

ECON3389 Machine Learning in Economics

Module 2 Classification

Alberto Cappello

Department of Economics, Boston College

Fall 2024

Overview

Agenda:

- Qualitative outcomes and classification problem.
- LPM, logit and probit models.
- Estimation, interpretation and accuracy measures.
- Extensions and other classification algorithms.

Readings:

- ISLR sections 4.1, 4.2, 4.3

Qualitative Outcomes

- So far our outcome variable Y has always been assumed to be quantitative, e.g. price, quantity, SAT score, etc.
- Qualitative variables (race, gender, geographical region, type of education, etc.) have only been discussed as predictors.
- But what if we want to quantify a relationship where the outcome Y is a qualitative variable?

Qualitative Outcomes

- So far our outcome variable Y has always been assumed to be quantitative, e.g. price, quantity, SAT score, etc.
- Qualitative variables (race, gender, geographical region, type of education, etc.) have only been discussed as predictors.
- But what if we want to quantify a relationship where the outcome Y is a qualitative variable?
 - A person arrives at an emergency room with symptoms that could be attributed to one of three medical conditions.

Qualitative Outcomes

- So far our outcome variable Y has always been assumed to be quantitative, e.g. price, quantity, SAT score, etc.
- Qualitative variables (race, gender, geographical region, type of education, etc.) have only been discussed as predictors.
- But what if we want to quantify a relationship where the outcome Y is a qualitative variable?
 - A person arrives at an emergency room with symptoms that could be attributed to one of three medical conditions.
 - Online banking service assesses whether a transaction being performed is fraudulent based on user's IP address, past transaction history, transaction amount, etc.

Qualitative Outcomes

- So far our outcome variable Y has always been assumed to be quantitative, e.g. price, quantity, SAT score, etc.
- Qualitative variables (race, gender, geographical region, type of education, etc.) have only been discussed as predictors.
- But what if we want to quantify a relationship where the outcome Y is a qualitative variable?
 - A person arrives at an emergency room with symptoms that could be attributed to one of three medical conditions.
 - Online banking service assesses whether a transaction being performed is fraudulent based on user's IP address, past transaction history, transaction amount, etc.
 - A researchers performs an analysis of socio-economic factors that affect whether a student graduates from college or drops out.

Basic Classification Problem

- Consider a qualitative variable Y that for every observation i takes a single value from a finite set of possible unordered values $\mathcal{C} = \{y_1, y_2, \dots, y_C\}$.
 - $Y = \text{eye color}$, $\mathcal{C} = \{\text{brown, blue, green}\}$
 - $Y = \text{medical diagnosis}$, $\mathcal{C} = \{\text{stroke, drug overdose, epileptic seizure}\}$
 - $Y = \text{transaction status}$, $\mathcal{C} = \{\text{fraudulent, non-fraudulent}\}$

Basic Classification Problem

- Consider a qualitative variable Y that for every observation i takes a single value from a finite set of possible unordered values $\mathcal{C} = \{y_1, y_2, \dots, y_C\}$.
 - $Y = \text{eye color}$, $\mathcal{C} = \{\text{brown, blue, green}\}$
 - $Y = \text{medical diagnosis}$, $\mathcal{C} = \{\text{stroke, drug overdose, epileptic seizure}\}$
 - $Y = \text{transaction status}$, $\mathcal{C} = \{\text{fraudulent, non-fraudulent}\}$
- Given a feature vector X and a qualitative response Y , the classification task is to build a function $C(X)$ that takes as input the feature vector X and predicts its value for Y , i.e. $C(X) \in \mathcal{C}$.
- In most cases we are interested in estimating the probabilities that Y belongs to a category in $\mathcal{C}$ given X , i.e.

$$\Pr(Y = y_c | X) \quad \forall y_c \in \mathcal{C}$$

Linear Probability Model: Problems

- Suppose for our medical condition classification we code

$$Y = \begin{cases} 1, & \text{if Stroke} \\ 2, & \text{if Drug Overdose} \\ 3, & \text{if Epileptic Seizure} \end{cases}$$

Can we use our standard regression model to predict the medical condition of a patient in the emergency room on the basis of her symptoms?

Linear Probability Model: Problems

- Suppose for our medical condition classification we code

$$Y = \begin{cases} 1, & \text{if Stroke} \\ 2, & \text{if Drug Overdose} \\ 3, & \text{if Epileptic Seizure} \end{cases}$$

Can we use our standard regression model to predict the medical condition of a patient in the emergency room on the basis of her symptoms?

- Does this coding imply an ordering of outcomes?

Linear Probability Model: Problems

- Suppose for our medical condition classification we code

$$Y = \begin{cases} 1, & \text{if Stroke} \\ 2, & \text{if Drug Overdose} \\ 3, & \text{if Epileptic Seizure} \end{cases}$$

Can we use our standard regression model to predict the medical condition of a patient in the emergency room on the basis of her symptoms?

- Does this coding imply an ordering of outcomes?
- In most cases, it is not possible for us to create a natural ordering in quantitative data

Linear Probability Model: Problems

- Suppose for our medical condition classification we code

$$Y = \begin{cases} 1, & \text{if Stroke} \\ 2, & \text{if Drug Overdose} \\ 3, & \text{if Epileptic Seizure} \end{cases}$$

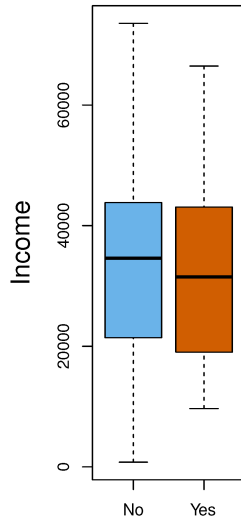
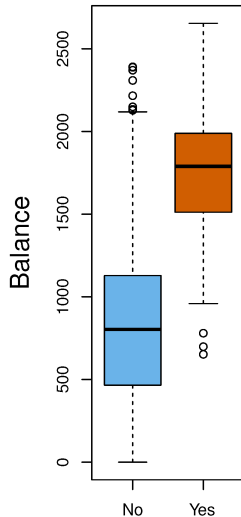
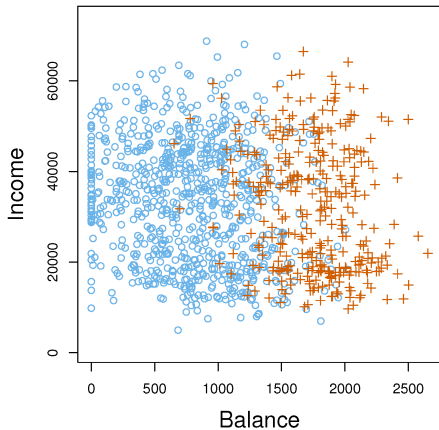
Can we use our standard regression model to predict the medical condition of a patient in the emergency room on the basis of her symptoms?

- Does this coding imply an ordering of outcomes?
- In most cases, it is not possible for us to create a natural ordering in quantitative data
- A different coding

$$Y = \begin{cases} 1, & \text{if Drug Overdose} \\ 2, & \text{if Stroke} \\ 3, & \text{if Epileptic Seizure} \end{cases}$$

Will generate a different model and different predictions

Example: Credit Card Defaults



Linear Probability Model

- Suppose for our credit card default classification we code

$$Y = \begin{cases} 0, & \text{if No} \\ 1, & \text{if Yes} \end{cases}$$

Can we use our standard regression model to estimate a regression of Y on X and classify outcome as **Yes** if $\hat{Y} > 0.5$?

Linear Probability Model

- Suppose for our credit card default classification we code

$$Y = \begin{cases} 0, & \text{if No} \\ 1, & \text{if Yes} \end{cases}$$

Can we use our standard regression model to estimate a regression of Y on X and classify outcome as **Yes** if $\hat{Y} > 0.5$?

- Given that Y is binary, we have

$$\mathbb{E}(Y|X) = ?$$

Linear Probability Model

- Suppose for our credit card default classification we code

$$Y = \begin{cases} 0, & \text{if No} \\ 1, & \text{if Yes} \end{cases}$$

Can we use our standard regression model to estimate a regression of Y on X and classify outcome as **Yes** if $\hat{Y} > 0.5$?

- Given that Y is binary, we have

$$\mathbb{E}(Y|X) = 1 \cdot \Pr(Y = 1|X) + 0 \cdot \Pr(Y = 0|X) = \Pr(Y = 1|X)$$

which means that in this case standard linear regression will estimate the probability of outcome $Y = 1$, hence the name *linear probability model* or *LPM*.

Linear Probability Model

- LPM retains all properties of linear regression, but the interpretation of the results is slightly different:

$$\Pr(Y = 1|X) = \beta_0 + \beta_1 X_1 + \beta_2 X_2 + \dots + \beta_p X_p + \epsilon$$

$$\beta_j = \frac{\partial \mathbb{E}[\Pr(Y = 1|X)]}{\partial X_j} = \frac{\mathbb{E}[\Delta \Pr(Y = 1|X)]}{\Delta X_j}$$

so regression coefficients now capture constant *marginal probabilities* of outcome $Y = 1$ given a change in X_j .

Linear Probability Model

- LPM retains all properties of linear regression, but the interpretation of the results is slightly different:

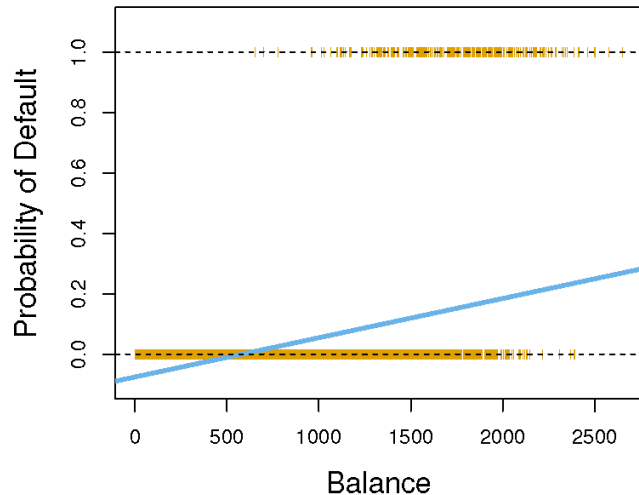
$$\Pr(Y = 1|X) = \beta_0 + \beta_1 X_1 + \beta_2 X_2 + \dots + \beta_p X_p + \epsilon$$

$$\beta_j = \frac{\partial \mathbb{E}[\Pr(Y = 1|X)]}{\partial X_j} = \frac{\mathbb{E}[\Delta \Pr(Y = 1|X)]}{\Delta X_j}$$

so regression coefficients now capture constant *marginal probabilities* of outcome $Y = 1$ given a change in X_j .

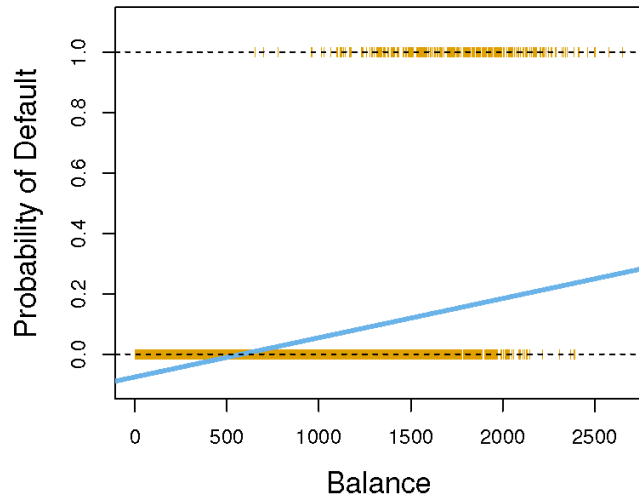
- LPM works very well in binary classification, especially if sample size is moderate to large.
- But it also has some inherent disadvantages, with the most common ones being unreasonable values of $\hat{\beta}_j$ and predictions for probability outside of $[0,1]$ interval.

Linear Probability Model



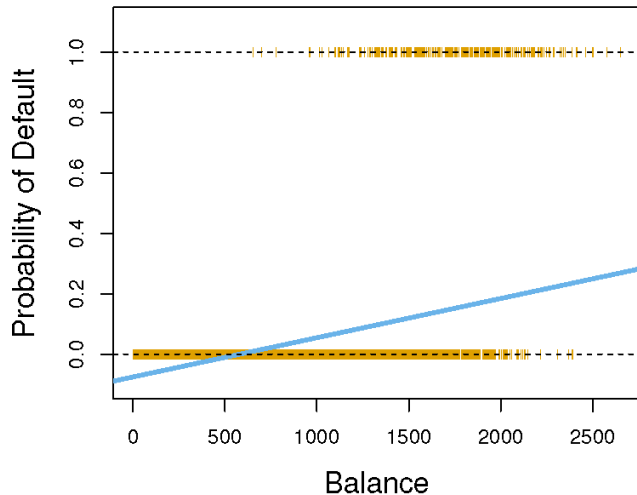
- The orange marks indicate the response Y , either 0 or 1.

Linear Probability Model



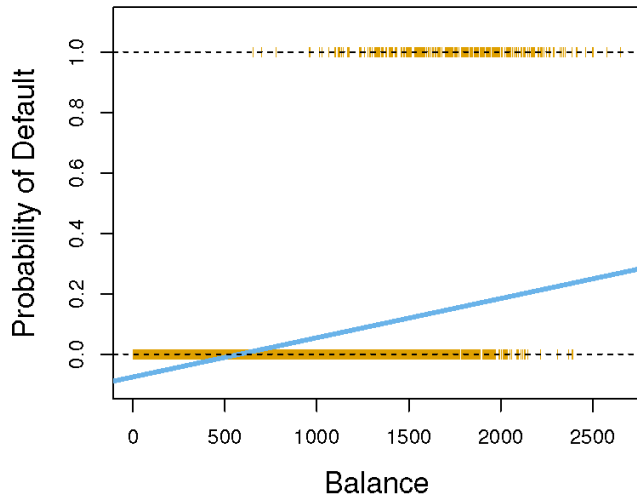
- The orange marks indicate the response Y , either 0 or 1.
- Blue line is estimated linear regression, which produces negative predictions for $\Pr(Y = 1|X)$ when **Balance** is less than 500.

Linear Probability Model



- The orange marks indicate the response Y , either 0 or 1.
- Blue line is estimated linear regression, which produces negative predictions for $\Pr(Y = 1|X)$ when **Balance** is less than 500.
- One way is to simply ignore the problem and bound any predictions from above and from below.

Linear Probability Model



- The orange marks indicate the response Y , either 0 or 1.
- Blue line is estimated linear regression, which produces negative predictions for $\Pr(Y = 1|X)$ when **Balance** is less than 500.
- One way is to simply ignore the problem and bound any predictions from above and from below.
- But can we do better?

Moving away from LPM

- The problem with prediction in LPM stems from the fact that we use linear combination of features X as the estimated probability itself.

Moving away from LPM

- The problem with prediction in LPM stems from the fact that we use linear combination of features X as the estimated probability itself.
- To avoid this problem, we can instead use a one-to-one mapping $F(\cdot)$ from $\mathbb{R}$ to $[0, 1]$ interval:

$$\hat{\Pr}(Y = 1|X) = F(\hat{\beta}_0 + \hat{\beta}_1 X_1 + \hat{\beta}_2 X_2 + \dots + \hat{\beta}_p X_p)$$

- What could serve as a function $F(\cdot)$? Infinitely many possibilities exist, but because this question was first addressed by statisticians, a natural choice was a *cumulative distribution function* (cdf) from some distribution.

Moving away from LPM

- The problem with prediction in LPM stems from the fact that we use linear combination of features X as the estimated probability itself.
- To avoid this problem, we can instead use a one-to-one mapping $F(\cdot)$ from $\mathbb{R}$ to $[0, 1]$ interval:

$$\hat{\Pr}(Y = 1|X) = F(\hat{\beta}_0 + \hat{\beta}_1 X_1 + \hat{\beta}_2 X_2 + \dots + \hat{\beta}_p X_p)$$

- What could serve as a function $F(\cdot)$? Infinitely many possibilities exist, but because this question was first addressed by statisticians, a natural choice was a *cumulative distribution function* (cdf) from some distribution.
- For any random variable Z its cdf $F_Z(\cdot)$ is by definition:

$$F_Z(a) = \Pr(Z \leq a)$$

- In classical statistical learning the two most common choice for $F(\cdot)$ are standard normal and logistic cdf.

Probit and Logit Models

- For simplicity, let's use the matrix notation:

$$\mathbf{X}\boldsymbol{\beta} = \beta_0 + \beta_1 X_1 + \beta_2 X_2 + \dots + \beta_p X_p$$

Probit and Logit Models

- For simplicity, let's use the matrix notation:

$$\mathbf{X}\boldsymbol{\beta} = \beta_0 + \beta_1 X_1 + \beta_2 X_2 + \dots + \beta_p X_p$$

- In *probit* model $F(\cdot)$ is assumed to be the cdf of a standard normal distribution:

$$F(\mathbf{X}\boldsymbol{\beta}) = \Phi(\mathbf{X}\boldsymbol{\beta}) = \int_{-\infty}^{\mathbf{X}\boldsymbol{\beta}} \frac{1}{\sqrt{2\pi}} \exp^{-\frac{z^2}{2}} dz$$

Probit and Logit Models

- For simplicity, let's use the matrix notation:

$$\mathbf{X}\beta = \beta_0 + \beta_1 X_1 + \beta_2 X_2 + \dots + \beta_p X_p$$

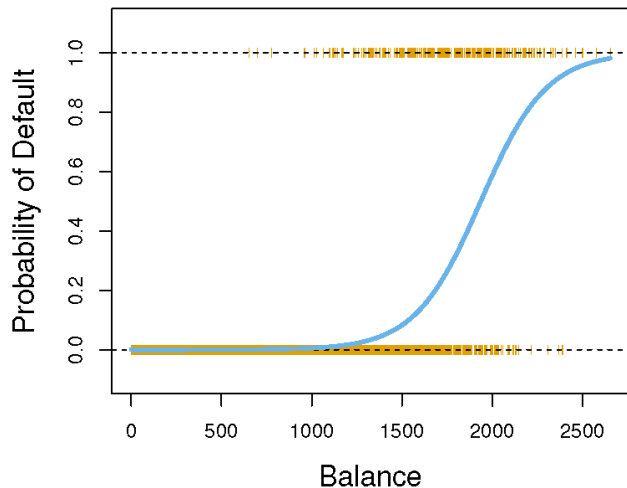
- In *probit* model $F(\cdot)$ is assumed to be the cdf of a standard normal distribution:

$$F(\mathbf{X}\beta) = \Phi(\mathbf{X}\beta) = \int_{-\infty}^{\mathbf{X}\beta} \frac{1}{\sqrt{2\pi}} \exp^{-\frac{z^2}{2}} dz$$

- In *logit* model $F(\cdot)$ is assumed to be the cdf of a logistic distribution:

$$F(\mathbf{X}\beta) = \Lambda(\mathbf{X}\beta) = \frac{\exp^{\mathbf{X}\beta}}{1 + \exp^{\mathbf{X}\beta}}$$

Probit and Logit Models



With either normal or logistic cdf as our $F(\cdot)$ function the estimated probability will, by definition, always lie in $[0, 1]$ interval.

But if that were the only advantage of logit/probit models, they would not become the de-facto standard econometric model for classification.

Generalized Linear Models

- Both logit and probit are special cases of the so-called *generalized linear models* (GLMs). Every GLM consists of three parts: the structural component, the link function and the response distribution.

Generalized Linear Models

- Both logit and probit are special cases of the so-called *generalized linear models* (GLMs). Every GLM consists of three parts: the structural component, the link function and the response distribution.
- *Structural component* is simply the linear combination of our predictors: $\mathbf{X}\beta$.

Generalized Linear Models

- Both logit and probit are special cases of the so-called *generalized linear models* (GLMs). Every GLM consists of three parts: the structural component, the link function and the response distribution.
- *Structural component* is simply the linear combination of our predictors: $\mathbf{X}\beta$.
- The *link function* $g(\mu)$ is such that its inverse gives us the (conditional) mean of our outcome Y as a function of the structural component:

$$g(\mu) = \mathbf{X}\beta \quad \text{or} \quad \mathbb{E}(Y|\mathbf{X}) = \mu = g^{-1}(\mathbf{X}\beta)$$

- The link function is the key to GLMs: since the distribution of the *response variable* Y is non-normal (in our simple example it is binomial), it's what lets us connect the structural component $\mathbf{X}\beta$ to the response Y — it 'links' them (hence the name).

Generalized Linear Models

- Because our outcome Y is binary, we have $\mathbb{E}(Y|\mathbf{X}) = \Pr(Y = 1|\mathbf{X})$, and thus our inverse link function g^{-1} is simply the function that defines conditional probability of $Y = 1$ given $\mathbf{X}$.
- For probit and logit models the corresponding cumulative distribution functions act as inverse link functions:

$$\text{Probit : } \mathbb{E}(Y|\mathbf{X}) = \mu_{Y|X} = \Pr(Y = 1|\mathbf{X}) = \Phi(\mathbf{X}\beta)$$

$$\text{Logit : } \mathbb{E}(Y|\mathbf{X}) = \mu_{Y|X} = \Pr(Y = 1|\mathbf{X}) = \Lambda(\mathbf{X}\beta)$$

Generalized Linear Models

- Because our outcome Y is binary, we have $\mathbb{E}(Y|\mathbf{X}) = \Pr(Y = 1|\mathbf{X})$, and thus our inverse link function g^{-1} is simply the function that defines conditional probability of $Y = 1$ given $\mathbf{X}$.
- For probit and logit models the corresponding cumulative distribution functions act as inverse link functions:

$$\text{Probit : } \mathbb{E}(Y|\mathbf{X}) = \mu_{Y|X} = \Pr(Y = 1|\mathbf{X}) = \Phi(\mathbf{X}\beta)$$

$$\text{Logit : } \mathbb{E}(Y|\mathbf{X}) = \mu_{Y|X} = \Pr(Y = 1|\mathbf{X}) = \Lambda(\mathbf{X}\beta)$$

- Note: standard MLR is also a special case of GLM with $g(\mu) = \mu = \mathbf{X}\beta$.

Generalized Linear Models

- Because our outcome Y is binary, we have $\mathbb{E}(Y|\mathbf{X}) = \Pr(Y = 1|\mathbf{X})$, and thus our inverse link function g^{-1} is simply the function that defines conditional probability of $Y = 1$ given $\mathbf{X}$.
- For probit and logit models the corresponding cumulative distribution functions act as inverse link functions:

$$\text{Probit : } \mathbb{E}(Y|\mathbf{X}) = \mu_{Y|X} = \Pr(Y = 1|\mathbf{X}) = \Phi(\mathbf{X}\beta)$$

$$\text{Logit : } \mathbb{E}(Y|\mathbf{X}) = \mu_{Y|X} = \Pr(Y = 1|\mathbf{X}) = \Lambda(\mathbf{X}\beta)$$

- Note: standard MLR is also a special case of GLM with $g(\mu) = \mu = \mathbf{X}\beta$.
- The two key differences of probit/logit models and usual MLR are estimation method and marginal effects calculation/interpretation.

Maximum Likelihood Estimation

- Because we no longer have a direct connection between Y and our structural component $\mathbf{X}\beta$, we need to specify our loss function in a different way. Using our link function, we can for every observation i write down the probability of observing a certain value of Y_i given values of $\mathbf{X}_i$

Maximum Likelihood Estimation

- Because we no longer have a direct connection between Y and our structural component $\mathbf{X}\beta$, we need to specify our loss function in a different way. Using our link function, we can for every observation i write down the probability of observing a certain value of Y_i given values of $\mathbf{X}_i$
- For example, for a logit model we have:

$$\Pr(Y = Y_i | \mathbf{X}_i) = \left(\frac{\exp^{\mathbf{X}_i \beta}}{1 + \exp^{\mathbf{X}_i \beta}} \right)^{Y_i} \left(1 - \frac{\exp^{\mathbf{X}_i \beta}}{1 + \exp^{\mathbf{X}_i \beta}} \right)^{1 - Y_i}$$

Maximum Likelihood Estimation

- Because we no longer have a direct connection between Y and our structural component $\mathbf{X}\beta$, we need to specify our loss function in a different way. Using our link function, we can for every observation i write down the probability of observing a certain value of Y_i given values of $\mathbf{X}_i$
- For example, for a logit model we have:

$$\Pr(Y = Y_i | \mathbf{X}_i) = \left(\frac{\exp^{X_i \beta}}{1 + \exp^{X_i \beta}} \right)^{Y_i} \left(1 - \frac{\exp^{X_i \beta}}{1 + \exp^{X_i \beta}} \right)^{1 - Y_i}$$

- With the default assumption of i.i.d. observations we can write down the joint probability or *likelihood function* of seeing our sample:

$$\ell(\beta) = \prod_{i=1}^n \Pr(Y = Y_i | \mathbf{X}_i)$$

Maximum Likelihood Estimation

- *Maximum likelihood* estimation (ML) is a method that chooses parameters β so as to minimize the loss function in form of the negative of the likelihood function:

$$\hat{\beta}_{ML} = \underset{\beta}{\operatorname{argmin}} -\ell(\beta)$$

Maximum Likelihood Estimation

- *Maximum likelihood* estimation (ML) is a method that chooses parameters β so as to minimize the loss function in form of the negative of the likelihood function:

$$\hat{\beta}_{ML} = \underset{\beta}{\operatorname{argmin}} -\ell(\beta)$$

- Under some general conditions $\hat{\beta}_{ML}$ is efficient, consistent and asymptotically normal, just like $\hat{\beta}_{OLS}$. In fact, one can show that standard OLS is a special case of ML if the error term ϵ in MLR is exactly normal.
- But unlike OLS, ML is a more general estimation procedure and allows one to recover structural parameters such as β in models that are far more flexible than standard MLR.

Marginal Effects in Probit/Logit

- The other key difference of logit/probit models from LPM is the fact that marginal effects are now calculated and interpreted in a different way:

$$\text{Probit : } \frac{\partial p(\mathbf{X})}{\partial X_j} = \frac{\partial \Phi(\mathbf{X}\beta)}{\partial X_j} \beta_j \neq \beta_j \quad \text{Logit : } \frac{\partial p(\mathbf{X})}{\partial X_j} = \frac{\partial \Lambda(\mathbf{X}\beta)}{\partial X_j} \beta_j \neq \beta_j$$

where $p(\mathbf{X}) = \Pr(Y = 1|\mathbf{X})$ for simplicity

Marginal Effects in Probit/Logit

- The other key difference of logit/probit models from LPM is the fact that marginal effects are now calculated and interpreted in a different way:

$$\text{Probit : } \frac{\partial p(\mathbf{X})}{\partial X_j} = \frac{\partial \Phi(\mathbf{X}\beta)}{\partial X_j} \beta_j \neq \beta_j \quad \text{Logit : } \frac{\partial p(\mathbf{X})}{\partial X_j} = \frac{\partial \Lambda(\mathbf{X}\beta)}{\partial X_j} \beta_j \neq \beta_j$$

where $p(\mathbf{X}) = \Pr(Y = 1|\mathbf{X})$ for simplicity

- The marginal effects now depend on values of all variables in $\mathbf{X}$, so we need to either estimate the marginal effects at a specific value of all our predictors (typically means or medians) or calculate their average over all values of $\mathbf{X}$ in our sample.
- Additionally, structural parameters β no longer have any direct interpretation on their own, with the exception of a few special cases.

Goodness-of-fit Measures

- While LPM can use the standard R^2 as a well-defined goodness-of-fit measure, it is not an option for logit/probit models due to the different loss function.
- In standard MLR an single R^2 value of 0.95 is an evidence of an excellent fit, but in classification problems we are often more interested in *class-specific* performance, especially in areas such as medicine or biology.

Goodness-of-fit Measures

- While LPM can use the standard R^2 as a well-defined goodness-of-fit measure, it is not an option for logit/probit models due to the different loss function.
- In standard MLR an single R^2 value of 0.95 is an evidence of an excellent fit, but in classification problems we are often more interested in *class-specific* performance, especially in areas such as medicine or biology.
- Consider the case where $Y = 1$ means a positive test result. Then our model's predictions fall into one of the 4 possible cases:

	$Y = 0$	$Y = 1$
$\hat{Y} = 0$		
$\hat{Y} = 1$		

Goodness-of-fit Measures

- While LPM can use the standard R^2 as a well-defined goodness-of-fit measure, it is not an option for logit/probit models due to the different loss function.
- In standard MLR an single R^2 value of 0.95 is an evidence of an excellent fit, but in classification problems we are often more interested in *class-specific* performance, especially in areas such as medicine or biology.
- Consider the case where $Y = 1$ means a positive test result. Then our model's predictions fall into one of the 4 possible cases:

	$Y = 0$	$Y = 1$
$\hat{Y} = 0$	True Negative	
$\hat{Y} = 1$		

Goodness-of-fit Measures

- While LPM can use the standard R^2 as a well-defined goodness-of-fit measure, it is not an option for logit/probit models due to the different loss function.
- In standard MLR an single R^2 value of 0.95 is an evidence of an excellent fit, but in classification problems we are often more interested in *class-specific* performance, especially in areas such as medicine or biology.
- Consider the case where $Y = 1$ means a positive test result. Then our model's predictions fall into one of the 4 possible cases:

	$Y = 0$	$Y = 1$
$\hat{Y} = 0$	True Negative	False Negative (Type II Error)
$\hat{Y} = 1$		

Goodness-of-fit Measures

- While LPM can use the standard R^2 as a well-defined goodness-of-fit measure, it is not an option for logit/probit models due to the different loss function.
- In standard MLR an single R^2 value of 0.95 is an evidence of an excellent fit, but in classification problems we are often more interested in *class-specific* performance, especially in areas such as medicine or biology.
- Consider the case where $Y = 1$ means a positive test result. Then our model's predictions fall into one of the 4 possible cases:

	$Y = 0$	$Y = 1$
$\hat{Y} = 0$	True Negative	False Negative (Type II Error)
$\hat{Y} = 1$	False Positive (Type I Error)	

Goodness-of-fit Measures

- While LPM can use the standard R^2 as a well-defined goodness-of-fit measure, it is not an option for logit/probit models due to the different loss function.
- In standard MLR an single R^2 value of 0.95 is an evidence of an excellent fit, but in classification problems we are often more interested in *class-specific* performance, especially in areas such as medicine or biology.
- Consider the case where $Y = 1$ means a positive test result. Then our model's predictions fall into one of the 4 possible cases:

	$Y = 0$	$Y = 1$
$\hat{Y} = 0$	True Negative	False Negative (Type II Error)
$\hat{Y} = 1$	False Positive (Type I Error)	True Positive

Goodness-of-fit Measures

- While LPM can use the standard R^2 as a well-defined goodness-of-fit measure, it is not an option for logit/probit models due to the different loss function.
- In standard MLR an single R^2 value of 0.95 is an evidence of an excellent fit, but in classification problems we are often more interested in *class-specific* performance, especially in areas such as medicine or biology.
- Consider the case where $Y = 1$ means a positive test result. Then our model's predictions fall into one of the 4 possible cases:

	$Y = 0$	$Y = 1$
$\hat{Y} = 0$	True Negative	False Negative (Type II Error)
$\hat{Y} = 1$	False Positive (Type I Error)	True Positive

Goodness-of-fit Measures

- While LPM can use the standard R^2 as a well-defined goodness-of-fit measure, it is not an option for logit/probit models due to the different loss function.
- In standard MLR an single R^2 value of 0.95 is an evidence of an excellent fit, but in classification problems we are often more interested in *class-specific* performance, especially in areas such as medicine or biology.
- Consider the case where $Y = 1$ means a positive test result. Then our model's predictions fall into one of the 4 possible cases:

	$Y = 0$	$Y = 1$
$\hat{Y} = 0$	True Negative	False Negative (Type II Error)
$\hat{Y} = 1$	False Positive (Type I Error)	True Positive

- For a COVID-19 test or cancer screening, we care more FN then about FP.

Goodness-of-fit Measures

- While LPM can use the standard R^2 as a well-defined goodness-of-fit measure, it is not an option for logit/probit models due to the different loss function.
- In standard MLR an single R^2 value of 0.95 is an evidence of an excellent fit, but in classification problems we are often more interested in *class-specific* performance, especially in areas such as medicine or biology.
- Consider the case where $Y = 1$ means a positive test result. Then our model's predictions fall into one of the 4 possible cases:

	$Y = 0$	$Y = 1$
$\hat{Y} = 0$	True Negative	False Negative (Type II Error)
$\hat{Y} = 1$	False Positive (Type I Error)	True Positive

- For a COVID-19 test or cancer screening, we care more FN then about FP.
- For city administration FPs in traffic cameras and speeding tickets are more important.

Goodness-of-fit Measures

- While LPM can use the standard R^2 as a well-defined goodness-of-fit measure, it is not an option for logit/probit models due to the different loss function.
- In standard MLR an single R^2 value of 0.95 is an evidence of an excellent fit, but in classification problems we are often more interested in *class-specific* performance, especially in areas such as medicine or biology.
- Consider the case where $Y = 1$ means a positive test result. Then our model's predictions fall into one of the 4 possible cases:

	$Y = 0$	$Y = 1$
$\hat{Y} = 0$	True Negative	False Negative (Type II Error)
$\hat{Y} = 1$	False Positive (Type I Error)	True Positive

- For a COVID-19 test or cancer screening, we care more FN then about FP.
- For city administration FPs in traffic cameras and speeding tickets are more important.
- In judicial system both FP and FN are equally important.

Goodness-of-fit Measures

- Consider the following *confusion matrix*, depicting the prediction results for the **Default** dataset, using 50% as a threshold for probability of default:

		True default status		
		No	Yes	Total
Predicted default status	No	9644	252	9896
	Yes	23	81	104
	Total	9667	333	10000

Goodness-of-fit Measures

- Consider the following *confusion matrix*, depicting the prediction results for the **Default** dataset, using 50% as a threshold for probability of default:

		True default status		
		No	Yes	Total
Predicted default status	No	9644	252	9896
	Yes	23	81	104
	Total	9667	333	10000

- If we simply look at pure prediction precision, then:
 - Our total error rate is $(23 + 252)/10000 = 2.75\%$, which seems low enough.

Goodness-of-fit Measures

- Consider the following *confusion matrix*, depicting the prediction results for the **Default** dataset, using 50% as a threshold for probability of default:

		True default status		
		No	Yes	Total
Predicted default status	No	9644	252	9896
	Yes	23	81	104
	Total	9667	333	10000

- If we simply look at pure prediction precision, then:
 - Our total error rate is $(23 + 252)/10000 = 2.75\%$, which seems low enough.
 - Out of 104 predicted defaults 81 ended up being classified correctly, which means only $23/9667 = 0.24\%$ of all non-defaults were classified incorrectly.

Goodness-of-fit Measures

- Consider the following *confusion matrix*, depicting the prediction results for the **Default** dataset, using 50% as a threshold for probability of default:

		True default status		
		No	Yes	Total
Predicted default status	No	9644	252	9896
	Yes	23	81	104
	Total	9667	333	10000

- If we simply look at pure prediction precision, then:
 - Our total error rate is $(23 + 252)/10000 = 2.75\%$, which seems low enough.
 - Out of 104 predicted defaults 81 ended up being classified correctly, which means only $23/9667 = 0.24\%$ of all non-defaults were classified incorrectly.
 - However, out of 333 true defaults we managed to miss $252/333 = 75.67\%$, which could be an unacceptably high error rate for this class.

Goodness-of-fit Measures

- This why in classification problems it is important to evaluate class-specific precision via the following four measures:

	$Y = 0$	$Y = 1$
$\hat{Y} = 0$	True Negative Rate (TNR) or <i>specificity</i> : $TNR = TN/N = 9644/9667 = 99.76\%$	
$\hat{Y} = 1$		

Goodness-of-fit Measures

- This why in classification problems it is important to evaluate class-specific precision via the following four measures:

	$Y = 0$	$Y = 1$
$\hat{Y} = 0$	True Negative Rate (TNR) or <i>specificity</i> : $TNR = TN/N = 9644/9667 = 99.76\%$	False Negative Rate (FNR): $FNR = FN/P = 252/333 = 75.67\%$
$\hat{Y} = 1$		

Goodness-of-fit Measures

- This why in classification problems it is important to evaluate class-specific precision via the following four measures:

	$Y = 0$	$Y = 1$
$\hat{Y} = 0$	True Negative Rate (TNR) or <i>specificity</i> : $TNR = TN/N = 9644/9667 = 99.76\%$	False Negative Rate (FNR): $FNR = FN/P = 252/333 = 75.67\%$
$\hat{Y} = 1$	False Positive Rate (FPR): $FPR = FP/N = 23/9667 = 0.24\%$	

Goodness-of-fit Measures

- This why in classification problems it is important to evaluate class-specific precision via the following four measures:

	$Y = 0$	$Y = 1$
$\hat{Y} = 0$	True Negative Rate (TNR) or <i>specificity</i> : $TNR = TN/N = 9644/9667 = 99.76\%$	False Negative Rate (FNR): $FNR = FN/P = 252/333 = 75.67\%$
$\hat{Y} = 1$	False Positive Rate (FPR): $FPR = FP/N = 23/9667 = 0.24\%$	True Positive Rate (TPR) or <i>sensitivity</i> : $TPR = TP/P = 81/333 = 24.33\%$

Goodness-of-fit Measures

- This why in classification problems it is important to evaluate class-specific precision via the following four measures:

	$Y = 0$	$Y = 1$
$\hat{Y} = 0$	True Negative Rate (TNR) or <i>specificity</i> : $TNR = TN/N = 9644/9667 = 99.76\%$	False Negative Rate (FNR): $FNR = FN/P = 252/333 = 75.67\%$
$\hat{Y} = 1$	False Positive Rate (FPR): $FPR = FP/N = 23/9667 = 0.24\%$	True Positive Rate (TPR) or <i>sensitivity</i> : $TPR = TP/P = 81/333 = 24.33\%$

Goodness-of-fit Measures

- This why in classification problems it is important to evaluate class-specific precision via the following four measures:

	$Y = 0$	$Y = 1$
$\hat{Y} = 0$	True Negative Rate (TNR) or <i>specificity</i> : $TNR = TN/N = 9644/9667 = 99.76\%$	False Negative Rate (FNR): $FNR = FN/P = 252/333 = 75.67\%$
$\hat{Y} = 1$	False Positive Rate (FPR): $FPR = FP/N = 23/9667 = 0.24\%$	True Positive Rate (TPR) or <i>sensitivity</i> : $TPR = TP/P = 81/333 = 24.33\%$

- As one can easily see, all four measures are related to each other. In particular, the following two identities must always hold:

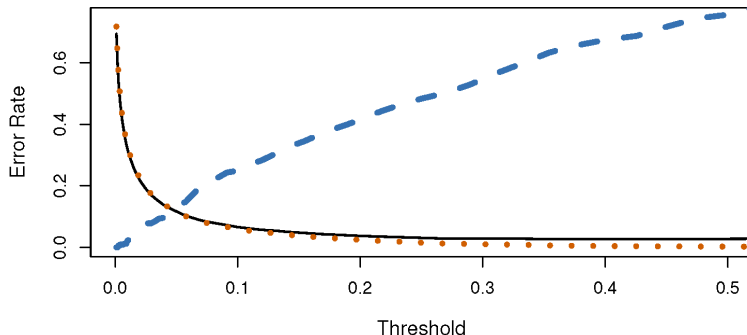
$$TNR + FPR = 100\% \quad \text{and} \quad TPR + FNR = 100\%$$

Goodness-of-fit Measures

The table on the previous slide was constructed using the rule $\hat{Y} = 1$ if $\hat{\Pr}(Y = 1|\mathbf{X}) > 0.5$, because 0.5 is the most common probability threshold used for classification predictions. However, the values of all 4 goodness-of-fit metrics will change if we change this threshold.

Goodness-of-fit Measures

The table on the previous slide was constructed using the rule $\hat{Y} = 1$ if $\hat{\Pr}(Y = 1|\mathbf{X}) > 0.5$, because 0.5 is the most common probability threshold used for classification predictions. However, the values of all 4 goodness-of-fit metrics will change if we change this threshold.



Black line: total error rate

Blue dashes: FNR

Orange dots: FPR

Based on this chart, we might want to set our threshold to 0.05 to achieve better error rate composition.

Logit vs Probit

While logit and probit models usually deliver very similar estimation results (especially on large datasets), modern statistical learning overwhelmingly prefers to use logistic regression. Why?

Logit vs Probit

While logit and probit models usually deliver very similar estimation results (especially on large datasets), modern statistical learning overwhelmingly prefers to use logistic regression. Why?

- **Coefficient interpretation.** In Economics we are interested in calculating and interpreting marginal effects, but in logit model one can also interpret the actual values of $\hat{\beta}_j$ themselves. This is because in logit model coefficient $\hat{\beta}_j$ shows how the log of odds ratio changes with changes in X_j :

$$\ln \left(\frac{p(\mathbf{X})}{1 - p(\mathbf{X})} \right) = \mathbf{X}\beta \quad \Rightarrow \quad \beta_j = \frac{\partial \ln \left(\frac{p(\mathbf{X})}{1 - p(\mathbf{X})} \right)}{\partial X_j}$$

Logit vs Probit

While logit and probit models usually deliver very similar estimation results (especially on large datasets), modern statistical learning overwhelmingly prefers to use logistic regression. Why?

- **Coefficient interpretation.** In Economics we are interested in calculating and interpreting marginal effects, but in logit model one can also interpret the actual values of $\hat{\beta}_j$ themselves. This is because in logit model coefficient $\hat{\beta}_j$ shows how the log of odds ratio changes with changes in X_j :

$$\ln \left(\frac{p(\mathbf{X})}{1 - p(\mathbf{X})} \right) = \mathbf{X}\beta \quad \Rightarrow \quad \beta_j = \frac{\partial \ln \left(\frac{p(\mathbf{X})}{1 - p(\mathbf{X})} \right)}{\partial X_j}$$

- **Random utility models.** Suppose consumer is choosing between two alternatives based on utility that is a function of observable product attributes $\mathbf{X}$ and a random utility shock ϵ . Then if ϵ follows Type I EV distribution, consumer's choice probabilities will take logit form (McFadden, D. (1973)).

Logit vs Probit

While logit and probit models usually deliver very similar estimation results (especially on large datasets), modern statistical learning overwhelmingly prefers to use logistic regression. Why?

- **Coefficient interpretation.** In Economics we are interested in calculating and interpreting marginal effects, but in logit model one can also interpret the actual values of $\hat{\beta}_j$ themselves. This is because in logit model coefficient $\hat{\beta}_j$ shows how the log of odds ratio changes with changes in X_j :

$$\ln \left(\frac{p(\mathbf{X})}{1 - p(\mathbf{X})} \right) = \mathbf{X}\beta \quad \Rightarrow \quad \beta_j = \frac{\partial \ln \left(\frac{p(\mathbf{X})}{1 - p(\mathbf{X})} \right)}{\partial X_j}$$

- **Random utility models.** Suppose consumer is choosing between two alternatives based on utility that is a function of observable product attributes $\mathbf{X}$ and a random utility shock ϵ . Then if ϵ follows Type I EV distribution, consumer's choice probabilities will take logit form (McFadden, D. (1973)).
- **Generalized choice models and information theory** (Matejka, F. and McKay, A. (2015)).

Multinomial and Ordered Outcomes

- In binary outcome case all three options (LPM, logit, probit) are applicable and most times yield similar results, especially in large samples. But once we move beyond basic binary case, LPM becomes structurally infeasible.

Multinomial and Ordered Outcomes

- In binary outcome case all three options (LPM, logit, probit) are applicable and most times yield similar results, especially in large samples. But once we move beyond basic binary case, LPM becomes structurally infeasible.
- Two hospitals use the following coding for incoming ER patients:

$$Y = \begin{cases} 1, & \text{if drug overdose} \\ 2, & \text{if stroke} \\ 3, & \text{if epileptic seizure} \end{cases} \quad \text{and} \quad Y = \begin{cases} 1, & \text{if stroke} \\ 5, & \text{if epileptic seizure} \\ 6, & \text{if drug overdose} \end{cases}$$

- This is an example of *multinomial unordered* classification. The fundamental difference with binary case is that here different coding will yield different results in LPM, but not in logit/probit.

Multinomial and Ordered Outcomes

- In binary outcome case all three options (LPM, logit, probit) are applicable and most times yield similar results, especially in large samples. But once we move beyond basic binary case, LPM becomes structurally infeasible.
- Two hospitals use the following coding for incoming ER patients:

$$Y = \begin{cases} 1, & \text{if drug overdose} \\ 2, & \text{if stroke} \\ 3, & \text{if epileptic seizure} \end{cases} \quad \text{and} \quad Y = \begin{cases} 1, & \text{if stroke} \\ 5, & \text{if epileptic seizure} \\ 6, & \text{if drug overdose} \end{cases}$$

- This is an example of *multinomial unordered* classification. The fundamental difference with binary case is that here different coding will yield different results in LPM, but not in logit/probit.
- Same thing happens with ordered outcomes, e.g. "strongly disagree, disagree, uncertain, agree, strongly agree", where we cannot impose a standardized numerical difference between outcomes (as opposed to something like number of kids in the family as the outcome).

Other SL and ML Classification Algorithms

- While research in econometrics over the past 50 years has developed a very wide range of discrete choice models, it wasn't just economics and casual inference that has been driving the development of statistical learning in classification problems.
- Areas such as genetics, biostatistics, pharmaceuticals and others have always had a need for statistical models that could precisely classify certain outcomes.

Other SL and ML Classification Algorithms

- While research in econometrics over the past 50 years has developed a very wide range of discrete choice models, it wasn't just economics and casual inference that has been driving the development of statistical learning in classification problems.
- Areas such as genetics, biostatistics, pharmaceuticals and others have always had a need for statistical models that could precisely classify certain outcomes.
- In recent decade huge advances in new variations of previously less used methods such as neural networks have been achieved due to ever-increasing demand for classification and prediction in modern data-dominated areas such as image and voice recognition.
- We will not be covering those methods in details (at least not till later in the course), because all of them are nearly completely irrelevant for causal inference. But because they have some other advantages, they deserve an honorary mentioning.

Discriminant analysis

- Instead of directly modeling $\Pr(Y|X)$, model the distribution of X in each of the C classes separately, and then use *Bayes theorem* to flip things around and obtain $\Pr(Y|X)$:

$$\Pr(Y = j|X = x) = \frac{\Pr(Y = j)\Pr(X = x|Y = j)}{\sum_{j=1}^C \Pr(Y = j)\Pr(X = x|Y = j)} = \frac{\pi_j f_j(x)}{\sum_{j=1}^C \pi_j f_j(x)}$$

where $f_j(x) = \Pr(X = x|Y = j)$ is the *density* for X in class j and $\pi_j = \Pr(Y = j)$ is the *prior* probability for class j .

Discriminant analysis

- Instead of directly modeling $\Pr(Y|X)$, model the distribution of X in each of the C classes separately, and then use *Bayes theorem* to flip things around and obtain $\Pr(Y|X)$:

$$\Pr(Y = j|X = x) = \frac{\Pr(Y = j)\Pr(X = x|Y = j)}{\sum_{j=1}^C \Pr(Y = j)\Pr(X = x|Y = j)} = \frac{\pi_j f_j(x)}{\sum_{j=1}^C \pi_j f_j(x)}$$

where $f_j(x) = \Pr(X = x|Y = j)$ is the *density* for X in class j and $\pi_j = \Pr(Y = j)$ is the *prior* probability for class j .

- Both $f_j(x)$ and π_j are estimated from the data, with the most common choice being normal density for $f_j(x)$.

Discriminant analysis

- Instead of directly modeling $\Pr(Y|X)$, model the distribution of X in each of the C classes separately, and then use *Bayes theorem* to flip things around and obtain $\Pr(Y|X)$:

$$\Pr(Y = j|X = x) = \frac{\Pr(Y = j)\Pr(X = x|Y = j)}{\sum_{j=1}^C \Pr(Y = j)\Pr(X = x|Y = j)} = \frac{\pi_j f_j(x)}{\sum_{j=1}^C \pi_j f_j(x)}$$

where $f_j(x) = \Pr(X = x|Y = j)$ is the *density* for X in class j and $\pi_j = \Pr(Y = j)$ is the *prior* probability for class j .

- Both $f_j(x)$ and π_j are estimated from the data, with the most common choice being normal density for $f_j(x)$.
- DA works especially well in small samples with well-separated classes and X variables that have approximately normal distributions.

K-nearest Neighbors

- For any given test observation x_0 and a positive integer K , first identify K points in the training data that are closest to x_0 , denoted as $\mathcal{N}_0$. Then the conditional probability for class j is calculated as

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

which is simply a fraction of points in $\mathcal{N}_0$ with response values equal to j .

K-nearest Neighbors

- For any given test observation x_0 and a positive integer K , first identify K points in the training data that are closest to x_0 , denoted as $\mathcal{N}_0$. Then the conditional probability for class j is calculated as

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

which is simply a fraction of points in $\mathcal{N}_0$ with response values equal to j .

- KNN is an example of non-parametric supervised learning algorithm and as such places almost no restrictions on the nature of the data.
- It quickly loses its potency when the number of features in X grows above 4-5 (too many points in high-dimensional space could be equally close to x_0). Thus it is often paired with other methods aimed at feature extraction and dimensionality reduction, such as *principal component analysis* (PCA).

Decision trees

- Create a sequence of binary splits that partitions (stratifies, segments) the feature space X (i.e. the set of possible values for $X_1, X_2, \dots, X_p$) into J distinct and non-overlapping regions $R_1, R_2, \dots, R_J$. Then predict that each observation belongs to the *most commonly occurring class* of training observations in the region to which it belongs.

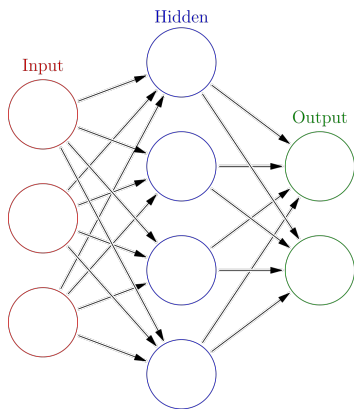
Decision trees

- Create a sequence of binary splits that partitions (stratifies, segments) the feature space X (i.e. the set of possible values for $X_1, X_2, \dots, X_p$) into J distinct and non-overlapping regions $R_1, R_2, \dots, R_J$. Then predict that each observation belongs to the *most commonly occurring class* of training observations in the region to which it belongs.
- Decision trees are very easy to explain and interpret, especially when they are small enough to be displayed graphically.
- Some believe that decision trees are a more natural way to model human decision making, as opposed to regressions and other methods.

Decision trees

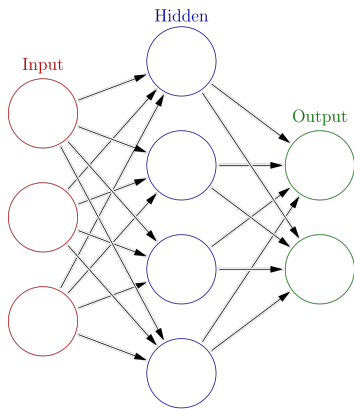
- Create a sequence of binary splits that partitions (stratifies, segments) the feature space X (i.e. the set of possible values for $X_1, X_2, \dots, X_p$) into J distinct and non-overlapping regions $R_1, R_2, \dots, R_J$. Then predict that each observation belongs to the *most commonly occurring class* of training observations in the region to which it belongs.
- Decision trees are very easy to explain and interpret, especially when they are small enough to be displayed graphically.
- Some believe that decision trees are a more natural way to model human decision making, as opposed to regressions and other methods.
- In their basic forms decision trees tend to have inferior levels of prediction accuracy compared to even basic LPM, but by employing modifications that allow aggregation and randomization of many trees into ensembles (forests) one can improve their accuracy drastically, although at the cost of interpretability.

Neural Networks



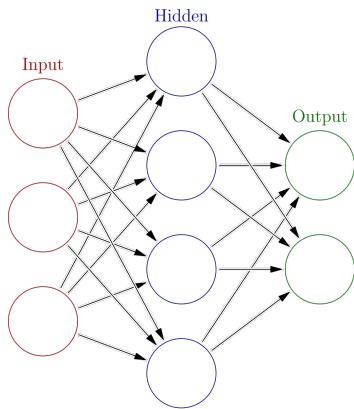
- Artificial neural networks (ANNs) is a learning/computing system that vaguely resembles neural systems in animal brains. ANNs consist of multiple layers (input layer, output layer and one or more hidden layers) of nodes called *neurons*, each capable of receiving and transmitting signals via connections to other neurons.

Neural Networks



- Artificial neural networks (ANNs) is a learning/computing system that vaguely resembles neural systems in animal brains. ANNs consist of multiple layers (input layer, output layer and one or more hidden layers) of nodes called *neurons*, each capable of receiving and transmitting signals via connections to other neurons.
- Each connection transferring the output of a neuron to the input of another neuron is assigned a weight. The *propagation function* computes the input to a neuron from the outputs of its predecessor neurons and their connections as a weighted sum.

Neural Networks



- Artificial neural networks (ANNs) is a learning/computing system that vaguely resembles neural systems in animal brains. ANNs consist of multiple layers (input layer, output layer and one or more hidden layers) of nodes called *neurons*, each capable of receiving and transmitting signals via connections to other neurons.
- Each connection transferring the output of a neuron to the input of another neuron is assigned a weight. The *propagation function* computes the input to a neuron from the outputs of its predecessor neurons and their connections as a weighted sum.
- ANN learns by adjusting its weighted associations according to a learning rule, using the error between the inputs and predicted outputs.