

ECON2228.05 Econometric Methods

Module 2 Binary Dependent Variable

Alberto Cappello

Department of Economics, Boston College

Spring 2025

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?

- This coding implies an ordering of outcomes. It also implies that the difference between Drug Overdose and Stroke is the same as the difference between Epileptic Seizure and Drug Overdose.

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?

- This coding implies an ordering of outcomes. It also implies that the difference between Drug Overdose and Stroke is the same as the difference between Epileptic Seizure and Drug Overdose.
- 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

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

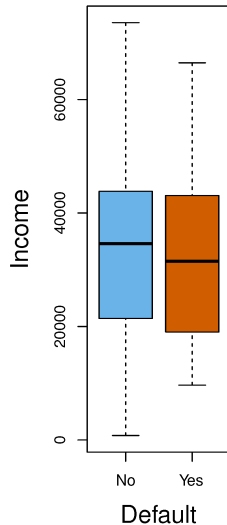
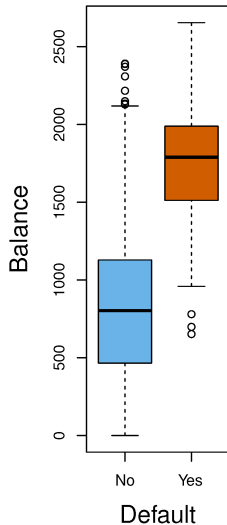
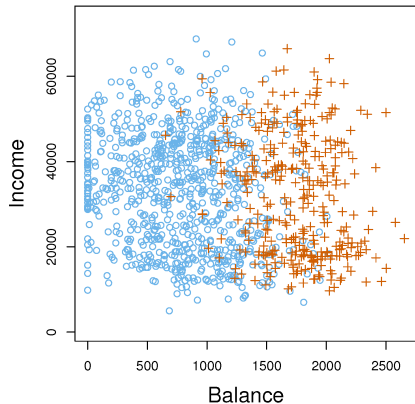
Can we use our standard regression model to predict the medical condition of a patient in the emergency room on the basis of her symptoms?

- This coding implies an ordering of outcomes. It also implies that the difference between Drug Overdose and Stroke is the same as the difference between Epileptic Seizure and Drug Overdose.
- 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:

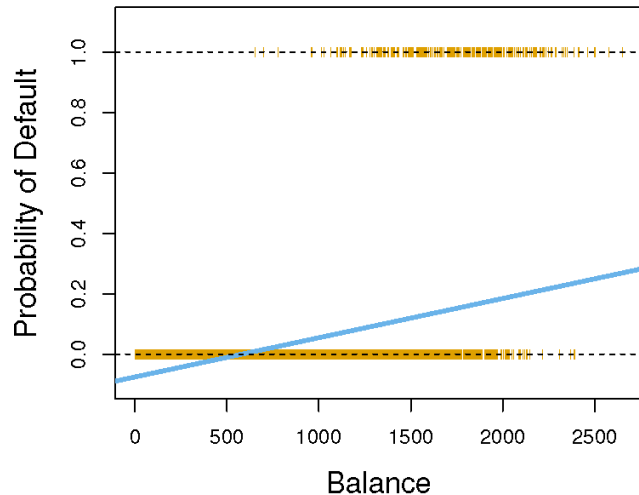
$$\Pr(Y = 1|X) = \beta_0 + \beta_1 X_1 + \beta_2 X_2 + \dots + \beta_p X_p + \epsilon$$

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

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

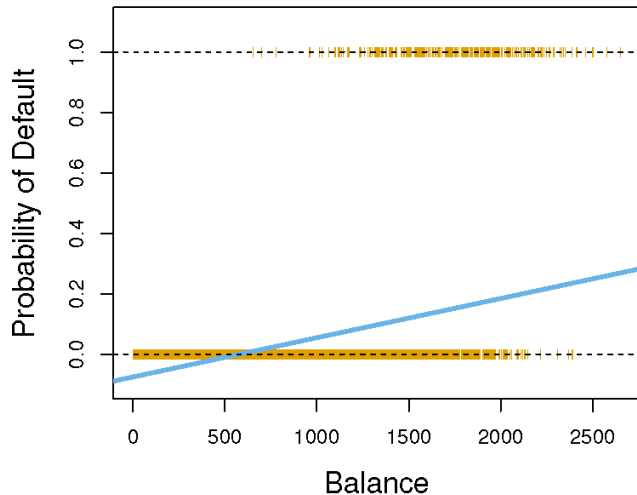
- LPM works very well in binary classification, especially if sample size is moderate to large.
- But it also has some inherent disadvantages

Linear Probability Model



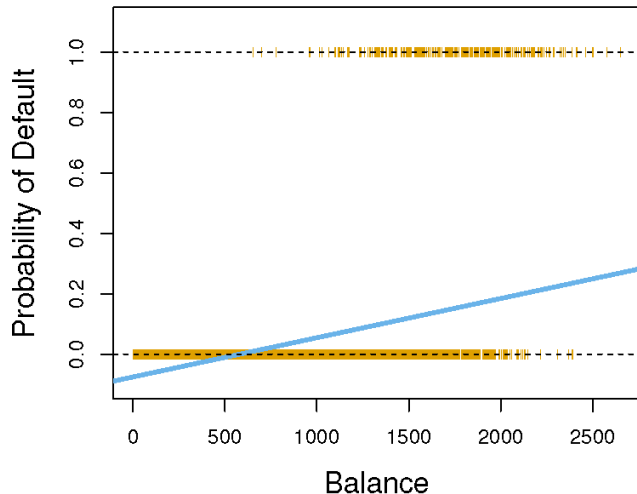
- The orange marks indicate the response Y , either 0 or 1.

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.
- 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:

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

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:

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

- 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:

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

- 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:

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

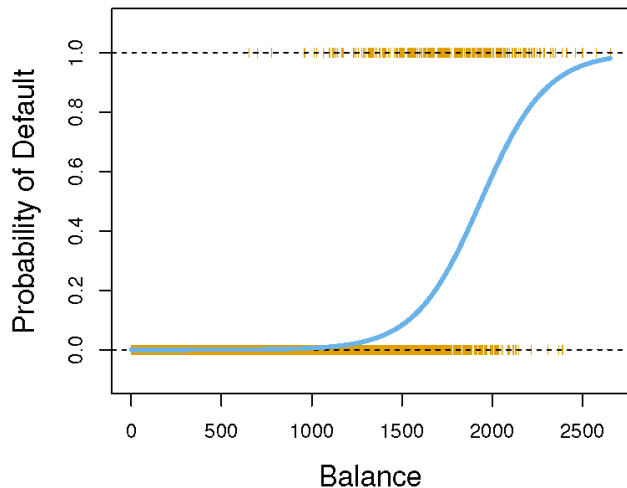
- In *probit* model $F(\cdot)$ is assumed to be the cdf of a standard normal distribution:

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

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

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

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.

Introduction to the Logit Model as a Latent Variable Model

- **Overview:** The logit model is commonly used in binary choice modeling, where the dependent variable takes values 0 or 1 (e.g., success or failure, yes or no). It can also be viewed as a latent variable model.
- **Latent Variable Framework:**
 - In the latent variable approach, the observed binary outcome is determined by an unobserved (latent) variable, which is modeled as a continuous variable.
 - The observed outcome is determined by whether this latent variable exceeds a certain threshold.
- **Key Idea:** The logit model can be interpreted as a probabilistic model where an individual's decision depends on the value of a latent variable (e.g., an underlying propensity to choose '1').

Latent Variable Representation in the Logit Model

- **Latent Variable Model:** The logit model assumes the existence of a latent variable y^* that is linearly related to the explanatory variables X and an error term ϵ :

$$y^* = X\beta + \epsilon$$

where:

- y^* is the latent variable,
 - X represents the explanatory variables,
 - β is the vector of parameters,
 - ϵ is the error term, typically assumed to follow a logistic distribution.
- **Thresholding:** The observed binary outcome y is determined by whether the latent variable exceeds a certain threshold, often assumed to be 0:

$$y = \begin{cases} 1 & \text{if } y^* > 0 \\ 0 & \text{if } y^* \leq 0 \end{cases}$$

This thresholding leads to the discrete binary outcome.

Derivation of the Probability Function – Step 1

- **Step 1: Latent Variable Model:** As mentioned earlier, the latent variable model is:

$$y^* = X\beta + \epsilon$$

where $\epsilon \sim \text{Logistic}(0, 1)$, meaning the error term follows a logistic distribution with a mean of 0 and a variance of $\frac{\pi^2}{3}$.

- **Step 2: Define the Probability of $y = 1$:** The observed outcome y is related to the latent variable y^* through the following thresholding rule:

$$y = \begin{cases} 1 & \text{if } y^* > 0 \\ 0 & \text{if } y^* \leq 0 \end{cases}$$

The probability of observing $y = 1$ is:

$$P(y = 1 \mid X) = P(y^* > 0) = P(X\beta + \epsilon > 0)$$

which simplifies to:

$$P(y = 1 \mid X) = P(\epsilon > -X\beta) = F_{\epsilon}(X\beta)$$

where F_{ϵ} is the CDF of the logistic distribution.

Derivation of the Probability Function – Step 2

- **Step 3: Using the Logistic CDF:** The CDF of the logistic distribution is:

$$F(\epsilon) = \frac{1}{1 + \exp(-\epsilon)}$$

Therefore, the probability of $y = 1$ becomes:

$$P(y = 1 | X) = \frac{1}{1 + \exp(-X\beta)}$$

- **Step 4: Probability of $y = 0$:** The probability of $y = 0$ is simply:

$$P(y = 0 | X) = 1 - P(y = 1 | X) = \frac{\exp(-X\beta)}{1 + \exp(-X\beta)}$$

Thus, we have derived the probabilities for the logit model as:

$$P(y = 1 | X) = \frac{1}{1 + \exp(-X\beta)} \quad \text{and} \quad P(y = 0 | X) = \frac{\exp(-X\beta)}{1 + \exp(-X\beta)}$$

Interpretation of β as a Change in Log Odds

- **Log Odds Interpretation:** The **log odds** of $y = 1$ (i.e., the log of the odds of success) is:

$$\text{log-odds}(y = 1 | X) = \log \left(\underbrace{\frac{P(y = 1 | X)}{P(y = 0 | X)}}_{\text{ODDS}} \right)$$

Substituting the probabilities:

$$\text{log-odds}(y = 1 | X) = X\beta$$

- **Change in Log Odds:** The coefficient β_j associated with an explanatory variable X_j represents the change in the log odds of $y = 1$ for a one-unit change in X_j , holding all other variables constant.
 - If $\beta_j > 0$, a one-unit increase in X_j increases the log odds of $y = 1$, making the outcome more likely.
 - If $\beta_j < 0$, a one-unit increase in X_j decreases the log odds of $y = 1$, making the outcome less likely.
- **Example:** If $\beta_j = 0.5$, then a one-unit increase in X_j increases the log odds of success (i.e., $y = 1$) by 0.5. This corresponds to multiplying the odds of success by $e^{0.5} \approx 1.65$, or an approximately 65% increase in the odds.

Maximum Likelihood Estimation: Idea

- Suppose $X_1, X_2, \dots, X_n$ form a random sample from a distribution for which the pdf is $f(x|\theta)$

Maximum Likelihood Estimation: Idea

- Suppose $X_1, X_2, \dots, X_n$ form a random sample from a distribution for which the pdf is $f(x|\theta)$
- For every observed vector $\mathbf{x} = (x_1, x_2, \dots, x_n)$, we can define the joint pdf as follows

$$L(\theta) = f_n(\mathbf{x}|\theta) = f(x_1|\theta).f(x_2|\theta).f(x_3|\theta)\dots f(x_n|\theta) \quad (1)$$

Maximum Likelihood Estimation: Idea

- Suppose $X_1, X_2, \dots, X_n$ form a random sample from a distribution for which the pdf is $f(x|\theta)$
- For every observed vector $\mathbf{x} = (x_1, x_2, \dots, x_n)$, we can define the joint pdf as follows

$$L(\theta) = f_n(\mathbf{x}|\theta) = f(x_1|\theta).f(x_2|\theta).f(x_3|\theta)\dots f(x_n|\theta) \quad (1)$$

- $L(\theta)$ is called the *Likelihood Function*

Maximum Likelihood Estimation: Idea

- Suppose $X_1, X_2, \dots, X_n$ form a random sample from a distribution for which the pdf is $f(x|\theta)$
- For every observed vector $\mathbf{x} = (x_1, x_2, \dots, x_n)$, we can define the joint pdf as follows

$$L(\theta) = f_n(\mathbf{x}|\theta) = f(x_1|\theta).f(x_2|\theta).f(x_3|\theta)\dots f(x_n|\theta) \quad (1)$$

- $L(\theta)$ is called the *Likelihood Function*
- The MLE estimator of θ will find the parameter values that maximize $L(\theta)$

Maximum Likelihood Estimation: Idea

- Suppose $X_1, X_2, \dots, X_n$ form a random sample from a distribution for which the pdf is $f(x|\theta)$
- For every observed vector $\mathbf{x} = (x_1, x_2, \dots, x_n)$, we can define the joint pdf as follows

$$L(\theta) = f_n(\mathbf{x}|\theta) = f(x_1|\theta).f(x_2|\theta).f(x_3|\theta)\dots f(x_n|\theta) \quad (1)$$

- $L(\theta)$ is called the *Likelihood Function*
- The MLE estimator of θ will find the parameter values that maximize $L(\theta)$
- In other words, it will find the parameter value that maximize the likelihood of the observed data being drawn from $f(x|\theta)$

Maximum Likelihood Estimation: Idea

- Suppose $X_1, X_2, \dots, X_n$ form a random sample from a distribution for which the pdf is $f(x|\theta)$
- For every observed vector $\mathbf{x} = (x_1, x_2, \dots, x_n)$, we can define the joint pdf as follows

$$L(\theta) = f_n(\mathbf{x}|\theta) = f(x_1|\theta).f(x_2|\theta).f(x_3|\theta)\dots f(x_n|\theta) \quad (1)$$

- $L(\theta)$ is called the *Likelihood Function*
- The MLE estimator of θ will find the parameter values that maximize $L(\theta)$
- In other words, it will find the parameter value that maximize the likelihood of the observed data being drawn from $f(x|\theta)$
- Suppose $x_1, x_2, \dots, x_n$ form a random sample from a normal distribution for which the mean μ is unknown and variance σ^2 is known

Maximum Likelihood Estimation: Idea

- Suppose $X_1, X_2, \dots, X_n$ form a random sample from a distribution for which the pdf is $f(x|\theta)$
- For every observed vector $\mathbf{x} = (x_1, x_2, \dots, x_n)$, we can define the joint pdf as follows

$$L(\theta) = f_n(\mathbf{x}|\theta) = f(x_1|\theta).f(x_2|\theta).f(x_3|\theta)\dots f(x_n|\theta) \quad (1)$$

- $L(\theta)$ is called the *Likelihood Function*
- The MLE estimator of θ will find the parameter values that maximize $L(\theta)$
- In other words, it will find the parameter value that maximize the likelihood of the observed data being drawn from $f(x|\theta)$
- Suppose $x_1, x_2, \dots, x_n$ form a random sample from a normal distribution for which the mean μ is unknown and variance σ^2 is known
- The likelihood function of μ is

$$L(\mu) = f_n(\mathbf{x}|\mu) = \prod_{i=1}^n \frac{1}{\sigma\sqrt{2\pi}} \exp^{-\frac{(x_i - \mu)^2}{2\sigma^2}} \quad (2)$$

Maximum Likelihood Estimation: Example 1

- **Example:** Suppose I toss a coin 100 times and get 56 heads. What is the MLE of the probability of heads in a single toss?

Maximum Likelihood Estimation: Example 1

- **Example:** Suppose I toss a coin 100 times and get 56 heads. What is the MLE of the probability of heads in a single toss?
- **Model:** $L(p) = L(p; n, x) = \binom{n}{x} p^x (1 - p)^{n-x} = \binom{100}{56} p^{56} (1 - p)^{44}$

Maximum Likelihood Estimation: Example 1

- **Example:** Suppose I toss a coin 100 times and get 56 heads. What is the MLE of the probability of heads in a single toss?
- **Model:** $L(p) = L(p; n, x) = \binom{n}{x} p^x (1 - p)^{n-x} = \binom{100}{56} p^{56} (1 - p)^{44}$
- $L(0.5) = 0.038$
- $L(0.52) = 0.058$
- $L(0.54) = 0.073$
- $L(0.56) = 0.081$
- $L(0.58) = 0.073$

Maximum Likelihood Estimation: Method

- **Method 1**

In the previous example, it is easy for us to write a simple equation that describes the likelihood surface that can be differentiated to find the MLE estimate

Maximum Likelihood Estimation: Method

- **Method 1**

In the previous example, it is easy for us to write a simple equation that describes the likelihood surface that can be differentiated to find the MLE estimate

- Step 1: Take the log

$$\ln L = \ln \binom{100}{56} + \ln(p^{56}(1-p)^{44}) \quad (3)$$

$$\ln L = 56 \ln(p) + 44 \ln(1-p) \quad (4)$$

(5)

Maximum Likelihood Estimation: Method

- **Method 1**

In the previous example, it is easy for us to write a simple equation that describes the likelihood surface that can be differentiated to find the MLE estimate

- Step 1: Take the log

$$\ln L = \ln \binom{100}{56} + \ln(p^{56}(1-p)^{44}) \quad (3)$$

$$\ln L = 56 \ln(p) + 44 \ln(1-p) \quad (4)$$

(5)

- Step 2: Differentiate the log likelihood to find the optimal parameter

$$\frac{56}{p} - \frac{44}{(1-p)} = 0 \quad (6)$$

$$56(1-p) - 44p = 0 \quad (7)$$

$$p = \frac{56}{100} \quad (8)$$

Maximum Likelihood Estimation: Method

- **Method 2**

In the previous example, we used the information that sum of bernoullis follows a binomial distribution to construct the overall likelihood surface

Maximum Likelihood Estimation: Method

- **Method 2**

In the previous example, we used the information that sum of bernoullis follows a binomial distribution to construct the overall likelihood surface

- In many cases this might not be possible because we are not working with such distributions and the model is at the level of individual coin tosses

Maximum Likelihood Estimation: Method

- **Method 2**

In the previous example, we used the information that sum of bernoullis follows a binomial distribution to construct the overall likelihood surface

- In many cases this might not be possible because we are not working with such distributions and the model is at the level of individual coin tosses
- In such cases we need to construct the overall likelihood surface using the individual likelihoods

Maximum Likelihood Estimation: Method

- **Method 2**

In the previous example, we used the information that sum of bernoullis follows a binomial distribution to construct the overall likelihood surface

- In many cases this might not be possible because we are not working with such distributions and the model is at the level of individual coin tosses
- In such cases we need to construct the overall likelihood surface using the individual likelihoods
- Each individual coin toss follows a Bernoulli distribution. Suppose $X = 1$ when heads and 0 otherwise

$$L(p; x) = p^x(1 - p)^{1-x}$$

Maximum Likelihood Estimation: Method

- **Method 2:** Remember the data is given to us

Observation	Outcome (x)	Likelihood of outcome
1	1	p
2	0	$1-p$
3	1	p
...		
99	0	$1-p$
100	1	p
Total	56	?

- What is the overall/joint likelihood of entries in the second column?

Maximum Likelihood Estimation: Method

- **Method 2:** Remember the data is given to us

Observation	Outcome (x)	Likelihood of outcome
1	1	p
2	0	1-p
3	1	p
...		
99	0	1-p
100	1	p
Total	56	?

- What is the overall/joint likelihood of entries in the second column?
- Each coin toss is independent

$$L(p) = p.p.p.p..(1-p)(1-p)...(1-p) \quad (9)$$

$$L(p) = p^{56}(1-p)^{44} \quad (10)$$

$$\ln L(p) = \ln(p^{56}(1-p)^{44}) \quad (11)$$

OLS as a special case of MLE

- Main assumption: The errors follow a normal distribution with mean 0 and variance σ^2

$$\epsilon \sim \mathcal{N}(0, \sigma^2)$$

$$Y_i = \beta_0 + \beta_1 X_i + \epsilon_i$$

$$\epsilon_i = Y_i - \beta_0 - \beta_1 X_i$$

OLS as a special case of MLE

- Main assumption: The errors follow a normal distribution with mean 0 and variance σ^2

$$\epsilon \sim \mathcal{N}(0, \sigma^2)$$

$$Y_i = \beta_0 + \beta_1 X_i + \epsilon_i$$

$$\epsilon_i = Y_i - \beta_0 - \beta_1 X_i$$

- The likelihood function of (β) is

$$L(\beta) = \prod_{i=1}^n \frac{1}{\sigma\sqrt{2\pi}} \exp^{-\frac{(y_i - \beta_0 - \beta_1 x_i)^2}{2\sigma^2}}$$

$$\ln(L(\beta)) = \sum_{i=1}^n \ln\left(\frac{1}{\sigma\sqrt{2\pi}}\right) - \sum_{i=1}^n \frac{(y_i - \beta_0 - \beta_1 x_i)^2}{2\sigma^2}$$

OLS as a special case of MLE

- Main assumption: The errors follow a normal distribution with mean 0 and variance σ^2

$$\epsilon \sim \mathcal{N}(0, \sigma^2)$$

$$Y_i = \beta_0 + \beta_1 X_i + \epsilon_i$$

$$\epsilon_i = Y_i - \beta_0 - \beta_1 X_i$$

- The likelihood function of (β) is

$$L(\beta) = \prod_{i=1}^n \frac{1}{\sigma\sqrt{2\pi}} \exp^{-\frac{(y_i - \beta_0 - \beta_1 x_i)^2}{2\sigma^2}}$$

$$\ln(L(\beta)) = \sum_{i=1}^n \ln\left(\frac{1}{\sigma\sqrt{2\pi}}\right) - \sum_{i=1}^n \frac{(y_i - \beta_0 - \beta_1 x_i)^2}{2\sigma^2}$$

- The first term does not depend on (β) and the second term has a constant σ^2 that we can bring outside the summation

$$\ln(L(\beta)) = -\frac{1}{2\sigma^2} \sum_{i=1}^n (y_i - \beta_0 - \beta_1 x_i)^2$$

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 log likelihood function:

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

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 log likelihood function:

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

- Under some general conditions $\hat{\beta}_{ML}$ is efficient, consistent and asymptotically normal, just like $\hat{\beta}_{OLS}$

	Coefficient	Std. error	z-statistic	p-value
Intercept	-10.6513	0.3612	-29.5	<0.0001
balance	0.0055	0.0002	24.9	<0.0001

- 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.

Logistic Regression

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

$$\Pr(Y = 1|\mathbf{X}) = p(\mathbf{X}) = \frac{\exp^{X\beta}}{1 + \exp^{X\beta}}$$

$$\frac{p(\mathbf{X})}{1 - p(\mathbf{X})} = \exp^{X\beta}$$

$$\ln \frac{p(\mathbf{X})}{1 - p(\mathbf{X})} = X\beta$$

- The quantity on the left is called *log-odds* or *logit*
- The logistic regression model has a logit that is linear in X
- In a logistic regression model, increasing X by one unit changes the log odds by its corresponding β

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.

Multinomial Logistic Regression

- It is possible to extend the two-class logistic regression approach to the setting of $K > 2$ classes
- Suppose we have K classes. First we need to define a base class (K^{th} one)

$$\Pr(Y = k|\mathbf{X}) = \frac{\exp^{X\beta_k}}{1 + \sum_{l=1}^{K-1} \exp^{X\beta_l}}$$
$$\Pr(Y = K|\mathbf{X}) = \frac{1}{1 + \sum_{l=1}^{K-1} \exp^{X\beta_l}}$$

- It can be shown that the above model implies

$$\frac{\Pr(Y = k|\mathbf{X})}{\Pr(Y = K|\mathbf{X})} = \exp^{X\beta_k}$$
$$\ln \frac{\Pr(Y = k|\mathbf{X})}{\Pr(Y = K|\mathbf{X})} = X\beta_k$$

- The interpretation of the β_k is with respect to the base category K

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:

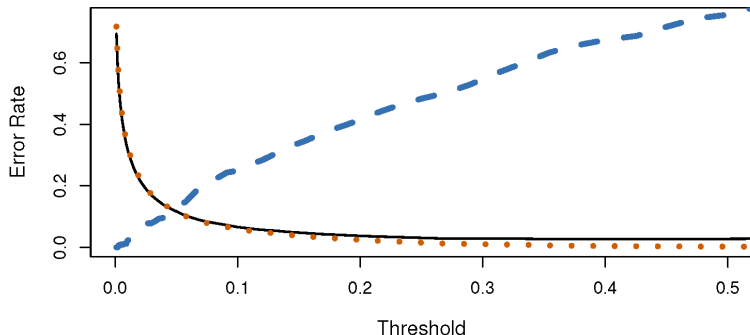
$$TNR + FPR = 100\% \quad \text{and} \quad TPR + FNR = 100\%$$

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.

Deviance and Model Fit

- Another way of assessing the goodness of fit the Logit model is based on the concept of Deviance.
- Intuitively, deviance measures the discrepancy of the fitted logistic model with respect to a perfect model for $P[Y = 1 \mid X_1 = x_1, \dots, X_k = x_k]$. This perfect model, known as the **saturated model**, represents an abstract model that fits the sample perfectly. That is, the model satisfies:

$$\hat{P}[Y = 1 \mid X_1 = X_{i1}, \dots, X_k = X_{ik}] = Y_i, \quad i = 1, \dots, n.$$

The saturated model assigns probabilities 0 or 1 to Y depending on the actual value of Y_i .

- More precisely, the deviance is defined as the difference in likelihoods between the fitted model and the saturated model:

$$D = -2 \log \text{lik}(\hat{\beta}) + 2 \log \text{lik}(\text{saturated model}).$$

Since the likelihood of the saturated model is exactly one, the deviance is simply another expression of the likelihood:

$$D = -2 \log \text{lik}(\hat{\beta}).$$

As a consequence, the deviance is always greater than or equal to zero, and it is equal to zero only if the fit is perfect.

- A benchmark for evaluating the magnitude of the deviance is the null deviance:

$$D_0 = -2 \log \text{lik}(\hat{\beta}_0),$$

where $\hat{\beta}_0$ represents the parameters of the model with only an intercept (no predictors).

- The null deviance, D_0 , represents the deviance of the worst model (the one fitted without any predictor) compared to the perfect model:

$$Y \mid (X_1 = x_1, \dots, X_k = x_k) \sim \text{Bernoulli}(\text{logistic}(\beta_0)),$$

where the parameter $\hat{\beta}_0$ is given by:

$$\hat{\beta}_0 = \text{logit}(m/n) = \log \frac{m/n}{1 - m/n},$$

and m is the number of 1's in $Y_1, \dots, Y_n$.

- The null deviance serves as a benchmark for comparing how much the model has improved by adding the predictors $X_1, \dots, X_k$. This comparison can be quantified using the R^2 statistic, which is a generalization of the coefficient of determination in multiple linear regression:

$$Pseudo - R^2 = 1 - \frac{\text{deviance fitted model } (D)}{\text{deviance null model } (D_0)} = 1 - \frac{\log \text{lik}(\hat{\beta})}{\log \text{lik}(\hat{\beta}_0)}$$

This global measure of fit shares some important properties with the R^2 in linear regression:

- It is a quantity between 0 and 1.
- If the fit is perfect, then $D = 0$ and $R^2 = 1$.
- If the predictors do not add anything to the regression, then $D = D_0$ and $R^2 = 0$.
- Notice that this R^2 is not the percentage of variance explained by the logistic model, but rather a ratio indicating how close is the fit to being perfect or the worst.