f=f_{m-1}}$$

residual / pseudo-residual

(b) Fit a regression tree to the targets r_{im} giving terminal regions R_{jm} , $j = 1, 2, \dots, J_m$.

(c) For $j = 1, 2, \dots, J_m$ compute

$$\rightarrow \gamma_{jm} = \arg \min_{\gamma} \sum_{x_i \in R_{jm}} L(y_i, f_{m-1}(x_i) + \gamma).$$

(d) Update $f_m(x) = f_{m-1}(x) + \sum_{j=1}^{J_m} \gamma_{jm} I(x \in R_{jm})$.

3. Output $\hat{f}(x) = f_M(x)$.

$$\textcircled{1} \quad f_1(x) = f_0(x) + \textcircled{dT}$$

$$f_3(x) = f_2(x) + dT_3$$

$$f_2(x) = f_1(x) + dT_2$$

$$f_1(x) + dT_2$$

$$\checkmark \quad f_0(x) + dT$$

$$f_0(x) + dT_1$$

→ recursion

$$f_4(x) = f_0(x) + \dots$$

$$f_1(x) \quad f_2(x)$$

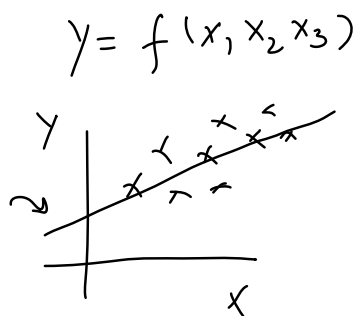
Additive Modelling

Monday, September 20, 2021 8:28 AM

$$\begin{array}{c} \overbrace{x}^{\curvearrowright} \\ \downarrow \\ y \end{array} \rightarrow f(\cdot) \Rightarrow y = f(x)$$

x, y

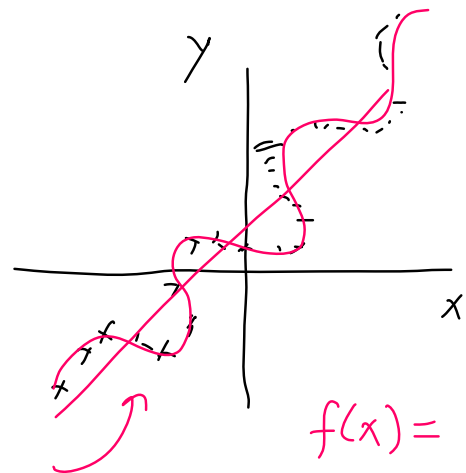
$x_1, x_2, x_3 | y$



additive

$$F(x) = f_0(x) + f_1(x) + f_2(x) + \dots$$

$\uparrow$ $\uparrow$
 D_T P_T



$$f(x) = x + \sin x$$

$$y = x \quad y = \sin(x)$$

Explanation

Monday, September 20, 2021 8:25 AM

$$L = \sum_{i=1}^n (y_i - \hat{y}_i)^2$$

$\uparrow$ actual $\downarrow$ pred

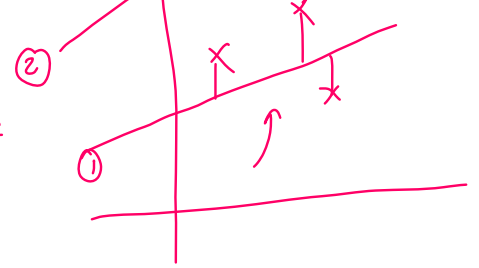
ls

$$L = \left(\frac{1}{2}\right) \sum_{i=1}^n (y_i - \hat{y}_i)^2 \rightarrow \text{diff}$$

$\frac{1}{2}$

$$L(y, F(x)) \rightarrow L(y, \hat{y})$$

$\uparrow$
 $\hat{y}$



$$\textcircled{1} = 10 = 5 \checkmark$$

$$\textcircled{2} = 20 = 10$$

$$y = f(x) \rightsquigarrow$$

$$f(x) = f_0(x) + \underbrace{f_1(x) + f_2(x) + \dots + f_n(x)}_{DT}$$

$$f_0(x) = \arg\min_{\gamma} \sum_{i=1}^n L(y_i, \gamma)$$

$$L = \frac{1}{2} \sum_{i=1}^n (y_i - \hat{y}_i)^2$$

$$f_0(x) = \arg\min_{\gamma} \left[\frac{1}{2} \sum_{i=1}^n (y_i - \gamma)^2 \right]$$

$\uparrow$ γ

$$\frac{d f_0(x)}{d \gamma} = \frac{d}{d \gamma} \left[\frac{1}{2} \sum_{i=1}^n (y_i - \gamma)^2 \right] = \frac{1}{2} \sum_{i=1}^n \frac{d}{d \gamma} (y_i - \gamma)^2$$

$$\sum_{i=1}^n (y_i - \gamma) \frac{d}{d \gamma} (y_i - \gamma) = - \sum_{i=1}^n (y_i - \gamma) = 0$$

$$\sum_{i=1}^n (\gamma - y_i) = 0$$

$$\sum_{i=1}^3 (\gamma - y_i) = 0 \Rightarrow (\gamma - 192) + (\gamma - 144) + (\gamma - 91) = 0$$

$$\sum_{i=1}^n (y - y_i) = 0 \Rightarrow (y - 192) + (y - 144) + \dots$$

$$3y = 192 + 144 + 91$$

mean

$F(x) = f_0(x)$

mean of output

$$y = \frac{192 + 144 + 91}{3}$$

$$F(x) = \underbrace{f_0(x)}_{\text{mean (leaf)}} + \underbrace{f_1(x)}_{b_1} + \underbrace{f_2(x)}_{b_2} + \dots + \underbrace{f_m(x)}_{b_m}$$

mean (leaf)

$$m = 1$$

$$\sigma_{im} = - \left[\frac{\partial L(y_i, f(x_i))}{\partial f(x_i)} \right]_{f=f_{m-1}}$$

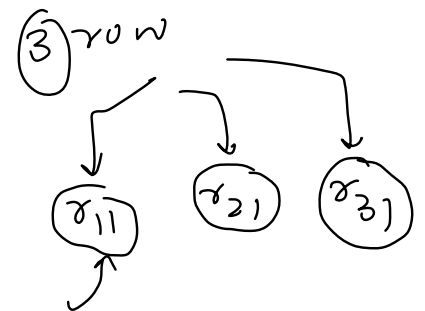
$$\sigma_{i1} = - \left[\frac{\partial L(y_i, f(x_i))}{\partial f(x_i)} \right]_{f=f_0}$$

$$\hat{y}_i = f(x_i)$$

$$\sigma_{i1} = - \left[\frac{\partial L(y_i, \hat{y}_i)}{\partial \hat{y}_i} \right]_{f=f_0}$$

$$\sigma_{i1} = - \left[\frac{\partial}{\partial \hat{y}_i} \left(\frac{1}{2} (y_i - \hat{y}_i)^2 \right) \right]_{f=f_0}$$

$$= \left[(y_i - \hat{y}_i) \right]_{f=f_0} = \left[(y_i - f(x_i)) \right]_{f=f_0}$$



$$L = \frac{1}{2} \sum_{i=1}^n (y_i - \hat{y}_i)^2$$

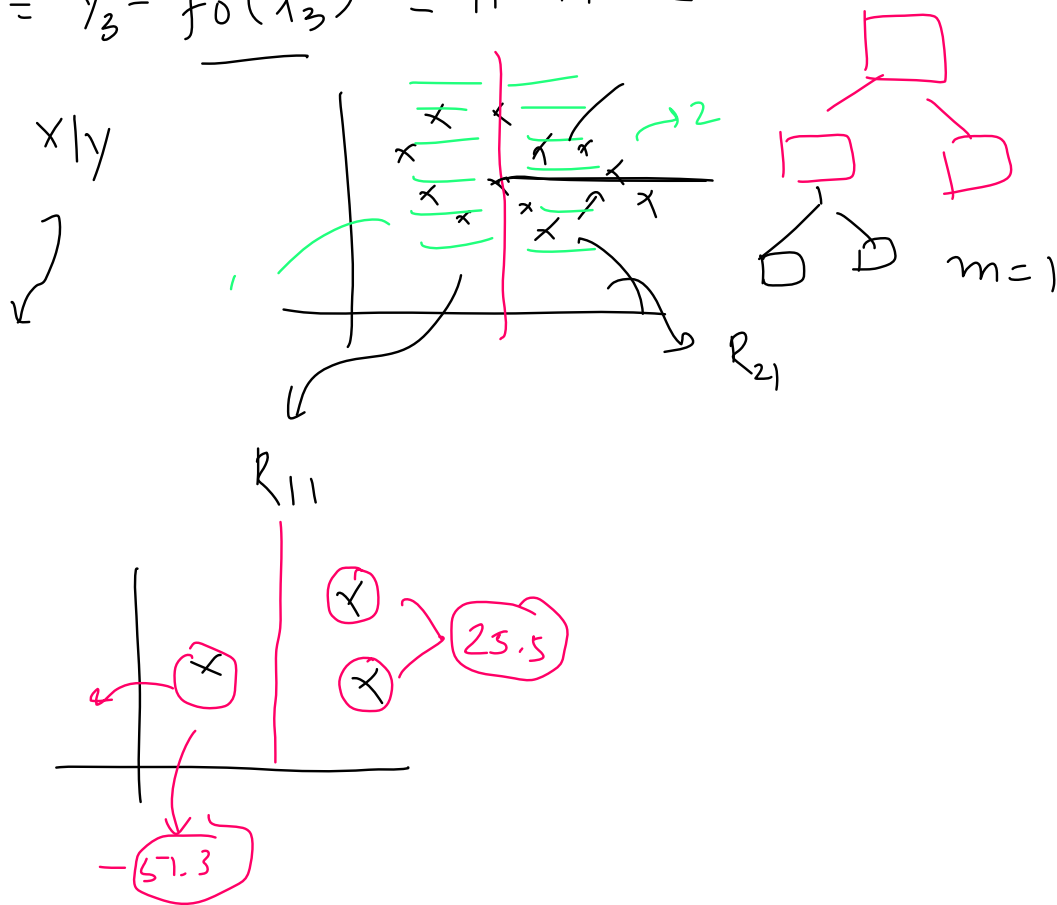
$$= \begin{bmatrix} (y_i - \hat{y}_i) \end{bmatrix} f = f_0 = \begin{bmatrix} (y_i - \hat{f}^{(1)}) \end{bmatrix} f = \underline{f_0}$$

$$r_{i1} = \underline{(y_i - f_0(x_i))}$$

$$r_{11} = y_1 - \underline{f_0(x_1)} = 192 - 142 =$$

$$r_{21} = y_2 - \underline{f_0(x_2)} = 144 - 142 =$$

$$r_{31} = y_3 - \underline{f_0(x_3)} = 91 - 142 =$$



$$\underline{\gamma_{jm}} = \underset{\gamma}{\operatorname{argmin}} \sum_{x_i \in R_{jm}} L(y_i, f_{m-1}(x_i) + \gamma)$$

$\gamma_{j1} \Rightarrow \underline{\gamma_{11}} = \underset{\textcircled{\gamma}}{\operatorname{argmin}} \sum_{x_i \in R_{11}} L(y_i, f_{\underline{m-1}}(x_i) + \gamma)$

Regions: γ_{11} , γ_{21}

↓
 γ_{11}

↓
 γ_{21}

⌈

↗ ↖

$$\gamma_{11} = \arg \min_{\gamma} \frac{1}{2} (\gamma_1 - (f_0(x_1) + \gamma))^2$$

$$\frac{dL}{d\gamma} = \frac{1}{2} \times 2 (\gamma_1 - (f_0(x) + \gamma)) \frac{d}{d\gamma} (\gamma_1 - f_0(x) - \gamma) = 0$$

$$\begin{aligned} &= -(\gamma_1 - f_0(x) - \gamma) = 0 \\ &= \gamma_1 - f_0(x) - \gamma = 0 \end{aligned}$$

$$\gamma_{11} = 91 - 142 - \gamma = 0$$

$$\boxed{\gamma = 91 - 142 = -51}$$

$$\gamma_{21} = \arg \min_{\gamma} \sum_{x_i \in \mathcal{R}_{21}} L(\gamma_i, f_0(x_i) + \gamma)$$

$$= \arg \min_{\gamma} \frac{1}{2} \sum_{i=1}^2 (\gamma_i - (f_0(x_i) + \gamma))^2$$

$$= - \sum_{i=1}^2 (\gamma_i - f_0(x_i) - \gamma) = 0$$

$$= \sum_{i=1}^2 (\gamma_i - f_0(x_i) - \gamma) = 0$$

$$= \gamma_1 - f_0(x_1) - \gamma + \gamma_2 - f_0(x_2) - \gamma = 0$$

336
284

$$= 192 - 142 - \gamma + 144 - 142 - \gamma = 0$$

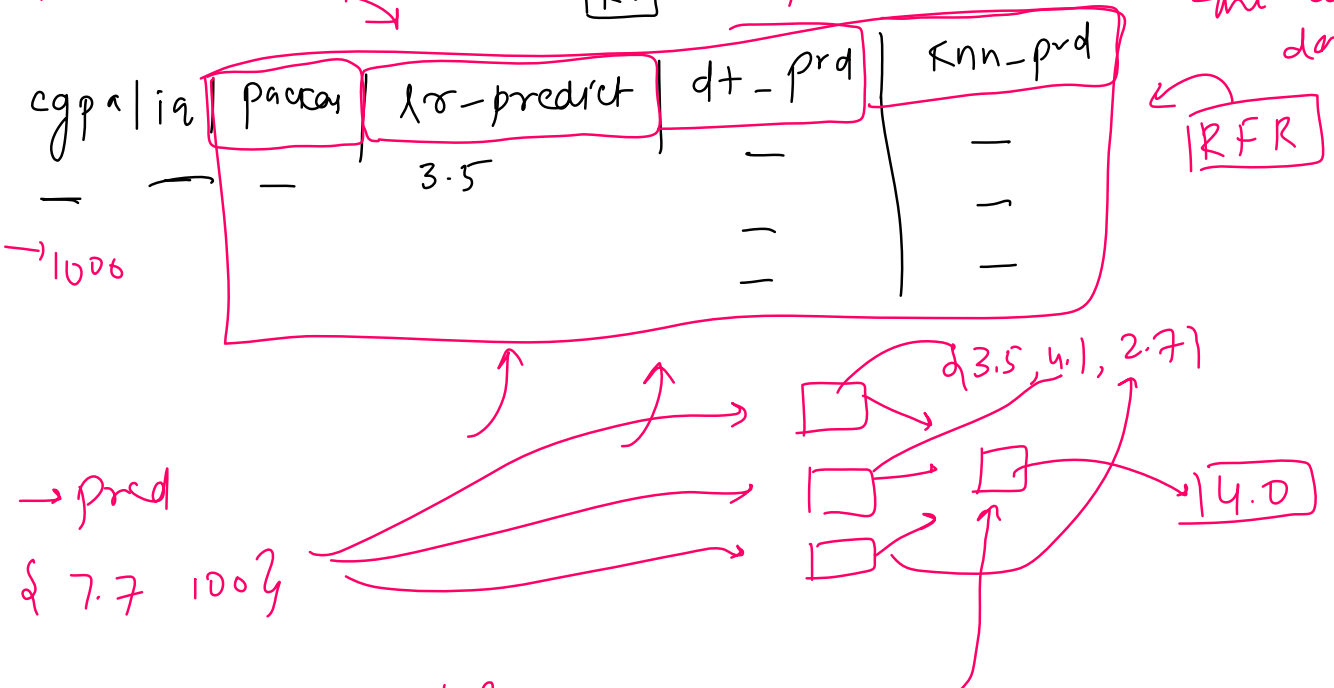
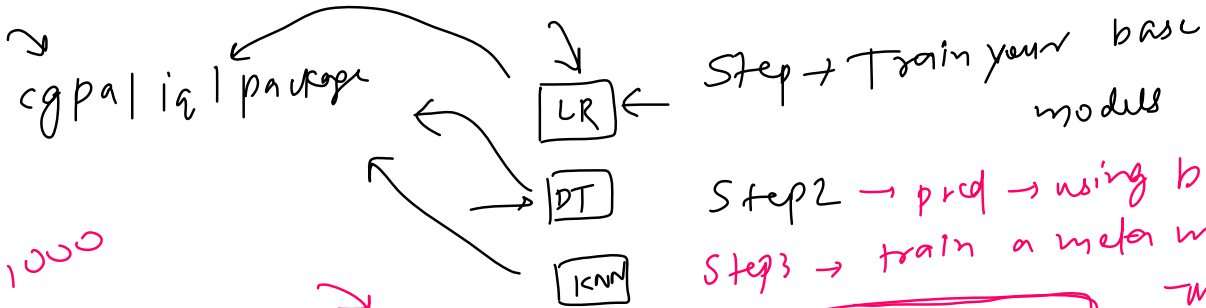
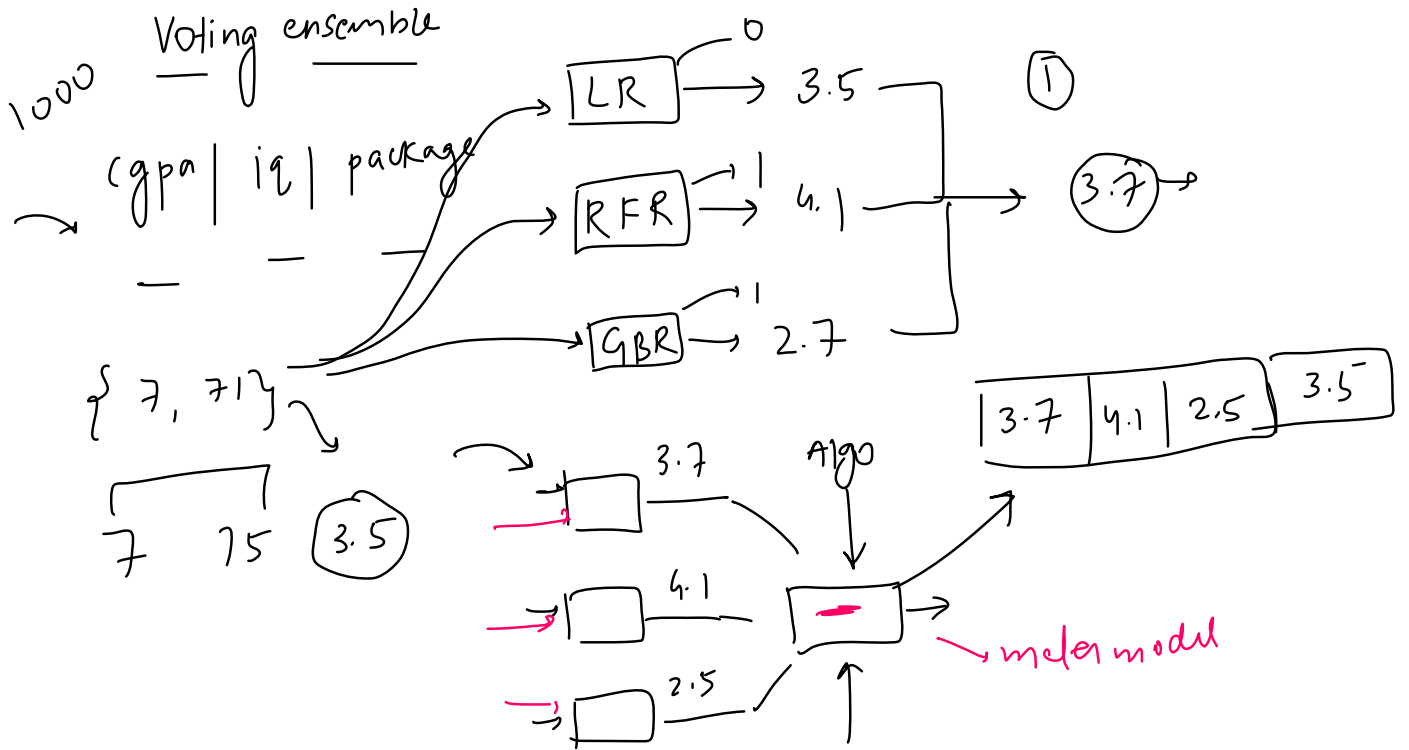
$$336 - 284$$

└───┘

$$52 - 2\gamma = 0$$

$$\gamma = \frac{52}{2} = 26$$

Stacking



Boosting / Bagging

→ base model

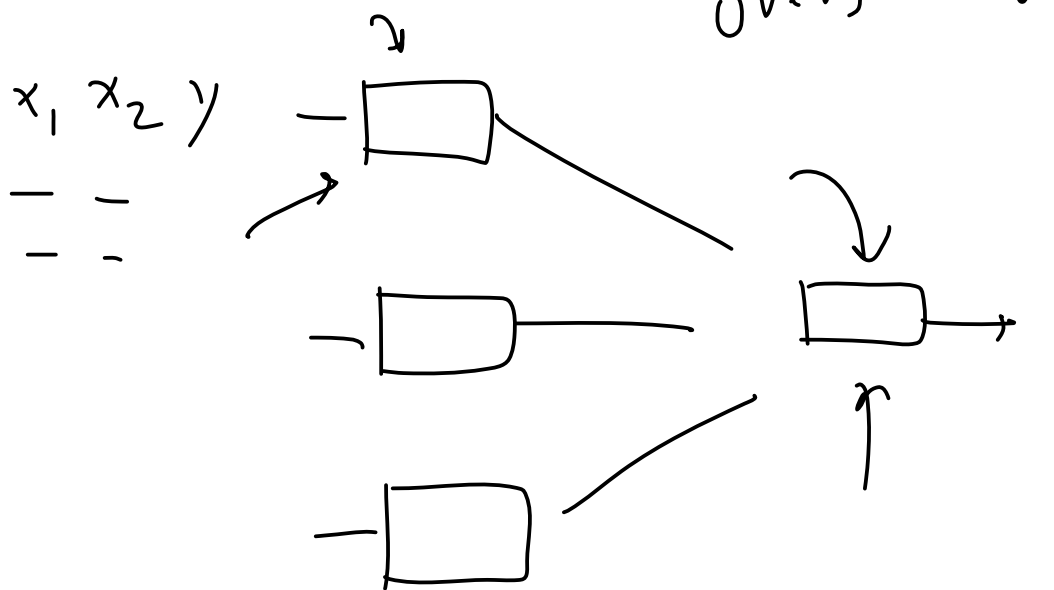
→ majority

— —
→ training

Problem

Thursday, September 30, 2021

3:48 PM

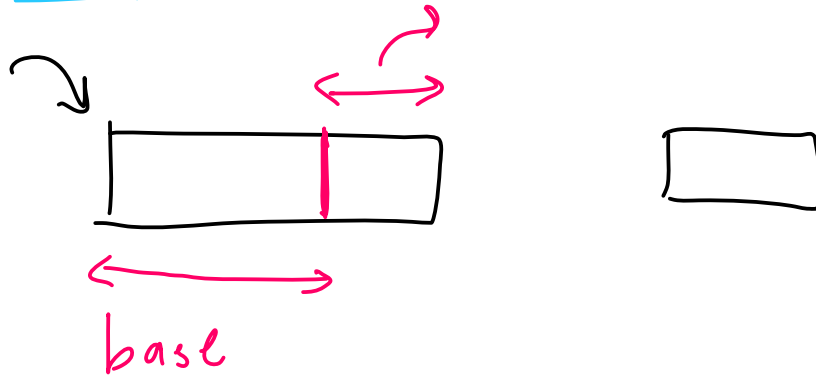


Solutions

Thursday, September 30, 2021

3:48 PM

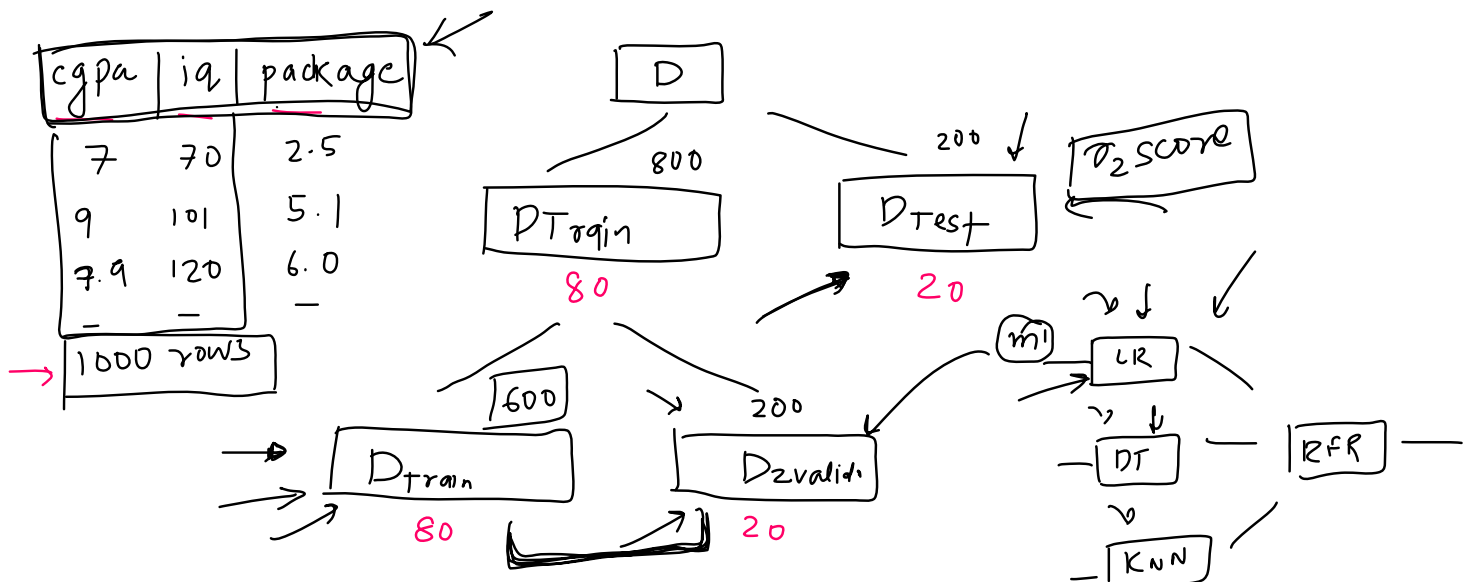
→ Hold out method → Blending



→ K-fold method → stacking

Hold Out Approach - Blending

Thursday, September 30, 2021 3:48 PM



Step 1 → train 3 base models (3 trained on D_{train})

Step 2 → You form new dataset

Step 3 → Train the meta-model on $\hookrightarrow$ meta-models

package	LR-pred	DT-pred	KNN-pred

(200, 4)

Blending

K Fold Approach - Stacking

Thursday, September 30, 2021 3:49 PM

$K = 5$

D

cgpa | iq | package

7	70	2.5
9	101	5.1
7.9	120	6.0

D_{train}

D_{test}

$K = 4$

Step 1

1000 rows

LR

DT

knn

RFR

200 200 200 200

12 models

$y \rightarrow l8$

$y \rightarrow dt$

$y \rightarrow knn$

lr-pred dt knn pack

400

600

RFR

(800, 4)

meta-model trained

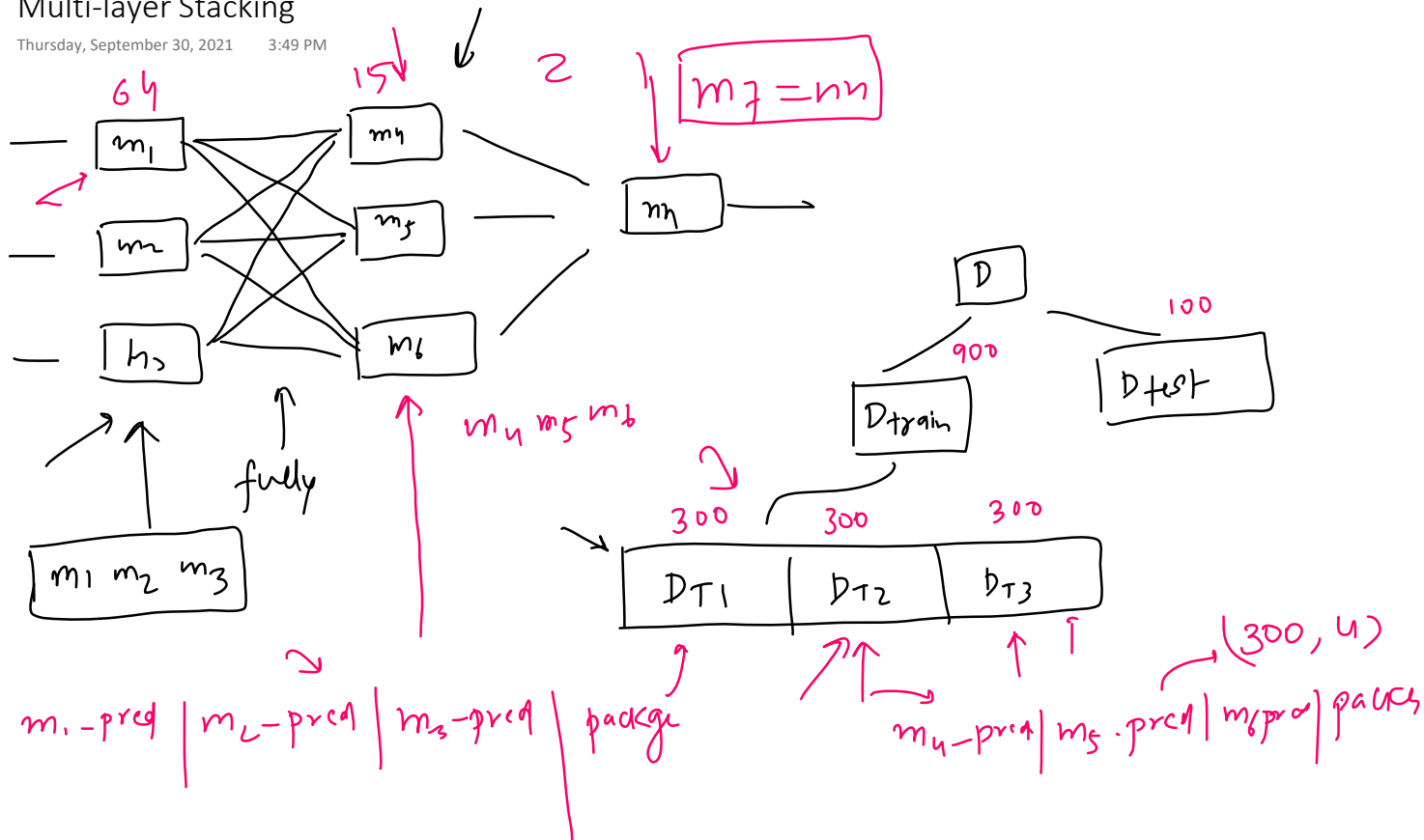
Step 3

retrain X_{train}, y_{train}
base model

x_2 -score

Multi-layer Stacking

Thursday, September 30, 2021 3:49 PM

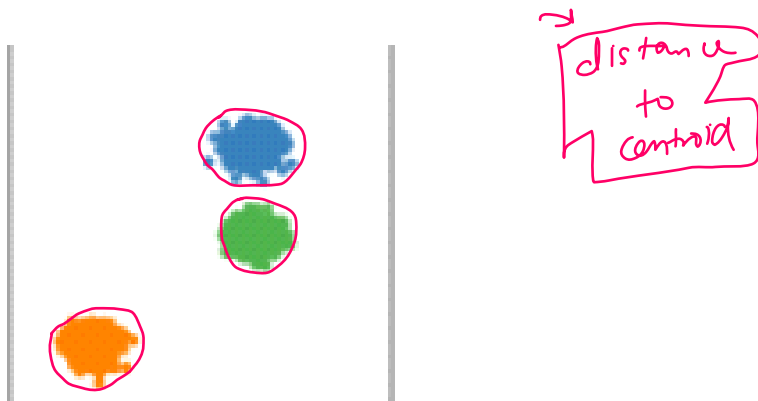


Demo

Thursday, September 30, 2021 3:49 PM

Need of Other Clustering Methods

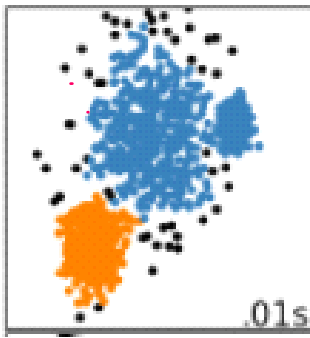
Saturday, November 6, 2021 7:38 AM



2 More clustering methods:

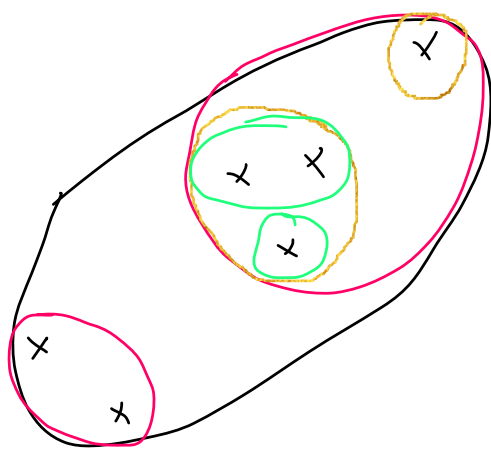
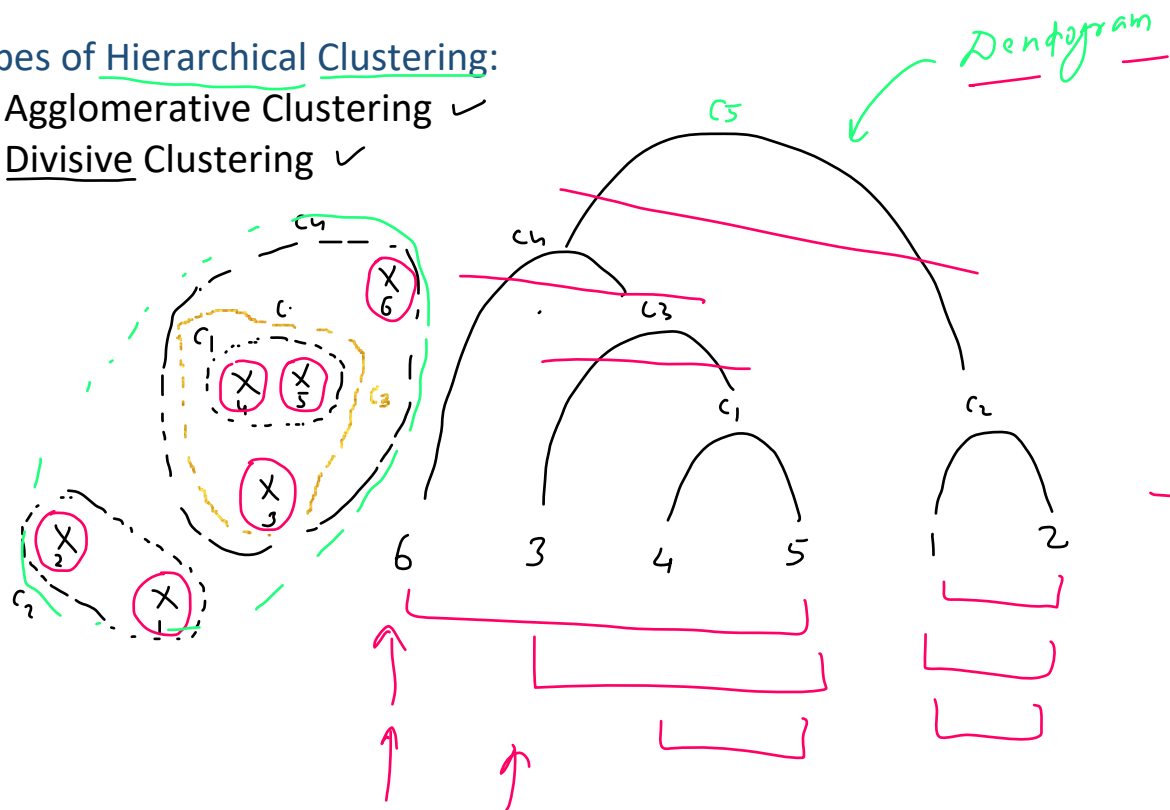
1. Hierarchical Clustering
2. Density Based Clustering DBSCAN





Types of Hierarchical Clustering:

1. Agglomerative Clustering ✓
2. Divisive Clustering ✓



1 cluster
↓
2 clusters
↓
3 clusters
↓
4 clusters
↓
6 clusters

Algorithm

Saturday, November 6, 2021 7:39 AM

1. Initialize the Proximity Matrix
2. Make each point a cluster
3. Inside a loop
 - a. Merge the 2 closest clusters
 - b. Update the Proximity Matrix
4. Until only one cluster is left

n points $n \times n$

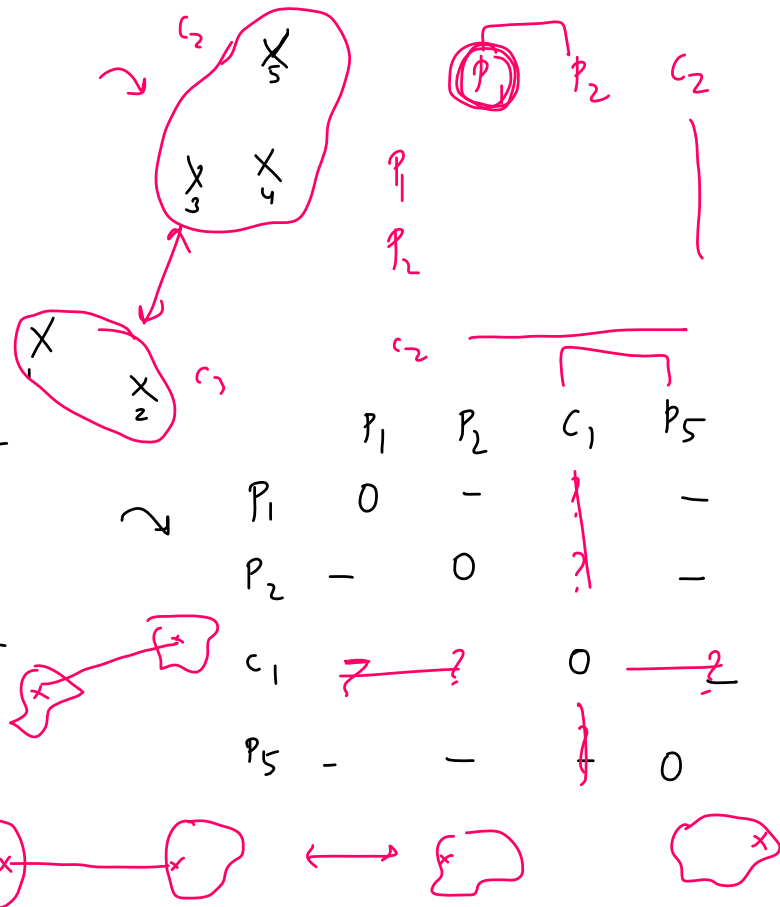
	p_1	p_2	p_3	p_4	p_5
p_1	0	-	-	-	-
p_2	-	0	-	-	-

p_3
 p_4

95

c_2 c_3

c_2
 c_3



Types

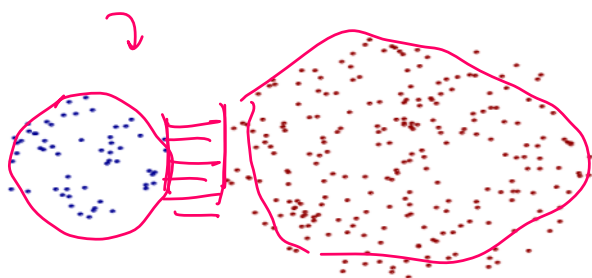
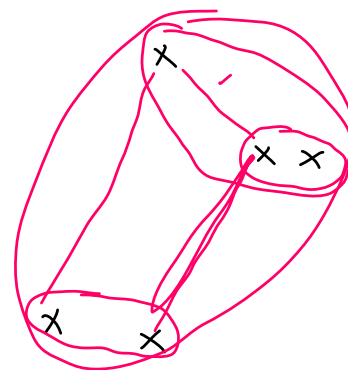
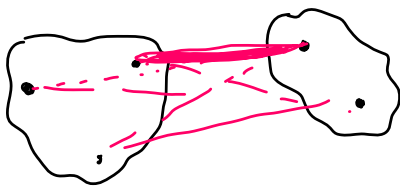
Saturday, November 6, 2021 7:39 AM

Types of Agglomerative Clustering

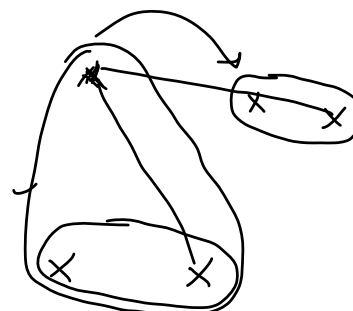
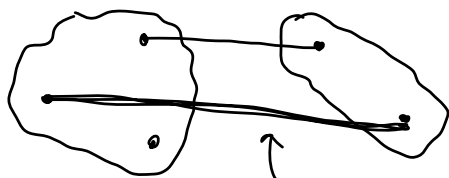
1. Min (Single-link)
2. Max (Complete Link)
3. Average
4. Ward

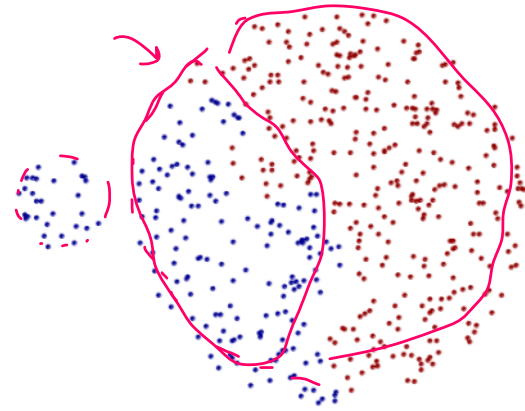
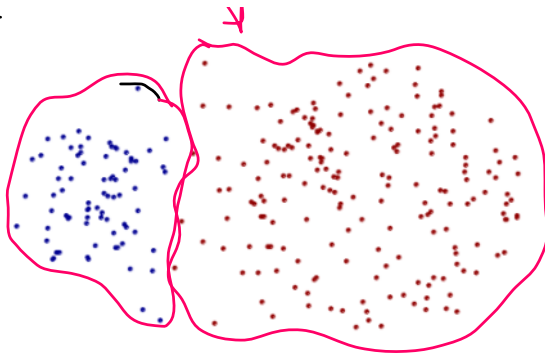
1. Single link

$$3 \times 2 = 6$$

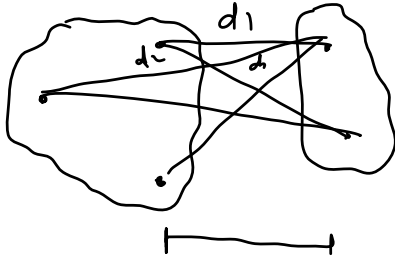


2. Complete link (Max)





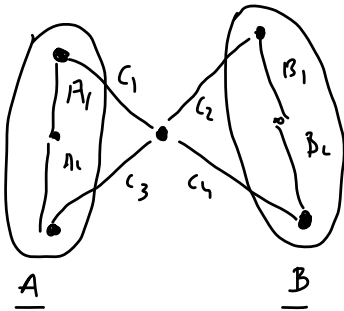
3. Group Average -



2x3

$$\frac{d_1 + d_2 + d_3 + \dots + d_l}{\boxed{2 \times 3}} =$$

4. Ward → SK Urm



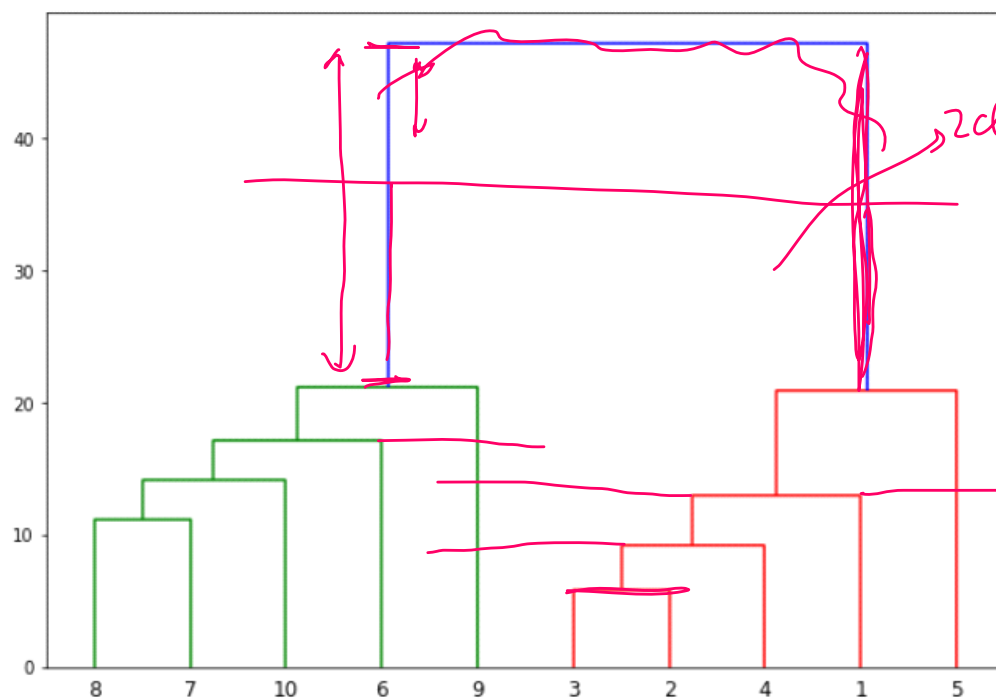
$$\text{dist}_{A \rightarrow B} = c_1^2 + c_2^2 + c_3^2 + c_4^2 - A_1^2 - A_2^2 - B_1^2 - B_2^2$$

W

Variance minimize

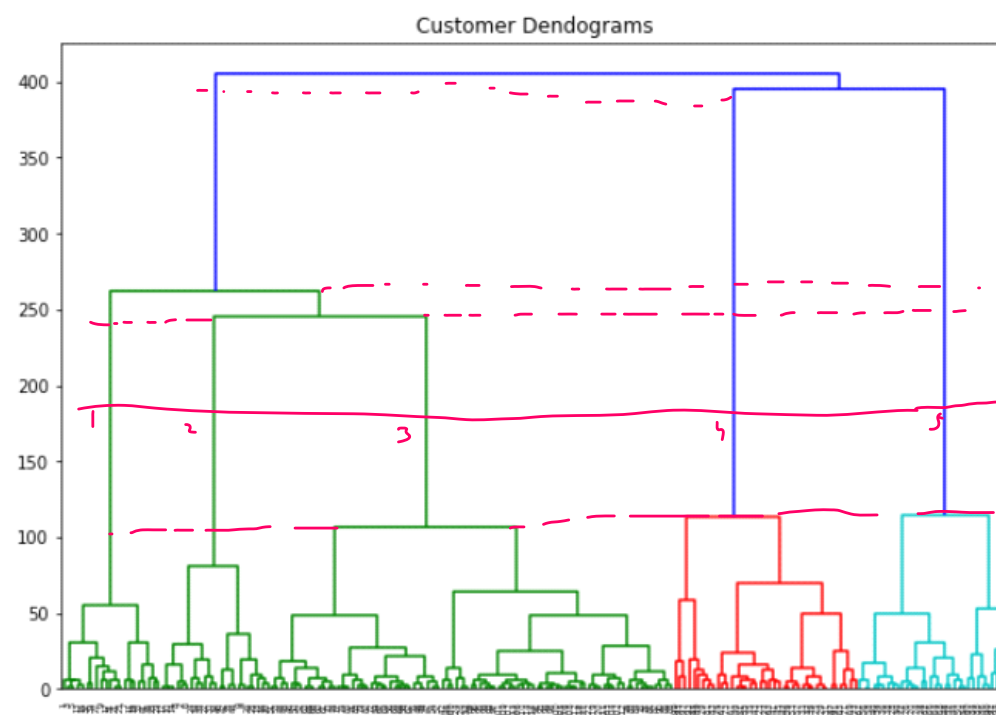
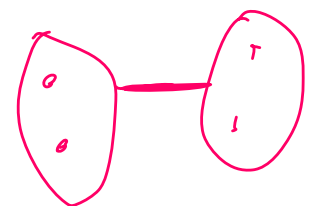
How to find the ideal number of clusters

Saturday, November 6, 2021 1:50 PM



Kmeans →
elbow method

inter cluster
similarity



5 clusters

Hyperparameter

Saturday, November 6, 2021 7:40 AM

Code Example

Saturday, November 6, 2021

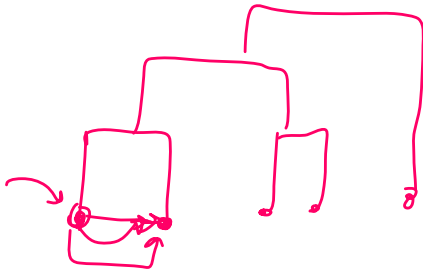
7:40 AM

Benefits/Limitations

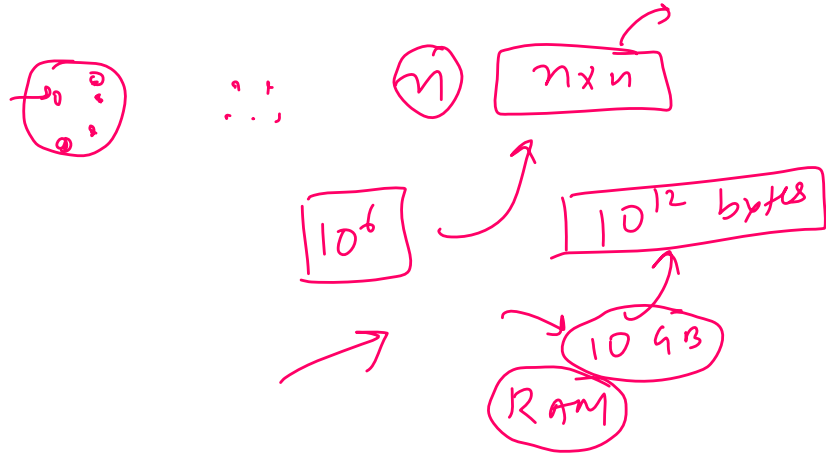
Saturday, November 6, 2021 7:39 AM

Benefits

- 1) Widely applicable
- 2) Dendrogram

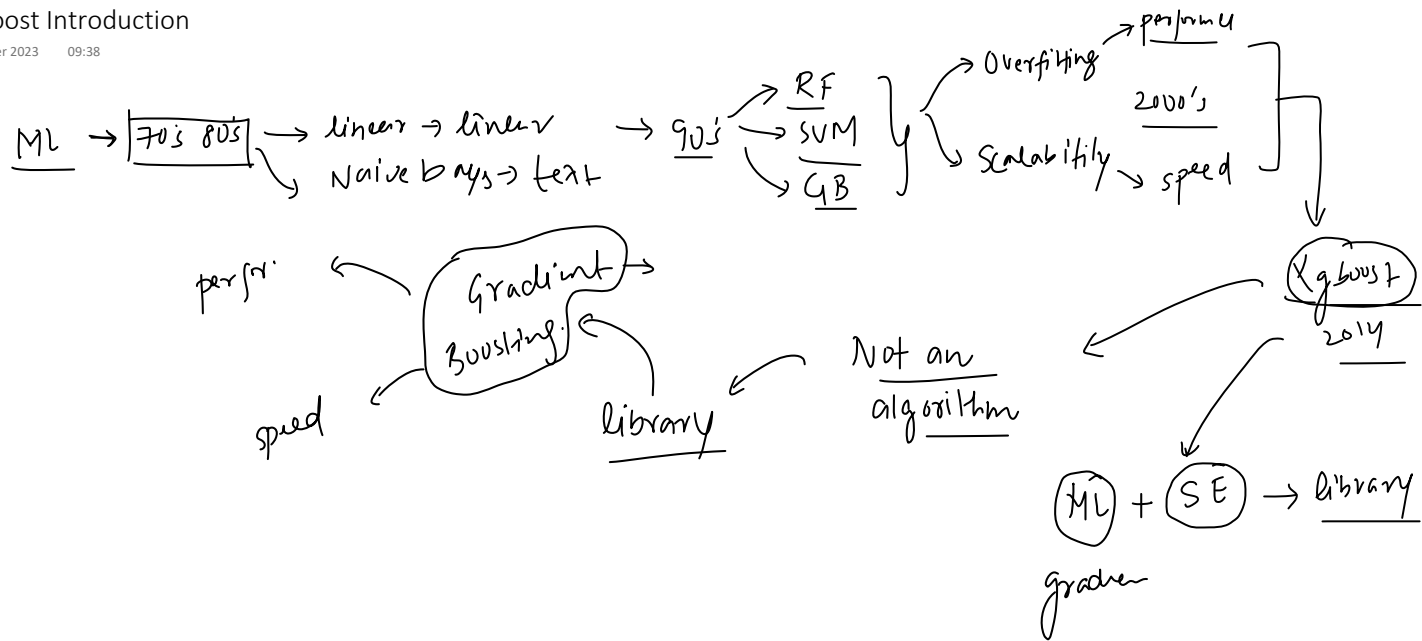


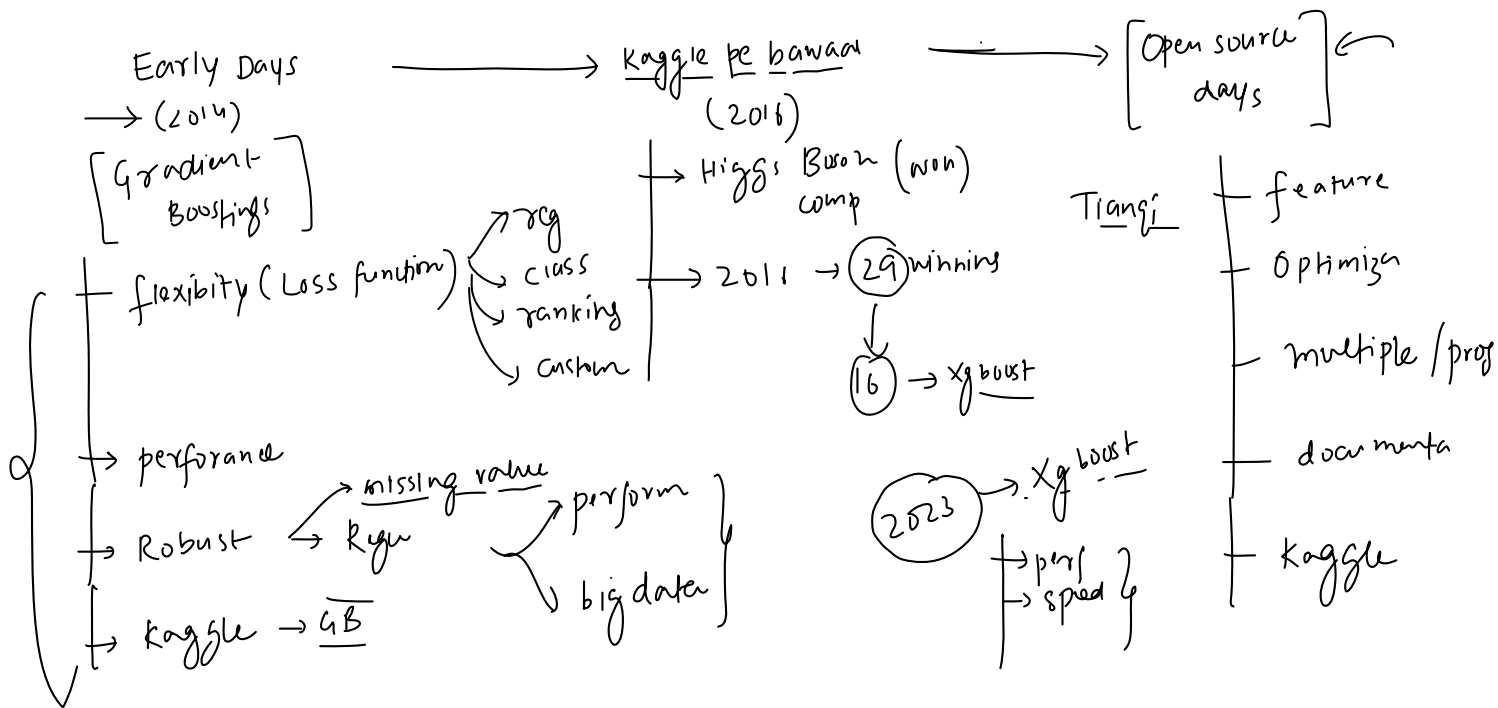
Limitation



XGBoost Introduction

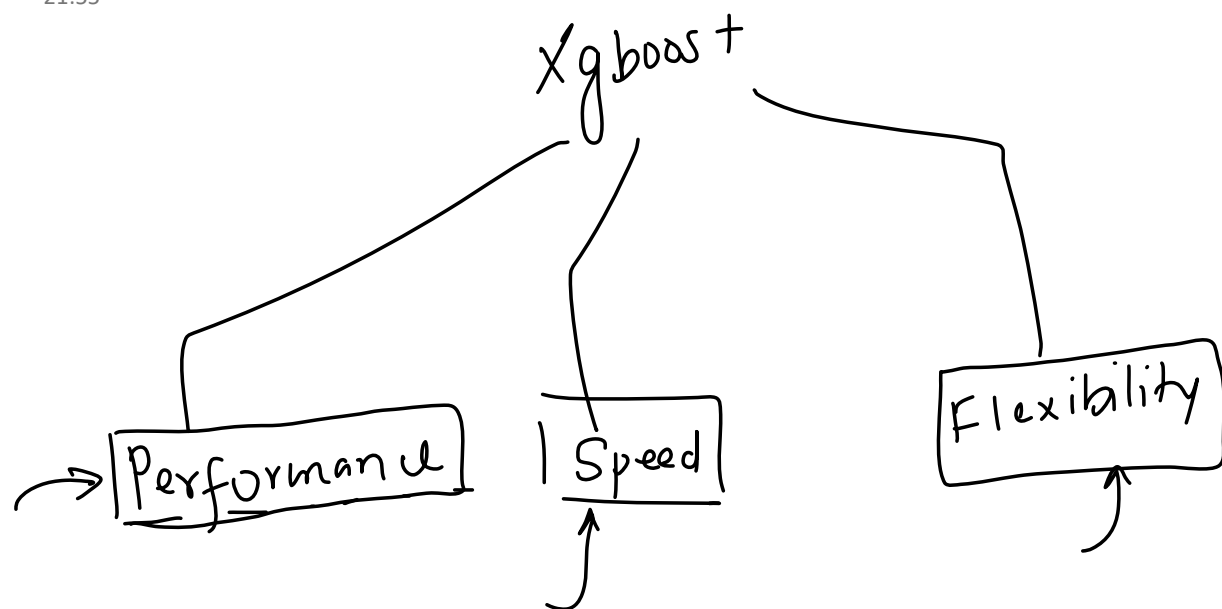
12 October 2023 09:38





XGBoost Features

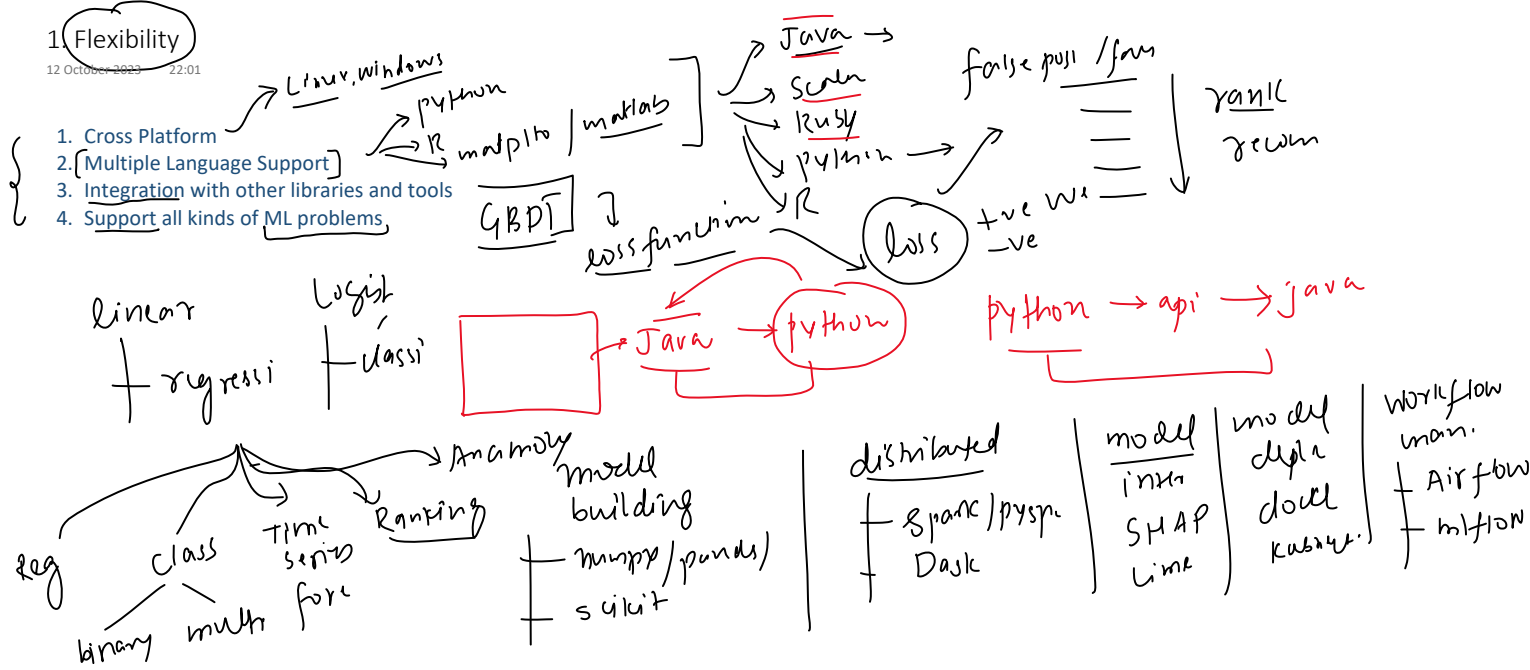
12 October 2023 21:55



1. Flexibility

12 October 2023 22:01

1. Cross Platform
2. Multiple Language Support
3. Integration with other libraries and tools
4. Support all kinds of ML problems



Python (Training and Saving the Model)

python

Copy code

```
import xgboost as xgb

# Assuming dtrain is your data in DMatrix format
model = xgb.train(params, dtrain, num_rounds)
model.save_model('xgboost_model.model')
```

python

Java (Loading and Using the Model)

java

Copy code

```
import ml.dmlc.xgboost4j.java.Booster;
import ml.dmlc.xgboost4j.java.DMatrix;
import ml.dmlc.xgboost4j.java.XGBoost;

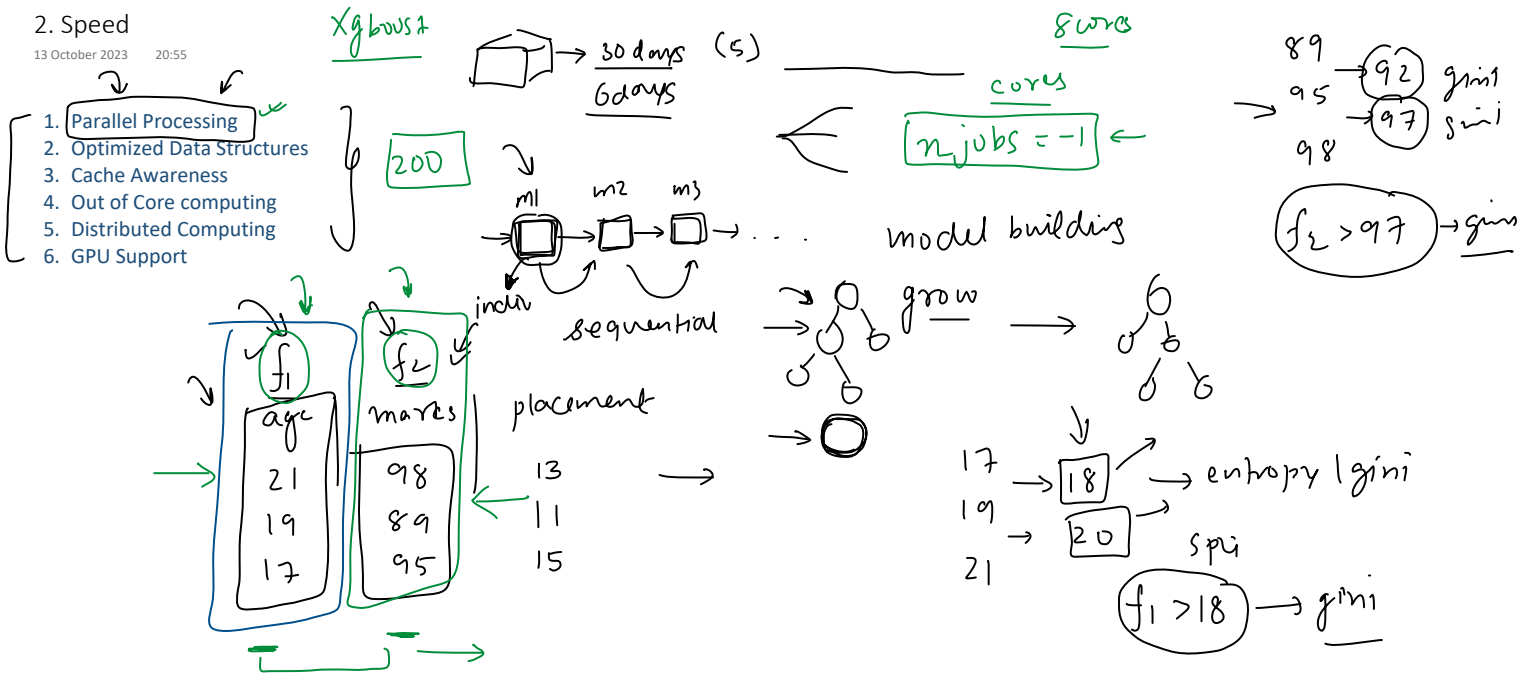
// Load model
Booster booster = XGBoost.loadModel("xgboost_model.model");

// Assuming data is your input data in DMatrix format
DMatrix data = new DMatrix("path_to_your_data");
float[][] predictions = booster.predict(data);
```

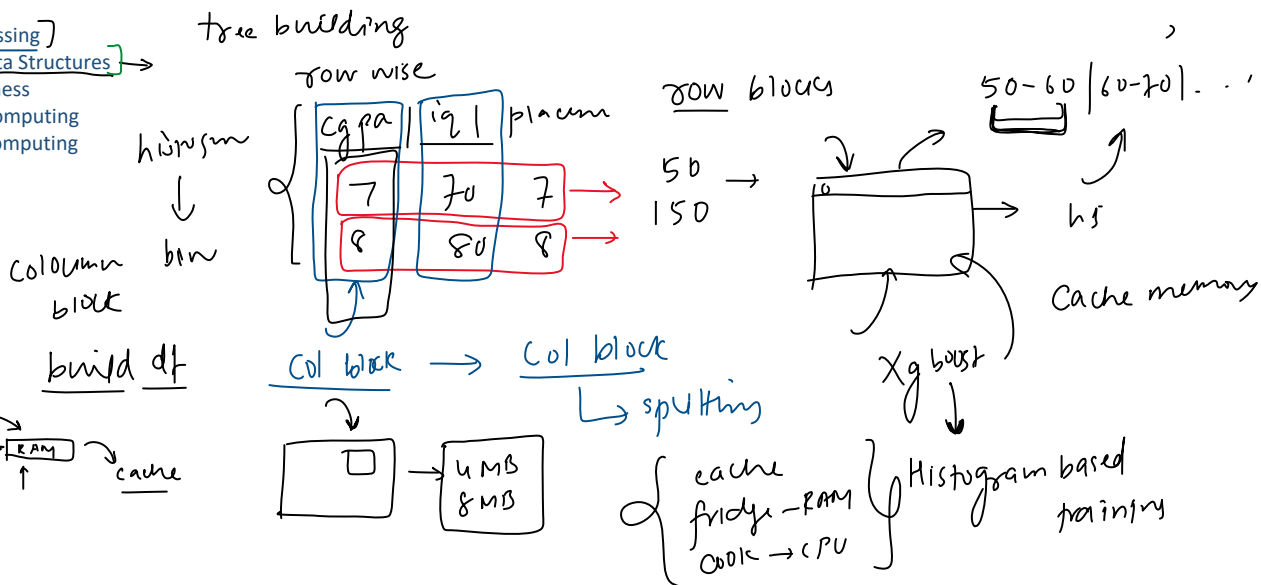
java

2. Speed

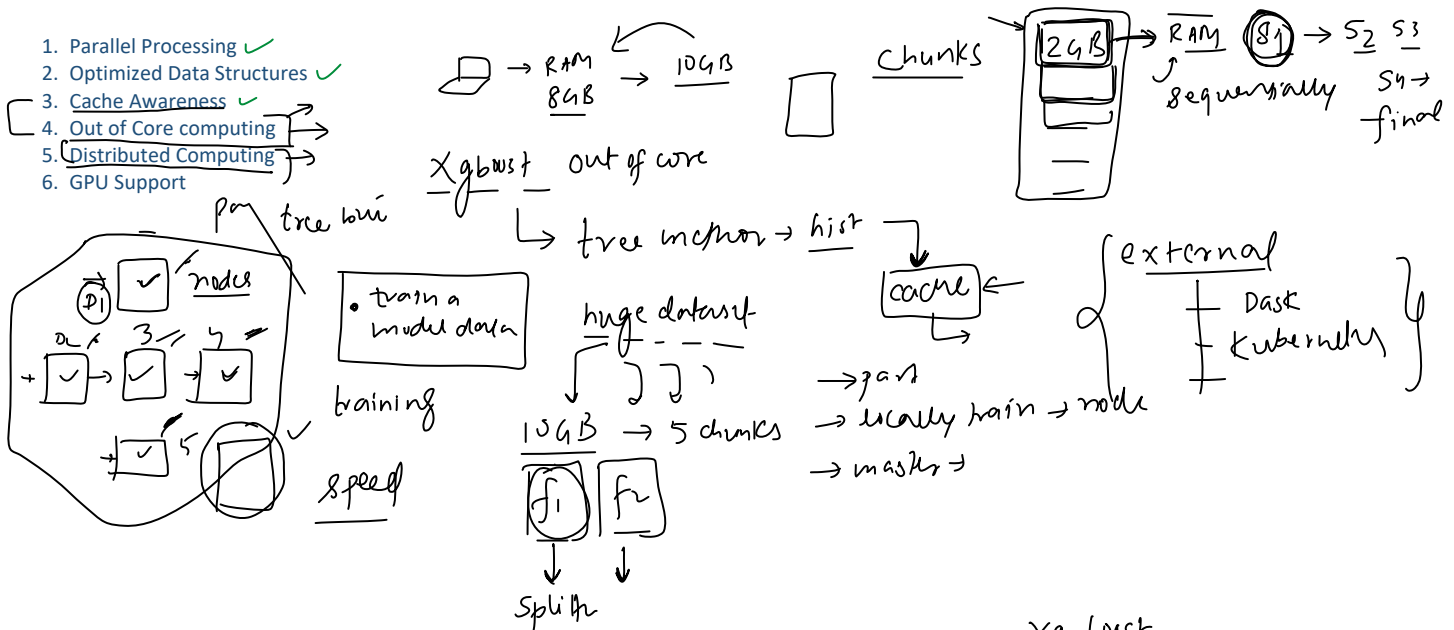
13 October 2023 20:55

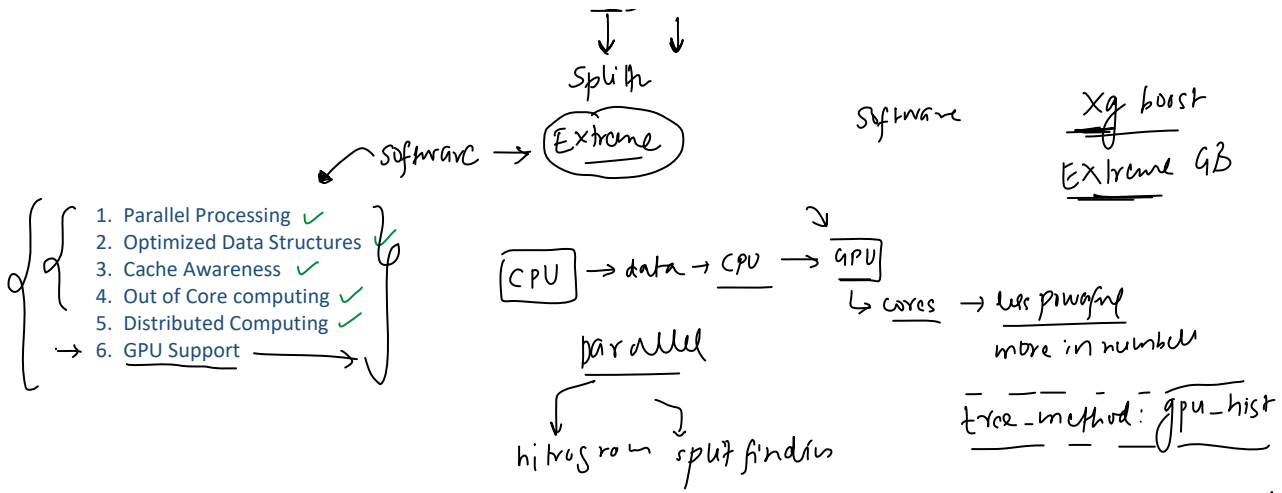


1. Parallel Processing
2. Optimized Data Structures
3. Cache Awareness
4. Out of Core computing
5. Distributed Computing
6. GPU Support



1. Parallel Processing
2. Optimized Data Structures
3. Cache Awareness
4. Out of Core computing
5. Distributed Computing
6. GPU Support





Software us
↓
ML → performance

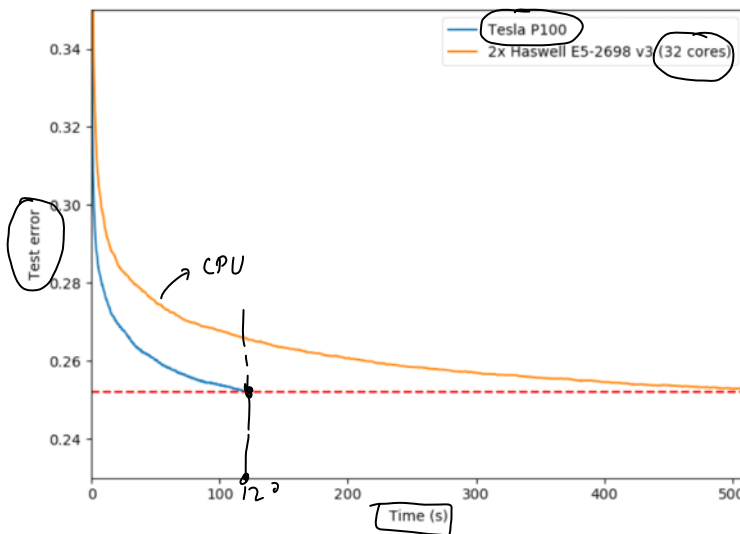


Figure 3. Test error over time for the Higgs dataset, 1000 boosting iterations.

Xgboost

Xgboost → Loss function → min

1. Regularized Learning Objective
2. Handling Missing values
3. Sparsity Aware Split Finding
4. Efficient Split Finding (Weighted Quantile Sketch + Approximate Tree Learning)
5. Tree Pruning

Linear reg
↓
Regular
(L1 & L2)

train → test
✓
overfitting

$L + \text{reg}$

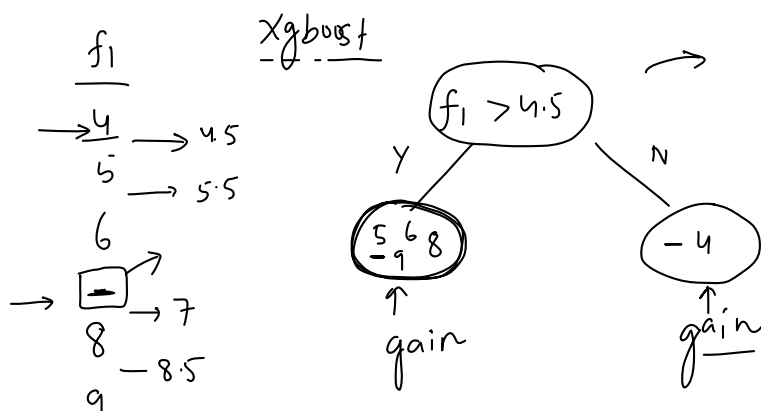
Gradient
(L) → reg

learning
rate

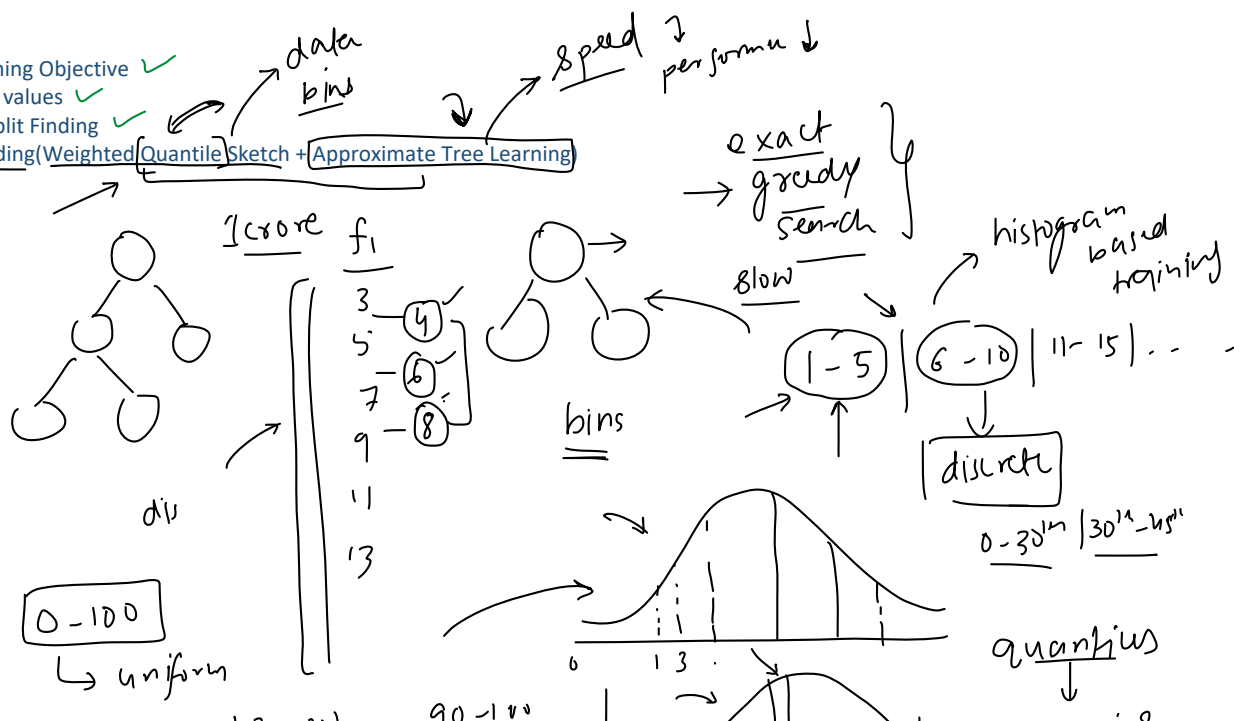
tree
construction

1. Regularized Learning Objective ✓
2. Handling Missing values
3. Sparsity Aware Split Finding →
4. Efficient Split Finding (Weighted Quantile Sketch + Approximate Tree Learning)
5. Tree Pruning

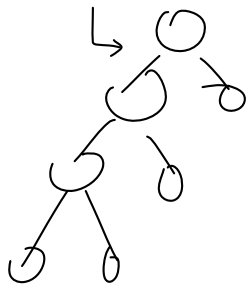
md → missing
↳ prepro
↳ impute



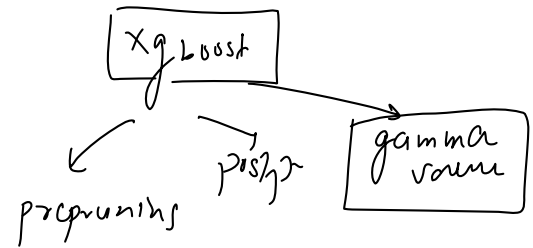
1. Regularized Learning Objective ✓
2. Handling Missing values ✓
3. Sparsity Aware Split Finding ✓
4. Efficient Split Finding (Weighted Quantile Sketch + Approximate Tree Learning)
5. Tree Pruning



- GB



→ reduce complexity
↓
overfitting



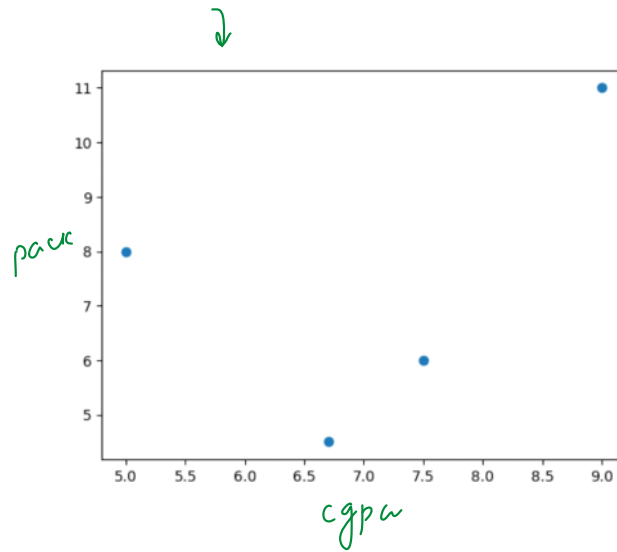
The Problem Statement

23 October 2023 15:05

lpa

cgpa	package
6.7	4.5
9.0	11.0
7.5	6.0
5.0	8.0

⑦



→ regression

new → cgpa → xgboost
↓
package

Solution Overview

23 October 2023 15:05

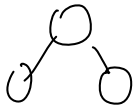
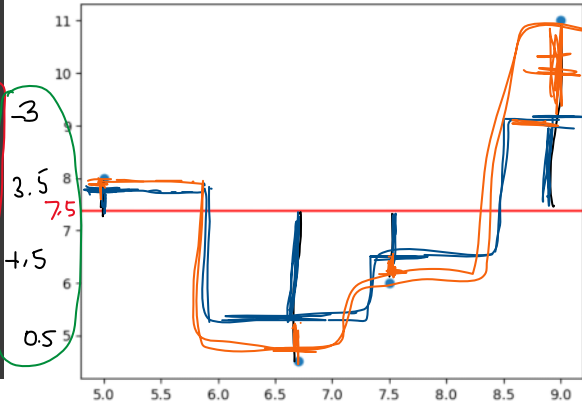
pseudo-residual

[perf / speed]

4.1, 9.1

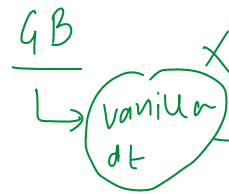
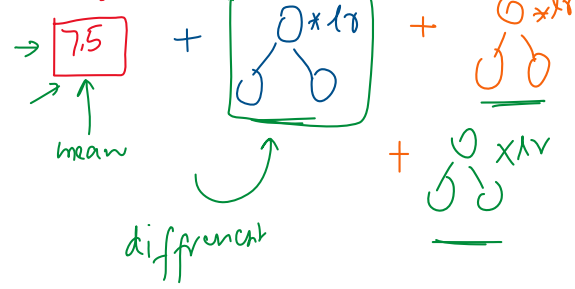
feature

cgpa	package
6.7	4.5
9.0	11.0
7.5	6.0
5.0	8.0



base estimator $\rightarrow$ mean

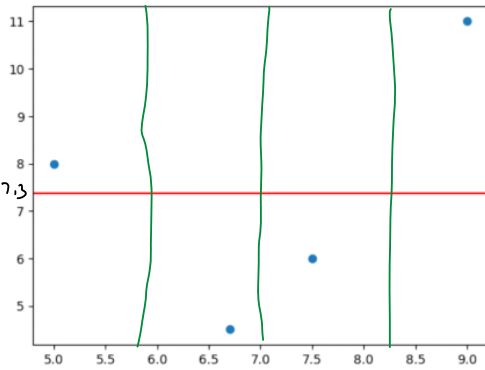
stage 1 [stage II]



gini / entropy $\rightarrow$ diff criteria

7.3

cgpa	package	model 1	residual 1
6.7	4.5	7.3	-2.8
9.0	11.0	7.3	3.7
7.5	6.0	7.3	-1.3
5.0	8.0	7.3	0.7



stage 1 → mean

cgpa | residual 1

decision tree

→ -2.8, 3.7, -1.3, 0.7 → Similarity score

$SS = \frac{(\text{sum of residuals})^2}{\# \text{residuals} + 1}$ → reg parameter

$\lambda = 0$

$\frac{(-2.8 + 3.7 - 1.3 + 0.7)^2}{4 + 0} = \frac{(0.3)^2}{4} = 0.02$

5.0 → 5.85
6.7 → 7.1
7.5 → 8.25
9.0 → 8.25

cgpa	residual 1
5.0	0.7
6.7	-2.8
7.5	-1.3
9.0	3.7

Splitting criteria ① → 5.85

→ cgpa < 5.85 ✓

0.7

$SS_L = \frac{(0.7)^2}{1 + 0} = 0.49$

-2.8, -1.3, 3.7

$SS_R = \frac{(-2.8 - 1.3 + 3.7)^2}{3} = \frac{0.16}{3} = 0.05$

gain = $(SS_L + SS_R) - SS_{\text{root}}$

= $0.49 + 0.05 - 0.02$

= 0.52 ✓

Splitting criteria ② → 7.1

→ cgpa < 7.1 → Similar

0.7, -2.8

$SS_L = \frac{(0.7 - 2.8)^2}{2} = \frac{4.41}{2} = 2.20$

-1.3, 3.7

$SS_R = \frac{(-1.3 + 3.7)^2}{2} = \frac{5.76}{2} = 2.88$

gain = $SS_L + SS_R - SS_{\text{root}}$

= $2.20 + 2.88 - 0.02$

= 5.06 ✓

Splitting criteria ③ → 8.25

cgpa < 8.25 ✓

0.7, -2.8, -1.3

$SS_L = \frac{(0.7 - 2.8 - 1.3)^2}{3} = \frac{11.56}{3} = 3.85$

3.7

$SS_R = \frac{(3.7)^2}{1} = 13.69$

gain = $3.85 + 13.69 - 0.02$

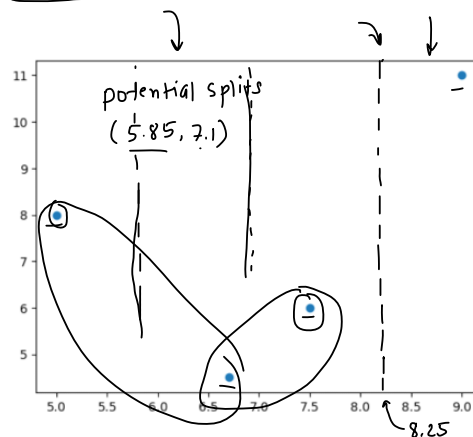
= 17.52

Stage 2 → dev

→ cgpa < 8.25

0.7, -2.8, -1.3

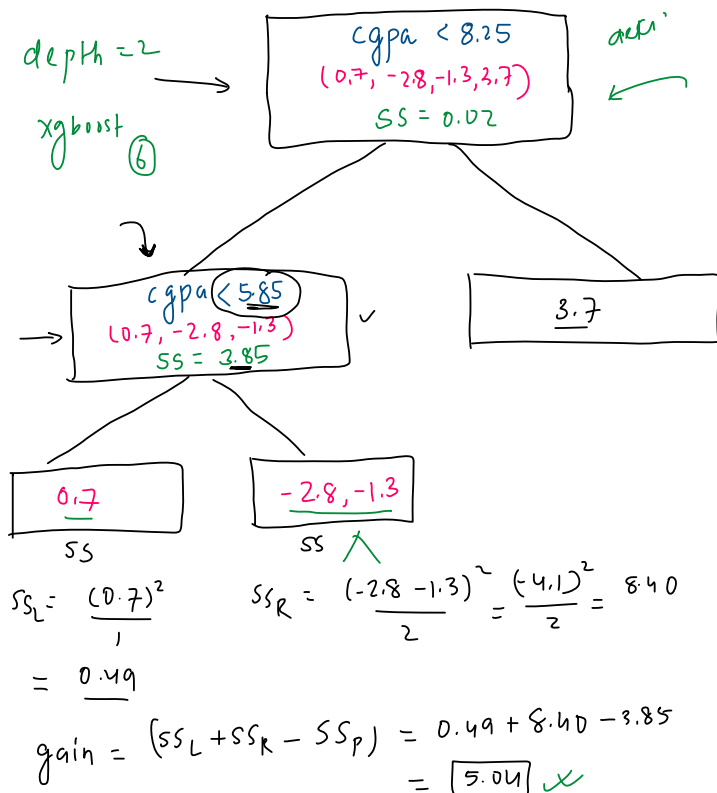
3.7



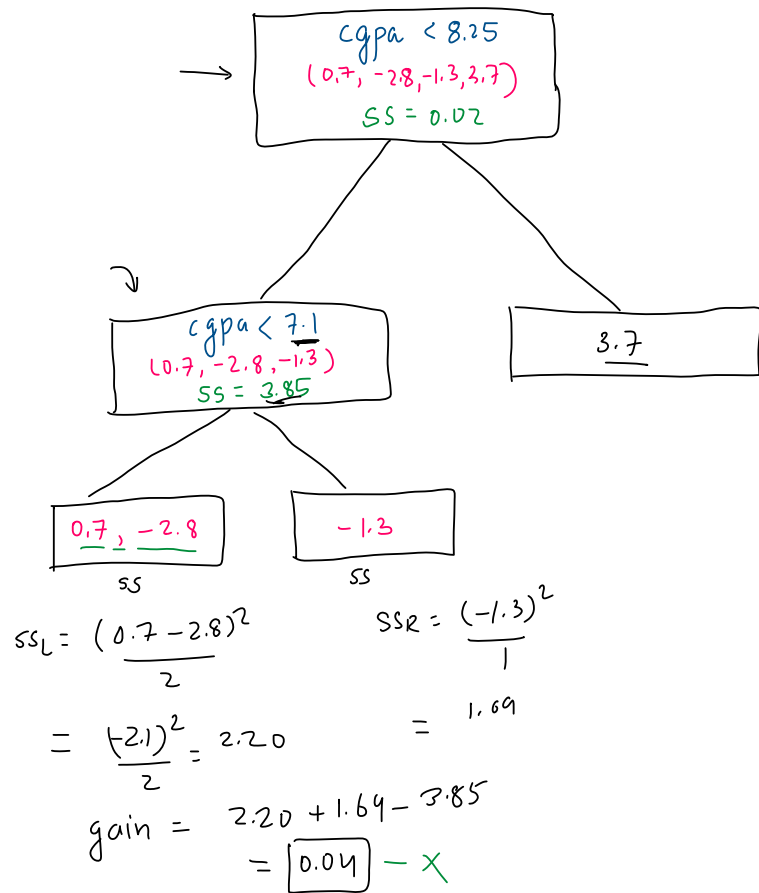
① Splitting Criteria $\rightarrow 5.85$

depth = 2

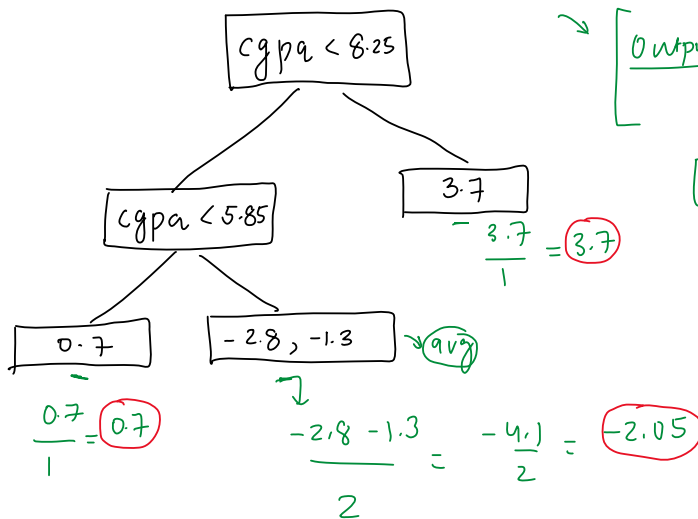
xgboost ⑥



② Splitting Criteria $\rightarrow 7.1$



The final decision tree



Output = $\frac{\text{sum of residuals}}{\# \text{residuals} + \lambda}$

$\lambda = 0 \rightarrow \lambda = 0$
 $\hookrightarrow \text{avg}$

cgpa	package	model 1	residual 1	model 2	residual 2
6.7	4.5	7.3	-2.8	6.69	-2.19
9.0	11.0	7.3	3.7	8.41	2.59
7.5	6.0	7.3	-1.3	6.69	-0.69
5.0	8.0	7.3	0.7	7.51	0.49

Stage 2: $\eta = 0.3$

Output: $7.3 + 0.3 \times (-2.05)$

Stage (3)
 $\boxed{cgpa / res 2}$

mean + $\lambda \times dt1 + \lambda \times dt2$

$7.3 + 0.3 \times (-2.05)$
 $7.3 + 0.3 \times 3.7$
 $7.3 + 0.3 \times (-2.05)$
 $7.3 + 0.3 \times 0.7$

$\boxed{100/200} \rightarrow \boxed{residual \rightarrow 0}$

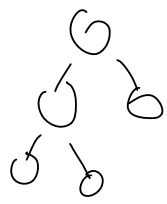
$cgpa / package \rightarrow \begin{bmatrix} f_1 \\ cgpa / 12^{th} \text{ mark} / package \end{bmatrix}$
 categorical $\rightarrow$ binary
 multipliers

$\left\{ \begin{array}{l} ss = \\ out \end{array} \right.$

$\lambda \rightarrow$

$\rightarrow$

xgboost



Prerequisite & Disclaimer

11 November 2023 09:29

{
→ Xgboost Reg
→ Gradient Boosting
→ log(odds)
} classifying

Disclaimers
→ optim
→ xgboost {
 perf
 speed
}
→ calculation → mistakes

Problem Statement

11 November 2023 09:29

cgpa	placed?
5.70	0
6.25	1
7.10	0
8.15	1
9.60	1

xgboost $\rightarrow$ classification

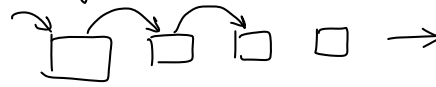
6.61 $\rightarrow$ 1,0

$\rightarrow$ overview
 $\rightarrow$ step by step

cgpa	placed?
5.70	0
6.25	1
7.10	0
8.15	1
9.60	1

xgboost $\rightarrow$ GBPT $\rightarrow$

stage wise add



base
model
(mean
 $\log(\text{odds})$)

m_1

+

$\log$

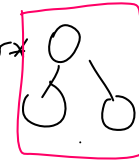


dt

m_L

+

$\log$



m_3

$\dots$

m_n

GBPT $\rightarrow$ vanilla $\left\{ \begin{array}{l} \text{entropy} \\ \text{gini} \end{array} \right\}$
 dt

xgboost $\rightarrow$ diff $\left\{ \begin{array}{l} \text{similarity} \\ \text{score} \end{array} \right\}$
 dt $\rightarrow$ xgboost $\log$

Stage 1 - $\square \rightarrow \log(\text{odds}) \rightarrow \text{prob}$

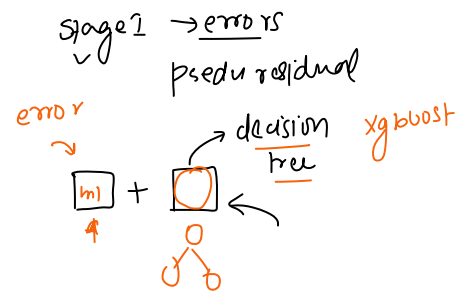
base estimator $\rightarrow \text{mean}$
 $\log(\text{odds}) \rightarrow \log\left(\frac{p}{1-p}\right)$ $p \rightarrow \text{prob +ve class (1)}$

$$\log\left(\frac{3/5}{2/5}\right) = \log\left(\frac{3}{2}\right) = 0.405$$

$$p = \frac{e^{\log \text{odds}}}{1 + e^{\log \text{odds}}} = \frac{e^{0.405}}{1 + e^{0.405}} = 0.60$$

cgpa	placed?	pred 1 (lo)
5.70	0	0.405 ✓
6.25	1	0.405
7.10	0	0.405
8.15	1	0.405
9.60	1	0.405

cgpa	placed?	pred 1 (lo)	pred 1 (prob)	res1
5.70	0	0.405 ✓	0.6	-0.6
6.25	1	0.405	0.6	0.4
7.10	0	0.405	0.6	-0.6
8.15	1	0.405	0.6	0.4
9.60	1	0.405	0.6	0.4



split - splitting criteria $\leftarrow$

cgpa	res1
5.70	-0.6 $\leftarrow$
6.25	0.4
7.10	-0.6
8.15	0.4
9.60	0.4

leaf node $\rightarrow$ $-0.6, 0.4, -0.6, 0.4, 0.4$ $\rightarrow$ similarity score

$SS = 0$ $\rightarrow$ maximize

GBDT classifier $\rightarrow$ $SS = \left(\sum \text{residual}_i\right)^2$ $\rightarrow$ $\square$

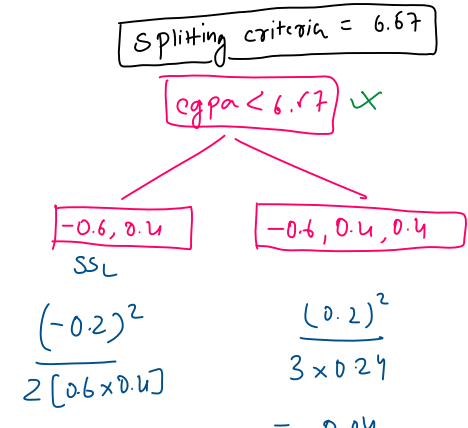
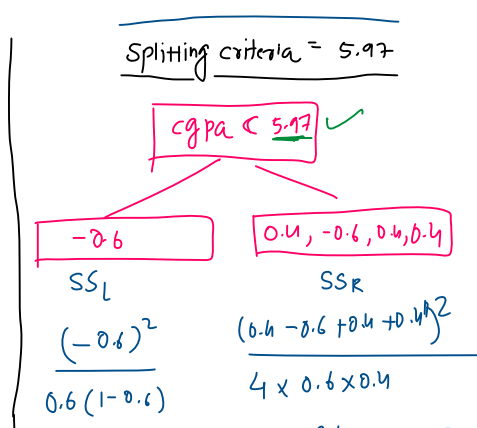
$\sum \text{prev_prob}_i (1 - \text{prev_prob}_i) + \lambda \rightarrow \text{reg } \lambda = 0$

$= (-0.6 + 0.4 - 0.6 + 0.4 + 0.4)^2 = 0$

$5 \times [0.6(1-0.6)] + 0$

gain $\rightarrow$ 1.17

cgpa	res1
5.70	-0.6 $\leftarrow$
6.25	0.4
7.10	-0.6
8.15	0.4
9.60	0.4



$$\begin{array}{c} \text{---} \\ 8.87 \end{array} \left\{ \begin{array}{l} 8.15 \\ 9.60 \end{array} \right\} \begin{array}{l} 0.4 \\ 0.4 \end{array}$$

$$\begin{array}{c} \text{---} \\ 0.6 \end{array} \left(\frac{-0.6}{1-0.6} \right) = \frac{0.36}{0.24} = 1.5$$

$$\begin{array}{c} \text{---} \\ 4 \times 0.6 \times 0.4 \end{array} = \frac{0.36}{4 \times 0.24} = \frac{1.5}{4} = 0.37$$

$$\text{gain} = (SS_L + SS_R) - SS_{\text{root}} = (1.5 + 0.37) - 0 = 1.87$$

$$\begin{array}{c} \text{---} \\ 2 \end{array} \left(\frac{0.6 \times 0.4}{0.6 \times 0.4} \right) = \frac{0.04}{2 \times 0.24} = \frac{0.04}{0.48} = 0.08$$

$$\begin{array}{c} \text{---} \\ 3 \times 0.24 \end{array} = \frac{0.04}{3 \times 0.24} = 0.05$$

$$\text{gain} = 0.08 + 0.05 - 0 = 0.13$$

splitting criteria = 7.62

splitting criteria = 8.87

cgpa	res1
5.70	-0.6
6.25	0.4
7.10	-0.6
8.15	0.4
9.60	0.4

$$\begin{array}{c} \text{cgpa} < 7.62 \\ \swarrow \quad \searrow \\ \boxed{-0.6, 0.4, -0.6} \quad \boxed{0.4, 0.4} \end{array}$$

$$\begin{array}{c} \frac{(-0.8)^2}{3 \times 0.24} = 0.88 \end{array}$$

$$\begin{array}{c} \frac{(0.8)^2}{2 \times 0.24} = 1.33 \end{array}$$

$$\text{gain} = 0.88 + 1.33 - 0 = 2.22$$

$$\begin{array}{c} \text{cgpa} < 8.87 \\ \swarrow \quad \searrow \\ \boxed{-0.6, 0.4, -0.6, 0.4} \quad \boxed{0.4} \end{array}$$

stage 2 model

$$\text{output} = \frac{\sum \text{residual}_i}{\sum \text{prev_prob}_i (1 - \text{prev_prob}_i) + \lambda} = 0$$

output

$$\begin{array}{c} \text{cgpa} < 7.62 \\ \swarrow \quad \searrow \\ \boxed{-0.6, 0.4, -0.6} \quad \boxed{0.4, 0.4} \end{array}$$

$$\downarrow (-1.11) \quad \downarrow (1.66)$$

$$\frac{-0.6 + 0.4 - 0.6}{3 \times 0.6 \times 0.4} = \frac{-0.8}{3 \times 0.24} = -1.11$$

$$\frac{0.4 + 0.4}{2 \times 0.24} = 1.66$$

stage 11 → prediction

cgpa	placed?	pred1(lo)	pred1(prob)	yes1	pred2(lo)	pred2(prob)	yes2
5.70	0	0.405	0.6	-0.6	0.072	0.518	-0.51
6.25	1	0.405	0.6	0.4	0.072	0.518	0.48
7.10	0	0.405	0.6	-0.6	0.072	0.518	-0.51
8.15	1	0.405	0.6	0.4	0.903	0.712	0.28
9.60	1	0.405	0.6	0.4	0.903	0.712	0.28

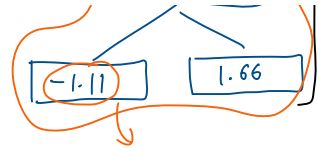
Stage 2 model

similar

$$\overline{m1} + 0.3 \times \overline{m2} \quad (\log \text{ odds}) \quad (\text{decision tree})$$

$$0.405 + 0.3 \times \begin{array}{c} \text{cgpa} < 7.62 \\ \swarrow \quad \searrow \\ \boxed{-1.11} \quad \boxed{1.66} \end{array}$$

→	8.15	1	0.405	0.6	{	0.4	0.903	0.712	0.28
→	9.60	1	0.405	0.6	}	0.4	0.903	0.712	0.28



$$0.405 + 0.3 \times (-1.11)$$

eta

$$0.405 + 0.3 \times (-1.11)$$

$$0.405 + 0.3 \times 1.66$$

log(odds)
→ prob

$$\frac{e^{\log odds}}{1 + e^{\log odds}}$$

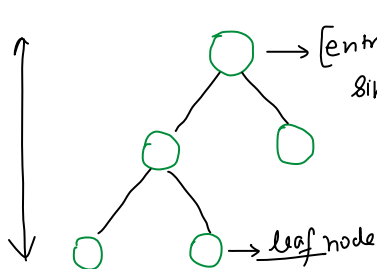
m3 → decision

same flow

Stage 3
 $[m1 + 0.3 \times m2 + 0.3 \times m3] \rightarrow \text{pred} \rightarrow \text{prob} \rightarrow \text{res3}$
 $\text{cgra/res3} \rightarrow m4 \rightarrow \text{decision}$

next video

→ maths → xgbost
 form



Regression

Classification

$$\left[\text{Similarity Score} = \frac{(\sum \text{Residuals})^2}{N + \lambda} \right] \textcircled{1}$$

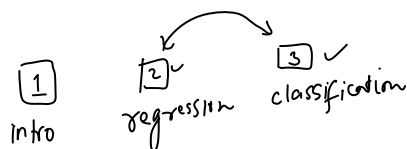
$$\left[\text{Similarity Score} = \frac{(\sum \text{Residuals})^2}{\sum [P (1 - P)] + \lambda} \right] \textcircled{2}$$

reg ✓

$$\left[\text{Output value} = \frac{\sum \text{Residuals}}{N + \lambda} \right] \textcircled{3}$$

✓

$$\left[\text{Output Value} = \frac{(\sum \text{Residuals})}{\sum [P (1 - P)] + \lambda} \right] \textcircled{4}$$

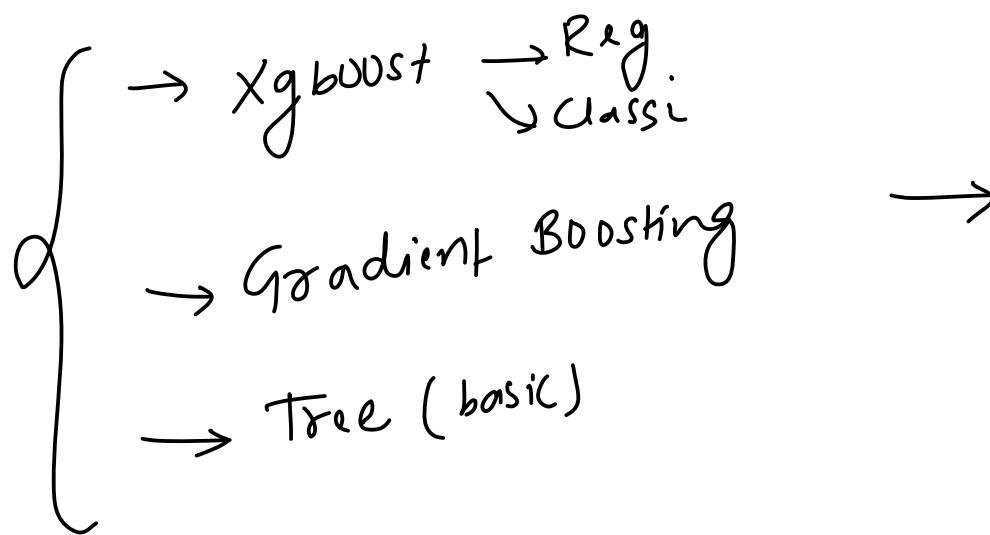


deep dive → goal

Prerequisite

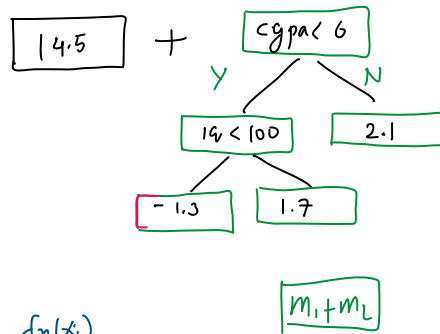
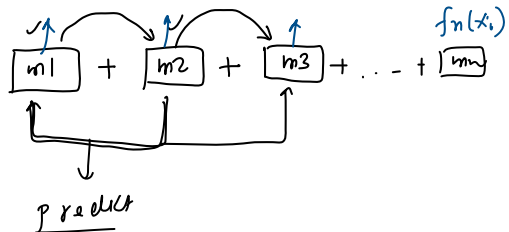
30 November 2023

10:44



Xgboost $\rightarrow$ 2 model

cgpa	ier	package
6.5	71	13
5	91	10
9.1	120	16



Xgboost

train

out leaf weight

df, 1023

$$\hat{y}_i = 14.5 + 2.1 = 16.6 \text{ lpa}$$

$$\hat{y}_i = f_1(x_i) + f_2(x_i) + \dots + f_n(x_i)$$

$$\hat{y}_i = \sum_{j=1}^n f_j(x_i) \rightarrow \text{Xgboost}$$

$$[\sigma_{reg}]$$

$$\text{output} =$$

$$\frac{\text{sum of residuals}}{N + \lambda}$$

origin story

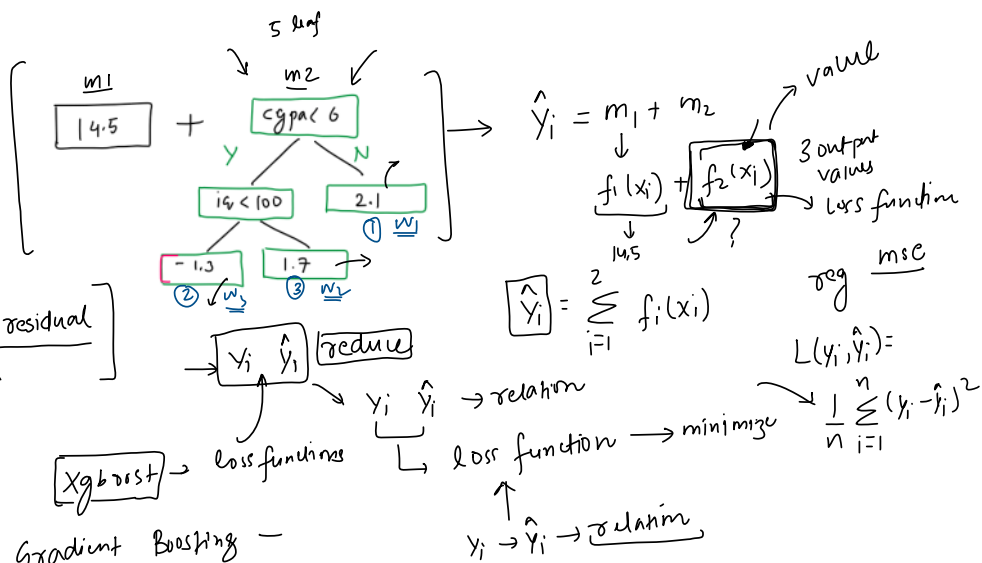
XGBoost Loss Function

30 November 2023 13:19

cgpa	iq	package
6.5	71	13
5	91	10
9.1	120	16

$$\left[\text{Output} = \frac{\sum \text{sum of residual}}{N + 1} \right]$$

number of residuals



Gradient Boosting -

flexible $\rightarrow$ loss function $\rightarrow$ differentiable

differentiable $\rightarrow$ loss-function $\rightarrow$ regularization parameter

$$[L(y_i, \hat{y}_i)]$$

mse log loss

$$m_1 + m_2$$

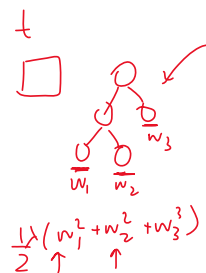
base dt

$$L = \sum_{i=1}^n L(y_i, \hat{y}_i) + \Omega(f_2(x_i))$$

minimize

$$\Omega(f) = \gamma T + \frac{1}{2} \lambda \|w\|^2$$

regularizer # leaf nodes regularizer

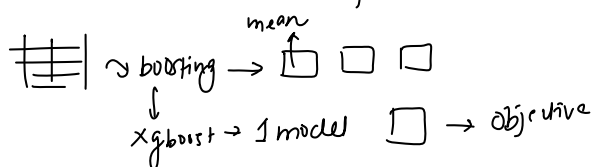


objective func

$$\mathcal{L} = \underbrace{\sum_{i=1}^n L(y_i, \hat{y}_i)}_{\text{loss function}} + \underbrace{\Omega(f_k)}_{\text{regular}}$$

$$\gamma \frac{T}{2} + \frac{1}{2} \lambda \|w\|^2$$

$\downarrow$ leaf $\downarrow$ leaf



Stage I

mean $\downarrow$

[m1] $\rightarrow$ $\mathcal{L}^{(1)} = \sum_{i=1}^n L(y_i, \hat{y}_i)$

$f_1(x_i)$

Stage II

mean $\downarrow$

[m1] $\rightarrow$ $\mathcal{L}^{(2)} = \sum_{i=1}^n L(y_i, f_1(x_i) + f_2(x_i)) + \Omega(f_2(x_i))$

$f_1(x_i)$ $f_2(x_i)$

decision $\rightarrow$

$\hat{y}_i = f_1(x_i) + f_2(x_i)$

Stage III

mean $\downarrow$

[m1] [m2] $\rightarrow$ $\mathcal{L}^{(3)} = \sum_{i=1}^n L(y_i, f_1(x_i) + f_2(x_i) + f_3(x_i)) + \Omega(f_3(x_i))$

$f_1(x_i)$ $f_2(x_i)$ $f_3(x_i)$

decision $\rightarrow$

$f_1(x_i) + f_2(x_i) + \dots + f_{t-1}(x_i) \rightarrow \hat{y}_i^{(t-1)}$

Stage t

mean $\downarrow$

[m1] [m2] ... [mt] $\rightarrow$ $\mathcal{L}^{(t)} = \sum_{i=1}^n L(y_i, f_1(x_i) + f_2(x_i) + \dots + f_t(x_i)) + \Omega(f_t(x_i))$

$f_1(x_i)$ $f_2(x_i)$... $f_t(x_i)$

last dec $\rightarrow$

$\mathcal{L}^{(t)} = \sum_{i=1}^n L(y_i, \hat{y}_i^{(t-1)} + f_t(x_i)) + \Omega(f_t)$

At any stage t

$$\mathcal{L}^{(t)} = \sum_{i=1}^n L(y_i, \hat{y}_i^{(t-1)} + f_t(x_i)) + \Omega(f_t(x_i))$$

Linear reg

$$\rightarrow L = \sum_{i=1}^n (y_i - \underset{\uparrow}{m} x_i - \underset{\uparrow}{b})^2$$

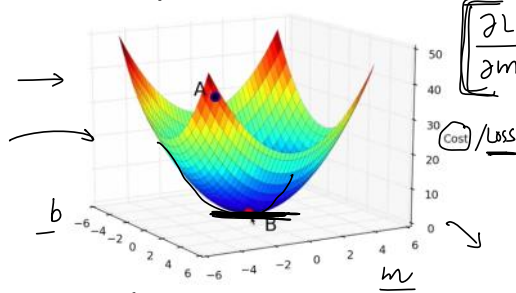
$$\underline{\mathcal{L}} = 0$$

m, b

91

grad desc

$$\left[\frac{\partial L}{\partial m} \right] \left[\frac{\partial L}{\partial b} \right] = 0$$



some
often $\rightarrow$ smooth
 $\rightarrow$ differentiable

$$\frac{\partial L}{\partial f_H(x_i)}$$

Taylor series



→ pericardial
constant

powerful approx techni

complex function $\rightarrow$ approx $\rightarrow$ polynomial e^x
accurate
approx

$$\longrightarrow \frac{3x^2 + 2}{1}$$

$$\uparrow$$

 $3x^2 + 3x^3 + x^4 + x^5 \dots$

Objective $\rightarrow$ usable
 $\downarrow$
Taylor series

How

 $e^x \rightarrow$ polynomial
 $\rightarrow$ derivatives

$$\boxed{e^x} \rightarrow e^x$$

 $f(x)$ point $\rightarrow \boxed{a}$

$$f(x) \simeq \underbrace{f(a)}_{1!} + \underbrace{f'(a)(x-a)}_{2!} + \underbrace{f''(a)(x-a)^2}_{3!} + \underbrace{f'''(a)(x-a)^3}_{4!} + \dots$$

$$f(x) = e^x \quad \boxed{a=0}$$

$$f(a) = e^a = 1$$

$$\frac{f'(a)}{1!} (x-a) = \frac{e^a}{1} (x-a) = e^0 (x-0) = x$$

$$\frac{f''(a)}{2!} (x-a)^2 = \frac{e^a}{2!} (x-a)^2 = \frac{e^0 (x)^2}{2!} = \frac{x^2}{2}$$

$$\frac{f'''(a)}{3!} (x-a)^3 = \frac{e^a}{3!} (x-a)^3 = \frac{1}{6} x^3 = \frac{x^3}{6}$$

$$e^x \simeq \boxed{1+x} + \frac{x^2}{2} + \frac{x^3}{6} + \dots$$

Applying Taylor Series

01 December 2023 00:51

$$f(x) \simeq \underbrace{f(a)} + \underbrace{f'(a)}(x-a) + \underbrace{\frac{f''(a)}{2!}}(x-a)^2 \dots \dots \dots]$$

$$\frac{\partial f}{\partial x} =$$

$$f(x) = x^2$$

$$f'(x) \frac{\partial f}{\partial x} = 2x$$

$$f'(a) \cdot 2a$$

$$\rightarrow \mathcal{L}^{<t>} = \left[\sum_{i=1}^n L(y_i, \hat{y}_i^{<t-1>} + f_t(x_i)) \right] + \underbrace{\Omega(f_t(x_i))}_{\text{2nd degree}}$$

$x \rightarrow \hat{y}_i^{<t-1>} + f_t(x_i) \quad a \rightarrow \hat{y}_i^{<t-1>}$

$$\mathcal{L}^{<t>} = \sum_{i=1}^n L(y_i, \hat{y}_i^{<t-1>}) + g_i f_t(x_i) + \frac{1}{2} h_i f_t^2(x_i) + \Omega(f_t(x_i))$$

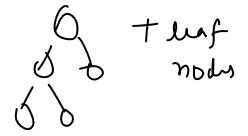
$$f^{(0)} = \sum_{i=1}^n L(y_i, \hat{y}_i^{<t-1>}) \rightarrow \text{gradient } g_i$$

$$\frac{f'(a)}{1} (x-a) = \sum_{i=1}^n \left[\frac{\partial L(y_i, \hat{y}_i^{<t-1>})}{\partial \hat{y}_i^{<t-1>}} \right] f_t(x_i) \rightarrow \text{hessian } h_i$$

$$\frac{f''(a)}{2} (x-a)^2 = \frac{1}{2} \sum_{i=1}^n \left[\frac{\partial^2 L(y_i, \hat{y}_i^{<t-1>})}{\partial \hat{y}_i^{<t-1>2}} \right] f_t^2(x_i)$$

$$\mathcal{L}^{(t)} \simeq \sum_{i=1}^n \left[\underbrace{\mathcal{L}(y_i, \hat{y}_i^{(t-1)})}_{\downarrow} + \underbrace{g_i f_t(x_i)}_{\uparrow} + \underbrace{\frac{1}{2} h_i f_t^2(x_i)}_{\checkmark} \right] + \underbrace{\mathcal{L}(f_t(x_i))}_{\checkmark}$$

$$\mathcal{L}^{(t)} = \sum_{i=1}^n \left[g_i f_t(x_i) + \frac{1}{2} h_i f_t^2(x_i) \right] + \boxed{\mathcal{L}(f_t(x_i))} \rightarrow \text{expand}$$



$$\mathcal{L}^{(t)} = \sum_{i=1}^n \left[g_i f_t(x_i) + \frac{1}{2} h_i f_t^2(x_i) \right] + \underbrace{\gamma T}_{\gamma T} + \underbrace{\frac{1}{2} \lambda \sum_{j=1}^T w_j^2}_{\substack{j=1 \rightarrow T \\ T \rightarrow \# \text{leaf nodes}}}$$



$$\mathcal{L}^{(t)} = \sum_{j=1}^T \left[\sum_{i \in I_j} g_i w_j + \frac{1}{2} \sum_{i \in I_j} h_i w_j^2 \right] + \gamma T + \frac{1}{2} \lambda \sum_{j=1}^T w_j^2$$

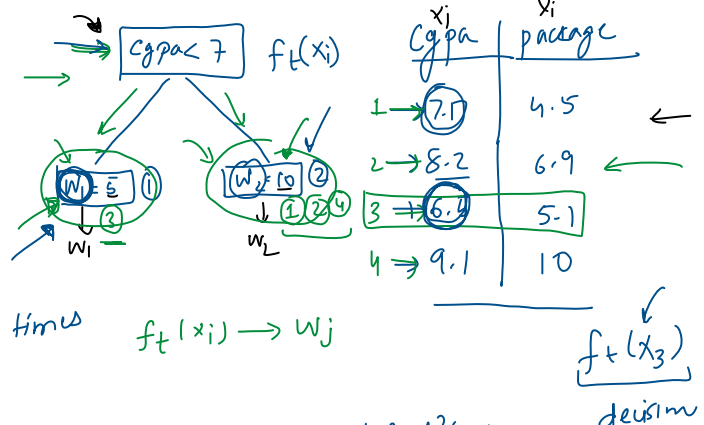
$$\left[\begin{array}{l} w_j \rightarrow f_t(x_i) \\ \sum_{i=1}^n \rightarrow \sum_{j=1}^T \sum_{i \in I_j} \end{array} \right]$$

$$\mathcal{L}^{(t)} = \sum_{i=1}^n \left[g_i f_t(x_i) + \frac{1}{2} h_i f_t^2(x_i) \right] + \gamma T + \frac{1}{2} \lambda \sum_{j=1}^T w_j^2$$

$$\mathcal{L}^{(t)} = \sum_{j=1}^T \left[\sum_{i \in I_j} g_i w_j + \frac{1}{2} \sum_{i \in I_j} h_i w_j^2 \right] + \gamma T + \frac{1}{2} \lambda \sum_{j=1}^T w_j^2$$

$i \in I_j \rightarrow$ instance set of leaf j

4 terms



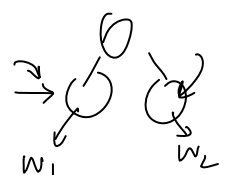
$$\rightarrow \left(g_1 f_t(x_1) + \frac{1}{2} h_1 f_t^2(x_1) \right) + \left(g_2 f_t(x_2) + \frac{1}{2} h_2 f_t^2(x_2) \right) + \dots + g_n f_t(x_n) + \frac{1}{2} h_n f_t^2(x_n)$$

$$\rightarrow \underbrace{g_3 f_t(x_3) + \frac{1}{2} h_3 f_t^2(x_3)}_{g_3 w_1 + \frac{1}{2} h_3 w_1^2} + \underbrace{g_1 f_t(x_1) + \frac{1}{2} h_1 f_t^2(x_1)}_{g_1 w_2 + \frac{1}{2} h_1 w_2^2} + g_2 f_t(x_2) + \frac{1}{2} h_2 f_t^2(x_2) + g_4 f_t(x_4) + \frac{1}{2} h_4 f_t^2(x_4)$$

$$\mathcal{L}^{(t)} = \sum_{j=1}^T \left[\sum_{i \in I_j} g_i w_j + \frac{1}{2} \sum_{i \in I_j} h_i w_j^2 \right] + \gamma T + \frac{1}{2} \lambda \sum_{j=1}^T w_j^2$$

$$\mathcal{L}^{(t)} = \sum_{j=1}^T \left[\sum_{i \in I_j} g_i w_j + \frac{1}{2} \sum_{i \in I_j} h_i w_j^2 + \frac{1}{2} \lambda w_j^2 \right] + \gamma T$$

$$\mathcal{L}^{(t)} = \sum_{j=1}^T \left[\sum_{i \in I_j} g_i w_j + \frac{1}{2} \left(\sum_{i \in I_j} h_i + \lambda \right) w_j^2 \right] + \gamma T$$

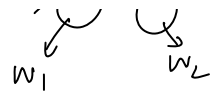


$$L^{(t)} = \sum_{j=1}^J \left[\underbrace{\sum_{i \in I_j} g_i w_j}_{\uparrow} + \frac{1}{2} \underbrace{\left(\sum_{i \in I_j} h_i + \lambda \right) w_j^2}_{\uparrow} \right] + \gamma T$$

$$\frac{\partial L^{(t)}}{\partial w_j} = \sum_{i \in I_j} g_i + \left(\sum_{i \in I_j} h_i + \lambda \right) w_j = 0$$

$$w_j = - \frac{\sum_{i \in I_j} g_i}{\sum_{i \in I_j} h_i + \lambda}$$

$g_i / h_i \rightarrow$ formula
the
output



$$w_j = \frac{- \sum_{i \in \mathcal{I}_j} g_i}{\sum_{i \in \mathcal{I}_j} h_i + \lambda}$$

gradient

hessian

$$g = \frac{\partial L(y_i, \hat{y}_i^{(t-1)})}{\partial \hat{y}_i^{(t-1)}}$$

$$h = \frac{\partial^2 L(y_i, \hat{y}_i^{(t-1)})}{\partial \hat{y}_i^{(t-1)2}}$$

Loss

reg

classi

Reg $\rightarrow$ mse $\rightarrow$

$$L = \frac{1}{2} \sum_{i=1}^n (y_i - \hat{y}_i)^2$$

residual for every point

$$g_i = \frac{\partial L}{\partial \hat{y}_i} = - \sum_{i=1}^n (y_i - \hat{y}_i) = - \sum_{i=1}^n (\hat{y}_i - y_i)$$

prev step

sum of residual

of residuals + 1

$$= \frac{\sum_{i \in \mathcal{I}_j} (y_i - \hat{y}_i^{(t-1)})}{\sum_{i \in \mathcal{I}_j} 1 + \lambda}$$

cgpa	ia	pa	pr	residual
8	80	8	8.5	8-8.5
9	90	9	8.5	9-8.5

$$\frac{\partial}{\partial \hat{y}_i} \frac{\partial L}{\partial \hat{y}_i} = \frac{\partial (y_i - \hat{y}_i)}{\partial \hat{y}_i} = -1$$

$$w_j = \frac{\text{sum of residual}}{\# \text{ residuals} + 1}$$

Output Value for Classification

01 December 2023 13:45

$$w_j = \frac{- \sum_{i \in I_j} g_i}{\sum_{i \in I_j} h_i + \lambda}$$

$$\frac{g}{f} = \frac{\partial L(y_i, \hat{y}_i^{(t-1)})}{\partial \hat{y}_i^{(t-1)}}$$

$$\underline{h} = \frac{\partial^2 L(y_i, \hat{y}_i^{(t-1)})}{\partial \hat{y}_i^{(t-1)2}}$$

loss - log loss $\rightarrow$ gradient boosting

$$\left[w_j = \frac{\text{Sum of residual}}{\sum_{i \in I_j} p_i (1 - p_i) + \lambda} \right]$$

$p_i \rightarrow$ pred prob of
prev timestep

Derivation of Similarity score

01 December 2023 13:45

$$\boxed{L^{(t)}} = \sum_{j=1}^T \left[\left(\sum_{i \in I_j} g_i \right) \underline{w_j} + \frac{1}{2} \left(\sum_{i \in I_j} h_i + \lambda \right) \underline{w_j^2} \right] + \gamma T$$

$$\boxed{w_j} = \frac{-\sum_{i \in I_j} g_i}{\sum_{i \in I_j} h_i + \lambda}$$

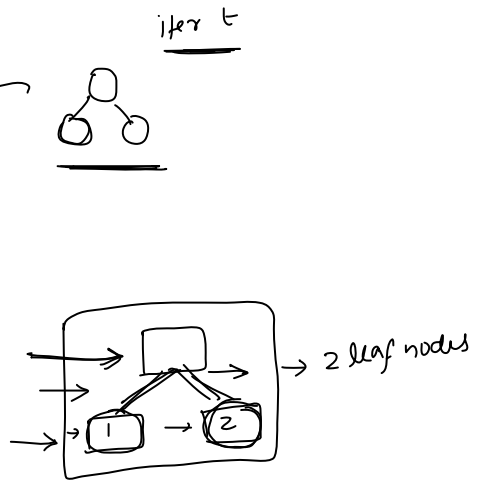
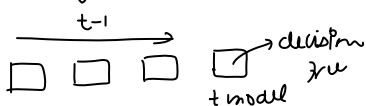
$$L^{(t)}(q) = \sum_{j=1}^T \left[\left(\sum_{i \in I_j} g_i \right) \frac{\left(-\sum_{i \in I_j} g_i \right)}{\sum_{i \in I_j} h_i + \lambda} + \frac{1}{2} \left(\sum_{i \in I_j} h_i + \lambda \right) \frac{\left(\sum_{i \in I_j} g_i \right)^2}{\left(\sum_{i \in I_j} h_i + \lambda \right)^2} \right] + \gamma T$$

$$L^{(t)}(q) = \sum_{j=1}^T \left[-\frac{\left(\sum_{i \in I_j} g_i \right)^2}{\sum_{i \in I_j} h_i + \lambda} + \frac{1}{2} \frac{\left(\sum_{i \in I_j} g_i \right)^2}{\sum_{i \in I_j} h_i + \lambda} \right] + \gamma T$$

$$\text{similarity score} = \sum_{j=1}^T \left[-\frac{1}{2} \frac{\left(\sum_{i \in I_j} g_i \right)^2}{\sum_{i \in I_j} h_i + \lambda} \right] + \gamma T$$

$$\boxed{L^{(t)}(q)} = -\frac{1}{2} \sum_{j=1}^T \frac{\left(\sum_{i \in I_j} g_i \right)^2}{\sum_{i \in I_j} h_i + \lambda} + \gamma T$$

gini/entropy



$$L^{(t)}(q_f) = -\frac{1}{2} \frac{\left(\sum_{i \in I_p} g_i \right)^2}{\sum_{i \in I_p} h_i + \lambda} + \gamma$$

$$\boxed{L^{(t)}(q_c)} = -\frac{1}{2} \left[\frac{\left(\sum_{i \in I_L} g_i \right)^2}{\sum_{i \in I_L} h_i + \lambda} + \frac{\left(\sum_{i \in I_R} g_i \right)^2}{\sum_{i \in I_R} h_i + \lambda} \right] + 2\gamma$$

$$L^{(t)}(q_f) \rightarrow L^{(t)}(q_c)$$

$$- \frac{1}{2} \left[\frac{\left(\sum_{i \in I_L} g_i \right)^2}{\sum_{i \in I_L} h_i + \lambda} + \frac{\left(\sum_{i \in I_R} g_i \right)^2}{\sum_{i \in I_R} h_i + \lambda} \right] + 2\gamma$$

$$-\frac{1}{2} \frac{(\sum_{i \in I_P} g_i)^2}{\sum_{i \in I_P} h_i + \lambda} + \gamma - \left(-\frac{1}{2} \left[\frac{(\sum_{i \in I_L} g_i)^2}{\sum_{i \in I_L} h_i + \lambda} + \frac{(\sum_{i \in I_R} g_i)^2}{\sum_{i \in I_R} h_i + \lambda} \right] \pm 2\gamma \right)$$

$$\left[-\frac{1}{2} \frac{(\sum_{i \in I_P} g_i)^2}{\sum_{i \in I_P} h_i + \lambda} + \frac{1}{2} \frac{(\sum_{i \in I_L} g_i)^2}{\sum_{i \in I_L} h_i + \lambda} + \frac{1}{2} \frac{(\sum_{i \in I_R} g_i)^2}{\sum_{i \in I_R} h_i + \lambda} \right] - \gamma$$

$$\frac{1}{2} \left[\frac{(\sum_{i \in I_L} g_i)^2}{\sum_{i \in I_L} h_i + \lambda} + \frac{(\sum_{i \in I_R} g_i)^2}{\sum_{i \in I_R} h_i + \lambda} - \frac{(\sum_{i \in I_P} g_i)^2}{\sum_{i \in I_P} h_i + \lambda} \right] - \gamma \quad \text{differ}$$

Final Calculation of Similarity

02 December 2023 01:27

$$\text{similarity score} = \sum_{j=1}^T \left[-\frac{1}{2} \frac{\left(\sum_{i \in I_j} g_i \right)^2}{\sum_{i \in I_j} h_i + \lambda} \right] + \gamma T$$

reg

$g \rightarrow$ sum of residuals

$h \rightarrow \sum 1 \rightarrow N$

Class

$g \rightarrow$ residual

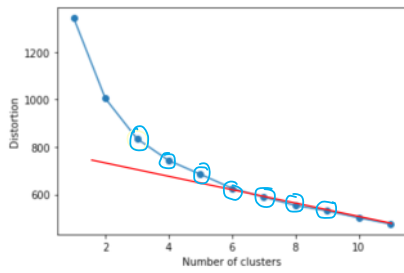
$h \rightarrow p_i(1-p_i)$

Why DBSCAN?

16 December 2023 15:10

Kmeans!

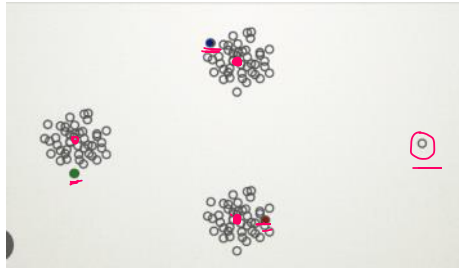
elbow



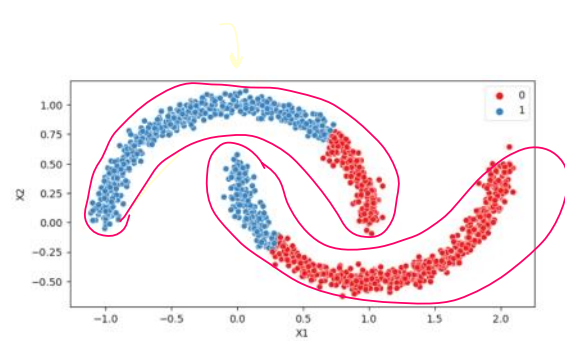
1) Specify clusters

↓

high $\rightarrow$ 10 dim $\rightarrow$ plot



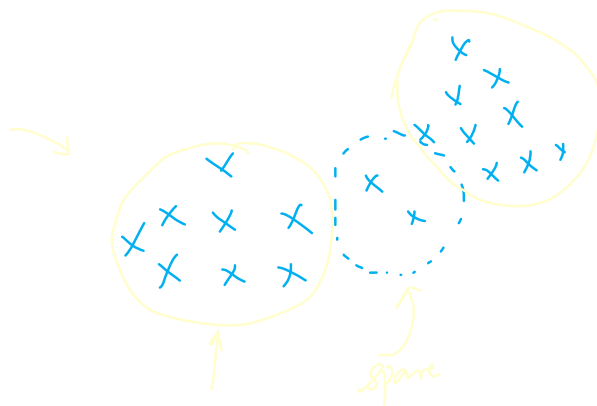
2) Sensitive to outliers



3)

What is Density Based Clustering

16 December 2023 15:12



Kmean
↓
→ centroid
clustering

DBSCAN → Density Based Spatial
Clustering of Applications
with Noise

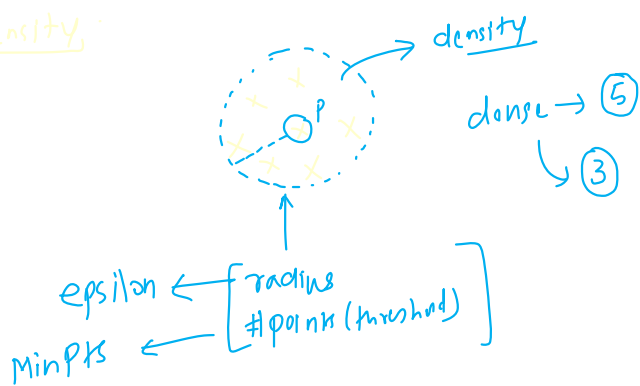
OPTICS

MinPts & Epsilon

16 December 2023 15:13

Density

density
hyperparameter



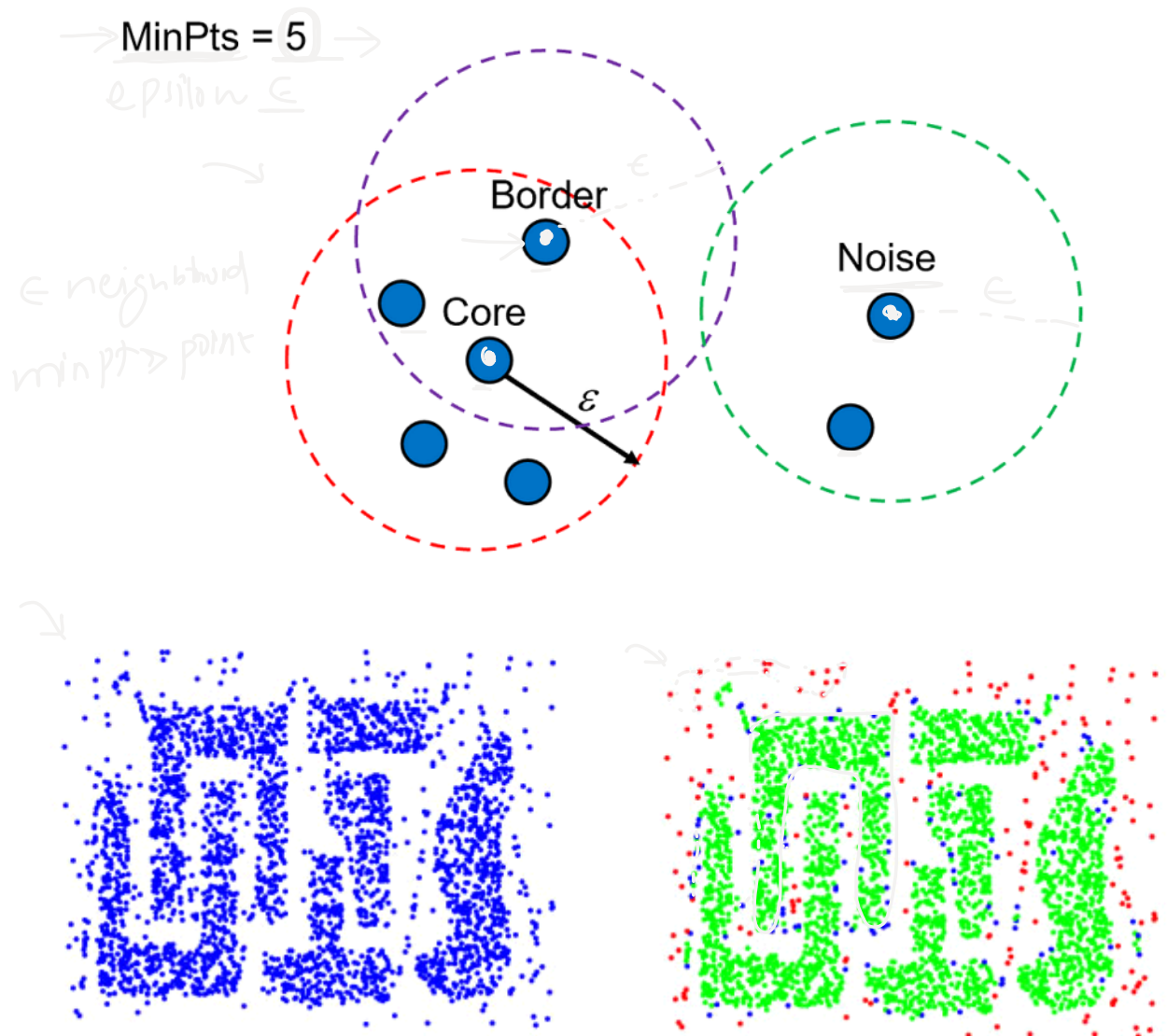
epsilon
 $\hookrightarrow$ radius of
the neighbor
minPts
 $\hookrightarrow$ threshold

epsilon $\rightarrow$ unit
minPts $\rightarrow$ 3



Core Points, Border Points & Noise Points

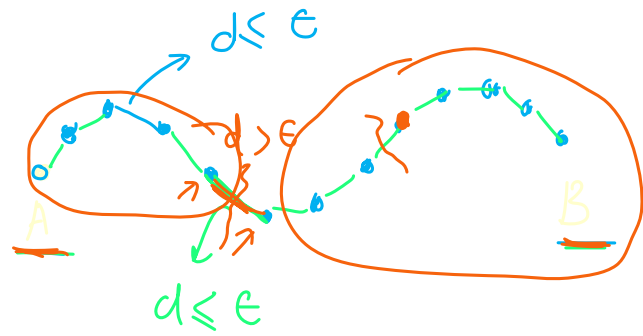
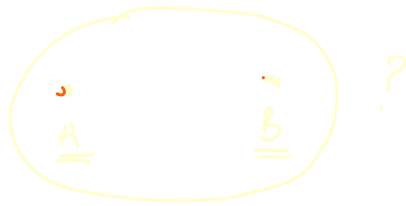
16 December 2023 15:13



Eps = 10, MinPts = 4

Density Connected Points

16 December 2023 15:13



DBSCAN Algorithm

16 December 2023 15:14

Step 0 $\rightarrow \frac{\text{minPts}}{\text{epsilon } \epsilon}$

Step 1 - Identify all points as either core point, border point or noise point

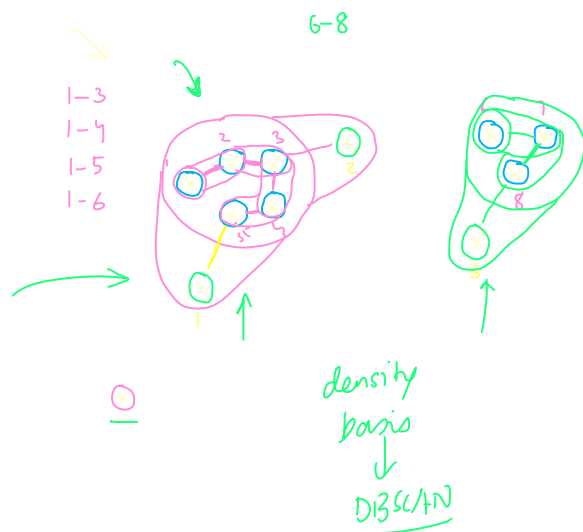
Step 2 - For all of the unclustered core points

Step 2a - Create a new cluster

Step 2b - add all the points that are unclustered and density connected to the current point into this cluster

→ **Step 3** - For each unclustered border point assign it to the cluster of nearest core point

Step 4 - Leave all the noise points as it is.



Code

16 December 2023

15:14

Limitations

16 December 2023 15:14

Advantages

1. Robust to outliers
2. No need to specify clusters
3. Can find arbitrary shaped clusters
4. Only 2 hyperparameters to tune

anomaly det
minPts
epsilon

1. Sensitivity to hyperparameters
2. Difficulty with varying density clusters
3. Does not predict

prediction predict



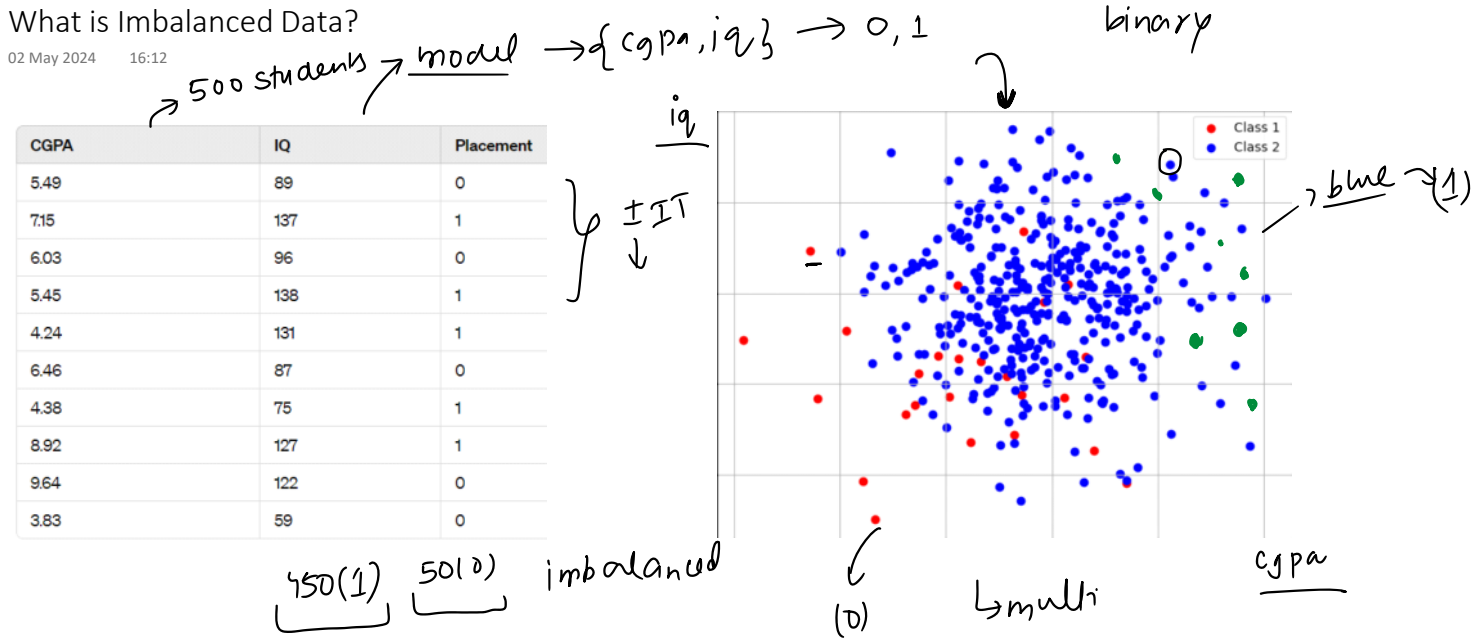
Visualization

16 December 2023

15:14

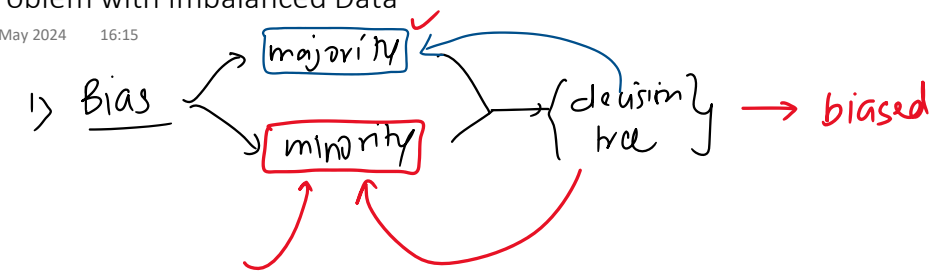
What is Imbalanced Data?

02 May 2024 16:12



Problem with Imbalanced Data

02 May 2024 16:15

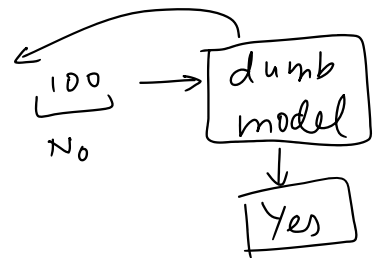


90% → Yes
100/900

2) metrics are not reliable

→ [accuracy] →

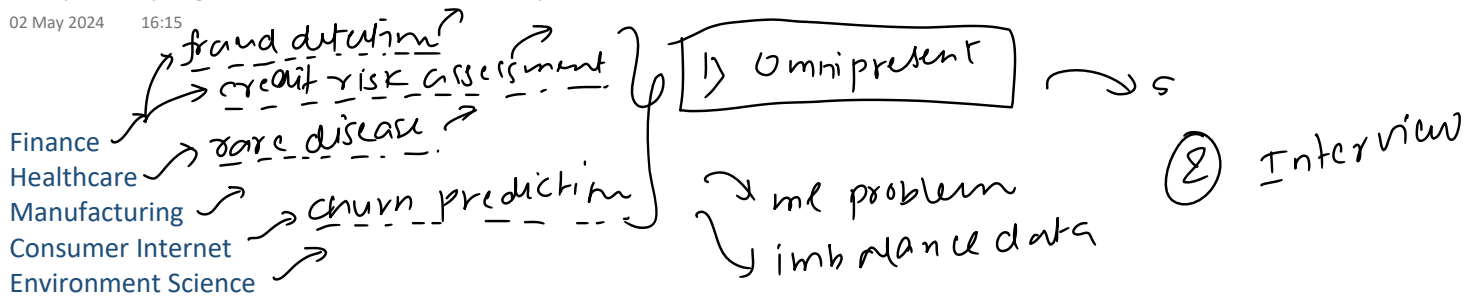
900
Yes



Why studying Imbalanced Data is important?

02 May 2024

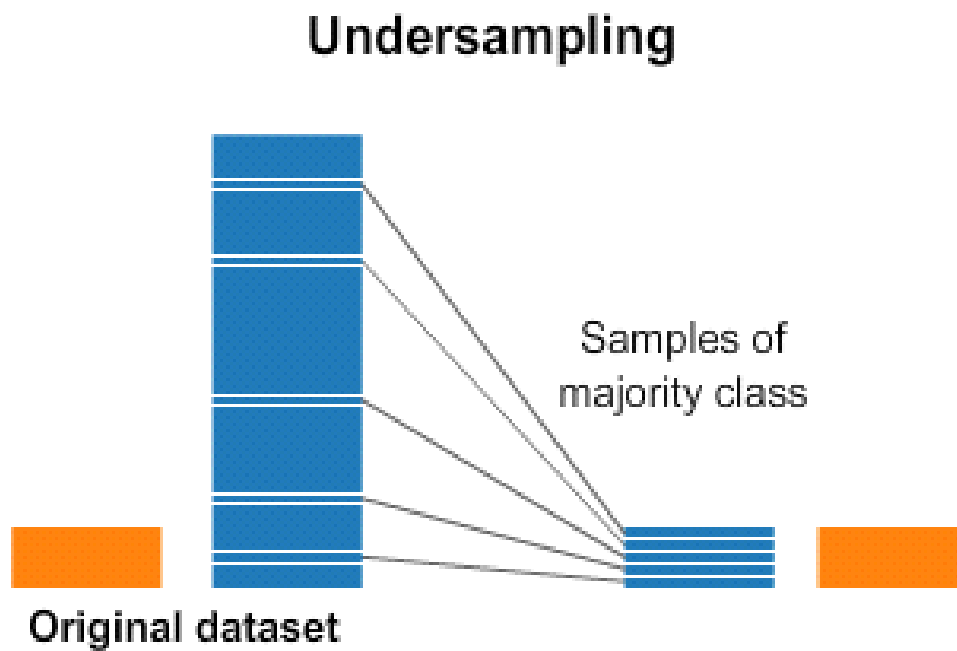
16:15



1. Undersampling

02 May 2024 16:17

2



Advantages

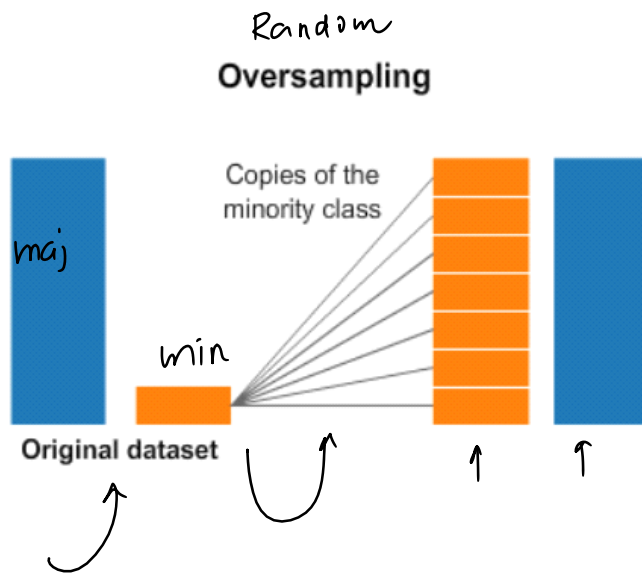
1. Reduction in bias
2. Faster training

Disadvantages

1. Information loss leading to underfitting
2. Sampling Bias

2. Oversampling

03 May 2024 01:13



maj 900
 min 100
 ↓
 900
 duplicate
 balance
 800 → duplication

Advantage

1. Reduced Bias ✓

Disadvantage

1. Increased size
2. Duplication of data may cause overfitting

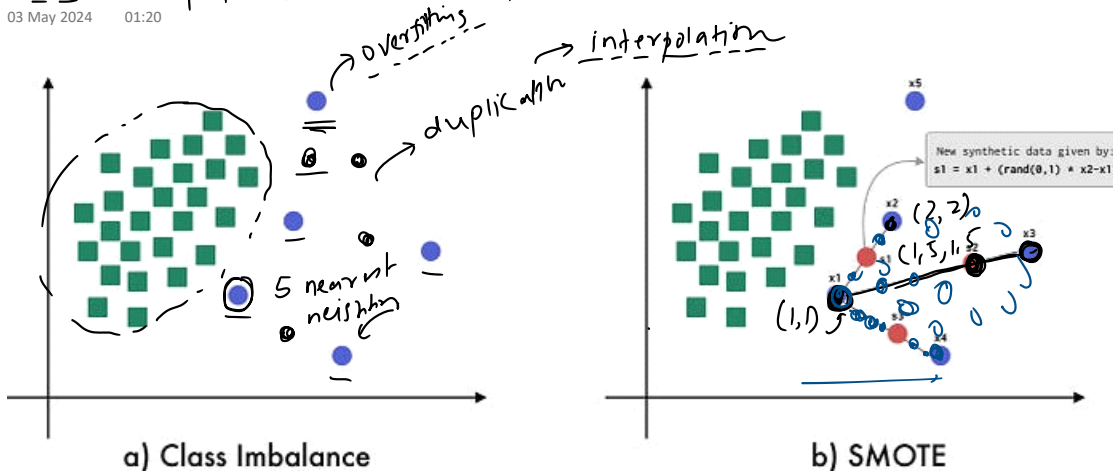
ml also
 → imp → overfitting

↓
 decision tree ROS → overfitting

3. SMOTE

03 May 2024 01:20

popular → Oversampling → min → upsample



a) Class Imbalance

b) SMOTE

0-1 → 0.5

sample - factor * [sample - neighbour]

$$\begin{aligned} [1,1] - 0.5 [(1,1) - (2,2)] \\ (1,1) - 0.5 [-1,-1] \\ (1,1) - (-0.5, -0.5) \\ = 2 [1.5, 1.5] \end{aligned}$$

- Train a KNN on minority class observations - find each observation's 5 closest neighbours

To create the new synthetic data:

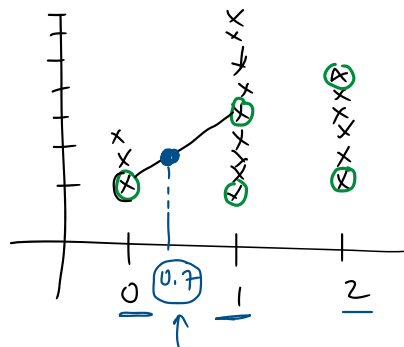
- Select examples from the minority class at random
- Select a neighbour of each example at random (for the interpolation)
- Extract a random number between 0 and 1
- Calculate the new examples as $\text{original sample} - \text{factor} * (\text{original sample} - \text{neighbour})$
- The final dataset consists of the original dataset + the newly created examples

$$0 \rightarrow K \rightarrow K=1$$

Disadvantages

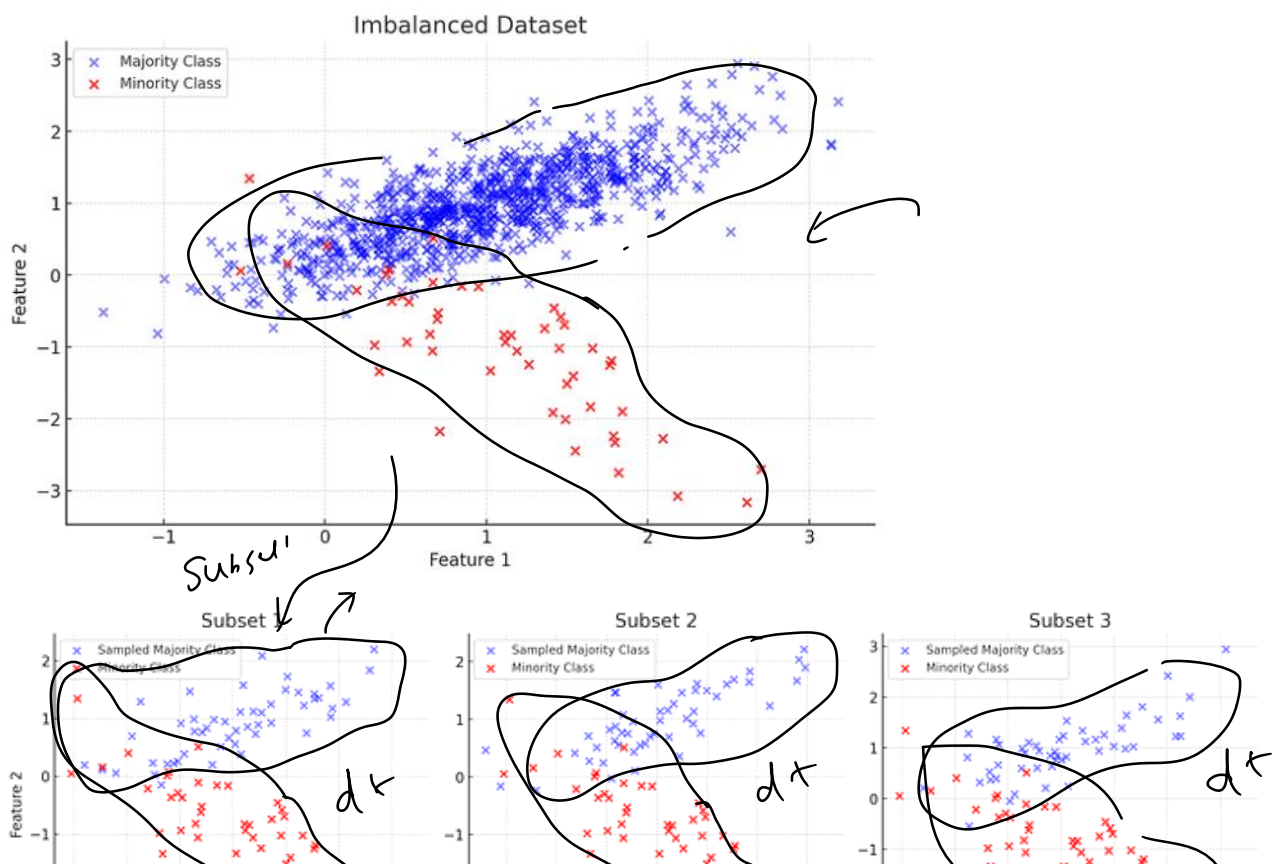
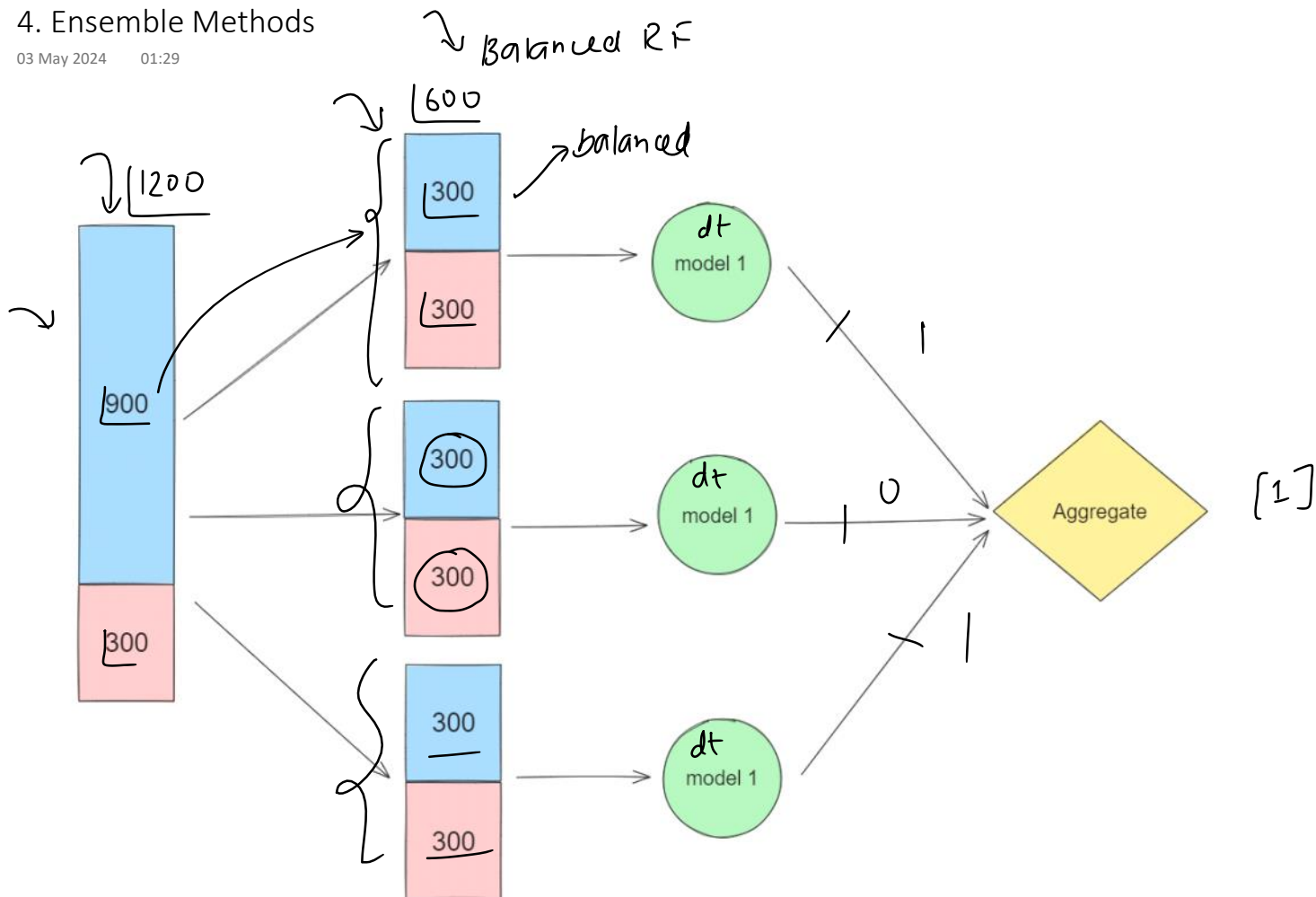
- Does Not Handle Categorical Data Well ✓
- Computational Complexity
- Dependency on the Choice of Neighbors
- Sensitive to Outliers
- Balance Achieved May Not Reflect True Nature

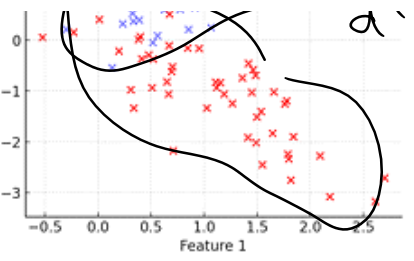
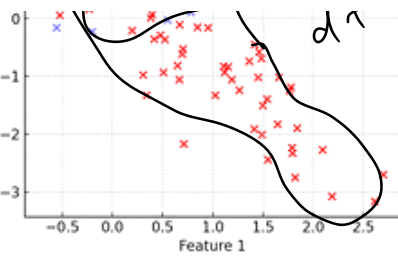
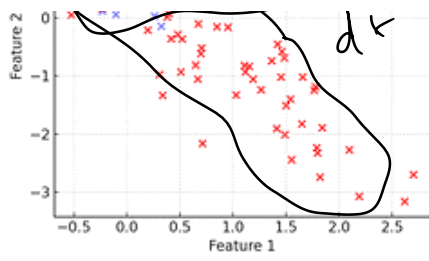
K=5



4. Ensemble Methods

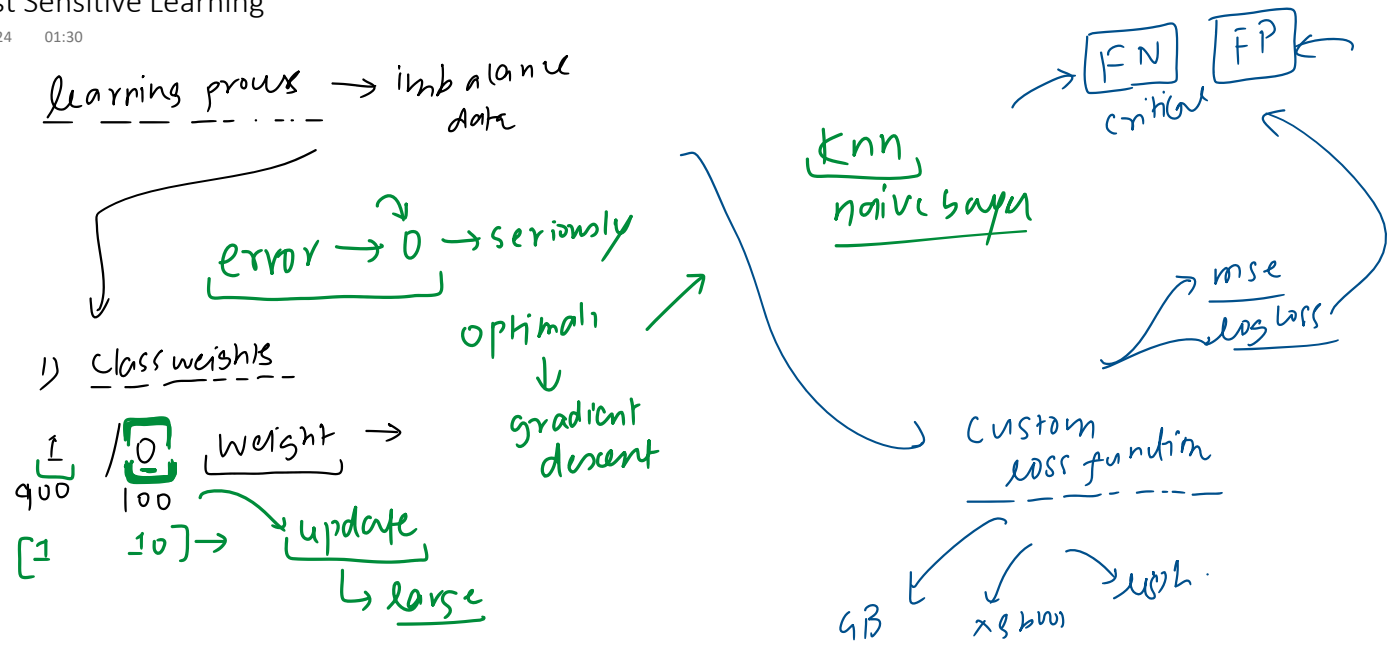
03 May 2024 01:29





5. Cost Sensitive Learning

03 May 2024 01:30



List of all the techniques

03 May 2024 15:06