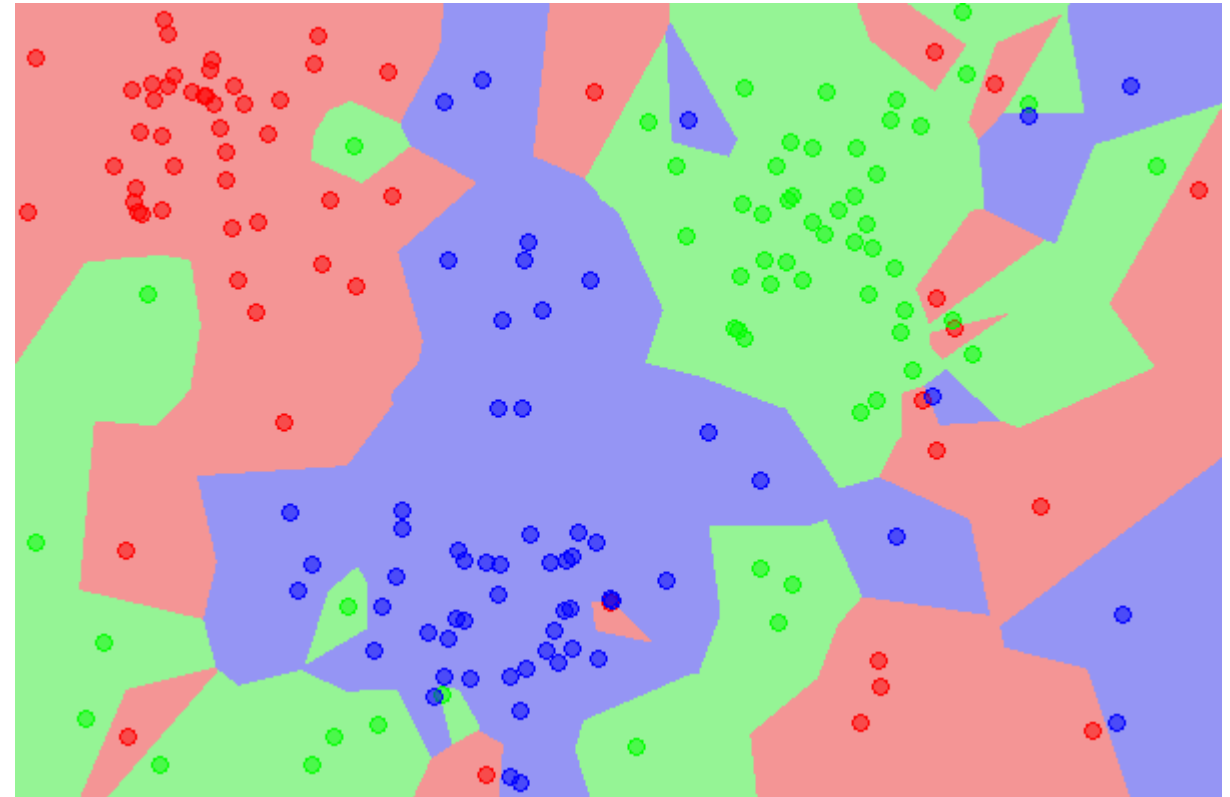




Intelligent Systems, Edition 2024

Contents

- Bias and variance of prediction models
- Bayes optimal classifier
- Simple regression models:
 - linear models, nearest neighbor, regression trees, regression rules
- Simple classification models:
 - nearest neighbor, naïve Bayes, decision trees, decision rules, logistic regression
- Biases in data

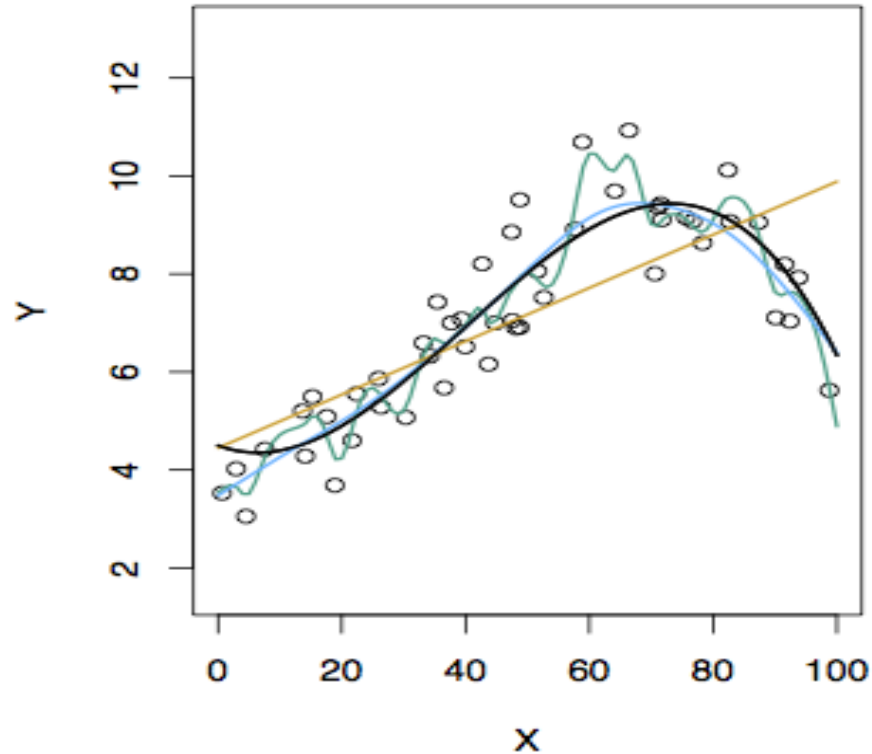
A generalization problem

- Our ML methods have generally been designed to make error small on the training data, e.g., with linear regression, we choose the line such that MSE is minimized.
- What we really care about is how well the method generalizes i.e. how well it works on new data. We call this new data “**Test data**”.
- There is no guarantee that the method with the smallest training error will have the smallest test (i.e. new data) error.
- One approach to address the problem is to reserve a portion of the training data to measure the generalization error, we this dataset “evaluation dataset”. We use this dataset during training to guide the learning process, e.g., to stop it. This approach is used especially with overparametrized methods such as neural networks and in learning settings with high likelihood of overfitting such as AutoML methods.

Training vs. test error

- In general the more flexible a method is the lower its training MSE will be, i.e. it will “fit” or explain the training data very well.
 - More flexible methods (such as splines) can generate a wider range of possible shapes to estimate f as compared to less flexible and more restrictive methods (such as linear regression). The less flexible the method, the easier to interpret the model. Thus, there is a trade-off between flexibility and model interpretability.
- However, the test MSE may in fact be higher for a more flexible method than for a simple approach like linear regression.

Different levels of flexibility: example 1



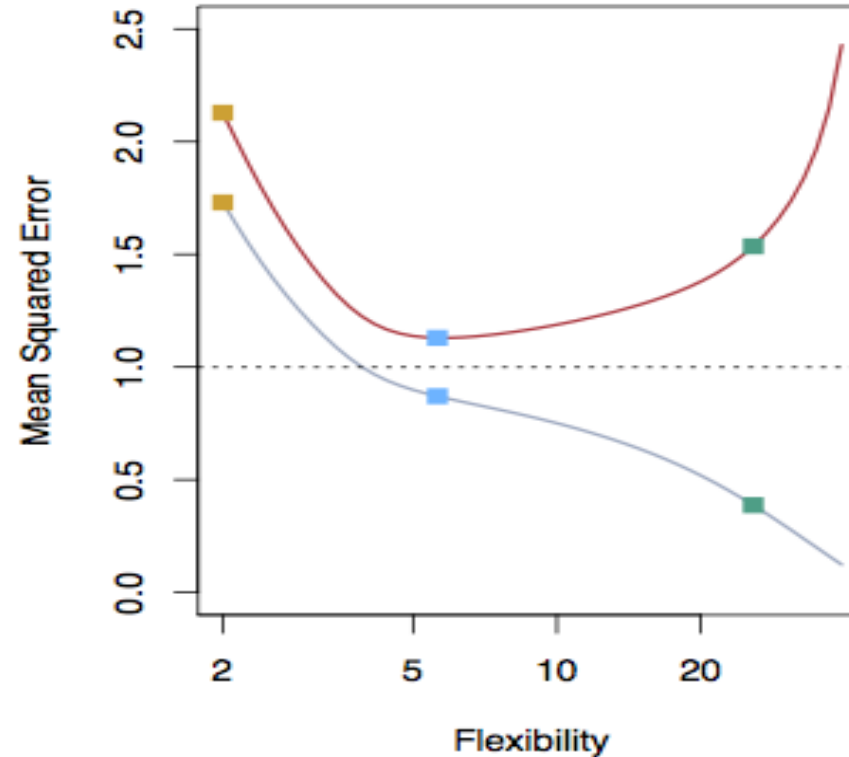
LEFT

Black: Truth

Orange: Linear Estimate

Blue: smoothing spline

Green: smoothing spline (more flexible)



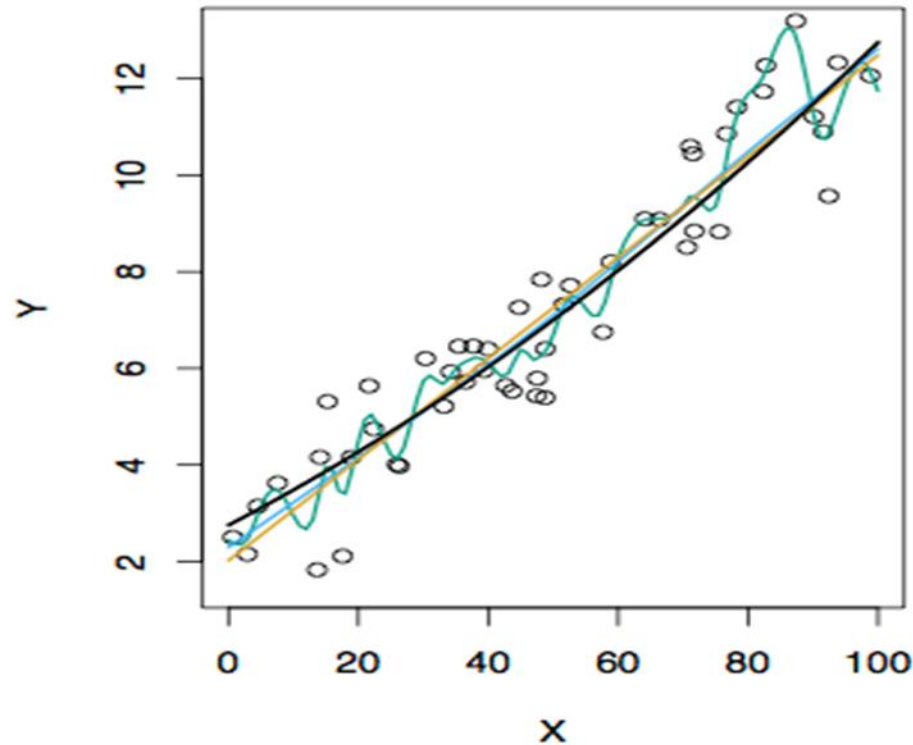
RIGHT

RED: Test MS

Grey: Training MSE

Dashed: Minimum possible test MSE (irreducible error)

Different levels of flexibility: example 2



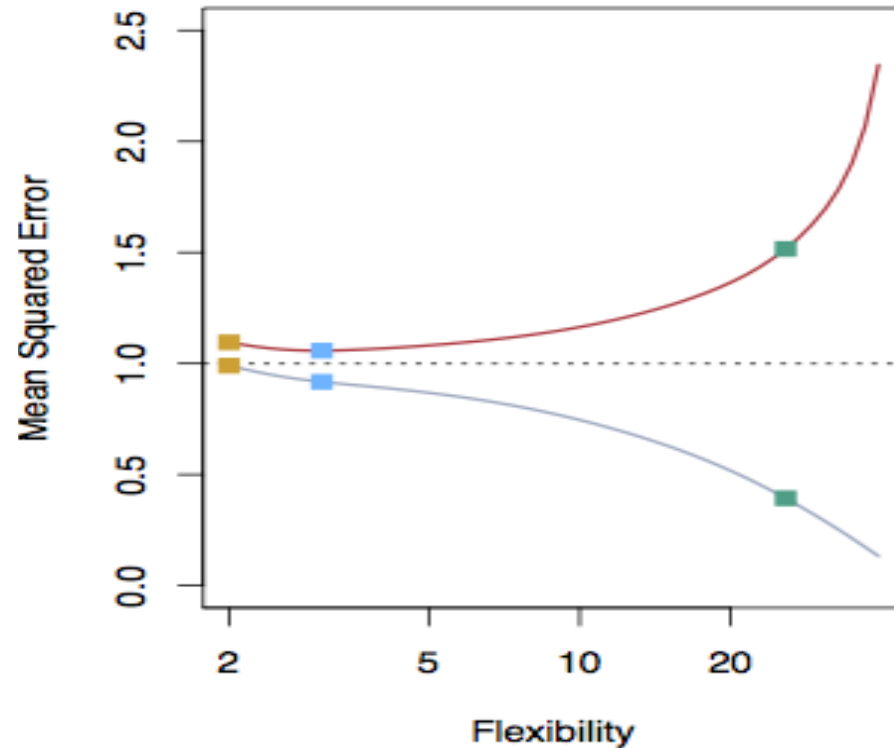
LEFT

Black: Truth

Orange: Linear Estimate

Blue: smoothing spline

Green: smoothing spline (more flexible)



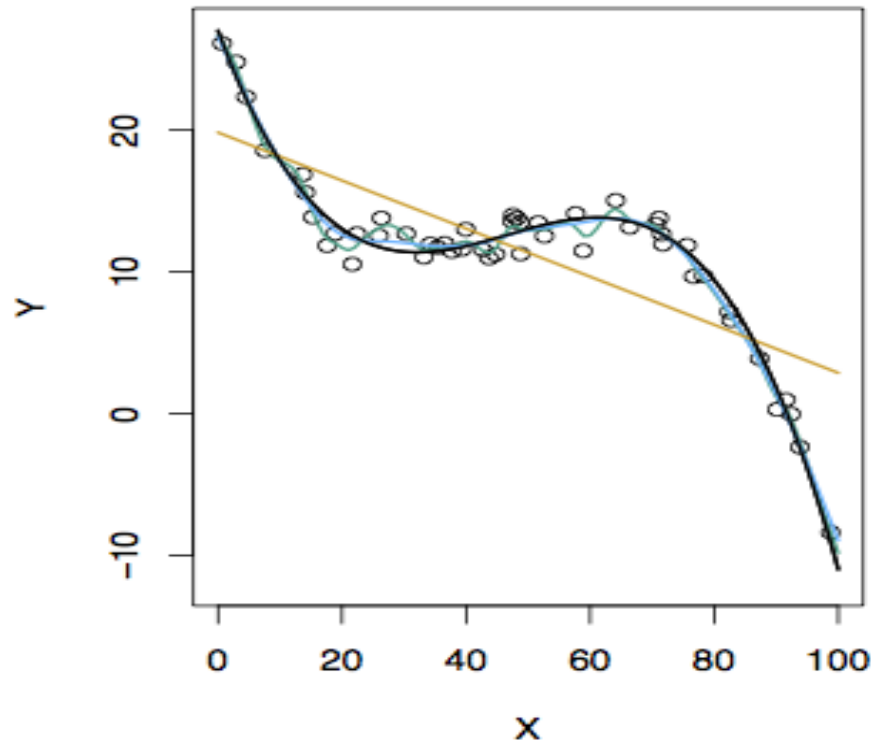
RIGHT

RED: Test MSE

Grey: Training MSE

Dashed: Minimum possible test MSE (irreducible error)

Different levels of flexibility: example 3



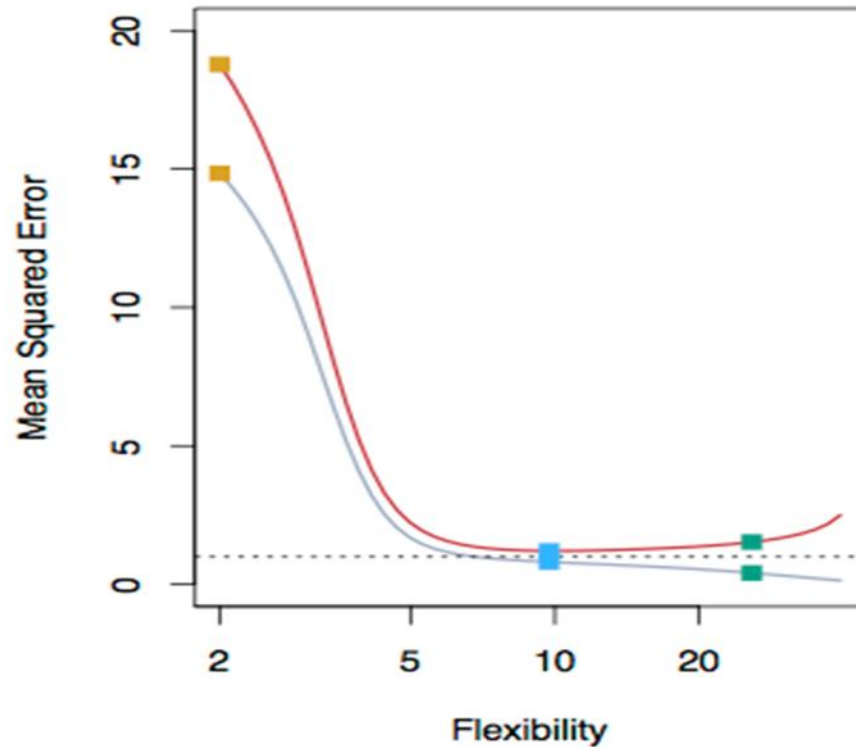
LEFT

Black: Truth

Orange: Linear Estimate

Blue: smoothing spline

Green: smoothing spline (more flexible)



RIGHT

RED: Test MSE

Grey: Training MSE

Dashed: Minimum possible test MSE
(irreducible error)

Bias - variance trade-off

- The previous graphs of test versus training MSE's illustrates a very important trade-off that governs the choice of statistical learning methods.
- There are always two competing forces that govern the choice of learning method, i.e. bias and variance.

Bias of learning methods

- Bias in general: inclination or prejudice for or against one person or group, especially in a way considered to be unfair.
- Bias in ML refers to the error that is introduced by modeling a real-life problem (that is usually extremely complicated) by a much simpler problem.
- A common definition of bias:

$$\text{Bias} = E[Y] - f(x)$$

- For example, linear regression assumes that there is a linear relationship between Y and X. It is unlikely that, in real life, the relationship is exactly linear so some bias will be present.
- The more flexible/complex a method is the less bias it will generally have.

Variance of learning methods

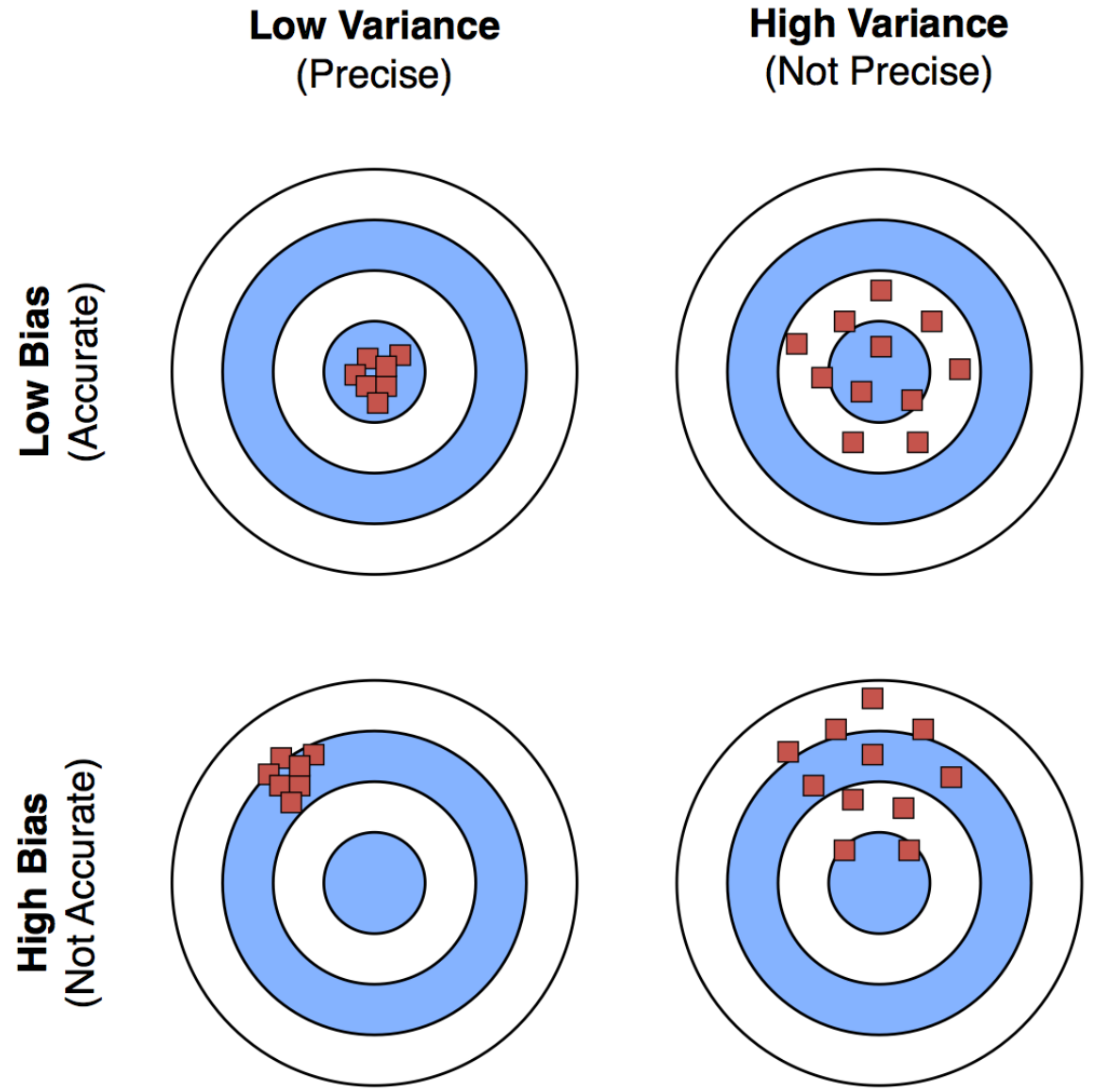
- Variance refers to how much your estimate for f would change if you had a different training data set.

- A common definition of variance:

$$\text{Var} = E[(Y - E[Y])^2]$$

- Generally, the more flexible a method is, the more variance it has.

Bias-variance illustration



The trade-off?

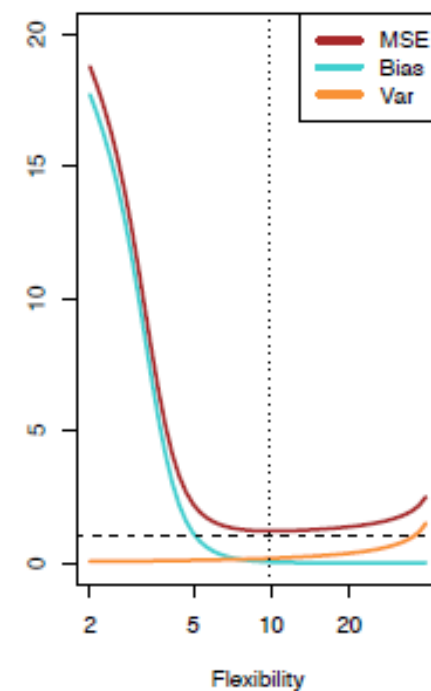
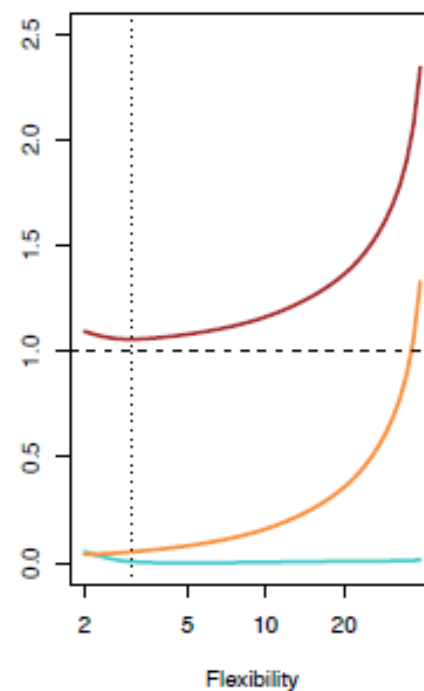
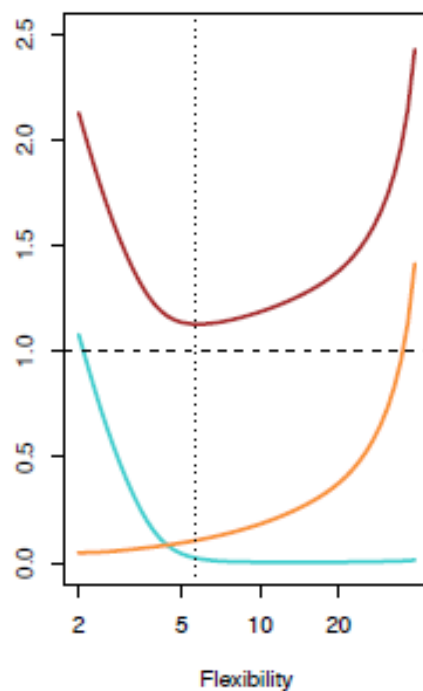
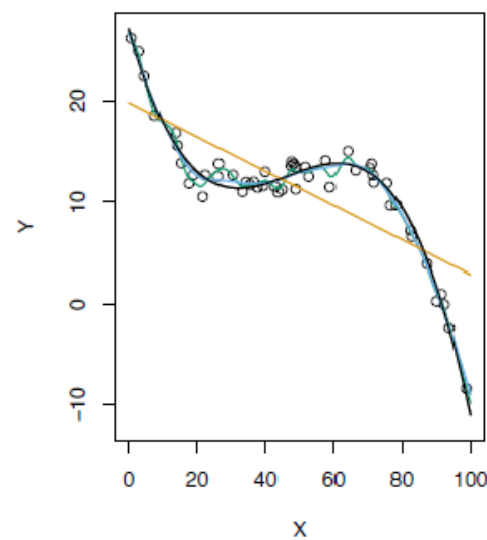
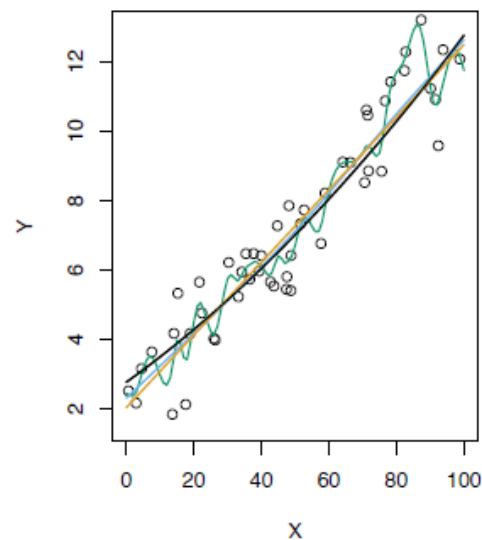
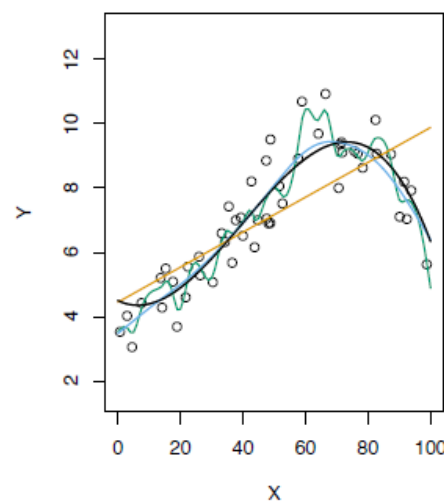
- It can be shown that for any given, $X=x_0$, the expected test MSE for a new Y at x_0 will be equal to

$$\text{Expected Test MSE} = E(Y - f(x_0))^2 = \text{Bias}^2 + \text{Var} + \underbrace{\sigma^2}_{\text{Irreducible Error}}$$

where $\text{Bias} = E[Y] - f(x)$ and $\text{Var} = E[(Y - E[Y])^2]$

- What this means is that as a method gets more complex
 - the bias will decrease and
 - the variance will likely increase
 - but expected test MSE may go up or down!
- The trade-off is only present if we assume fixed error!
- For some models there may be no trade-off!

Test MSE, bias and variance



Bayes classifier

- In classification, the optimal classification for an instance (x_0, y_0) can be obtained by selecting the class j which maximizes the probability

$$P(Y = j | X = x_0)$$

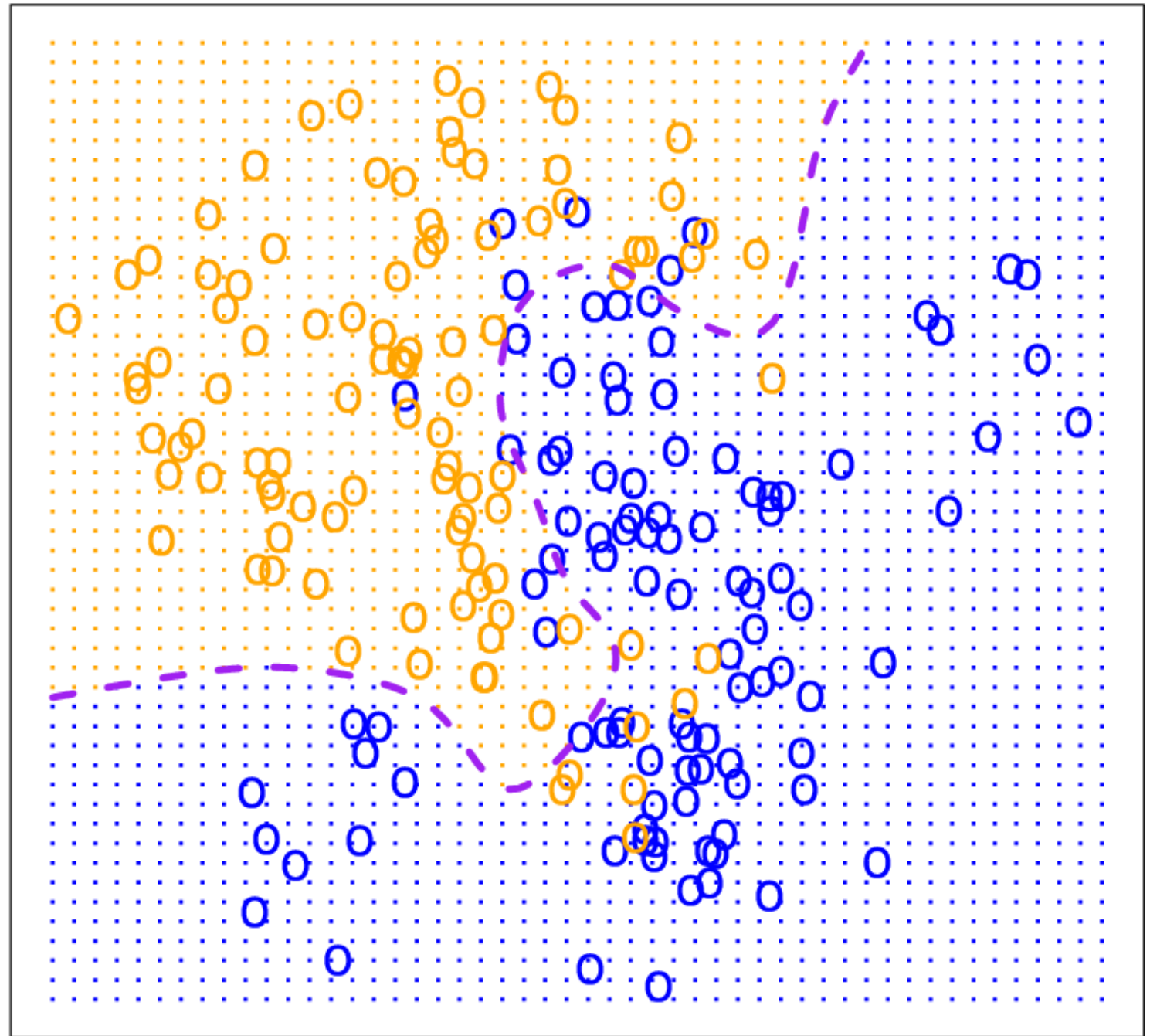
- This classifier is called the Bayes (optimal) classifier
- It implies that learning is actually an estimation of the conditional data distribution

Bayes error rate

- The Bayes error rate refers to the lowest possible error rate that could be achieved if somehow we knew exactly what the “true” probability distribution of the data looked like.
- On test data, no classifier (or statistical learning method) can get lower error rates than the Bayes error rate.
- Of course, in real life problems, the Bayes error rate can't be calculated exactly (why not?) but it is useful to think about it

Bayes optimal classifier

- for new x_0 returns the maximally probable prediction value $P(Y=y \mid X=x_0)$
- this means to select the class j with $\arg \max_j P(Y=y_j \mid X=x_0)$
- Why is this probability not easy to estimate?
- The dotted line show the Bayes decision boundary, where $P(Y=y \mid X=x_0) = 0.5$



Bayes classifier approximations

- Two models can be viewed as directly approximating the Bayes classifier $P(Y = j | X = x_0)$
- Naive Bayesian classifier
 - uses Bayesian formula to get inverse conditional probabilities
 - assuming conditional independence between features

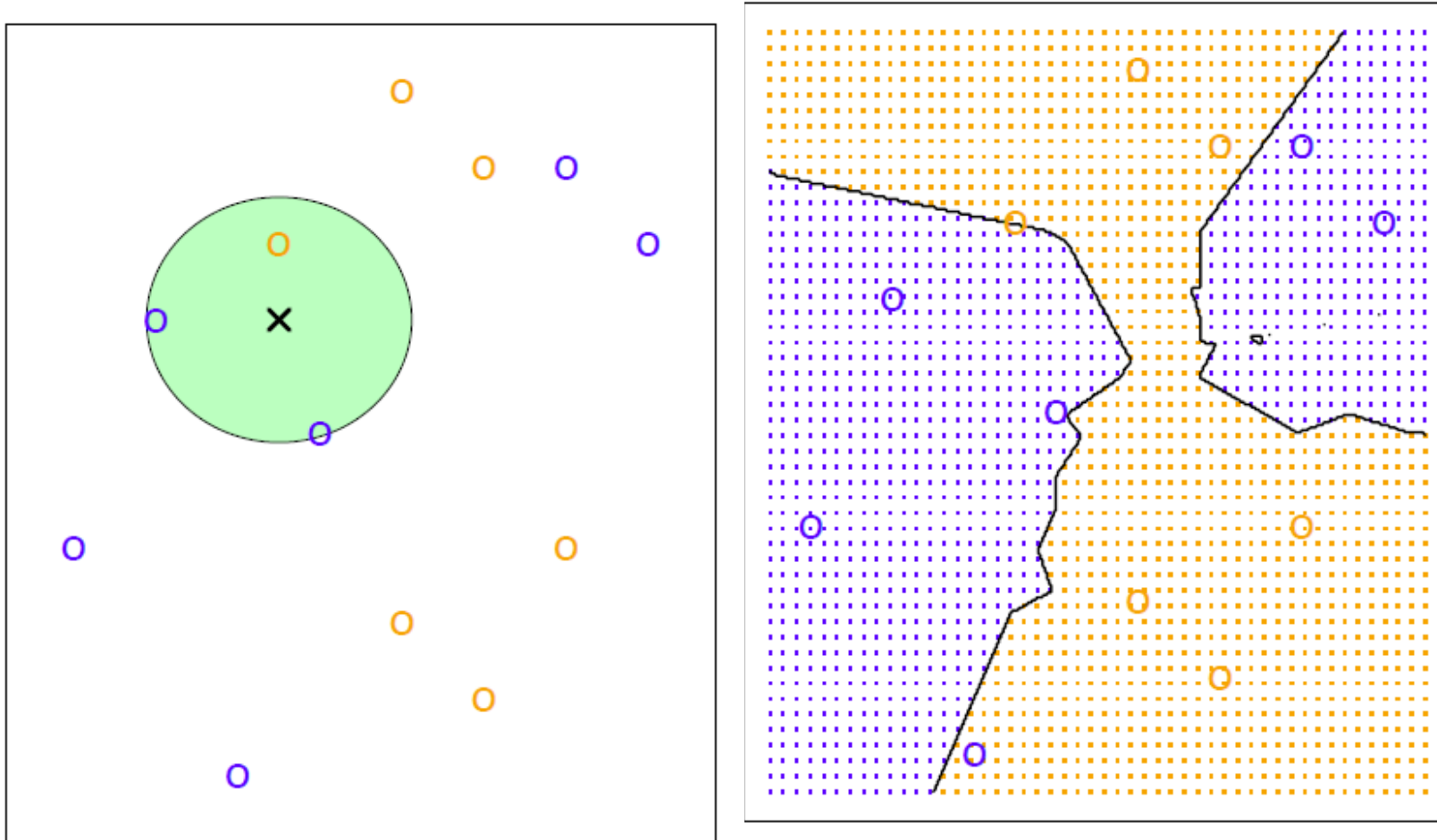
$$P(C|X_1X_2 \dots X_n) = \frac{P(C) \cdot P(X_1X_2 \dots X_n|C)}{P(X_1X_2 \dots X_n)} \approx \frac{P(C) \cdot \prod_i P(X_i|C)}{\prod_i P(X_i)}$$

- Nearest neighbour classifier
 - directly estimates the conditional probability using instances near to x_0

K-Nearest Neighbors (KNN)

- k Nearest Neighbors is a flexible approach to estimate the Bayes classifier.
- For any given x we find the k closest neighbors to x in the training data, and examine their corresponding y .
- If the majority of the y 's are orange, we predict the orange label otherwise the blue label.
- The smaller that k is the more flexible the method will be.

KNN example with $k = 3$



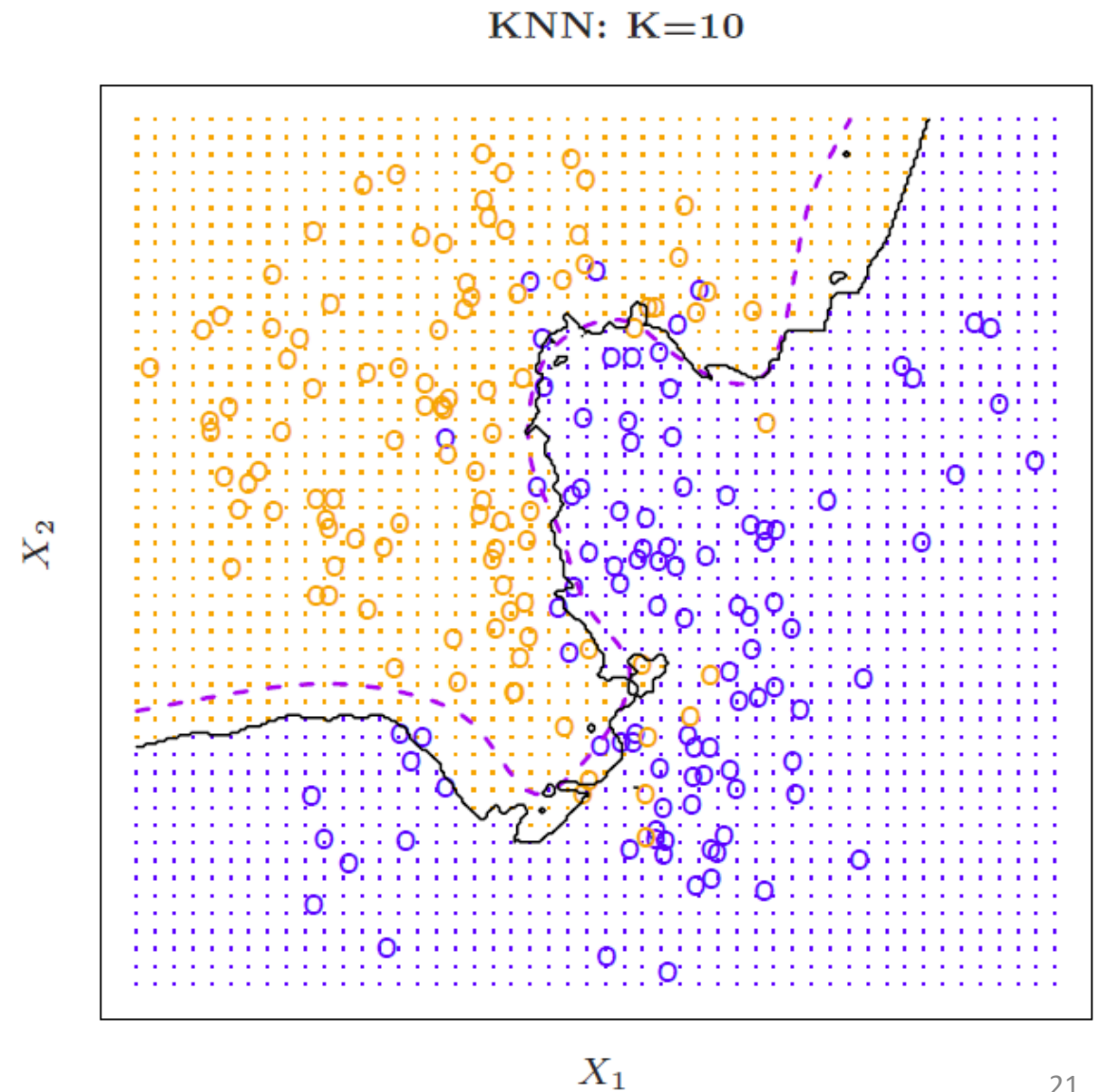
K-NN classifier

- Given a positive integer K and a test observation x_0 , the KNN classifier first identifies the K points in the training data that are closest to x_0 , represented by the set $\mathcal{N}_0$.
- It then estimates the conditional probability for class j as the fraction of points in $\mathcal{N}_0$ whose response values equal j :

$$\Pr(Y = j | X = x_0) = \frac{1}{K} \sum_{i \in \mathcal{N}_0} I(y_i = j).$$

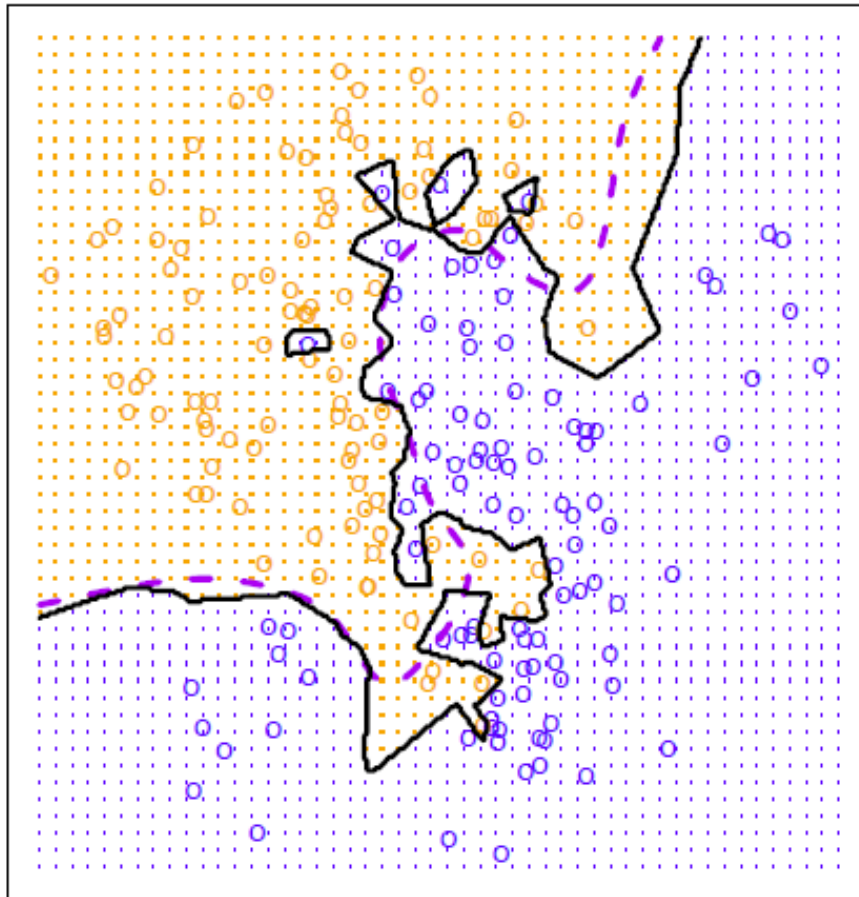
- applies Bayes rule and classifies the test observation x_0 to the class with the largest probability.

Simulated data:
 $K = 10$

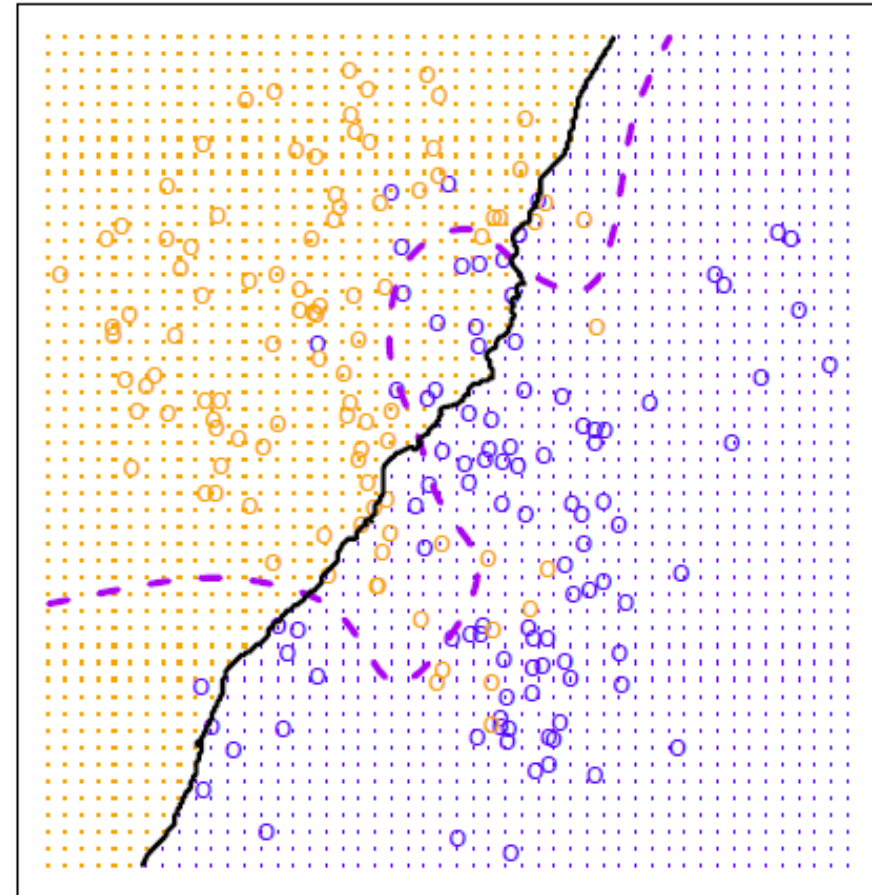


$K = 1$ and $K = 100$

KNN: $K=1$

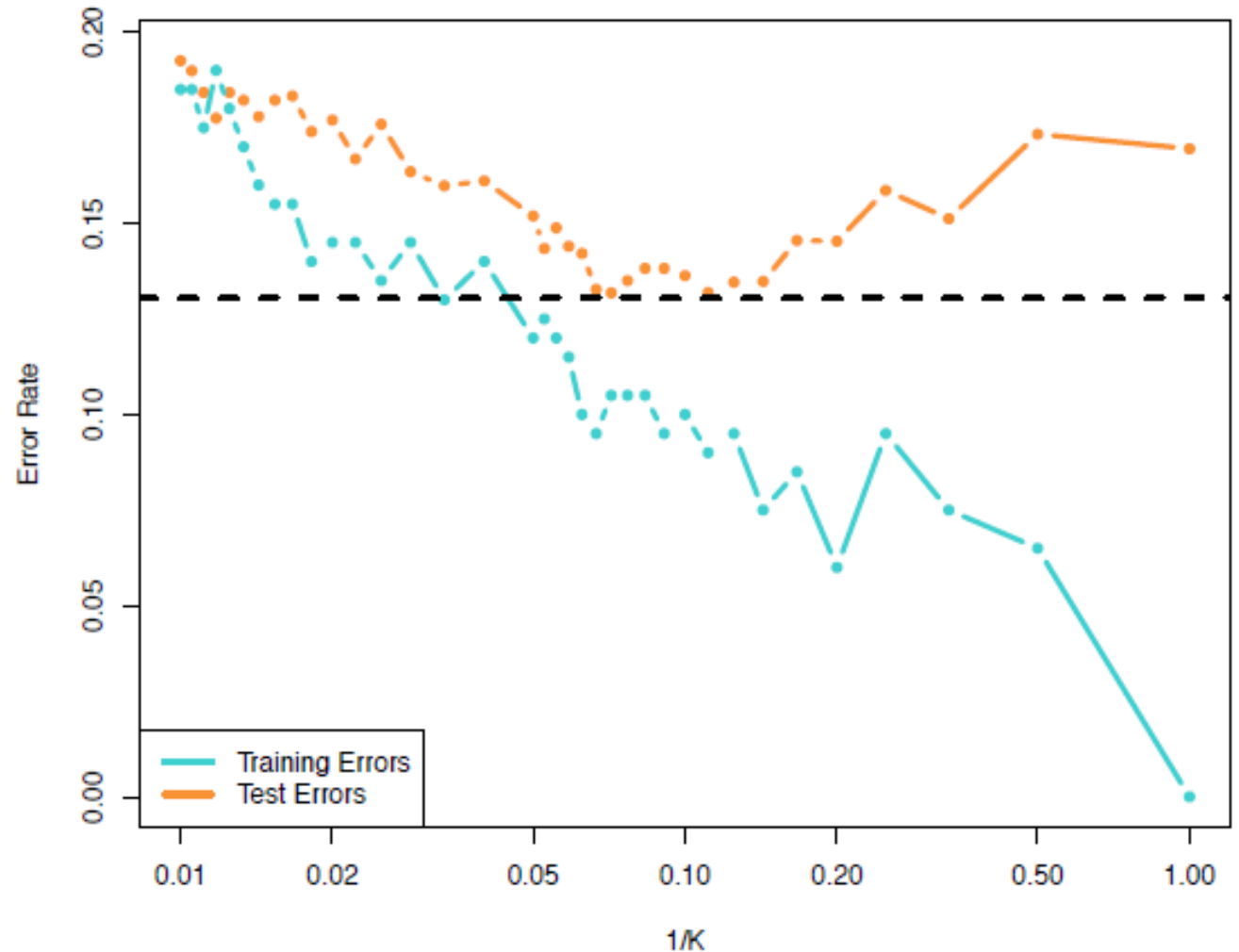


KNN: $K=100$



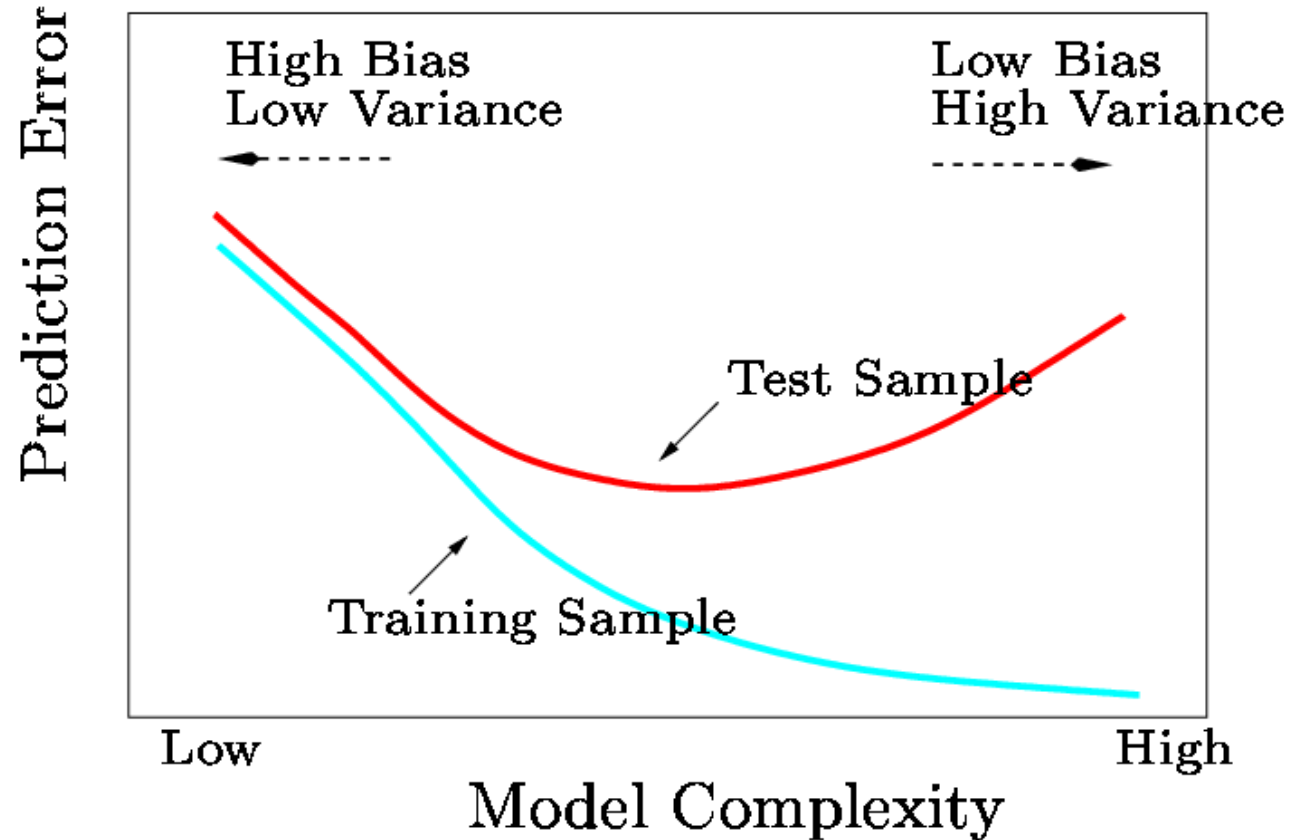
Training vs. test error rates on the simulated data

- Notice that training error rates keep going down as k decreases or equivalently as the flexibility increases.
- However, the test error rate at first decreases but then starts to increase again.



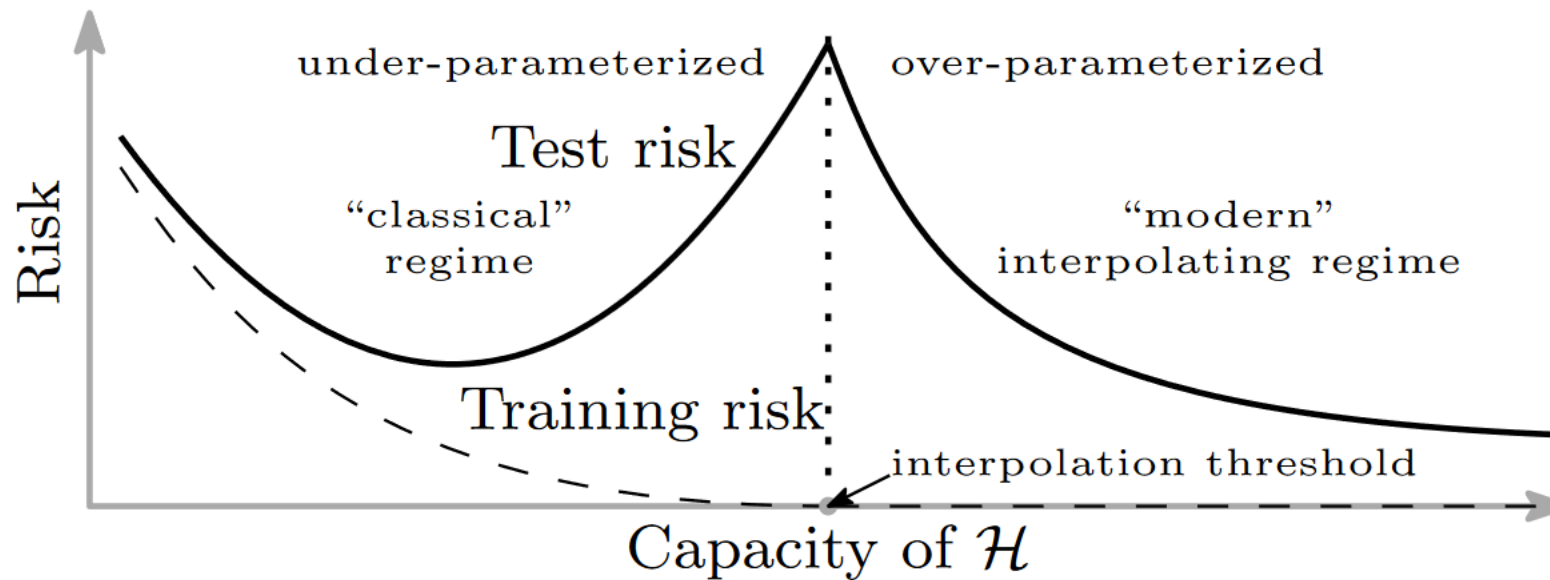
A fundamental picture

- In general training errors will always decline.
- However, test errors will decline at first (as reductions in bias dominate) but will then start to increase again (as increases in variance dominate).
- This is a conventional wisdom, but it is not true for all methods and all training regimes.

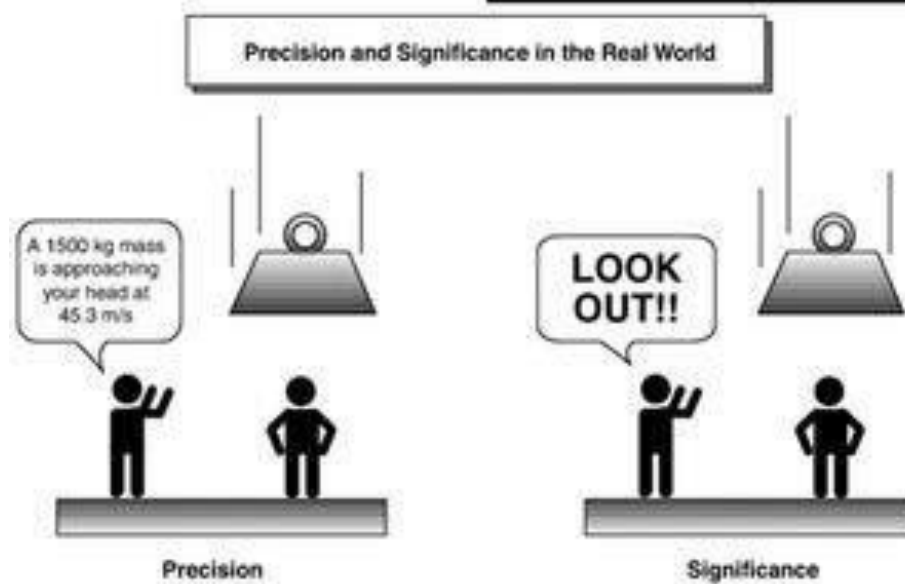


The double descent curve

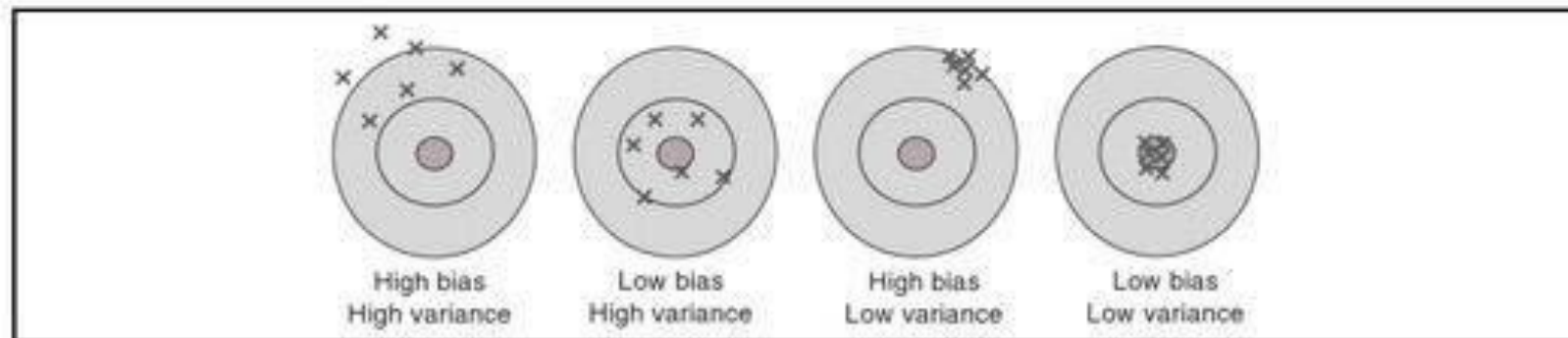
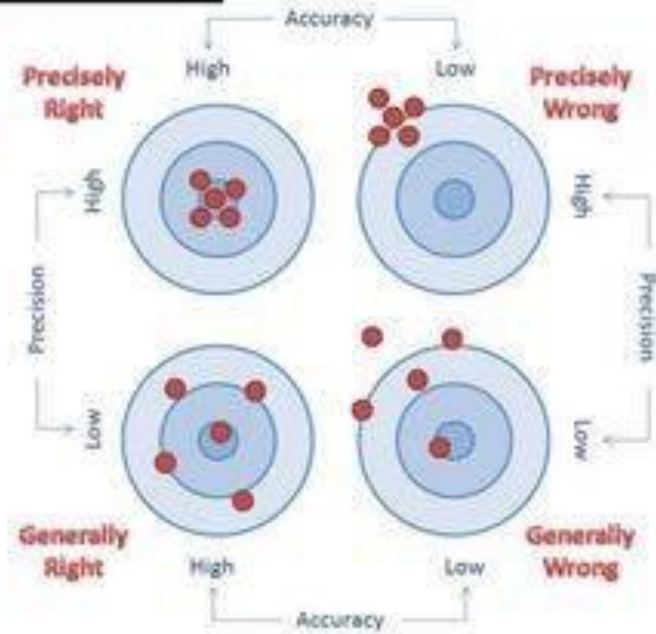
- While for some models, like kNN, there seem to be a trade-off between bias and variance, this is not a universal phenomenon
- E.g., overparametrization in neural networks produce double descent curve (similar evidence for random forests)



Precision vs. Significance Accuracy vs. Precision Bias vs. Variance



**Accuracy
and
Precision**

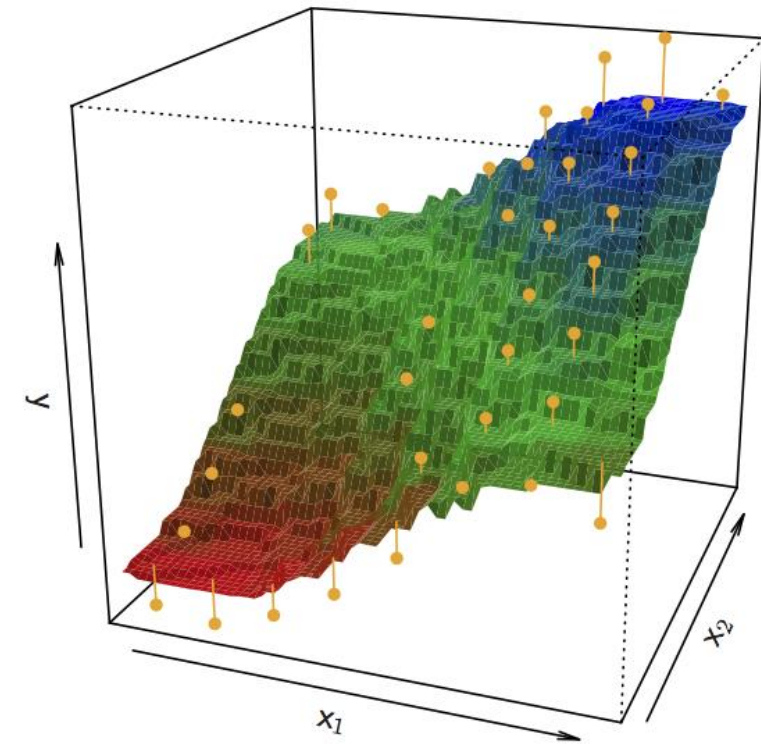
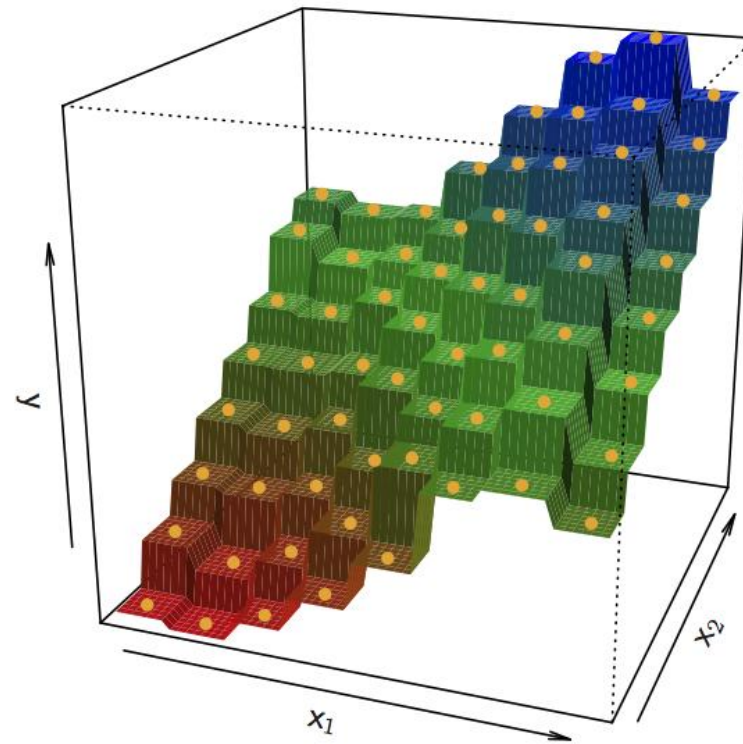


K-nearest neighbor for regression

- kNN regression is similar to the kNN classifier.
- given a set of instances (x_i, y_i)
- To predict y for a given value of x , consider k closest points to x in training data $N_k(x)$ and take the average of the responses. i.e.

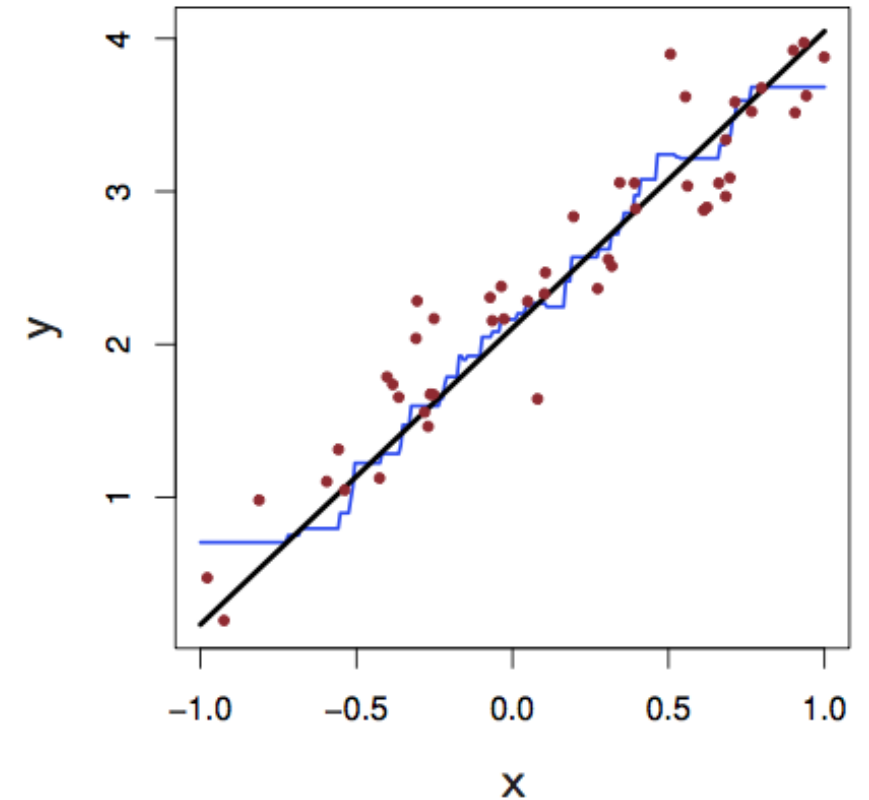
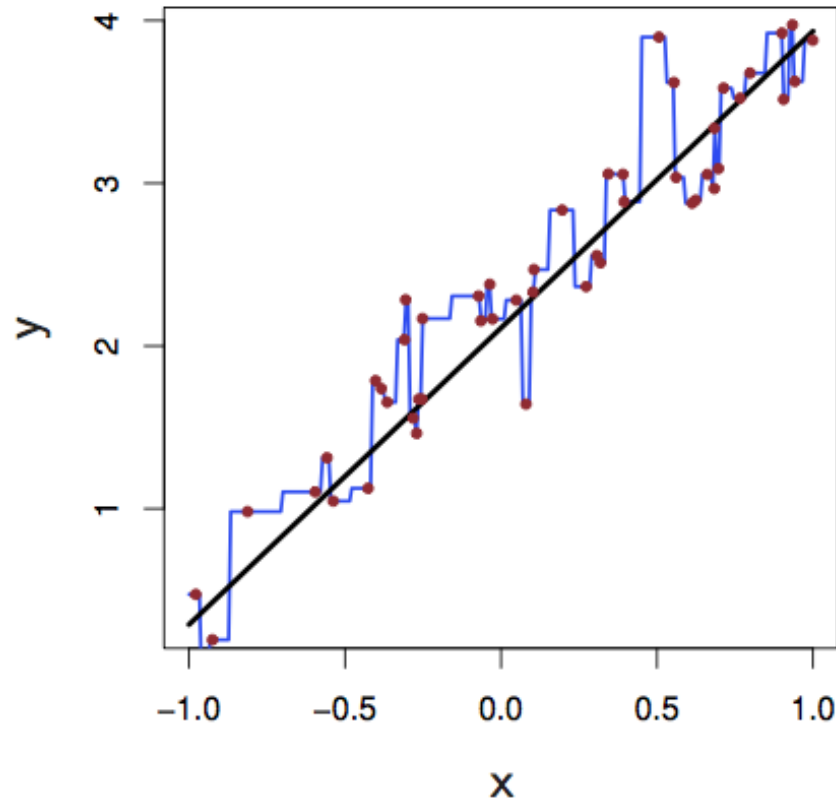
$$f(x) = \frac{1}{k} \sum_{x_i \in N_k(x)} y_i$$

KNN Fits for $k=1$ and $k=9$



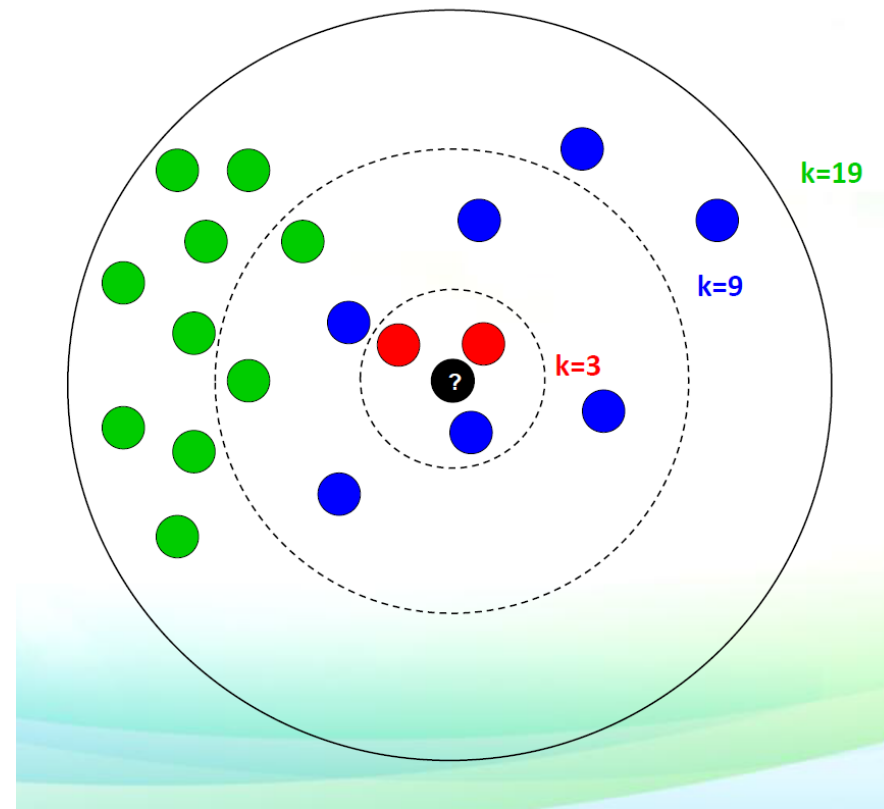
KNN fits in one dimension ($k=1$ and $k=9$)

- black line: actual function,
- blue line: regressional kNN



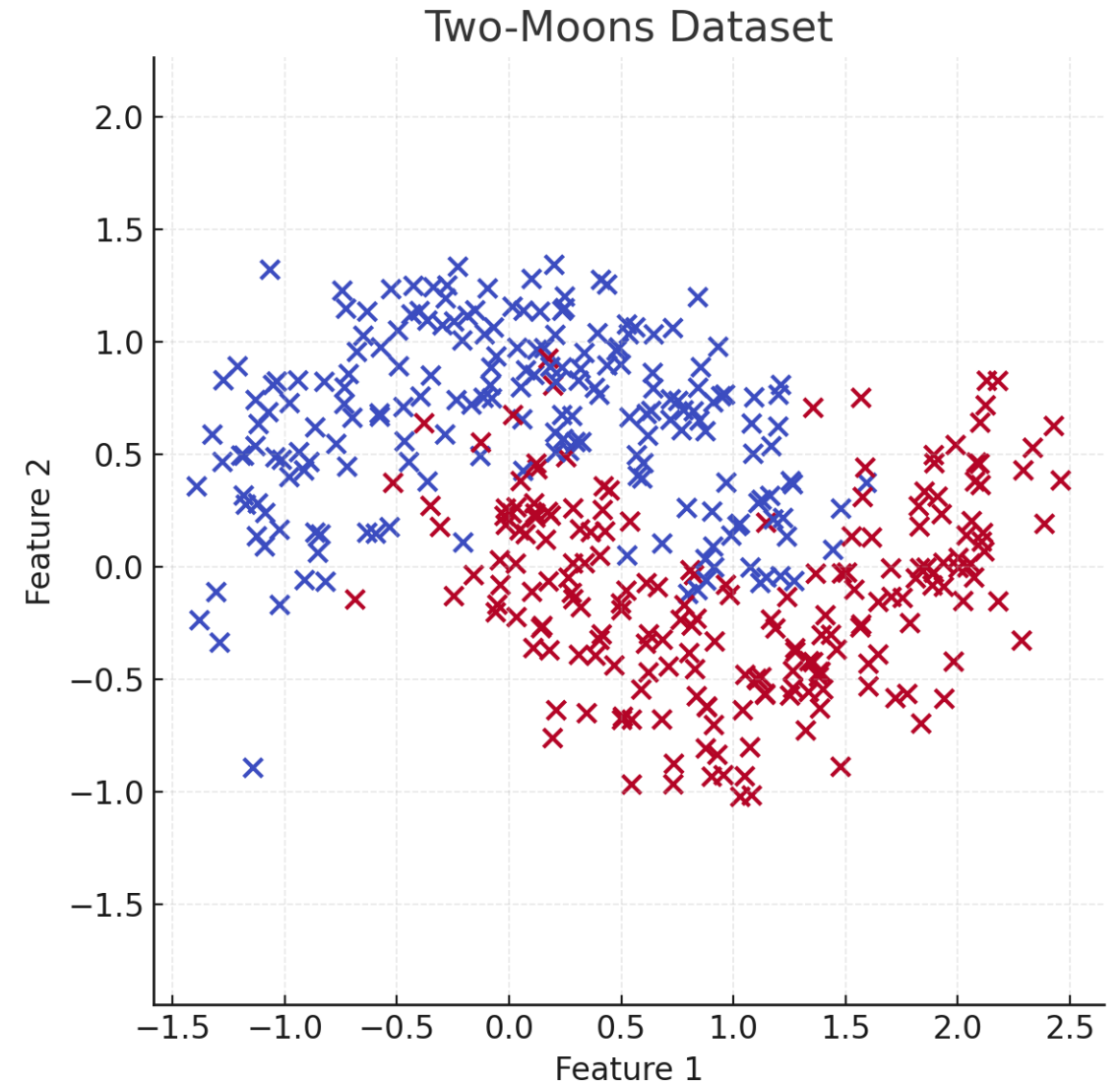
Choice of k in KNN

- If k is small, kNN is much more flexible than linear regression.
- Is that better?
- The results may be highly dependent on the choice of k .

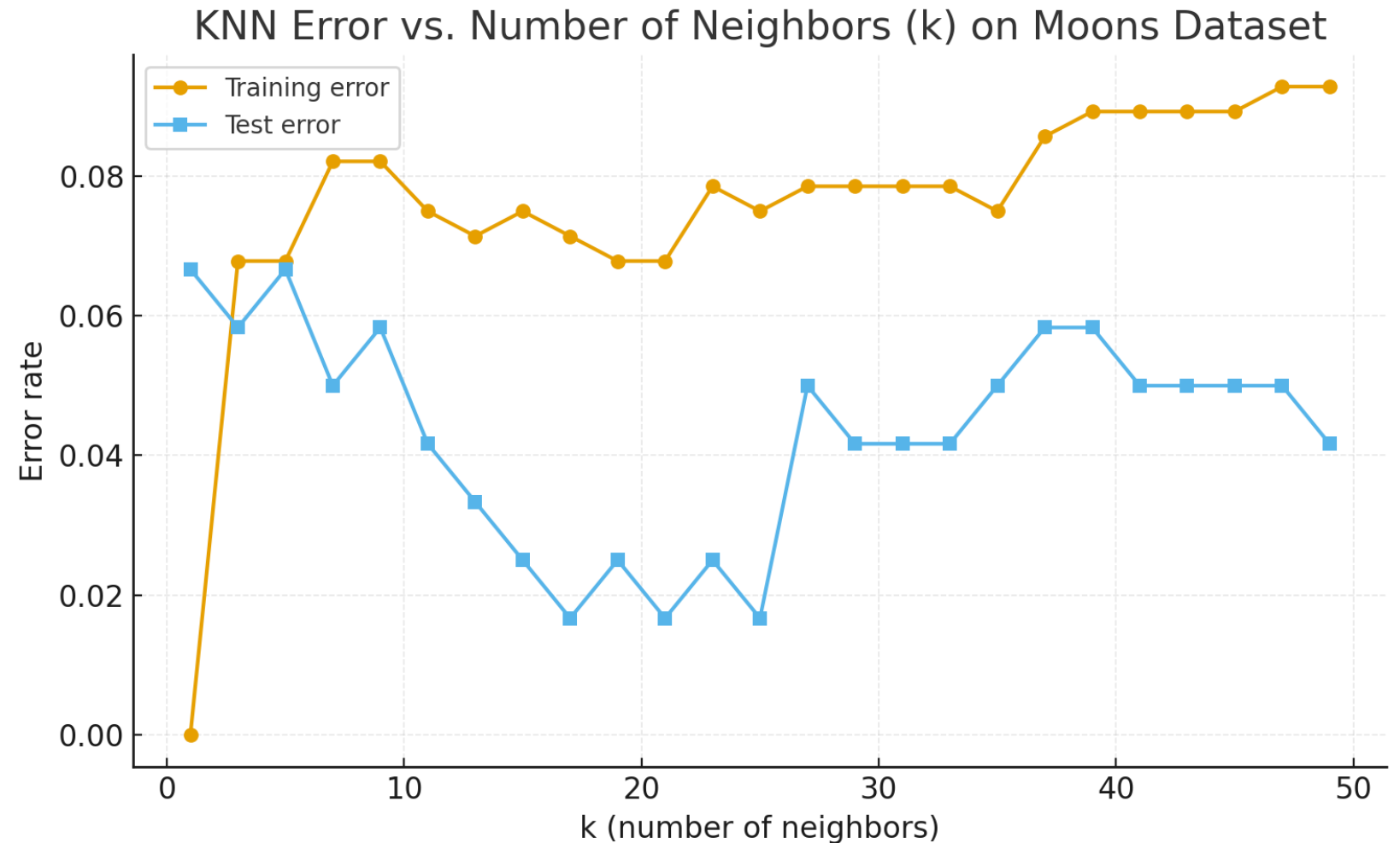


Example: two moons dataset

- Non-linear decision boundary



KNN is not so good in high dimensional situations



Speeding up KNN algorithm

- precondition: normalization of dimensions, e.g., to $[0, 1]$
- naive search for nearest neighbors: $O(n \cdot d \cdot t)$
 - n is number of instances
 - d is number of dimensions
 - t is number of nearest neighbors
- exact search for low dimensional spaces
 - k-d trees (d is around 10)
 - quad-trees ($d=2$), octrees ($d=3$)
 - R-tree (rectangular tree, also R^+ , R^* , ...), $d=2$ or 3
- approximate search
 - RKD-tree (random k-d tree)
 - locally sensitive hashing (LSH),
 - hierarchical k-means

Word bias have several meanings

1. General / Everyday Meaning

- **Preference or prejudice** toward or against something or someone, often in an unfair way.
 - *Example:* “The judge must avoid bias in court.”
 - → *Synonyms:* partiality, favoritism, prejudice.
- **Tendency** to lean in a certain direction, consciously or unconsciously.
 - *Example:* “She has a bias toward traditional art styles.”

2. Social and Psychological Context

- **Cognitive bias:** A systematic pattern of deviation from rational judgment or objective standards.
 - *Example:* *confirmation bias* (favoring information that confirms one’s beliefs).
- **Social bias:** Prejudice or discrimination based on group characteristics (e.g., gender, race, culture).
 - *Example:* “Hiring practices should be checked for gender bias.”

3. Statistics & Machine Learning

- **Statistical bias:** Systematic error that leads an estimator or model to deviate from the true value.
 - *Example:* “The sample mean is an unbiased estimator of the population mean.”
- **Bias–variance tradeoff:** In machine learning, bias refers to the error introduced by simplifying assumptions in a model.
 - High bias → model underfits the data.
 - Low bias → model can better capture complexity.

Ethical consideration of social and cognitive bias in ML models

- The word bias is ambiguous even within ML
- bias in models:
 - characteristic of models,
 - affects error,
 - unlikely to be ethically problematic
 - when it can be problematic?
- bias in data:
 - data unrepresentative of the true population,
 - might be ethically problematic



Biases in the data

- Machine learning models are not inherently objective. Engineers train models by feeding them a data set of training examples, and human involvement in the provision and curation of this data can make a model's predictions susceptible to bias.
- When building models, it's important to be aware of common human biases that can manifest in your data, so you can take proactive steps to mitigate their effects.
- The biases listed next provide just a small selection of biases that are often uncovered in machine learning data sets; this list *is not intended to be exhaustive*. Wikipedia's [catalog of cognitive biases](#) enumerates over 100 different types of human bias that can affect our judgment. When auditing your data, you should be on the lookout for any and all potential sources of bias that might skew your model's predictions.

Reporting bias

- **Reporting bias** occurs when the frequency of events, properties, and/or outcomes captured in a data set does not accurately reflect their real-world frequency. This bias can arise because people tend to focus on documenting circumstances that are unusual or especially memorable, assuming that the ordinary can "go without saying."
 - **EXAMPLE:** A sentiment-analysis model is trained to predict whether book reviews are positive or negative based on a corpus of user submissions to a popular website. The majority of reviews in the training data set reflect extreme opinions (reviewers who either loved or hated a book), because people were less likely to submit a review of a book if they did not respond to it strongly. As a result, the model is less able to correctly predict sentiment of reviews that use more subtle language to describe a book.

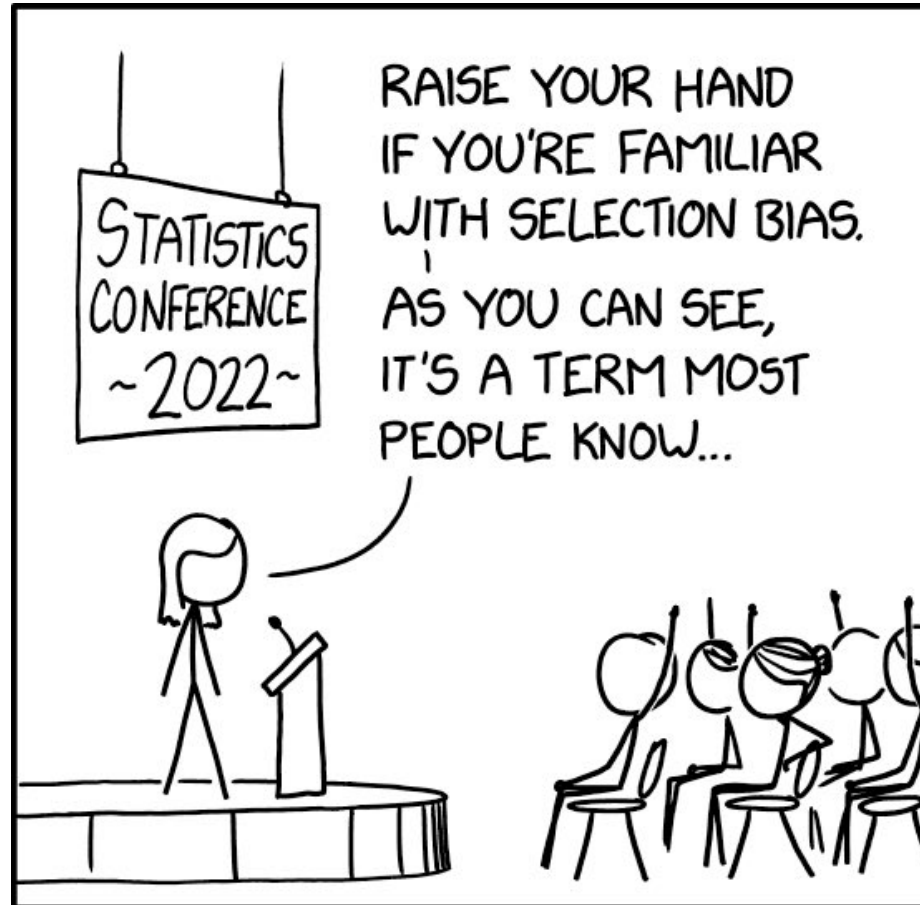
Automation bias

- **Automation bias** is a tendency to favor results generated by automated systems over those generated by non-automated systems, irrespective of the error rates of each (we also have the opposite bias)
 - **EXAMPLE:** Software engineers working for a sprocket manufacturer were eager to deploy the new "groundbreaking" model they trained to identify tooth defects, until the factory supervisor pointed out that the model's precision and recall rates were both 15% lower than those of human inspectors.



Selection bias

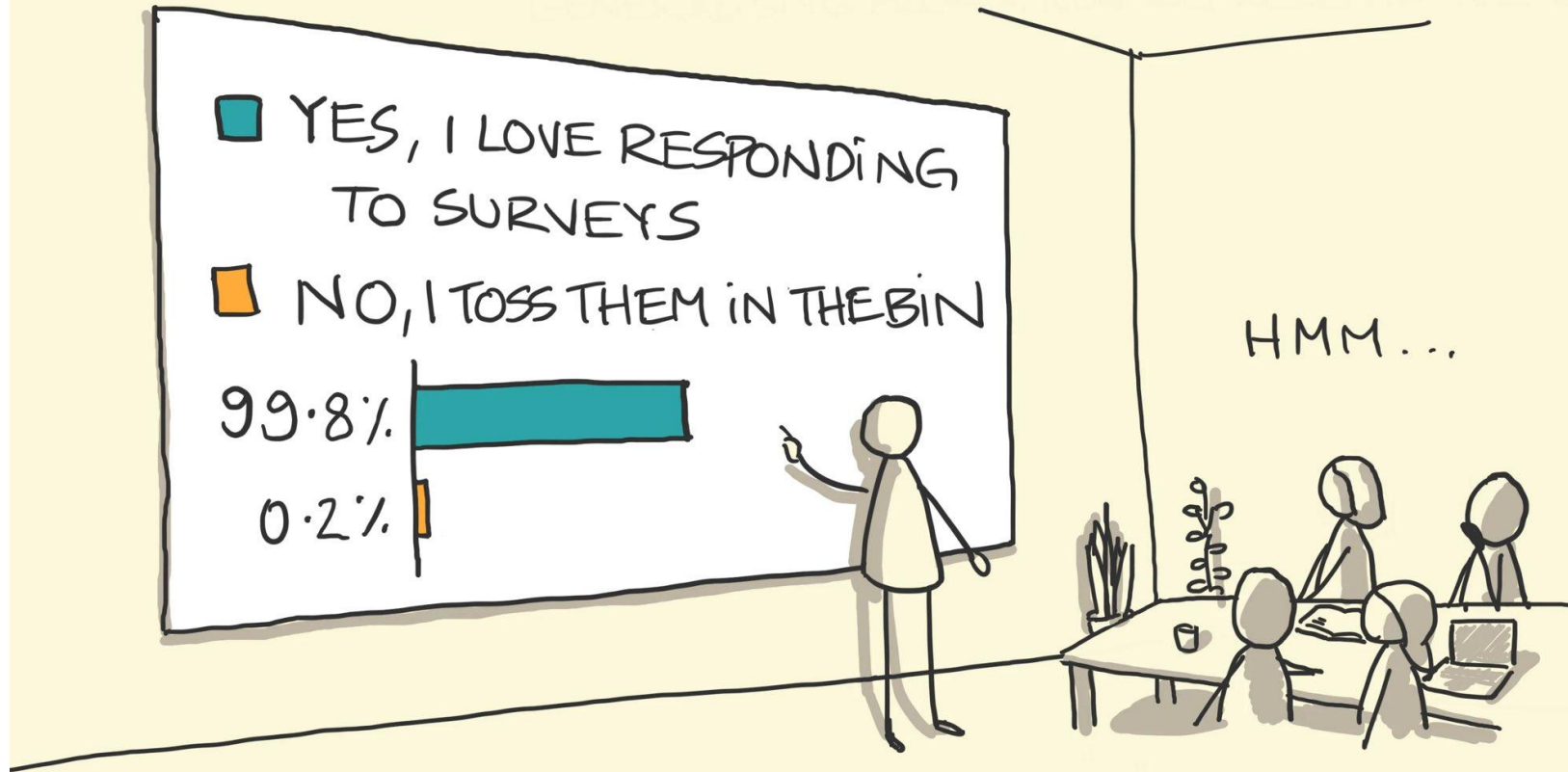
- **Selection bias** occurs if a data set's examples are chosen in a way that is not reflective of their real-world distribution.



Selection bias variants

- Selection bias can take many different forms:
 - **Coverage bias:** Data is not selected in a representative fashion.
 - **EXAMPLE:** A model is trained to predict future sales of a new product based on phone surveys conducted with a sample of consumers who bought the product. Consumers who instead opted to buy a competing product were not surveyed, and as a result, this group of people was not represented in the training data.
 - **Non-response bias (or participation bias):** Data ends up being unrepresentative due to participation gaps in the data-collection process.
 - **EXAMPLE:** A model is trained to predict future sales of a new product based on phone surveys conducted with a sample of consumers who bought the product and with a sample of consumers who bought a competing product. Consumers who bought the competing product were 80% more likely to refuse to complete the survey, and their data was underrepresented in the sample.
 - **Sampling bias:** Proper randomization is not used during data collection.
 - **EXAMPLE:** A model is trained to predict future sales of a new product based on phone surveys conducted with a sample of consumers who bought the product and with a sample of consumers who bought a competing product. Instead of randomly targeting consumers, the surveyor chose the first 200 consumers that responded to an email, who might have been more enthusiastic about the product than average purchasers.

SAMPLING BIAS



" WE RECEIVED 500 RESPONSES AND
FOUND THAT PEOPLE LOVE RESPONDING
TO SURVEYS "

sketchplanations

Group attribution bias

- **Group attribution bias** is a tendency to generalize what is true of individuals to an entire group to which they belong. Two key manifestations of this bias are:
 - **In-group bias:** A preference for members of a group to which *you also belong*, or for characteristics that you also share.
 - **EXAMPLE:** Two engineers training a resume-screening model for software developers are predisposed to believe that applicants who attended the same computer-science academy as they both did are more qualified for the role.
 - **Out-group homogeneity bias:** A tendency to stereotype individual members of a group to which *you do not belong*, or to see their characteristics as more uniform.
 - **EXAMPLE:** Two engineers training a resume-screening model for software developers are predisposed to believe that all applicants who did not attend a computer-science academy do not have sufficient expertise for the role.

Implicit bias

- **Implicit bias** occurs when assumptions are made based on one's own mental models and personal experiences that do not necessarily apply more generally. We are often not aware of these biases, and some may be contrary to our conscious beliefs.
 - **EXAMPLE:** An engineer training a gesture-recognition model uses a [head shake](#) as a feature to indicate a person is communicating the word “no.” However, in some regions of the world, a head shake actually signifies “yes.” A common form of implicit bias is **confirmation bias**, where model builders unconsciously process data in ways that affirm preexisting beliefs and hypotheses. In some cases, a model builder may actually keep training a model until it produces a result that aligns with their original hypothesis; this is called an **experimenter's bias**.
 - **EXAMPLE:** An engineer is building a model that predicts aggressiveness in dogs based on a variety of features (height, weight, breed, environment). The engineer had an unpleasant encounter with a hyperactive toy poodle as a child, and ever since has associated the breed with aggression. When the trained model predicted most toy poodles to be relatively docile, the engineer retrained the model several more times until it produced a result showing smaller poodles to be more violent.

Implicit Bias is...

Attitudes, Stereotypes, & Beliefs
that can affect how we treat others

based on categorizations such as...

Race



Ability



Gender



Culture



Language



Implicit bias runs contrary to our stated beliefs. We can say that we believe in equity (and truly believe it). But then unintentionally behave in ways that are biased and discriminatory.

Good: out-of-sample predictions. **Bias** and **variance** of a machine learning method on a specific dataset is a **core diagnostic**.

Here's a clear and practical breakdown of how to measure **bias** and **variance** in this context:

$\hat{y}_m(x_i) = f_m(x_i)$,
Where:

Bias measures how far the **average prediction** of the model is from the true value.
for each data point (x_i) and model (m).

Variance measures how much **predictions fluctuate** when the model is trained on different samples of the dataset.

Irreducible loss function should match the task:

`import numpy as np from sklearn.metrics import mean_squared_error M = 50 # number of bootstrap samples predictions = [] for _ in range(M): X_train, y_train = resample(X, y) model`

In practice, to approximate these quantities by **resampling the dataset**, training multiple models, and analyzing predictions.

• **Cross-entropy loss** for probabilistic classifiers.
• **Use a held-out test set** for evaluation to avoid optimistic bias.

• **Bias and variance** are typically computed per class or per example and then averaged.

✓ **Interpreting** Training time estimates can be expensive — use smaller subsamples or Monte Carlo dropout for an approximate.

Bias-variance trade-off can be visualized by plotting **bias², variance, and total error** as a function of model complexity.

Low bias, low variance anytime fit (ideal region).

$\text{Variance}(x_i) = \frac{1}{M} \sum_{m=1}^M (\hat{y}_m(x_i) - \bar{y}(x_i))^2$
2. Making multiple models across model complexity helps with **model selection** and **regularization tuning**.

3. Compute mean predictions, squared bias, and variance at each test point.

Step 5: Average across all data points

4. **Would you like me to adapt this method specifically for your use case** (e.g. regression vs classification, neural nets vs trees).
5. Interpret results to guide model design and complexity.

Revision question:
How to measure bias and
variance of an ML method?

Background

- For a given data point x_0 , a model's prediction can be decomposed as:
- $\mathbb{E}[(\hat{y}(x_0) - y(x_0))^2] = \text{Bias}^2(x_0) + \text{Variance}(x_0) + \text{Irreducible Error}.$
- Bias: error due to simplifying assumptions
- Variance: sensitivity to training set fluctuations
- Key idea: use resampling or cross-validation to estimate bias and variance

Practical steps 1/3

- Generate M training sets (bootstrap or k-fold)
- Train models on each set
- Collect predictions on a fixed test set for all models:

$$\hat{y}_m(x_i) = f_m(x_i),$$

for each data point x_i and model m .

Practical steps 2/3

- Compute mean prediction, bias² and variance per point

$$\bar{y}(x_i) = \frac{1}{M} \sum_{m=1}^M \hat{y}_m(x_i)$$

$$\text{Bias}^2(x_i) = (\bar{y}(x_i) - y(x_i))^2$$

$$\text{Variance}(x_i) = \frac{1}{M} \sum_{m=1}^M (\hat{y}_m(x_i) - \bar{y}(x_i))^2$$

Practical steps 3/3

- Average across all test data points

$$\text{Bias}^2 = \frac{1}{N} \sum_{i=1}^N \text{Bias}^2 (x_i)$$

$$\text{Variance} = \frac{1}{N} \sum_{i=1}^N \text{Variance} (x_i)$$

-

Implementation

- Use bootstrapping or multilevel cross-validation to train multiple models; compute bias and variance
- Use $M = 30\text{--}100$ for stable estimates

Interpretation

- High bias, low variance $\rightarrow$ underfitting
- Low bias, high variance $\rightarrow$ overfitting
- Low bias, low variance $\rightarrow$ good fit
- Plot bias^2 and variance vs model complexity for insights