

Linear Regression. Logistic Regression. Overfitting and Regularization.

COMP5318/COMP4318 Machine Learning and Data Mining
semester 1, 2026, week 3

Lecturer: Tilda Yu

Some slides prepared by Irena Koprinska

Reference: Witten ch.4: 128-131, Müller & Guido: ch.2: 28-31, 47-63,
Geron: ch.4 132-137, 149-161



- Linear regression
- Logistic regression
- Overfitting and regularization
- Ridge and Lasso regression

- **Linear regression** is a prediction method used for *regression* tasks
 - Regression tasks – the predicted variables is numeric
 - Examples: predict the exchange rate of AU\$ based on economic indicators, predict the sales of a company based on the amount spent for advertising
- **Logistic regression** is an extension of linear regression for *classification* tasks
 - Classification tasks – the predicted variable is nominal
- Both linear regression and logistic regression are very popular in statistics



Linear Regression

- Given: a dataset with 2 continuous variables:
 - feature x (also called independent variable)
 - predicted variable y (also called target variable or dependent variable)
- Goal: Approximate the relationship between these variables with a straight line for the given dataset
 - **Prediction (typical task in DM)**: Given a new value of independent variable, use the line to predict the value of the dependent variable
 - **Descriptive analysis (typical task in psychology, health and social sciences)**: assess the strength of the relationship between x and y

- Contains nutritional information for 77 breakfast cereals
- 14 features
 - cereal manufacturer, type (hot or cold), calories, protein [g], fat [g], sodium [mg], fiber [g], carbohydrates [g], **sugar [g]**, potassium [mg], %recommended daily vitamins, weight of 1 serving, number of cups per serving, shelf location (bottom, middle or top)
- Class variable (numeric): **nutritional rating**
- Task: Predict the nutritional rating of a cereal based on its sugar content
 1. Use this data to build the model
 2. Given the sugar content of a new cereal, use the model to predict its nutritional rating
 - New cereal = cereal not used for building of the model

Task: Predict the nutritional rating of a cereal based on its sugar content

1. Use this data to build the model
2. Given the sugar content of a new cereal, use the model to predict its nutritional rating

Cereal Name	Manuf.	Sugars	Calories	Protein	Fat	Sodium	Rating
100% Bran	N	6	70	4	1	130	68.4030
100% Natural Bran	Q	8	120	3	5	15	33.9837
All-Bran	K	5	70	4	1	260	59.4255
All-Bran Extra Fiber	K	0	50	4	0	140	93.7049
Almond Delight	R	8	110	2	2	200	34.3848
Apple Cinnamon Cheerios	G	10	110	2	2	180	29.5095
Apple Jacks	K	14	110	2	0	125	33.1741
Basic 4	G	8	130	3	2	210	37.0386
Bran Chex	R	6	90	2	1	200	49.1203
Bran Flakes	P	5	90	3	0	210	53.3138
Cap'n crunch	Q	12	120	1	2	220	18.0429
Cheerios	G	1	110	6	2	290	50.7650
Cinnamon Toast Crunch	G	9	120	1	3	210	19.8236
Clusters	G	7	110	3	2	140	40.4002
Cocoa Puffs	G	13	110	1	1	180	22.7364

Dependent variable?

rating

Independent variable?

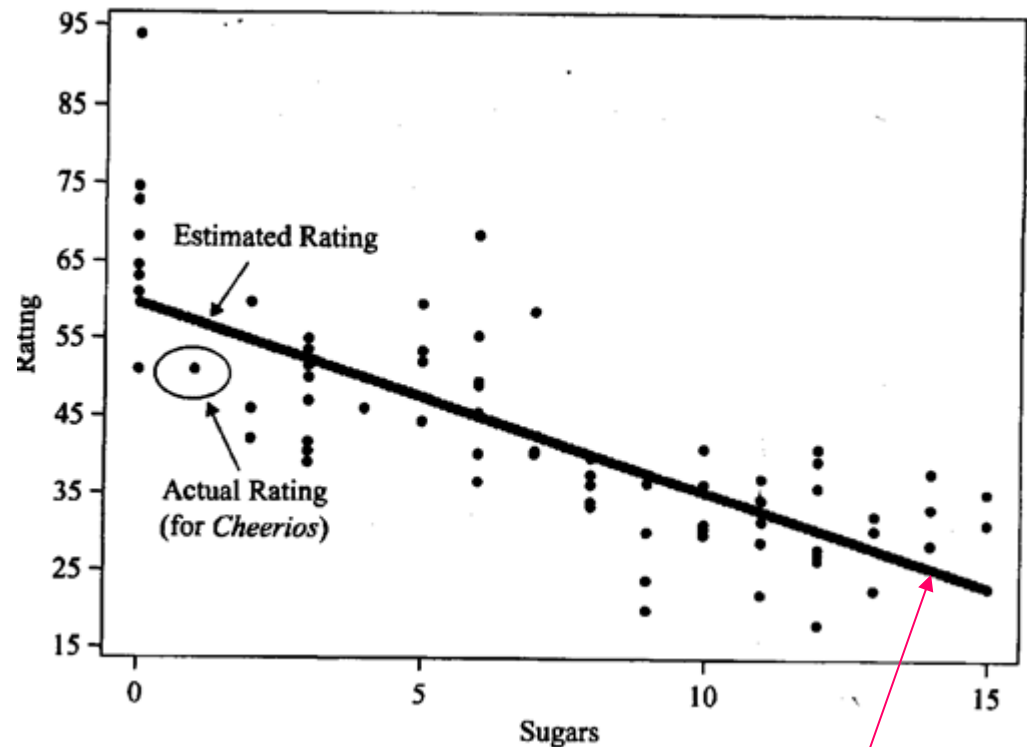
sugars

Example from D. Larose, Data Mining: Methods and Models, 2006, Wiley



- The relationship between sugars and rating is modeled by a line
- The line is used to make predictions

Cereal Name	Manuf.	Sugars	Calories	Protein	Fat	Sodium	Rating
100% Bran	N	6	70	4	1	130	68.4030
100% Natural Bran	Q	8	120	3	5	15	33.9837
All-Bran	K	5	70	4	1	260	59.4255
All-Bran Extra Fiber	K	0	50	4	0	140	93.7049
Almond Delight	R	8	110	2	2	200	34.3848
Apple Cinnamon Cheerios	G	10	110	2	2	180	29.5095
Apple Jacks	K	14	110	2	0	125	33.1741
Basic 4	G	8	130	3	2	210	37.0386
Bran Chex	R	6	90	2	1	200	49.1203
Bran Flakes	P	5	90	3	0	210	53.3138
Cap'n crunch	Q	12	120	1	2	220	18.0429
Cheerios	G	1	110	6	2	290	50.7650
Cinnamon Toast Crunch	G	9	120	1	3	210	19.8236
Clusters	G	7	110	3	2	140	40.4002
Cocoa Puffs	G	13	110	1	1	180	22.7364



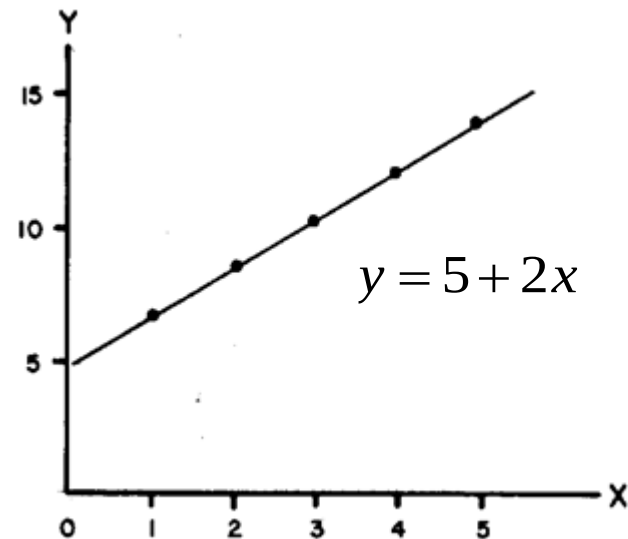
model (regression line)



Equation of a line

$$y = b_0 + b_1x$$

intercept slope

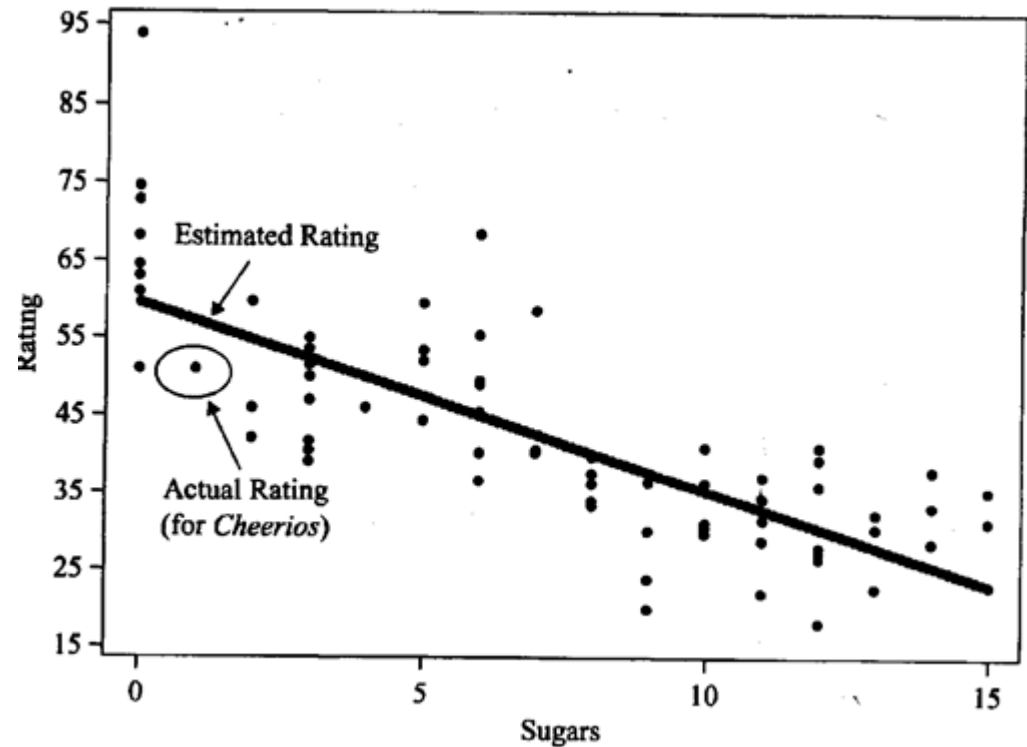


Equation of a regression line

$$\hat{y} = b_0 + b_1x$$

$\hat{y}$ Estimated (predicted)
value of y from the
regression line

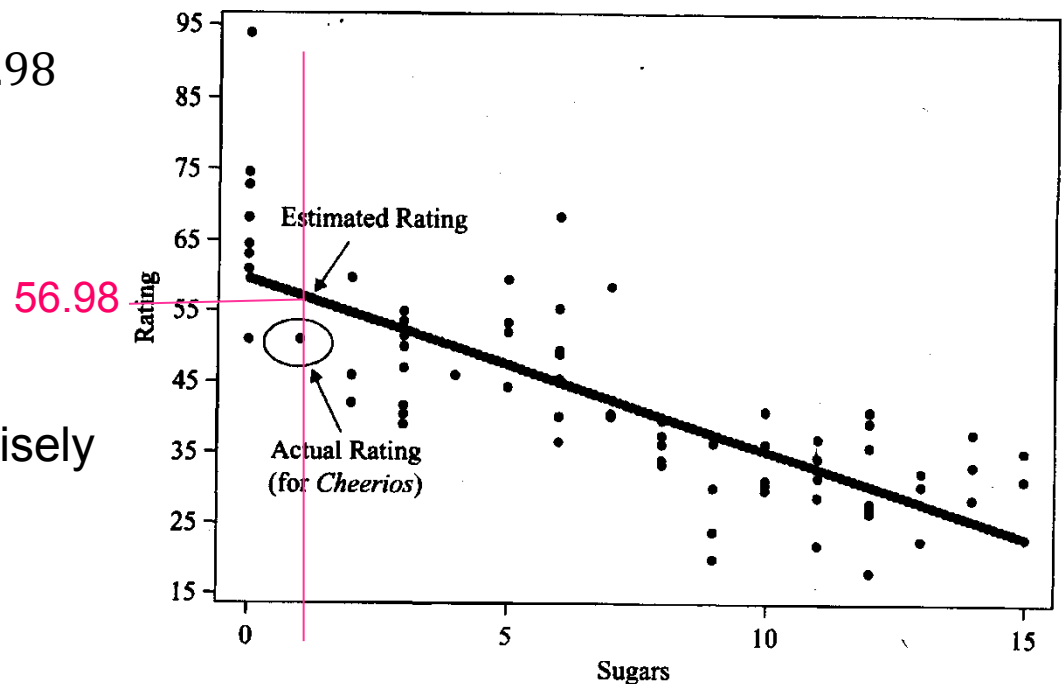
b_0 and b_1 Regression
coefficients



- In our case the computed regression line (model) is
$$\hat{y} = 59.4 - 2.42x$$
- It can be used to make predictions
 - e.g. predict the nutritional rating of a new cereal type (not in the original data) that contains $x=1\text{g}$ sugar

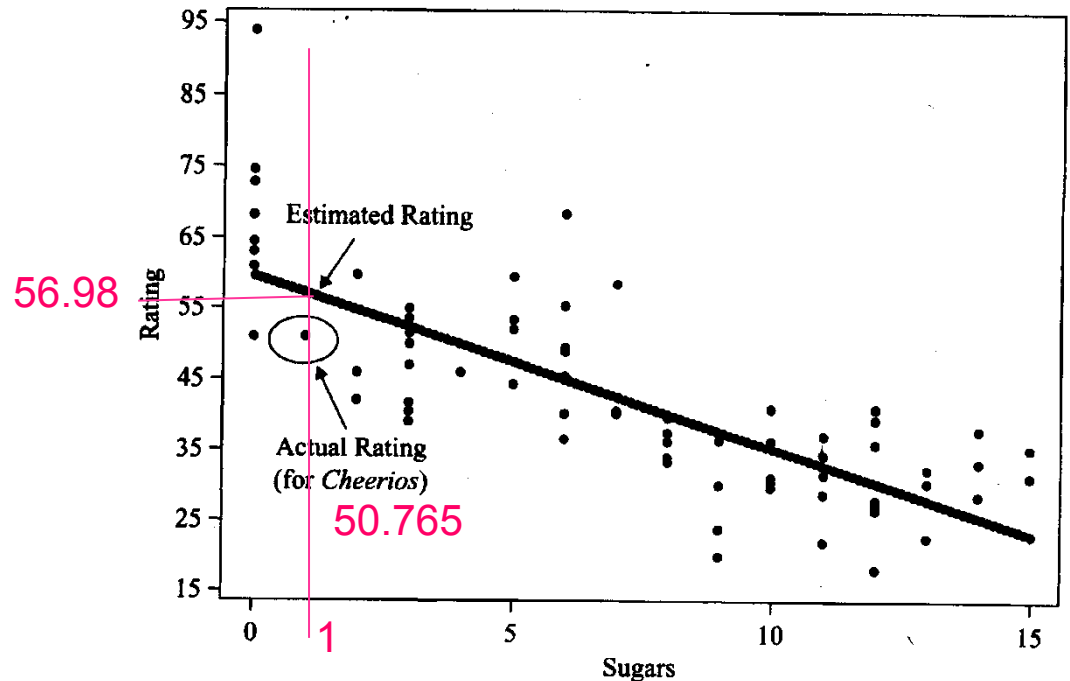
$$\hat{y} = 59.4 - 2.42 * 1 = 56.98$$

- The predicted value lies precisely on the regression line



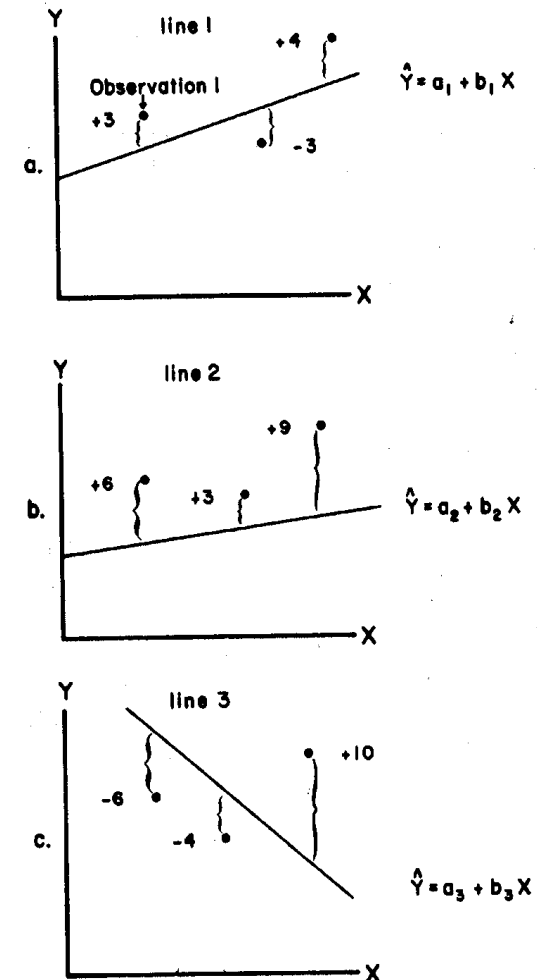
How to make predictions (2)

- We have a cereal type in our dataset with sugar = 1g: Cheerios
- Its nutritional rating is: 50.765 (actual value) not 56.98 (predicted)
- The difference is called **prediction error** or **residual**



- There are many lines that can be fitted to the given dataset. Which one is the best one?
 - The one “closest” to the data
 - Mathematically:
 - Prediction error (residual) =
actual value - predicted value:

$$\varepsilon = y_i - \hat{y}_i$$
- Performance index: sum of squared prediction errors (SSE): $SSE = \sum_i (y_i - \hat{y}_i)^2$
- Our goal: select the line which minimizes SSE
- Can be solved using the *method of the least squares*



$$b_1 = \frac{\sum x_i y_i - [(\sum x_i) (\sum y_i)] / n}{\sum x_i^2 - (\sum x_i)^2 / n}$$

$$b_0 = \bar{y} - b_1 \bar{x}$$

$\bar{x}$ - mean value of x

$\bar{y}$ - mean value of y

n – number of training examples (= data points, observations)

- This solution is obtained by minimizing SSE using differential calculus
- If you are interested to see how this was done, please see Appendix 1 at the end

- The least squares method finds the best fit to the data but doesn't tell us how good this fit is
 - E.g. $SSE=12$; is this large or small?
- R^2 measures the *goodness of fit* of the regression line found by the least squares method:

$$R^2 = \frac{SSR}{SST}$$

- Values between 0 and 1; the higher the better
 - = 1: the regression line fits perfectly the training data
 - close to 0: poor fit
- What are SSR and SST?

- 1. SSE - Sum of squared *prediction* errors

$$SSE = \sum_i^n (y_i - \hat{y}_i)^2 \quad = \text{actual value} - \text{predicted value}$$

- 2. SST - Sum of squared *total* errors

$$SST = \sum_{i=1}^n (y_i - \bar{y})^2 \quad = \text{actual value} - \text{mean value}$$

- Hence, SST measures the prediction error when the predicted value is the mean value
- SST is a function of the variance of y (variance = standard deviation²) => SST is a measure of the variability of y, without considering x

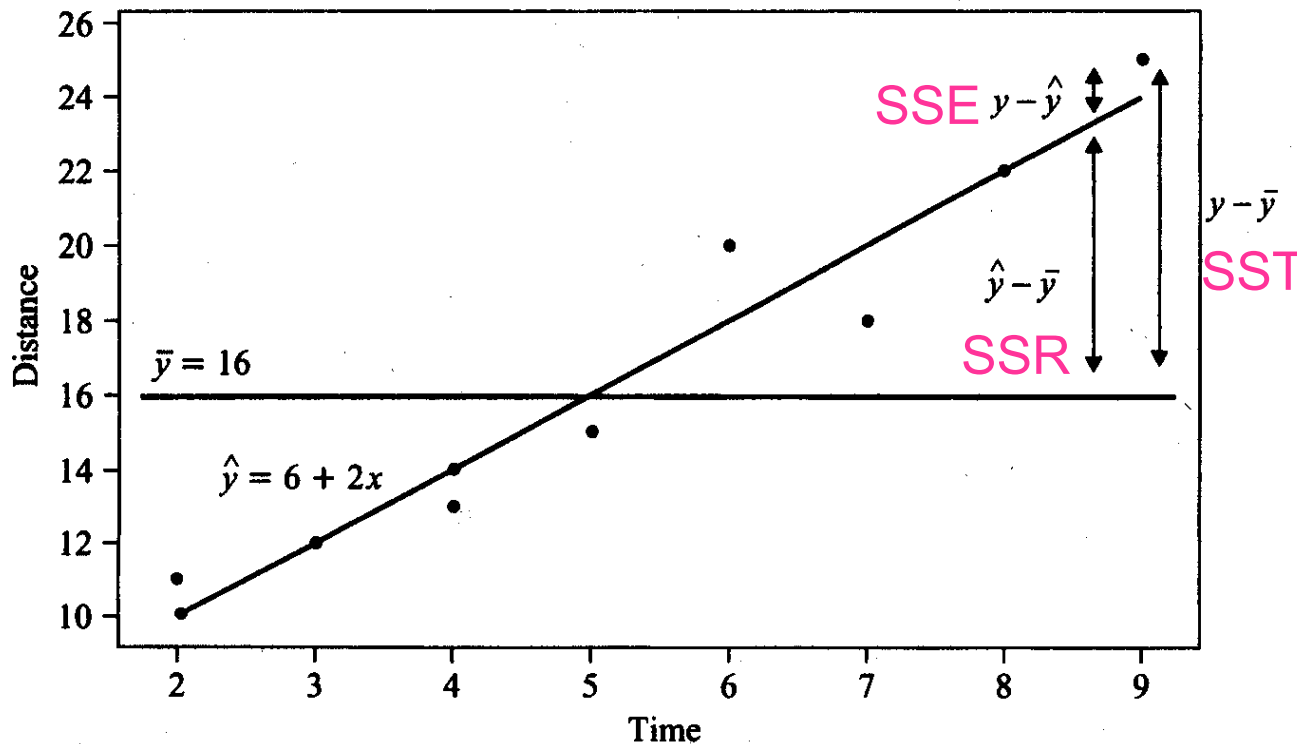
$$SST = \sum_i^n (y_i - \bar{y})^2 = (n-1) \text{var}(y)$$

Can be used as a baseline - predicting y without knowing x

Three types of errors (2)

- 3. SSR - Sum of squared *regression* errors = predicted value – mean value

$$SSR = \sum_i^n (\hat{y}_i - \bar{y})^2$$



Ex.: Distance travelled for a number of hours

Subject	Time, x (hours)	Distance, y (km)	S
1	2	10	
2	2	11	
3	3	12	
4	4	13	
5	4	14	
6	5	15	
7	6	20	
8	7	18	
9	8	22	
10	9	25	

Relation between SST, SSR and SSE

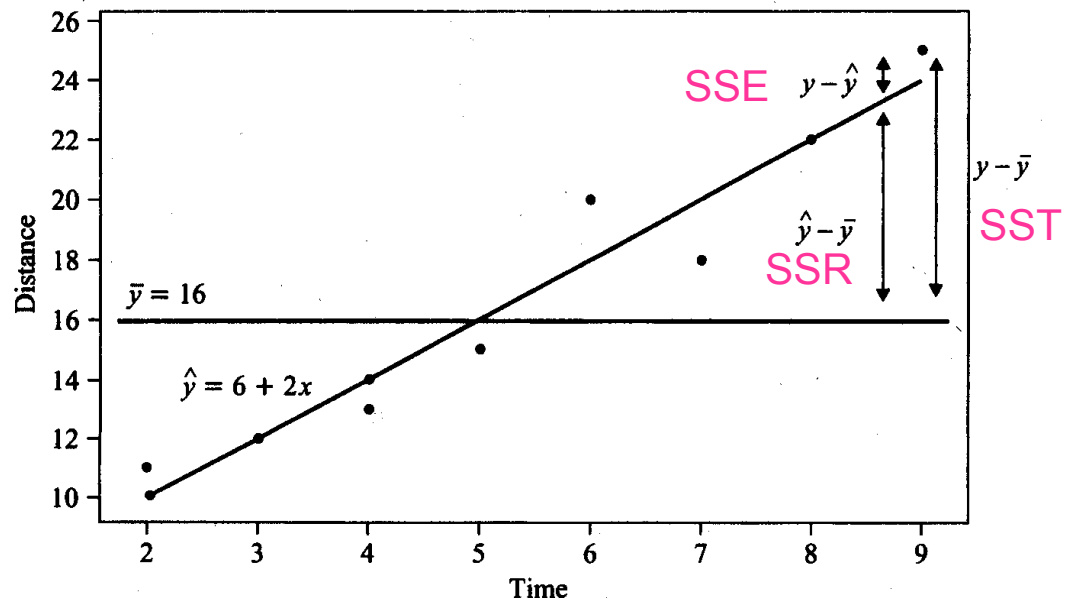
- From the graph: $y_i - \bar{y}_i = (\hat{y}_i - \bar{y}_i) + (y_i - \hat{y}_i)$

- It can be shown that $SST = SSR + SSE$

(For the interested students:
How? By squaring each side: $\sum_{i=1}^n (y_i - \bar{y}_i)^2 = \sum_{i=1}^n (\hat{y}_i - \bar{y}_i)^2 + \sum_{i=1}^n (y_i - \hat{y}_i)^2$

The cross product cancels out as shown in this book:

N. Draper and H. Smith, Applied Regression Analysis, Wiley, 1998)



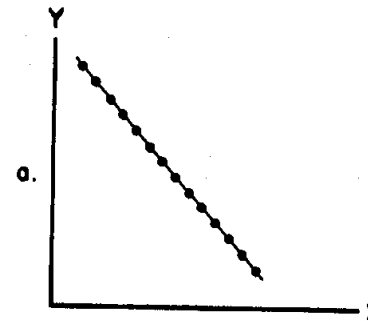
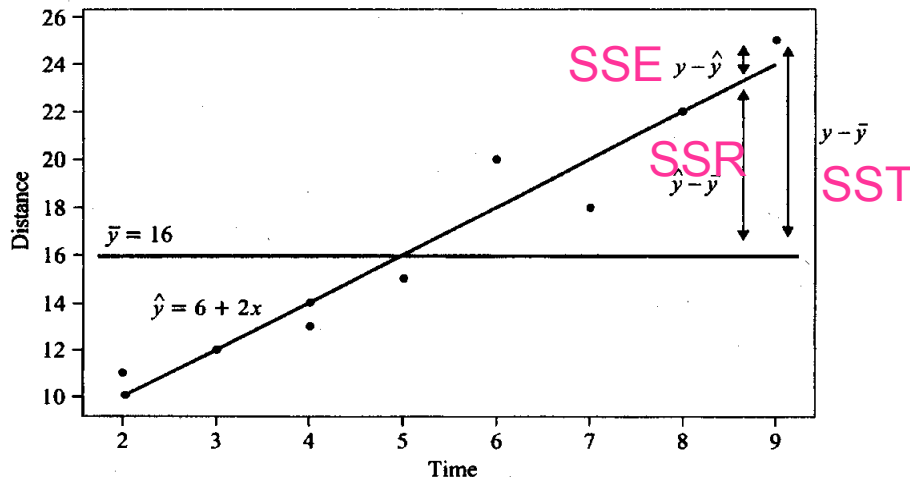
Coefficient of determination R^2 - again

$$R^2 = \frac{SSR}{SST}$$

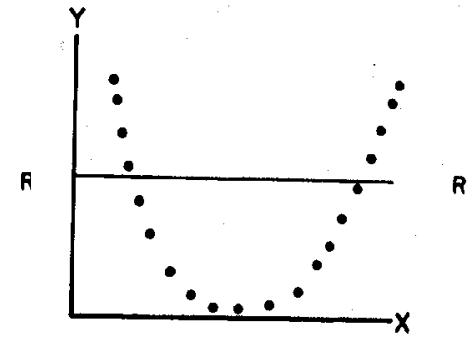
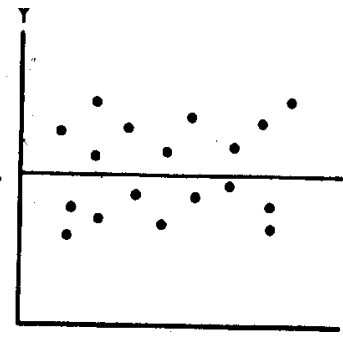
- Measures the goodness of fit of the regression line to the training data
- Values between 0 and 1; the higher the better
 - 1: perfect fit, $SSE=0$; Why is it 1 when $SSE=0$?
 - 0: x is not helpful for predicting y, $SSR=0$

If $SSE=0$

$$R^2 = \frac{SSR}{SST} = \frac{SST - SSE}{SST} = \frac{SST}{SST} = 1$$



Is R^2 high or low?



- r - correlation coefficient; measures linear relationship between 2 vectors $\mathbf{x}$ and $\mathbf{y}$ (see slides for week 1b):

$$r = \text{corr}(\mathbf{x}, \mathbf{y}) = \frac{\text{covar}(\mathbf{x}, \mathbf{y})}{\text{std}(\mathbf{x})\text{std}(\mathbf{y})} = \frac{\text{covar}(\mathbf{x}, \mathbf{y})}{\sqrt{\text{var}(\mathbf{x})\text{var}(\mathbf{y})}}$$

- R^2 – coefficient of determination; measures how well the regression line represents the data: $R^2 = \frac{SSR}{SST}$

- It can be shown that $r = \sqrt{R^2}$

Except for the sign of r , which depends on the direction of the relationship, positive or negative, so:

$$r = \pm\sqrt{R^2}$$

- MAE, MSE and RMSE are other performance measures for evaluating:
 - how good the model is (performance on training data) and
 - how well it works on new data (performance on test data)
- They are widely used in ML and DM

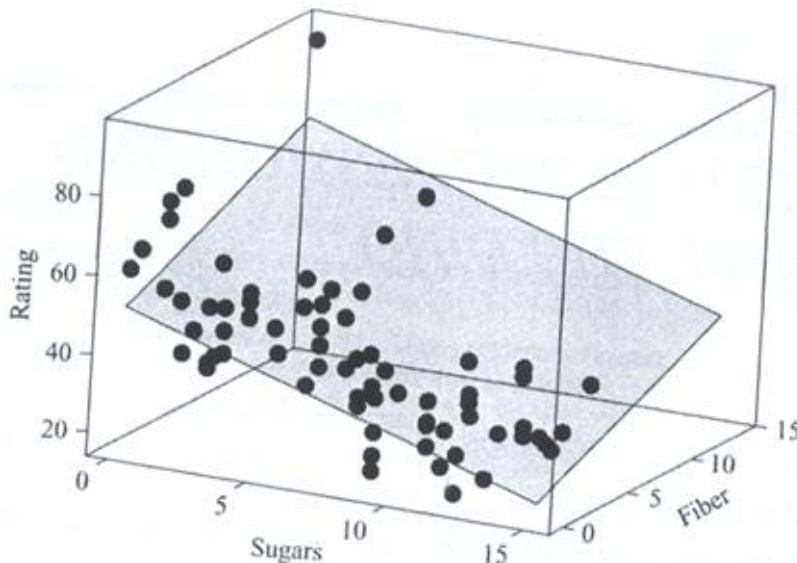
- Mean Absolute Error (MAE): $MAE = \frac{1}{n} \sum_{i=1}^n |\hat{y}_i - y_i|$

- Mean Squared Error (MSE): $MSE = \frac{1}{n} \sum_{i=1}^n (\hat{y}_i - y_i)^2$

- Root Mean Squared Error (RMSE):

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

- Simple regression: 1 feature
- Multiple regression: more than 1 feature



- The line becomes a plane in 2-dim. space and a hyperplane in >2 -dim. space
- R^2 is similarly defined, called **multiple coefficient of determination**

- True or False?
- 1) The regression line minimizes the sum of the residuals
- 2) If all residuals are 0, $SST=SSR$
- 3) If the value of the correlation coefficient is negative, this indicates that the 2 variables are negatively correlated
- 4) The value of the correlation coefficient can be calculated given the value of R^2
- 5) SSR represents an overall measure of the prediction error on the training set by using the regression line

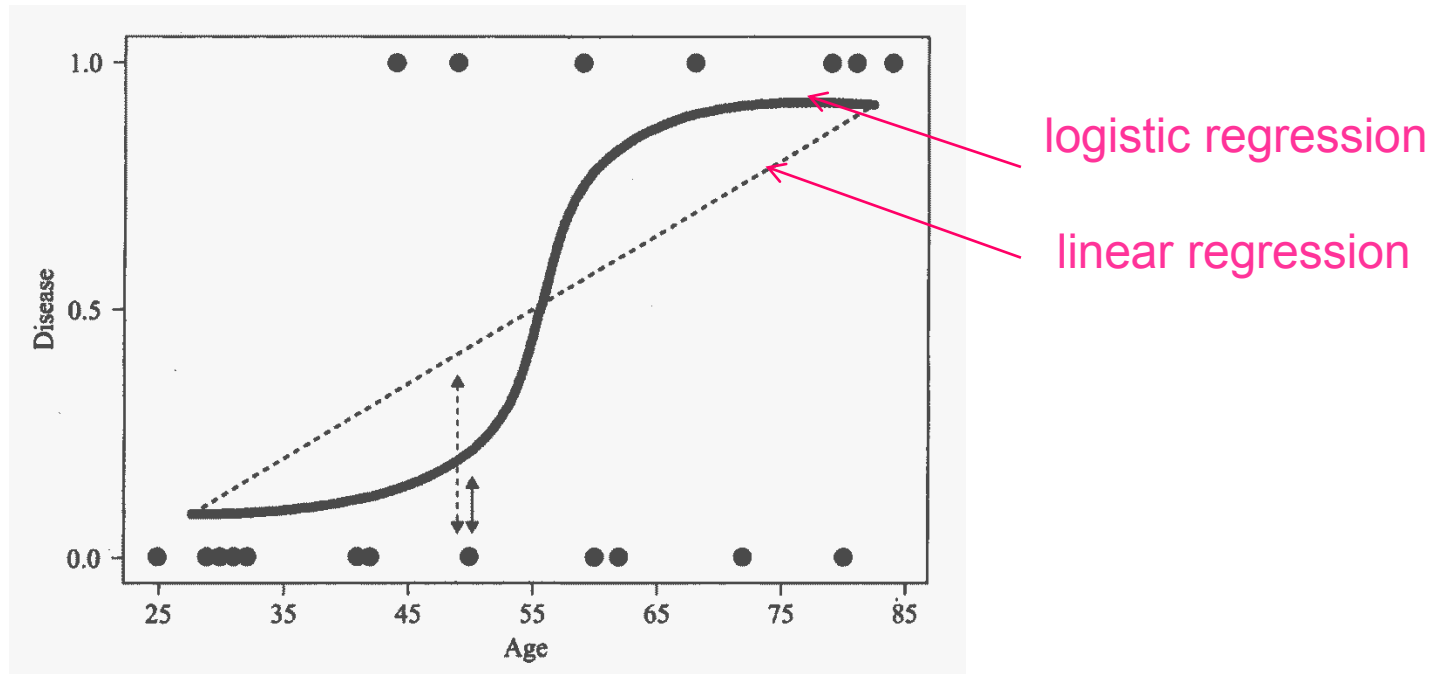
- True or False?
- 1) The regression line minimizes the sum of the residuals **False**
No, the sum of squared residuals $SSE = \sum_i (y_i - \hat{y}_i)^2$
- 2) If all residuals are 0, SST=SSR **True**
If the residuals are 0 => SSE will be 0; SST=SSR+SSE => SST=SSR
- 3) If the value of the correlation coefficient is negative, this indicates that the 2 variables are negatively correlated **True**
- 4) The value of the correlation coefficient can be calculated given the value of R^2 **False. Only the value can be calculated, but not the sign.**
 $r = \pm\sqrt{R^2}$
- 5) SSR represents an overall measure of the prediction error on the training set by using the regression line **False**
No, this is SSE, R^2 or other measures such as MAE, MSE, RMSE

- Note that if the LR model is fitted on one dataset but tested on another dataset, then it is possible that the R^2 value is negative
- We will see such case during the tutorial – a LR model trained on the training set and tested on the test set
 - R^2 on the training set: 0.69
 - R^2 on the test set: -0.73 (negative)
- Negative value means a poor fit
 - R^2 on the training set: 0.69 - good fit on the training data
 - R^2 on the test set: -0.73 - poor fit on the test data
 - => overfitting
- If the model is trained and tested on the same dataset, R^2 is always between 0 and 1



Logistic Regression

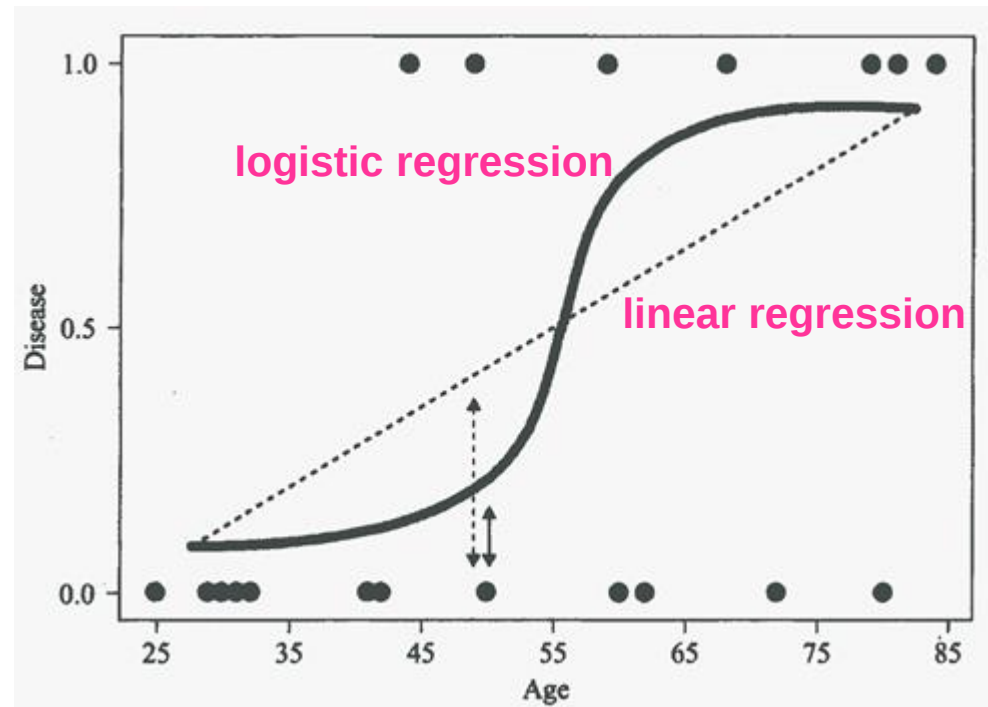
- Used for classification tasks
- Two classes: 0 and 1 (there are extensions for more than 2 classes)
- Fits the data to a logistic (sigmoidal) curve instead of fitting it to a straight line
 - => assumes that the relationship between the feature and class variable is *nonlinear*



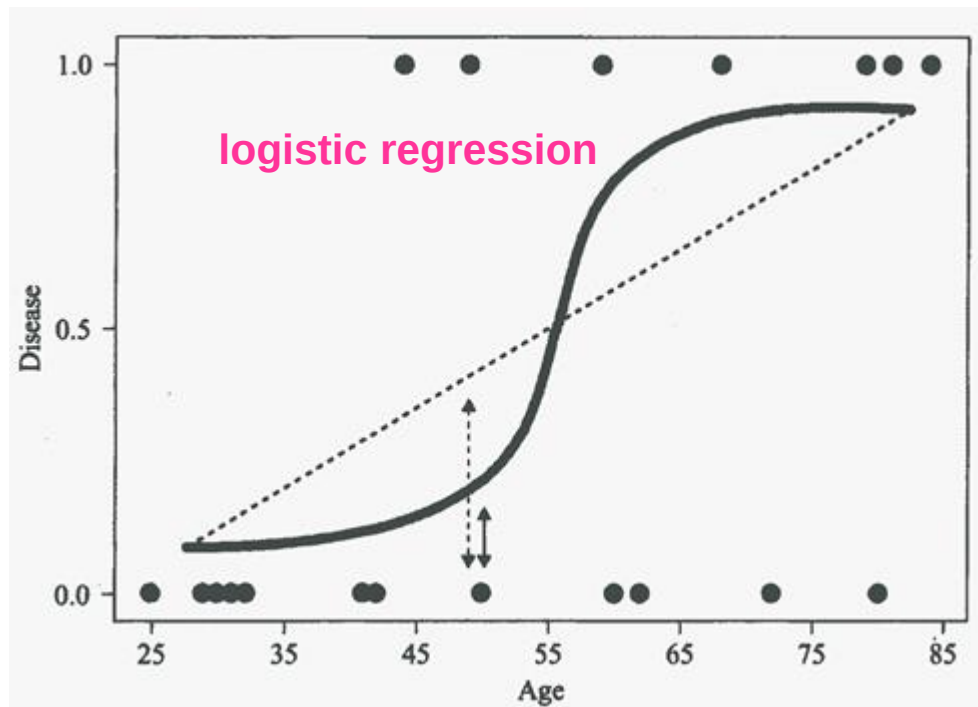
Simple (bivariate) logistic regression

- Example: Predicting the presence (class=1) or absence (class=0) of a particular disease, given the patient's age

ID	age	disease	ID	age	disease
1	25	0	11	50	0
2	29	0	12	59	1
3	30	0	13	60	0
4	31	0	14	62	0
5	32	0	15	68	1
6	41	0	16	72	0
7	41	0	17	79	1
8	42	0	18	80	0
9	44	1	19	81	1
10	49	1	20	84	1



- What will be the prediction of Logistic Regression for patient 11 from the training data (age=50, disease=0)?



- The equation of the logistic (sigmoidal) curve is:

$$p = \frac{e^{b_0 + b_1 x}}{1 + e^{b_0 + b_1 x}}$$

- It gives a value between 0 and 1 that is interpreted as the probability for class membership:
- p is the probability for class 1 and $1-p$ is the probability for class 0
- It uses the *maximum likelihood method* to find the parameters b_0 and b_1 - the curve that best fits the data

- The logistic regression produced $b_0 = -4.372$, $b_1 = 0.06696$
- $\Rightarrow$ the probability for a patient aged 50 (training example 11) to have the disease:

$$p = \frac{e^{b_0 + b_1 x}}{1 + e^{b_0 + b_1 x}} = \frac{e^{-4.372 + 0.06696 \cdot \text{age}}}{1 + e^{-4.372 + 0.06696 \cdot \text{age}}} = 0.26$$

- $\Rightarrow$ 26% to have the disease and 74% not to have the disease
- We can use the probability directly or convert it into 0/1 answer required for classification tasks, e.g. 0 if $p < 0.5$ and 1 if $p \geq 0.5$
- $\Rightarrow$ We predict class 0 for this patient
 - Other thresholds (not 0.5) are also possible depending on domain knowledge
- The class for new examples can be predicted similarly – e.g. make a prediction for a patient aged 45

$$p = \frac{e^{b_0 + b_1 x}}{1 + e^{b_0 + b_1 x}}$$

- It also follows that: How can this be shown? See Appendix 2 at the end.

$$b_0 + b_1 x = \ln \frac{p}{1 - p}$$

$$\ln \frac{p}{1 - p} = b_0 + b_1 x$$

linear calculation, as in linear regression

called odds ratio for the default class (class 1)

$$\ln(odds) = b_0 + b_1 x \quad \Rightarrow \quad odds = e^{(b_0 + b_1 x)}$$

Compare:

- Logistic regression: $\ln(odds) = b_0 + b_1 x$
- Linear regression: $\hat{y} = b_0 + b_1 x$

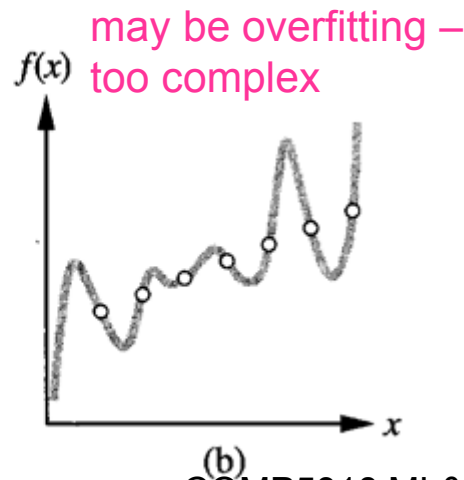
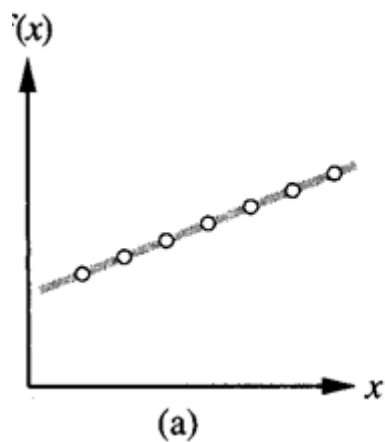
The model is still a linear combination of the input features, but this combination determines the log odds of the class not directly the predicted value



Overfitting and Regularization

- Overfitting:
 - Small error on the training set but high error on test set (new examples)
 - The classifier has memorized the training examples but has not learned to generalize to new examples!
- It occurs when
 - we fit a model too closely to the particularities of the training set – the resulting model is too specific, works well on the training data but doesn't work well on new data

Ex.1



Ex.2

Rule1: may be overfitting – too specific
If age>45, income>100K, has_children=3,
divorced=no -> buy_boat=yes

Rule2:
If age>45, income>100K -> buy_boat=yes

- Various reasons, e.g.
 - Issues with the data
 - Noise in the training data
 - Training data does not contain enough representative examples – too small
 - Training data very different from test data – not representative enough
 - How the algorithm operates
 - Some algorithms are more susceptible to overfitting than others
 - Different algorithms have different strategies to deal with overfitting, e.g.
 - Decision tree – prune the tree
 - Neural networks – early stopping of the training
 - ...

- The model is too simple and doesn't capture all important aspects of the data
- It performs badly on both training and test data

Rule1: **may be overfitting – too specific**

If age>45, income>100K, has_children=3,
divorced=no -> buy_boat=yes

Rule2:

If age>45, income>100K -> buy_boat=yes

Rule3: **may be underfitting – too general**

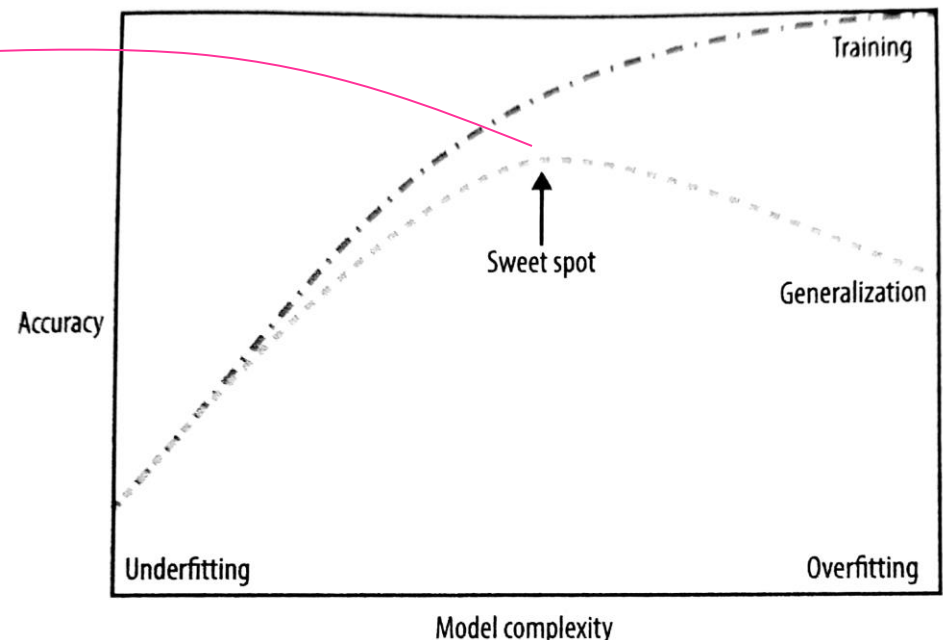
If owns_hourse=yes -> buy_boat=yes

Trade-off between model complexity and generalization performance

- Generalization performance = accuracy on test set
- Usually, the more complex we allow the model to be, the better it will predict on the training data
- However, if it becomes too complex, it will start focusing too much on each individual data point, and will not generalize well on new data

Image from A. Mueller and S. Guido, Introduction to ML with Python

- There is **point** in between, which will yield the best test accuracy
- This is the model we want to find



- Regularization means explicitly restricting a model to avoid overfitting
- It is used in some regression models (e.g. Ridge and Lasso regression) and in some neural networks



Ridge and Lasso Regression

- A regularized version of the standard Linear Regression (LR)
- Also called Tikhonov regularization
- Uses the same equation as LR to make predictions
- However, the regression coefficients w are chosen so that they not only fit well the training data (as in LR) but also satisfy an additional constraint:
 - the magnitude of the coefficients is as small as possible, i.e. close to 0
- Small values of the coefficients means
 - each feature will have little effect on the outcome
 - small slope of the regression line
- Rationale: a more restricted model (less complex) is less likely to overfit
- Ridge regression uses the so called L2 regularization (L2 norm of the weight vector)

- Minimizes the following cost function:

$$\underbrace{\frac{1}{n} \sum_{i=1}^n (\hat{y}_i - y_i)^2}_{\text{MSE}} + \underbrace{\alpha \sum_{i=1}^m w_i^2}_{\text{regularization term}}$$

n - number of training examples

m – number of regression coefficients (weights)

MSE

regularization term

Goal: high accuracy
on training data (low
MSE)

low complexity

model – w close to 0

- Parameter α controls the trade-off between the performance on the training set and model complexity



Ridge regression (3)

$$\underbrace{\frac{1}{n} \sum_{i=1}^n (\hat{y}_i - y_i)^2}_{\text{MSE}} + \alpha \underbrace{\sum_{i=1}^m w_i^2}_{\text{regularization term (L2 norm)}}$$

- α controls the trade-off between the performance on the training set and model complexity
 - Increasing α makes the coefficients smaller (close to 0); this typically decreases the performance on the training set but may improve the performance on the test set
 - Decreasing α means less restricted coefficients. For very small α , Ridge Regression will behave similarly to LR

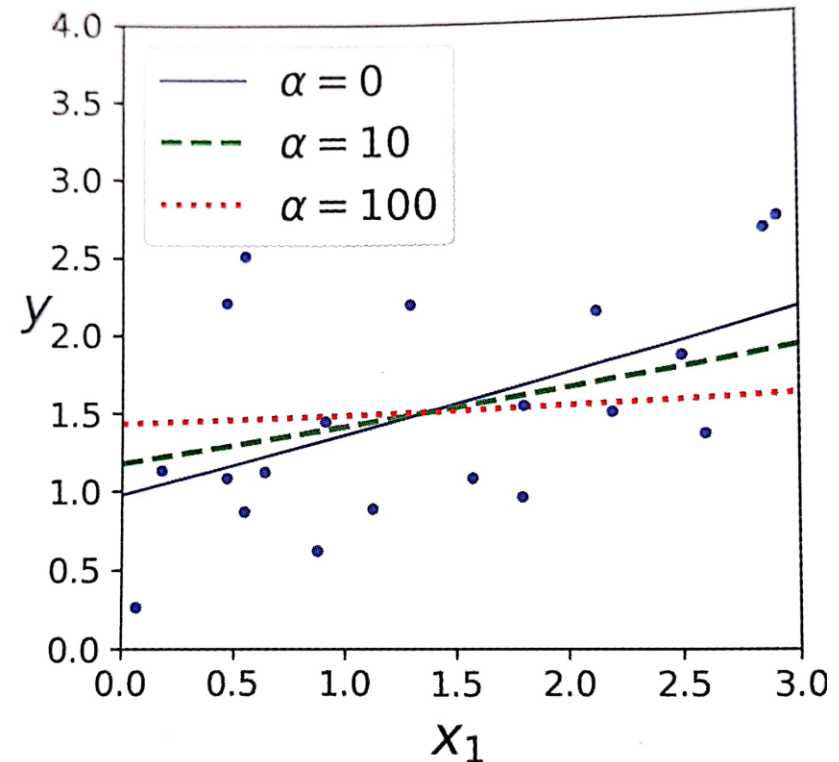


Image from A. Geron, Hands-on ML with Scikit-learn, Keras & TensorFlow

- Another regularized version of the standard LR
- LASSO = Least Absolute Shrinkage and Selection Operator Regression
- As Ridge Regression, it adds a regularization term to the cost function but it uses the L1 norm of the regression coefficient vector w

$$\underbrace{\frac{1}{n} \sum_{i=1}^n (\hat{y}_i - y_i)^2}_{\text{MSE}} + \underbrace{\alpha \sum_{i=1}^m |w_i|}_{\text{regularization term (L1 norm)}}$$

Goal: high accuracy
on training data (low
MSE)

low complexity model

- Consequence of using L1 – some w will become **exactly 0**
- => some features will be completely ignored by the model – a form of automatic feature selection
- Less features – simpler model, easier to interpret



Lasso regression (2)

$$\underbrace{\frac{1}{n} \sum_{i=1}^n (\hat{y}_i - y_i)^2}_{\text{MSE}} + \underbrace{\alpha \sum_{i=1}^m |w_i|}_{\text{regularization term (L1 norm)}}$$

- As in Ridge Regression:
 - α controls the trade-off between the performance on the training set and model complexity
 - Increasing/decreasing α - similar reasoning as before

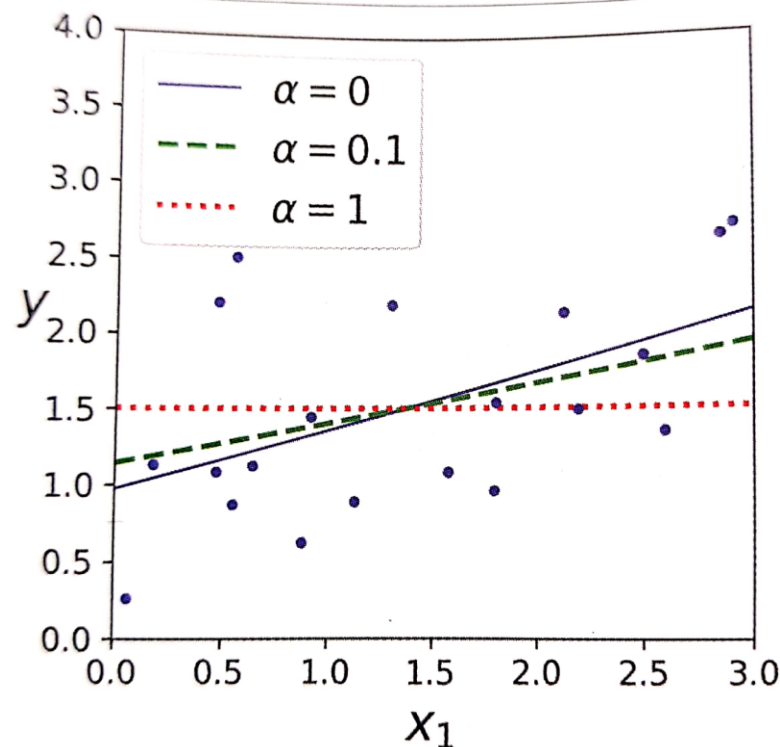


Image from A. Geron, Hands-on ML with Scikit-learn, Keras & TensorFlow

- Linear regression
 - Simple (bivariate) - a line is used to approximate the relationship between 2 continuous variables (feature x and class variable y)
 - Multiple – more than 1 feature; the line becomes a hyperplane
 - The least-square method is used to find the line (hyperplane) which best fit the given data (training data)
 - “Best fit”: minimizes the sum of the squared errors (SSE) between the actual and predicted values of y , over all data points
 - R^2 = coefficient of determination = SSR/SST – how well the line fits the data; the higher the better
 - MAE, MSE and RMSE – widely used accuracy measures in ML (can be measured on both training and test data)

- Logistic regression
 - Simple (bivariate) - a sigmoidal curve is used to approximate the relationship between the feature x and class variable y
 - $\Rightarrow$ assumes the relationship between the feature and class variable is nonlinear
 - Multiple – more than 1 feature; the sigmoidal curve becomes a sigmoidal hyperplane
 - Uses the maximum likelihood method to find the curve (hyperplane) which best fit the given data (training data)
- Overfitting and regularization
 - Overfitting - high accuracy on training data but low accuracy on test data (low generalization)
 - High model complexity $\rightarrow$ low generalization
 - Regularization is a method to avoid overfitting – it makes the model more restrictive (less complex)
 - Ridge and Lasso regression are regularized linear regression models

- M. Lewis-Beck, Applied statistics, SAGE University Paper Series on Quantitative Analysis.
- D. Larose, Data Mining: Methods and Models, 2006, Wiley.

- For interested students; not examinable
- From D. Larose, Data Mining: Methods and Models, 2006, Wiley; p.36-37

The least-squares line is that line which minimizes the population sum of squared errors, $SSE_p = \sum_{i=1}^n \varepsilon_i^2$. First, we reexpress the population sum of squared errors as

$$SSE_p = \sum_{i=1}^n \varepsilon_i^2 = \sum_{i=1}^n (y_i - \beta_0 - \beta_1 x_i)^2 \quad (2.2)$$

Then, recalling our differential calculus, we may find the values of β_0 and β_1 that minimize $\sum_{i=1}^n \varepsilon_i^2$ by differentiating equation (2.2) with respect to β_0 and β_1 and setting the results equal to zero. The partial derivatives of equation (2.2) with respect to β_0 and β_1 are, respectively,

$$\begin{aligned} \frac{\partial SSE_p}{\partial \beta_0} &= -2 \sum_{i=1}^n (y_i - \beta_0 - \beta_1 x_i) \\ \frac{\partial SSE_p}{\partial \beta_1} &= -2 \sum_{i=1}^n x_i (y_i - \beta_0 - \beta_1 x_i) \end{aligned} \quad (2.3)$$

We are interested in the values for the estimates b_0 and b_1 , so setting equations (2.3) equal to zero, we have

$$\begin{aligned} \sum_{i=1}^n (y_i - b_0 - b_1 x_i) &= 0 \\ \sum_{i=1}^n x_i (y_i - b_0 - b_1 x_i) &= 0 \end{aligned}$$



Appendix 1: Minimizing SSE (2)

Distributing the summation gives us

$$\begin{aligned}\sum_{i=1}^n y_i - nb_0 - b_1 \sum_{i=1}^n x_i &= 0 \\ \sum_{i=1}^n x_i y_i - b_0 \sum_{i=1}^n x_i - b_1 \sum_{i=1}^n x_i^2 &= 0\end{aligned}$$

which is reexpressed as

$$\begin{aligned}b_0 n + b_1 \sum_{i=1}^n x_i &= \sum_{i=1}^n y_i \\ b_0 \sum_{i=1}^n x_i + b_1 \sum_{i=1}^n x_i^2 &= \sum_{i=1}^n x_i y_i\end{aligned}\tag{2.4}$$

Solving equations (2.4) for b_1 and b_0 , we have

$$b_1 = \frac{\sum x_i y_i - [(\sum x_i)(\sum y_i)]/n}{\sum x_i^2 - (\sum x_i)^2/n}\tag{2.5}$$

$$b_0 = \bar{y} - b_1 \bar{x}\tag{2.6}$$

where n is the total number of observations, $\bar{x}$ the mean value for the predictor variable, $\bar{y}$ the mean value for the response variable, and the summations are $i = 1$ to n . Equations (2.5) and (2.6) are therefore the least-squares estimates for β_0 and β_1 , the values that minimize the sum of squared errors.



Appendix 2: Logistic regression

Given: $p = \frac{e^{b_0 + b_1 x}}{1 + e^{b_0 + b_1 x}}$

Show that: $b_0 + b_1 x = \ln \frac{p}{1-p}$

Solution:

Let $b_0 + b_1 x = a$

Re-formulate problem:

Given: $p = \frac{e^a}{1 + e^a}$

$$\Rightarrow \ln p = \underbrace{\ln e^a}_a - \ln(1 + e^a) =$$

$$\Rightarrow \boxed{\ln p = \frac{a + \ln(1 + e^a)^{-1}}{1}} \quad (1)$$

Show that:

$$\ln \frac{p}{1-p} = a$$

$$\Rightarrow \ln p - \ln(1-p) = a$$

$$\Rightarrow \boxed{\ln p = a + \ln(1-p)} \quad (2)$$

Compare (1) & (2) - we need to show that:

$$\ln(1-p) = \ln(1 + e^a)^{-1}$$

$$= \ln\left(\frac{1 - e^a}{1 + e^a}\right) = \ln\left(\frac{1 + e^a - e^a}{1 + e^a}\right) = \ln \frac{1}{1 + e^a} = \underbrace{\ln(1 + e^a)^{-1}}_{\text{Done!}} \quad \square$$

- For interested students; not examinable