

Lecture 7 - Modelling and Analysis

Topics in Econometrics - M2 ENS Lyon

Vincent Bagilet

2025-10-14

Introduction

Short feedback form



<https://forms.gle/wxwZNFXyPxprLELT7>

Steps of an Econometrics Analysis

- **Design:** decisions of data collection and measurement
 - eg, decisions related to sample size and ensuring exogeneity of the treatment
- **Modelling:** define statistical models
 - eg selecting variables, functional forms, etc
- **Analysis:** estimation and questions of statistical inference
 - eg standard errors, hypothesis tests, and estimator properties

Goal of the session

- Main focus in causal inference is often **identification**
- So far, we have: a good question and a convincing quasi-random allocation of the treatment
- How do we make inference for the population?
- We make causal claims based on significance $\Rightarrow$ need **modelling assumptions to hold** and **reliable SEs**
- Aim to give intuition about *some* key points and provide you with resources to learn more about them

Modelling

Modelling assumptions matter

- OLS valid under a set of assumptions: the Gauss–Markov conditions
- If these assumptions, or any modelling one, do not hold we cannot make reliable inference
- Let's see why!
- **Modelling matters**: specification choices affect the results of our studies and our ability to make reliable inference

Gauss-Markov conditions

Assumption	Idea	If violated
1. Linearity	Model linear in its parameters	Estimates misspecified $\Rightarrow$ biased/inconsistent
2. No perfect collinearity	$(X'X)^{-1}$ exists	Coefficients not identifiable
3. Exogeneity	$E[u_i X_i] = 0$	Estimator biased
4.a. Independent errors	$Cov(u_i, u_j X) = 0$ for $i \neq j$ (eg no autocorrelation)	Invalid inference
4.b. Homoskedasticity	$Var(u_i X_i) = \sigma^2 = \text{cst}$	Inefficient

- 4.a + 4.b = **spherical errors**

Implications

- Finite sample properties:
 - 1 → 3: $\hat{\beta}_{OLS}$ unbiased
 - 1 → 4: $\hat{\beta}_{OLS}$ efficient (Best **Linear** Unbiased Estimator)
- Asymptotically:
 - Unbiased
 - Normally distributed
 - Efficient

Normally distributed estimator

- **Required** for making inference (eg computing confidence intervals or p-values)
- Additional assumption: **normal errors** $\Rightarrow \hat{\beta}_{OLS}$ normally distributed
- If errors non-normal
 - Alternative way to compute SE (eg bootstrap)
 - If n is large enough, Central Limit Theorem (CLT) + Weak Law of Large Numbers (WLLN)
 $\Rightarrow \hat{\beta}_{OLS}$ approximately normal

Exercise

- Generate fake data and analyse the impact of violations of some of these assumptions
 1. Non-linearity
 2. Perfect collinearity
 3. Endogeneity
 4. Autocorrelation
 5. Heteroskedasticity
 6. Non-normal errors

Non-linear

Code

Regression table

Graph

```
1 n <- 1000
2 alpha <- 10
3 beta <- 2
4 mu_x <- 2
5 sigma_x <- 1
6 sigma_u <- 2
7
8 data_non_linear <- tibble(
9   x = rnorm(n, mu_x, sigma_x),
10  u = rnorm(n, 0, sigma_u),
11  y = alpha + beta*x^2 + u
12 )
13
14 reg_non_linear <- lm(y ~ x, data = data_non_linear)
15
16 # list("Non-linear" = reg_non_linear) |>
17 #   modelsummary(gof_omit = "IC|Adj|F|RMSE|Log")
```

Collinearity

Code

Perfect collinearity

Almost perfect collinearity

```
1 gamma <- 0.2
2
3 data_collin <- tibble(
4   x = rnorm(n, mu_x, sigma_x),
5   w = 0.3*x,
6   u = rnorm(n, 0, sigma_u),
7   y = alpha + beta*x + gamma*w + u
8 )
9
10 reg_collin <- lm(y ~ x + w, data = data_collin)
11
12 # list("Perfect collin." = reg_collin) |>
13 #   modelsummary(gof_omit = "IC|Adj|F|RMSE|Log")
14
15 data_almost_collin <- tibble(
16   x = rnorm(n, mu_x, sigma_x),
17   w = 0.3*x + rnorm(n, 0, 0.01),
18   u = rnorm(n, 0, sigma_u),
19   y = alpha + beta*x + gamma*w + u
20 )
21
22 reg_almost_collin <- lm(y ~ x + w, data = data_almost_collin)
23
24 # list("Almost perfect collin." = reg_almost_collin) |>
25 #   modelsummary(gof_omit = "IC|Adj|F|RMSE|Log")
```

Endogeneity

Code

Regression table

```
1 data_endog <- tibble(  
2   x = rnorm(n, mu_x, sigma_x),  
3   u = 0.5 * x + rnorm(n, 0, sigma_u),  
4   y = alpha + beta*x + u  
5 )  
6  
7 reg_endog <- lm(y ~ x, data = data_endog)  
8  
9 # list("Endogeneity" = reg_endog) |>  
10 #   modelsummary(gof_omit = "IC|Adj|F|RMSE|Log")
```

Autocorrelation

Code

Regression table

```
1 u_auto <- numeric(n)
2 for (i in 2:n) u_auto[i] <- 0.9*u_auto[i-1] + rnorm(1, 0, sigma_u)
3
4 data_auto <- tibble(
5   x = rnorm(n, mu_x, sigma_x),
6   u = u_auto,
7   y = alpha + beta*x + u
8 )
9
10 reg_auto <- lm(y ~ x, data = data_auto)
11
12 # reg_auto |>
13 #   modelsummary(gof_omit = "IC|Adj|F|RMSE|Log", vcov = c("classical", "HAC"))
```

Heteroskedasticity

Code

Regression table

Graph

```
1 data_heterosked <- tibble(  
2   x = rnorm(n, mu_x, sigma_x),  
3   u = rnorm(n, 0, sigma_u + x^2),  
4   y = alpha + beta*x + u  
5 )  
6  
7 reg_heterosked <- lm(y ~ x, data = data_heterosked)  
8  
9 # reg_heterosked |>  
10 #   modelsummary(gof_omit = "IC|Adj|F|RMSE|Log", vcov = c("classical", "robust"))
```


Non-normal errors

Code

Regression table

```
1 df <- 2
2
3 data_non_normal <- tibble(
4   x = rnorm(n, mu_x, sigma_x),
5   u = rt(n, df), #fatter tails
6   y = alpha + beta*x + u
7 )
8
9 reg_non_normal <- lm(y ~ x, data = data_non_normal)
10
11 # reg_non_normal |>
12 #   modelsummary(gof_omit = "IC|Adj|F|RMSE|Log", vcov = c("classical", "bootstrap"))
```

Limited outcome models

- Often, y is limited: **binary**, **categorical**, **censored**, etc
- Linear regression is not appropriate:
 - It does not take the constraints on y into account
 - *ie* wrongly assumes linearity and errors non-normal
- Use **Generalized Linear Models** (GLMs):
 - Idea: uses an invertible link function g to transform a limited y into a continuous variable
 - $g(\mathbb{E}[y|X]) = \beta_0 + \beta_1 X_1 + \dots + \beta_k X_k$

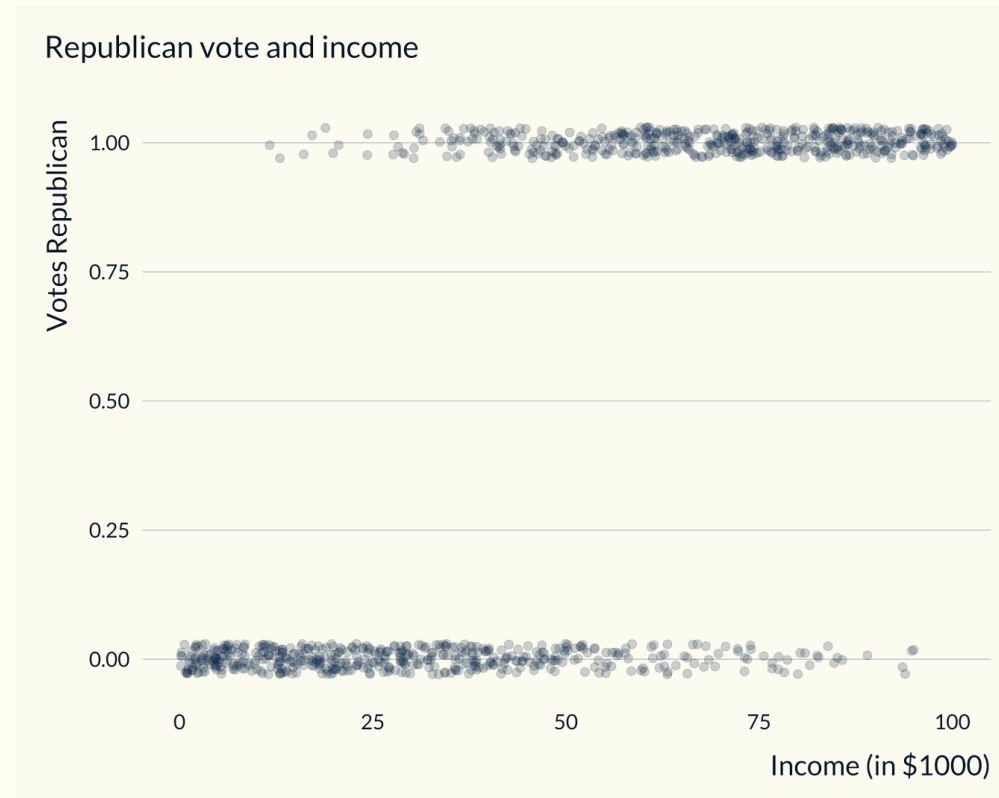
Limited outcome models

Limited y	Example	Regression model
Binary	$y \in \{0, 1\}$	Probit, logit, ...
Count	$y \in \{0, 1, 2, 3, \dots\}$	Poisson, negative binomial, ...
Censored	eg $y = \max(0, y^*)$	Censored regression models (eg tobit)

Binary outcome

Example

- The outcome follows a Bernoulli distribution: $y|X \sim \text{Ber}(\pi)$
- A regression model expresses the conditional probability $\pi = P[y = 1|X]$ as a function of X and β : $\pi = g^{-1}(X'\beta)$



Binary outcome

Models

- **Linear Probability Model**

- $g: x \mapsto x$
- Almost always yields biased and inconsistent estimates

- **Logistic regression model**

- $g: x \mapsto \log\left(\frac{x}{1-x}\right)$, the logit function

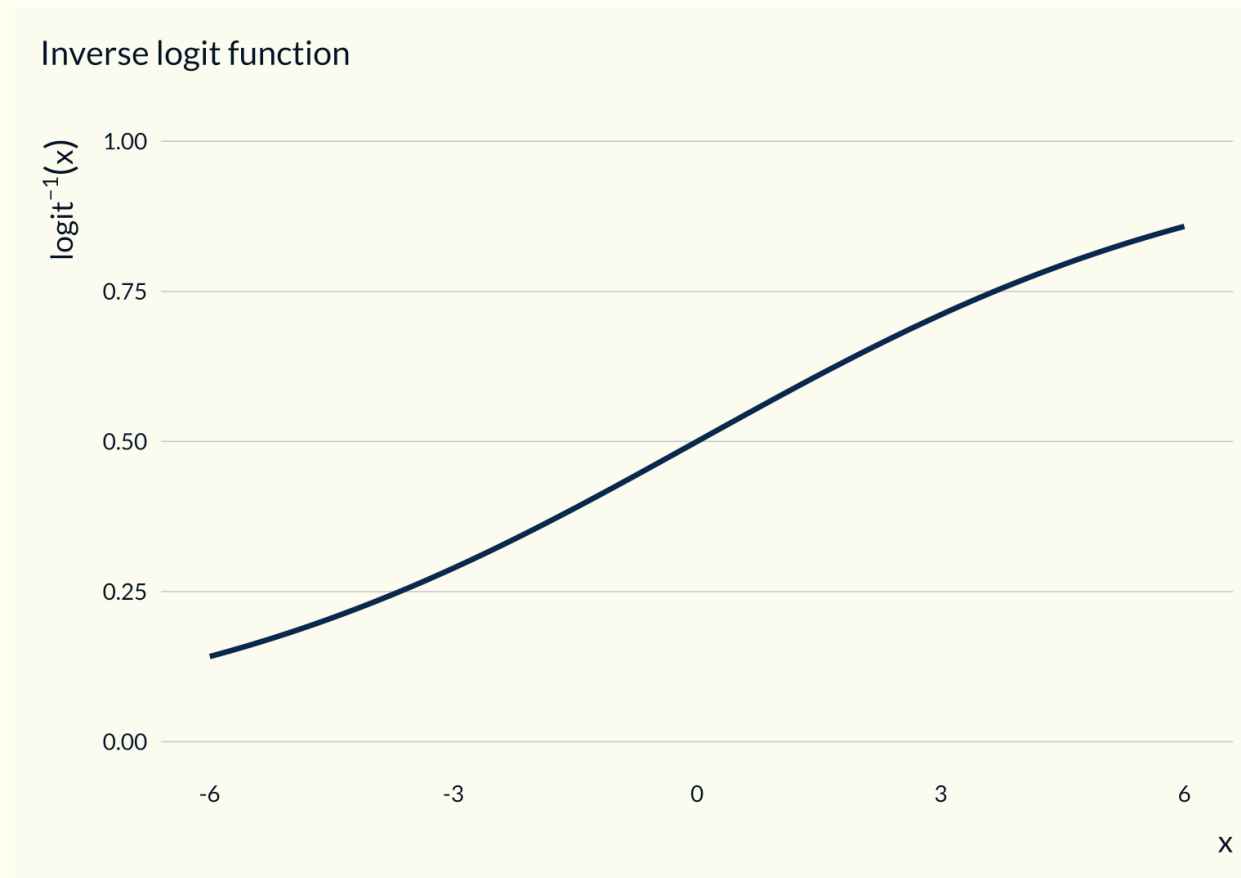
- **Probit regression model**

- g : probit, ie the quantile function of the standard normal distribution (and g^{-1} is the CDF of the normal)
- Rule of thumb: coefs $\simeq$ logit coefs divided by 1.6
- Very similar to logit; sometimes easier to implement

Binary outcome models

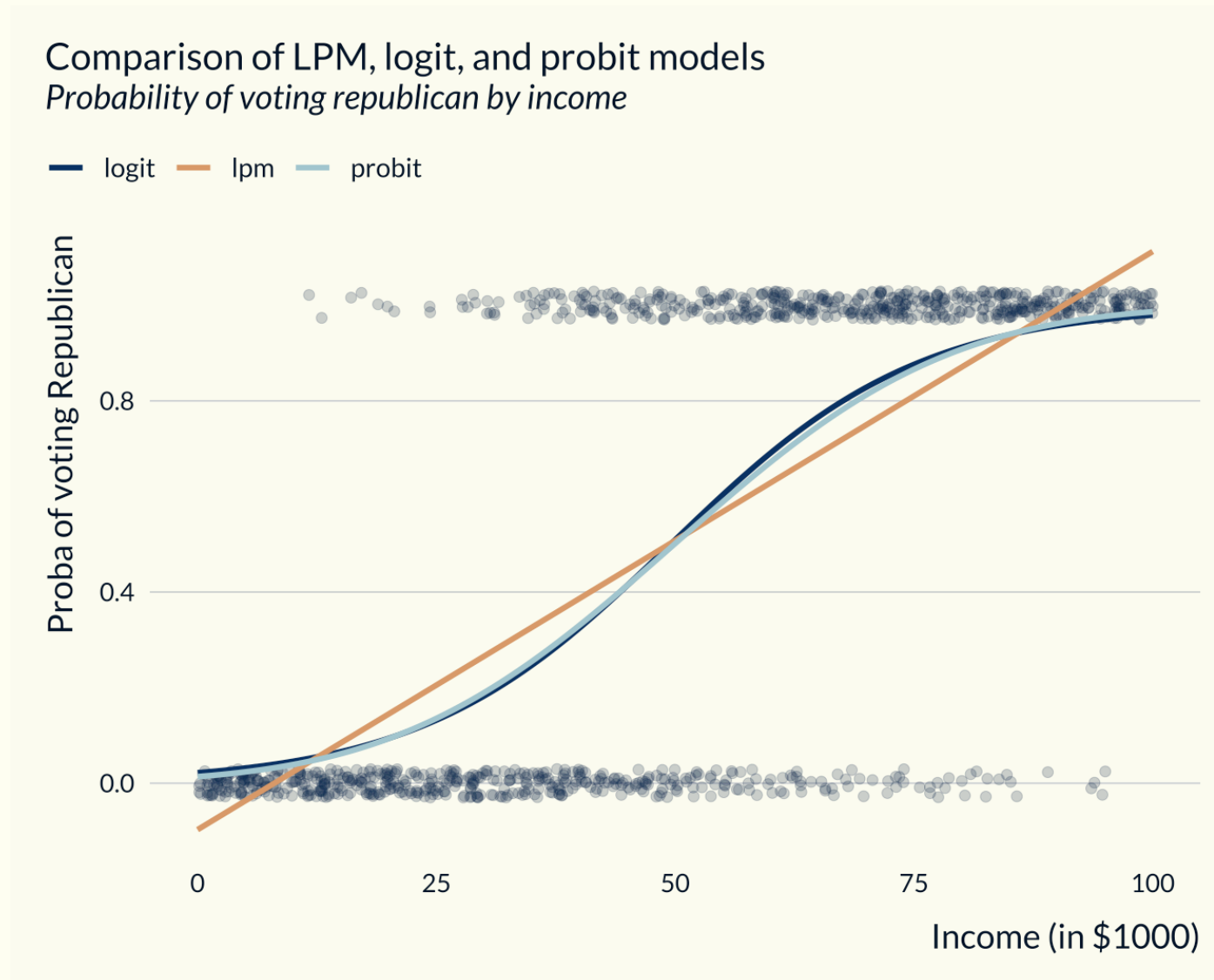
Interpreting coefficients

- β gives the **direction** BUT not the magnitude
- The **marginal effect** depends on X : effects largest in the middle of the distribution



Binary outcome models

Example fits



Count data models

- **Poisson regression model**

- The Poisson distribution $\text{Pois}(\lambda)$ models the number of events occurring in a fixed interval if when events occur independently and at a constant mean rate
- Choice function: $\ln$
- Imposes $\mathbb{V}[y|X] = \mathbb{E}[y|X]$

- **Negative binomial model**

- The negative binomial distribution $NB(p, r)$ models the number of successes in a sequence of iid $\text{Ber}(p)$ trials before r failures occur
- Allows $\mathbb{V}[y|X] \neq \mathbb{E}[y|X]$

A note on controls

- Overall two reasons to include controls:
 1. **To ensure random allocation** of the treatment
 - Necessary for identification
 2. **To improve precision**
 - To better explain y : $\nearrow R^2 (= 1 - FUV)$ and $\searrow \sigma_u^2$
- Adjusting for pre-treatment covariates may
 - Reduce the variation in y , increases precision 👍
 - Reduce the variation in x , decreases precision 👎

Bad controls

Do not adjust for post-treatment variables that may be affected by the treatment

Simulating bad controls

Code

Table

```
1 #reuse the parameters from above
2 n <- 10000 #to limit sampling variation
3 mu_b <- 3
4 sigma_b <- 1
5 gamma <- 5
6 kappa <- 4
7 sigma_a <- 2
8 delta <- 0
9
10 data_bad_control <- tibble(
11   x = rnorm(n, mu_x, sigma_x),
12   u = rnorm(n, 0, sigma_u),
13   b = rnorm(n, mu_b, sigma_b),
14   a = kappa*x + rnorm(n, 0, sigma_a),
15   y = alpha + beta*x + gamma*b + delta*a + u
16 )
17
18 reg_short <- lm(data = data_bad_control, y ~ x)
19 reg_pre <- lm(data = data_bad_control, y ~ x + b)
20 reg_post <- lm(data = data_bad_control, y ~ x + a)
21 reg_pre_post <- lm(data = data_bad_control, y ~ x + b + a)
```

Analysis

SEs: what are they?

- Standard errors = estimate of **sampling variability**
- Tell us how precise estimates are, how much $\hat{\beta}$ would vary across repeated samples
- They determine confidence intervals and p-values
- An example will better illustrate

Illustration of sampling variability

Relationship between education and wage in fake data

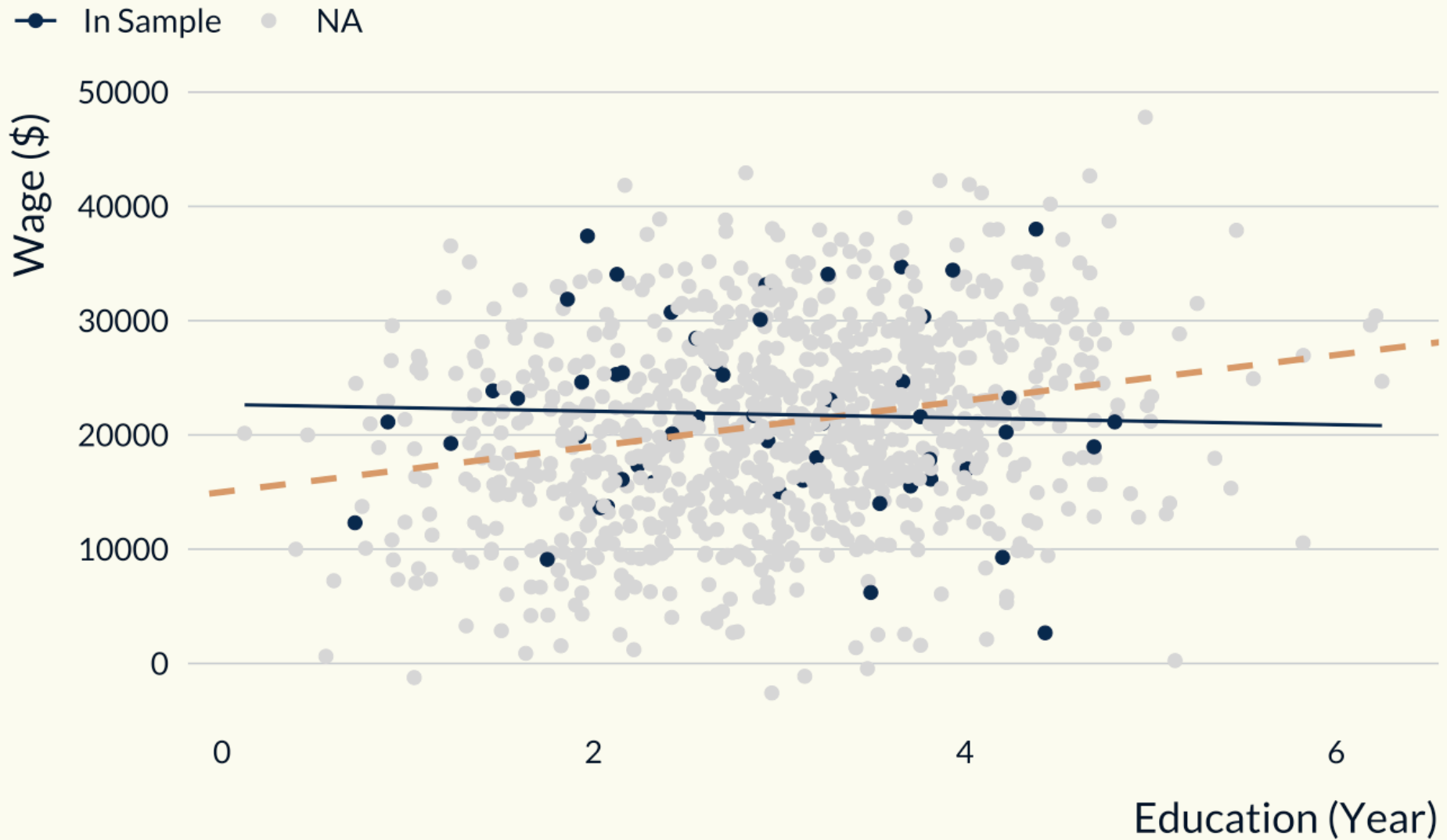


Illustration of sampling variability

Relationship between education and wage in fake data

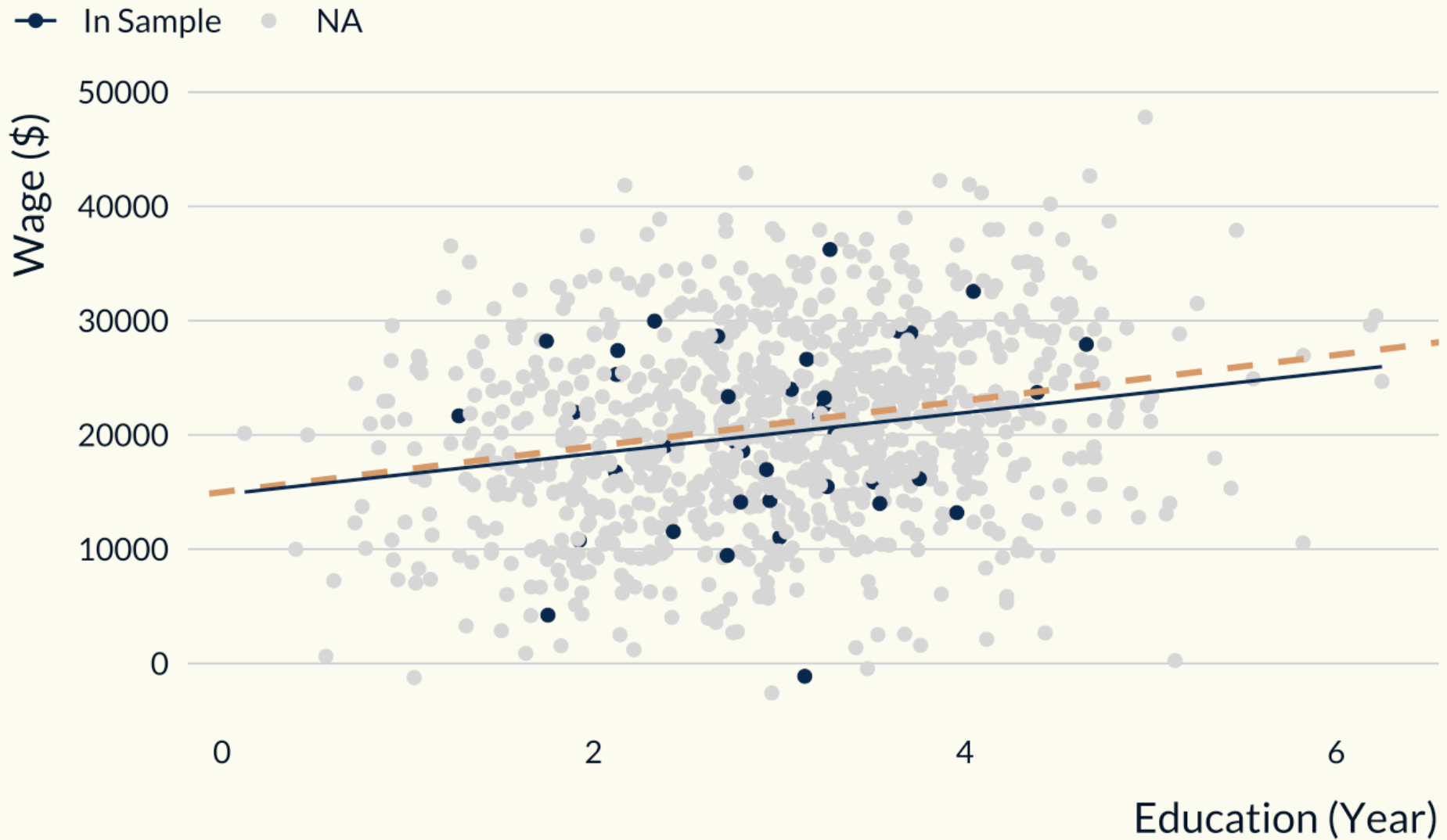


Illustration of sampling variability

Relationship between education and wage in fake data

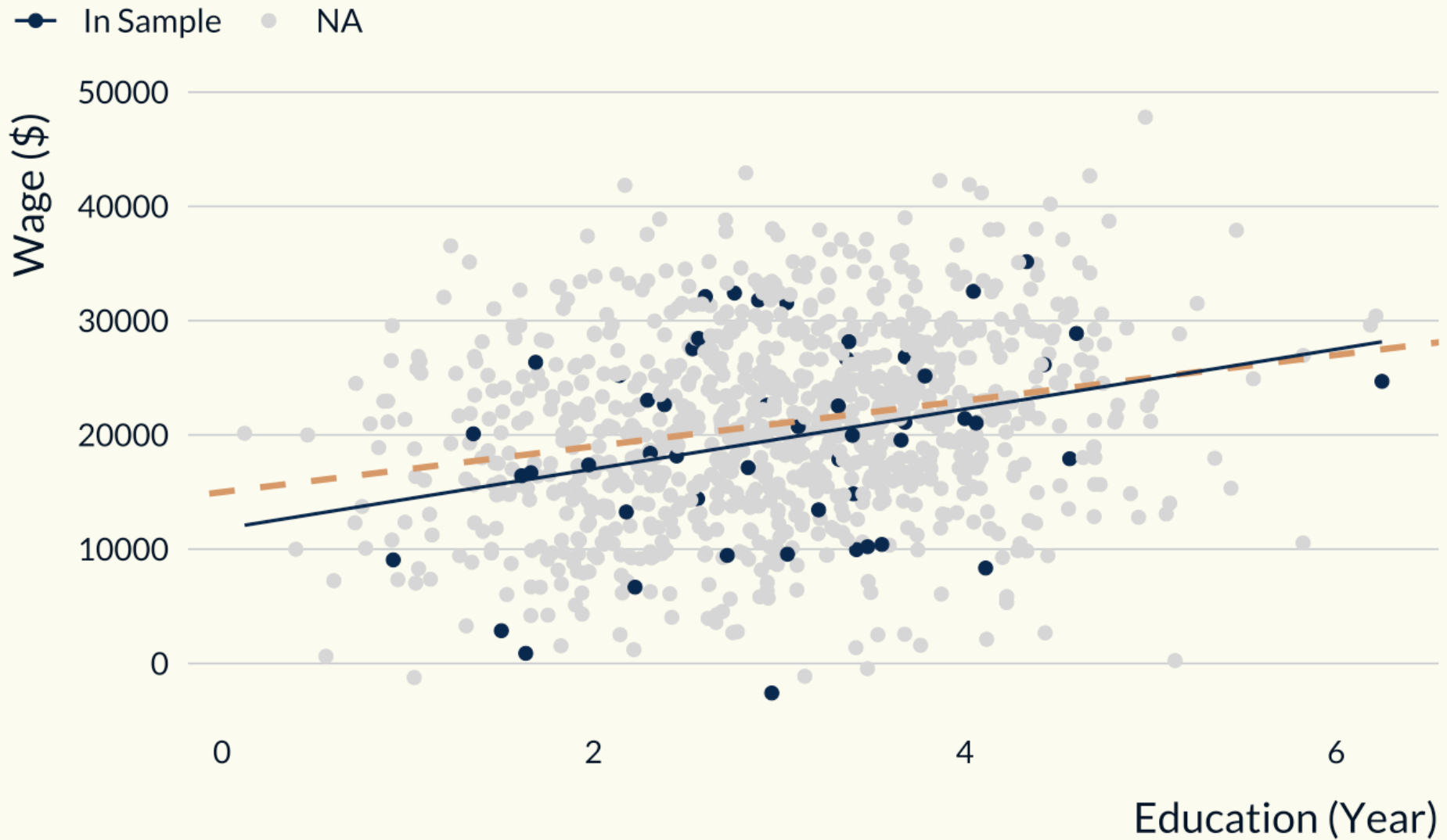
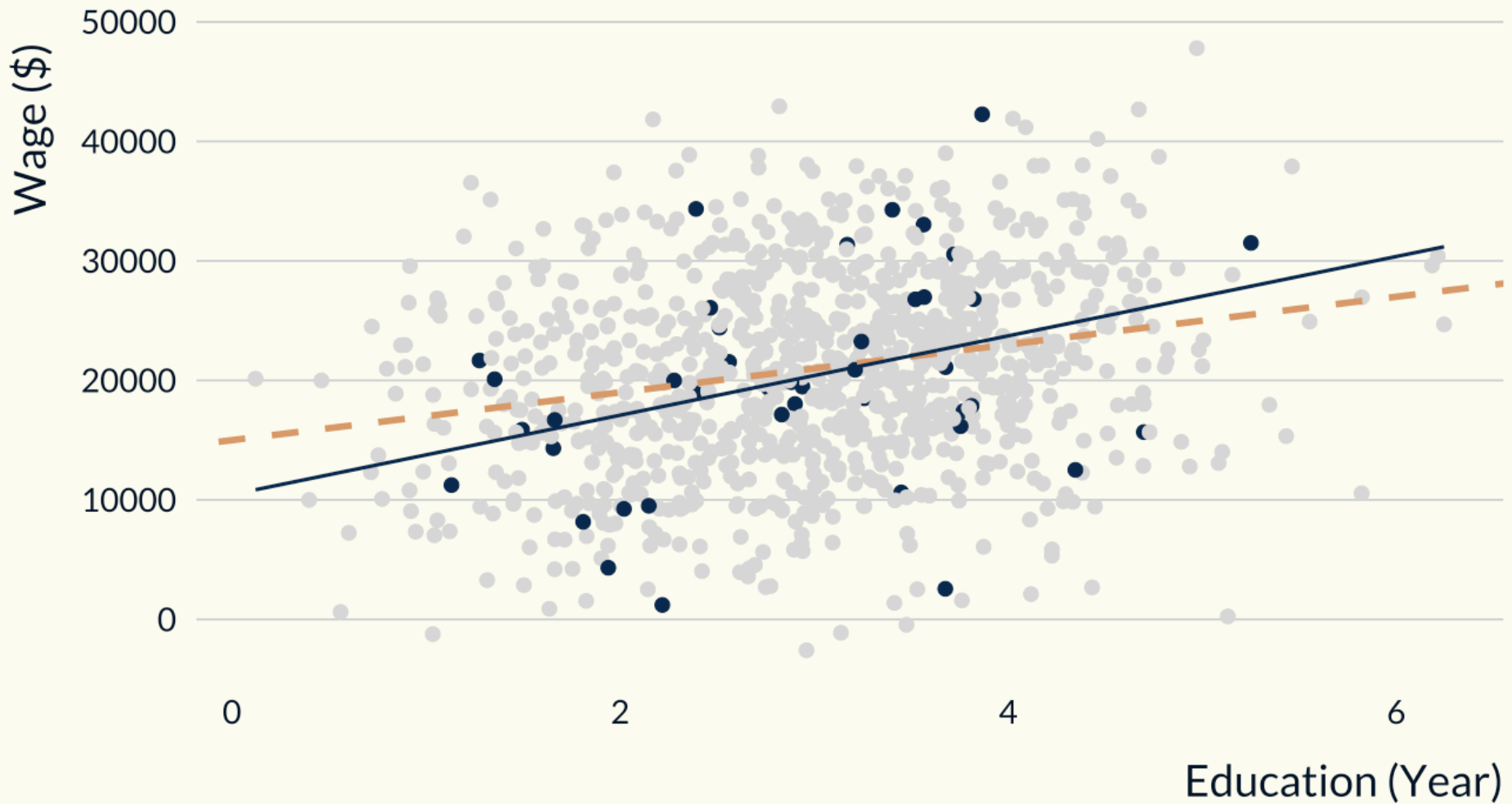


Illustration of sampling variability

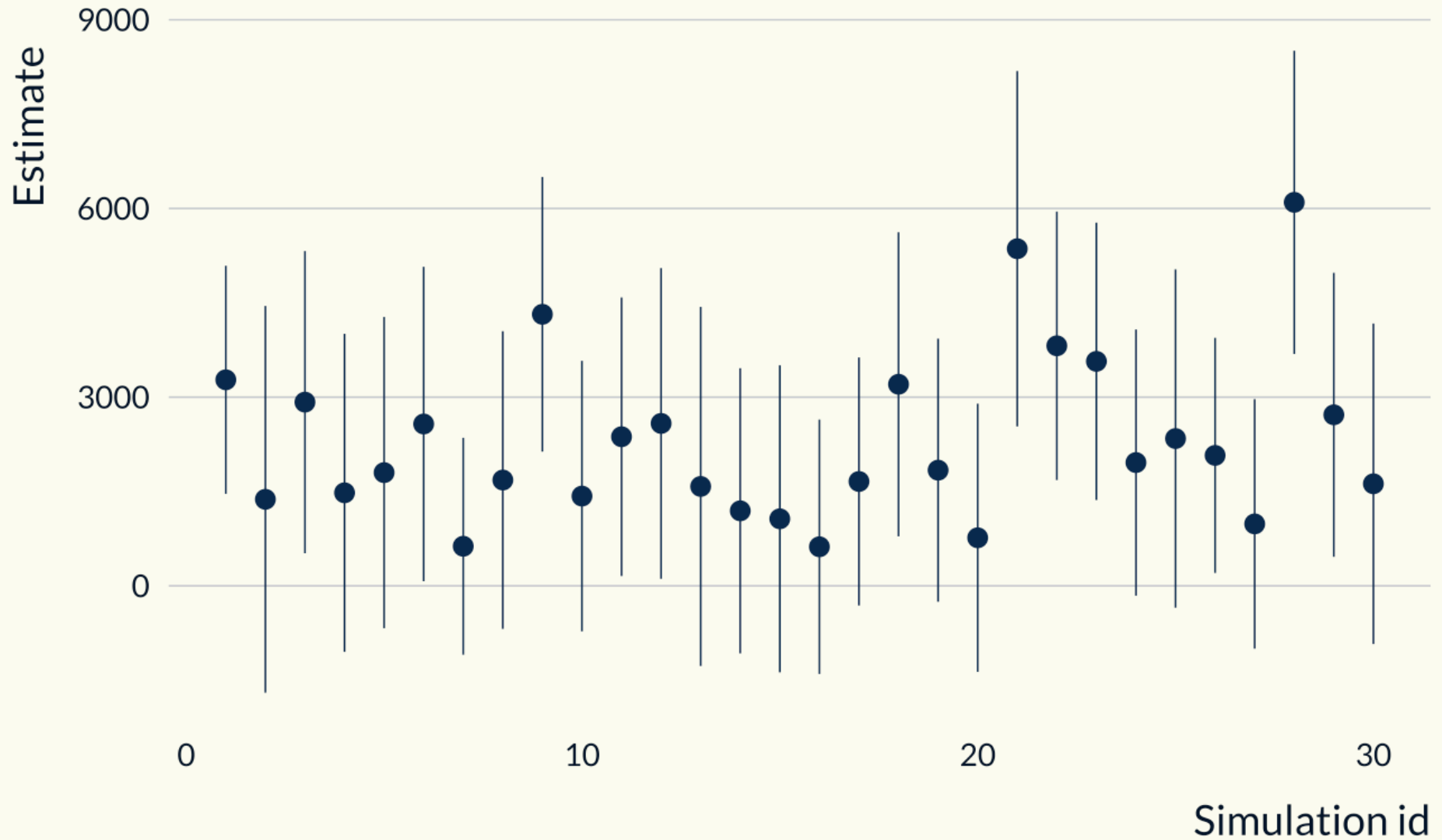
Relationship between education and wage in fake data

● In Sample ● NA



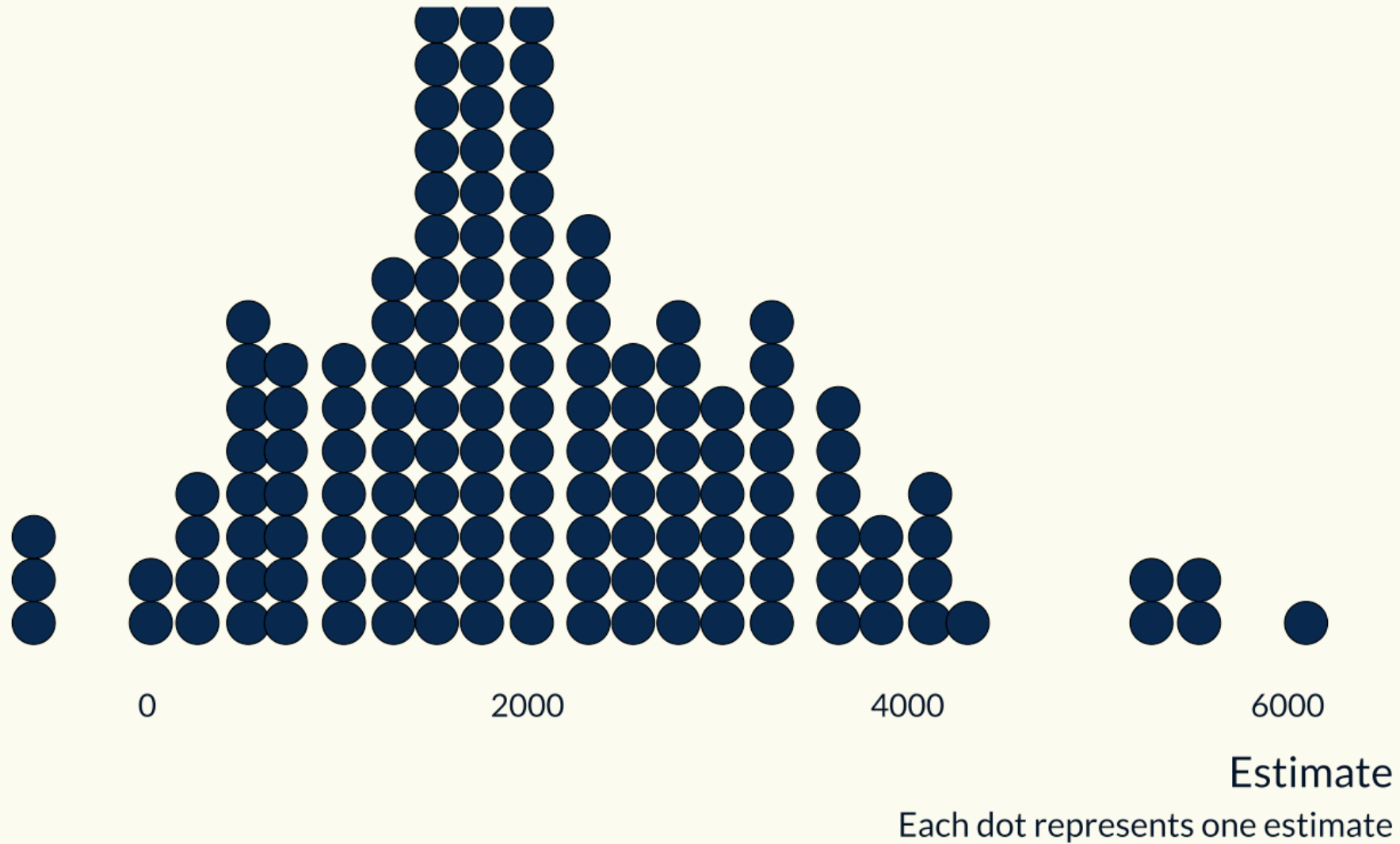
Estimates of the parameter of interest

Computed on 30 different data sets



Distribution of 150 estimates

Computed on 150 different samples



Why care about them?

- Violations of classical assumptions:
 - **Heteroskedasticity**: variance depends on X
 - **Non-independence**: errors correlated within groups
- Implications:
 - t-tests and CIs misleading
 - In general, SEs too small $\Rightarrow$ inference too optimistic

(Non-)spherical errors

- Under assumptions 1 $\rightarrow$ 3 we saw earlier, the asymptotic distribution of $\widehat{\beta}_{OLS}$ is

$$\widehat{\beta}_{OLS} \overset{a}{\sim} \mathcal{N}(\beta_0, (X'X)^{-1}X'\Sigma X(X'X)^{-1})$$

- If spherical errors (ie $\Sigma = \sigma^2 I$), use unbiased sample variance
- If non-spherical errors, need a covariance matrix estimator that is consistent under this misspecification
- $\Rightarrow$ use **sandwich estimators** of the variance $(\underbrace{(X'X)^{-1}}_{bread} \underbrace{X'\widehat{\Sigma}X}_{meat} \underbrace{(X'X)^{-1}}_{bread})$
- Heteroskedasticity: compute **White SEs**
- Autocorrelation: compute **Newey-West/Conley SEs** if correlated in time/space

Why clustering SEs

- When errors are correlated within groups (e.g. individuals, firms, regions), clustering adjust for this

$$\mathbb{E}[u_i u_j | X] \neq 0 \text{ for } i, j \text{ in the same cluster}$$

- **Examples**

- Panel data: repeated measures of same unit
- Group-level treatment, eg policy at regional level with many individuals per region

- **Intuition**

- We do not have N independent observations but G clusters
- eg 10 classrooms of 30 students does not correspond to 300 independent observations
- The real sample size is closer to the number of clusters, not observations

What clustering does

- Do not affect point estimates
- Allows for intra-cluster correlation and adjusts the effective number of independent observations
- Increases SEs to reflect within-group dependence $\Rightarrow$ wider CIs

$$\widehat{\text{Var}}_{CR}(\widehat{\beta}) = (X'X)^{-1} \left(\frac{1}{n-p} \sum_{g \in G} X_g' r_g r_g' X_g \right) (X'X)^{-1}$$

where X_g and r_g are data and residuals for cluster g

- Intuition: each cluster provides one independent piece of information about how residuals co-evolve with regressors

Which level to cluster at?

- Clustering accounts for correlation in residuals
- $\Rightarrow$ the level of clustering depends on **where correlation comes from**
- Rule of thumb:
 - Cluster at the level of the shock or treatment variation
 - If unsure, cluster higher rather than lower
- If level of clustering too low, SEs too small, overconfidence in results
- If level of clustering too large, SEs too large, under reject

When is clustering difficult to implement?

- When **few clusters** (eg < 30) $\Rightarrow$ asymptotic results unreliable
- When complex dependence or unclear correlation structure
- When model residuals violate independence in unknown ways
- Solution
 - **Resampling** methods (eg bootstrap)
 - Hierarchical clustering

Bootstrap

- **Intuition:** simulate the sampling variability by resampling from our data
- **Steps:**
 1. Draw samples (**with replacement**)
 2. Estimate $\hat{\beta}$ on each sample
 3. Compute the SD of $\hat{\beta}$ across replications
 4. That SD = bootstrap SE
- When resampling, respect the correlation structure: **sample clusters**

Summaries

Summary of today

- Regression models and OLS estimation rest on a set of **assumptions**
- Ensure estimates are unbiased, efficient, and statistically valid
- When fail, estimates are biased and standard errors misleading
- We reviewed these modelling assumptions and discuss **what happens when they fail**
- Limited outcome models or clustering help **overcome** the issues

Goal of the whole course

Give us a deeper understanding of:

- How regression works “**under the hood**”: **intuition**
- Causal identification strategies and their assumptions
- How **design**, **modeling**, and **analysis** choices shape empirical results
- Common pitfalls and challenges in empirical work
- How to use **simulations** to explore estimator behavior and diagnose potential problems specific to your own cases
- Existing **references** and where to find additional information on a specific topic

To mention when pitching your analysis

1. Research question

- What causal effect of interest are you trying to estimate?

2. Ideal experiment

- What ideal experiment would capture the causal effect?

3. Identification strategy

- How are the observational data used to make comparisons that approximate such an experiment?

4. Estimation method (including assumptions made when constructing standard errors)

5. Falsification tests that support the identifying assumptions

To also mention in your pitch

- **Motivation**
 - Why is your research question important?
- **Contributions to the literature**
- **Methodological contributions**
- **Internal validity**
 - Are the identifying assumptions plausible?
 - Are there unexplained results?
- **External validity**
 - Gap between policy questions and the analyses performed?
 - Generalization to other populations and settings?

Summary of the entire course

Take away messages

- Your **research question** and underlying **theory** are crucial
- Think about what is the **identifying variation** in your model:
 - What are you estimating exactly with your model?
 - Which observations contribute to identification?
- Always wonder what you are **comparing**
- Use **simulations** to explore and understand points

References