

Handling Large Scale Unit Level Data Using STATA

Professor Pratap C. Mohanty

Department of Humanities and Social Sciences

Indian Institute of Technology, Roorkee

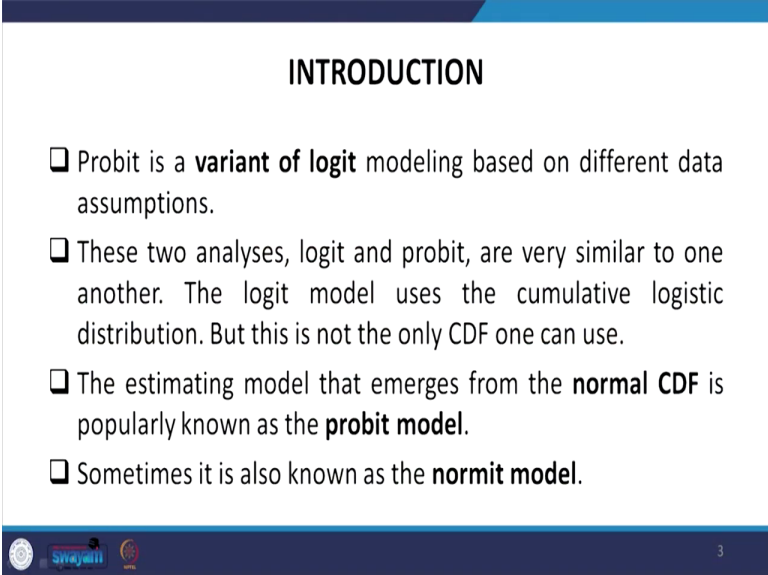
Lecture 35

Binary Response Models - IV

Welcome once again to this module of Handling Large Scale Data Using STATA. So, far we have been trying to understand ways to go into the deeper understanding of quantitative analysis of qualitative data. So, in the last 3 lectures we have been explaining the qualitative variables, qualitative dependent variables. And LPM and Logit has already been discussed but now onwards we are going to discuss, in this particular lecture we are going to discuss Probit and Tobit model.

Probit where we wanted to understand whether the error distribution or the distribution function follows normality or not. In Tobit somewhere we are going to discuss whether there is censoring of that particular data. and we will also compare all those 3 models together; whether LPM, Logit and Probit, which one is better, why it is better?

(Refer Slide Time: 01:34)



INTRODUCTION

- ❑ Probit is a **variant of logit** modeling based on different data assumptions.
- ❑ These two analyses, logit and probit, are very similar to one another. The logit model uses the cumulative logistic distribution. But this is not the only CDF one can use.
- ❑ The estimating model that emerges from the **normal CDF** is popularly known as the **probit model**.
- ❑ Sometimes it is also known as the **normit model**.

3

So, coming to the understanding of Probit regression so I will take a quick wrap up of the understanding since many aspects of Logit and distribution function has already been explained. So, Probit is a variant of Logit, modeling based on different data assumptions. These two analyses that is Logit and Probit are very similar to one another but the Logit uses a cumulative

distribution function whereas the Probit uses normal density function. So, coming to the estimates of the Probit, since it follows a normal density function we wanted to mention that Probit model is also called as a normit model.

(Refer Slide Time: 2:34)

- ☐ The term 'probit' was coined in the **1934** by **Chester Bliss** and stands for **probability unit**. and a fast method for computing maximum likelihood estimates for them was proposed by **Ronald Fisher** in an **appendix to Bliss 1935**.
- ☐ The name is from **probability + unit**.
- ☐ The purpose of the model is to estimate the probability.
- ☐ The parameters of interest in probit model model can be estimated through **maximum likelihood** method.
- ☐ Probit is significantly more sensitive to outliers than logistic sigmoid.

- ☐ Probit model is used to model dichotomous or binary outcome variables. In the probit model, the **inverse standard normal distribution of the probability** is modeled as a linear combination of the predictors.

This was initially coined by Chester Bliss in 1934. Then the Probit stands for probability unit. And basically the first maximum likelihood estimates were initially proposed by Ronald Fisher in an appendix to Bliss of 1935. So, the name is from probability plus unit. That is why it is called Probit, in short it is called probability unit.

The purpose of this model is to estimate the probability value. The parameter of interest in Probit model can be estimated through maximum likelihood method also. Probit is significantly more sensitive to outliers than that of Logit because distribution of the probability change in the outliers is very less in case of logistic distribution whereas the probability change in case of the outlier in case of Probit distribution very sharp.

The Probit model is used to model dichotomous and binary outcome variables. In the Probit model the inverse standard normal distribution of the probability's model is a linear combination of the predictors. So, inverse of the standard normal can be modeled in terms of linear combination.

(Refer Slide Time: 04:15)

THE PROBIT REGRESSION MODEL

- ☐ Response variable Y binary, can take only two possible outcomes 1 and 0.
- ☐ Here, Y represents presence (1) and absence (0) of certain condition say, person in labour force (1) and not in labour force (0).
- ☐ Y is influenced by the set of regressors X .
- ☐ We assume that the model takes form:
$$P(Y = 1|X) = F(X\beta)$$

6

□ Where,

P denotes probability,

F denotes cumulative distribution function (CDF) of the standard normal distribution.

Parameters β 's are typically estimated by maximum likelihood method.

□ It is also possible to motivate the probit model as a **latent variable model**.

□ Suppose there exist a latent variable (auxiliary random variable).

□ $Y^* = X\beta + \varepsilon$

Where Y^* is unobserved and $\varepsilon \sim N(0, \sigma^2)$

□ The Y can be viewed as an indicator for whether this latent variable is positive:

$$Y = 1, \text{ If } Y^* > 0 = X\beta + \varepsilon > 0 \text{ i.e. } \varepsilon > -X\beta$$

0 otherwise.

$$P(Y^* > 0 | X) = P(Y = 1 | X) = P(\varepsilon > -X\beta)$$

$$= P\left(\frac{\varepsilon}{\sigma} > \frac{-X\beta}{\sigma}\right)$$

Probit regression model especially for the response variable Y binary can take only 2 possible outcomes that is 1 and 0. Here Y represents 1 as presence and absence with 0 of certain conditions such as person in labor force to be coded as 1 and not in the labor force to be 0. And Y is influenced by set of regressors that is X.

We assume now the model takes the form like Y with the probability of success given X is nothing but the probability density function or the probability function, cumulative function of X beta where P denotes probability and F denotes the cumulative distribution function of the standard normal distribution. Parameter betas are typically estimated by maximum likelihood

method. Here we also wanted to mention that it is possible to motivate the Probit model as a latent variable model since probabilities of the variables are estimated the probabilities and its expected values could be defined as not the exact values rather representative values called latent variable model.

Suppose there exists a latent variable or auxiliary random variable then Probit can be also most fitted. Then the model looks like this. Y^* is the unobserved variables with error distribution of epsilon with 0 and Sigma square as standard variance. Then Y_i 's viewed as an indicator of whether this is a latent variable is positive if Y is equal to 1. And if Y^* is greater than 0, that is nothing but this is greater than 0, if and only if that is the case then the error will be greater than that of the negative value of the estimated values. if it is 0 then in that case the probability of error values could be estimated till that of the XB limit.

So, basically in the density function or the probability density function, if you can estimate till that particular time is more important. Based on this equation if you simply divide the standard deviation in both the side the probability limit is defined here. If any specific value till that value if you are calculating and their probabilities you are estimating the probability limit of this error is estimated through this. We are going to mention.

(Refer Slide Time: 07:43)

$$= 1 - F\left(\frac{-X\beta}{\sigma}\right) = F\left(\frac{X\beta}{\sigma}\right) \text{ if } F \text{ is symmetric}$$

$$P(Y = 1 | X) = F\left(\frac{X\beta}{\sigma}\right) = F\left(\frac{\beta_0 + \beta_1 X}{\sigma}\right)$$

Where, F is the standard normal cumulative density function for ϵ . We can't estimate both β and σ , since they enter the equation as a ratio, so, we set $\sigma=1$, making the distribution on ϵ a standard normal density.

$$P(Y = 1) = F(\beta_0 + \beta_1 X) = \int_{-\infty}^{\beta_0 + \beta_1 X} \frac{1}{\sqrt{2\pi}} e^{-t^2/2} dt$$

Where, t is the standard normal variable i.e., $t \sim N(0, \sigma^2)$

So, in the total density function, cumulative density function since the entire probability is 1, so 1 minus till is nothing but the cumulative probabilities of X beta upon sigma. And if these two are equal that means the distribution is perfectly symmetric. If the both sides are equal, this side as well as this side, basically probability of success in one side as compared to another side if they are equal that means the distribution is bell-shaped and symmetric.

So, probability of success given the explanatory variables is nothing but the density function or the cumulative distribution function. Denominator with F nothing but X beta upon sigma. So, that is basically beta not plus beta 1 X , we are offering here, divided by sigma. So, this is nothing but, F is nothing standard normal cumulative density function for the level of error values.

We cannot estimate both beta and sigma since they enter the equation as a ratio. So, here they are coming as a ratio. It is not individually separated, alright. So, we set sigma is equal to 1 making the distribution on epsilon or error a standard normal density function. When this is 1 so basically we are referring to a standard normal density function.

If the probability value with the success of occurrence in the dependent variable is 1, probability of success is 1 that means we end up with this estimation. So, we are basically estimating till that level this estimated value with the Z value of minus infinity to infinity, probability values from 0 to 1.

So, we are estimating on the diagram till this level which is equivalent to our estimated value in the model. Since it follows a normal density function we are referring to a normal density function equation. This is a normal density function and since integration we are using so the delta change in that distribution is also attached. So, 1 upon square root of 2π e to the power minus, usually it is to be Z^2 upon half of, Z^2 upon 2 but here T distribution is considered because of the limiting in the sample. So, the density function equation is clearly given for the estimation, where t is the standard normal variable with its distribution follows 0 and sigma square.

(Refer Slide Time: 10:51)

- Now we can obtain information on probability $P(Y=1)$ and β 's:

$$F^{-1}(P) = \beta_0 + \beta_1 X$$

The inverse transformation is called the **probit**, which gives the linear predictor as a function of the probability.

We can obtain information on probability that is of occurrence 1. That is the inverse of the probability is nothing but the linear function of it. That can be estimated. The inverse of the probabilities will give a linear equation and the interpretation whatever the probabilities values you get, if we take inverse value of it is basically called beta 1 or beta not. The inverse transformation is called the Probit value. Basically the P value is nothing but the inverse transformation which gives the linear predictor of a function of the probability.

(Refer Slide Time: 11:35)

LOGIT Vs. PROBIT

- Neither the logit model nor the probit model are linear, which makes things difficult. To make the model linear, a transformation is done on the dependent variable.
- Both methods use maximum likelihood, and so require more cases than a similar OLS model. Unlike logit models, we don't get odds ratios with probit models.
- Probit regression is an alternative log-linear approach to handling categorical dependent variables.

- ❑ Logit and probit models generally give similar results.
- ❑ The main difference between the two models is that the logistic distribution has slightly **fatter tails**. The conditional probability P_i approaches 0 or 1 at a **slower rate** in logit than in probit.
- ❑ In practice there is no compelling reason to choose one over the other. Many researchers choose the logit over the probit because of its comparative **mathematical simplicity**.
- ❑ Another difference is when the error term follows **logistic distribution function**, **logit model** is appropriate. If it follows **standard normal distribution**, **probit model** is appropriate.

Let us compare logit, probit and LPM together, let us first compare logit model or probit model both are not linear. So, which makes things difficult to make the model linear we make some form of transformation. So, log transformation or inverse transformation in either of the case that is on the dependent variable is going to make it linear.

So, both methods use maximum likelihood estimator this is one, so linearity aspect I discussed, transformation is important, so we require more cases than the similar OLS model. In order to make likelihood estimator, it usually goes by iteration, iteration requires more number of cases. OLS less number of cases can derive a result a better result but both the models require, more number of cases that is another important aspect to be mentioned.

Logit model do not mention online logit models we do not get odds ratio with probit models. So, probit never gives odds ratio it gives the probability value. Probit regression is an alternative log linear approach to handling categorical dependent variables. Question always comes which one to be taken, we have it clearly, that there is no such major difference except the distribution function. Logit and probit models give similar results.

The main difference between the two models is that the logistic distribution has slightly fatter tails. So, Logit has fatter tails. The tails are more fatter whereas in case of Probit the tails are very sharp, and when it approaches to 0 or 1. The end is at a very slower rate. The conditional probability P approaches 0 or 1 at a very slower rate in Logit than that of Probit. Since the tails are flatter approaching very fatter at the end the slope is very very less that is the reason why the

change the end nearby 0 and 1 in case of Logit it is very slow whereas in case of Probit there is a sharp change.

In practice there is no completing reason to choose over the other. That is important. And do not get confused. In practice there is no such hard and fast rule to go by any single model many researchers choose the Logit over Probit because of its comparative mathematical simplicity. Logit is preferred. Another difference is when the error term follows logistic distribution, Logit model is appropriate whereas if the distribution is of standard normal function Probit is most appropriate.

(Refer Slide Time: 14:44)

MARGINAL EFFECT OF THE PROBIT MODEL

- ❑ The drawback of probit model is that the coefficients are more difficult to interpret, hence they are less used.
- ❑ The marginal effect of changing X on $\hat{P}$, the probability of getting $Y = 1$.
- ❑ For the probit model,

$$P(Y = 1|X) = F(\hat{\beta}_0 + \hat{\beta}_1 X)$$
$$\frac{\partial P(Y = 1|X)}{\partial X} = F(\hat{\beta}_0 + \hat{\beta}_1 X) \hat{\beta}_1$$

13

- The effect of the change in X on $P(Y = 1|X)$ depends on the value of X. In practice, we usually evaluate the *marginal effect* at the sample average $\bar{X}$. i.e. The marginal effect is

$$F(\hat{\beta}_0 + \hat{\beta}_1 \bar{X}) \hat{\beta}_1.$$

- When X is binary, it is not clear what does the sample average mean.
- The marginal effect then measures the probability difference between $X = 1$ and $X = 0$.

$$P(Y = 1 | X = 1) - P(Y = 1 | X = 0)$$

$$F(\hat{\beta}_0 + \hat{\beta}_1) - F(\hat{\beta}_0)$$

Marginal effect of the Probit model and Logit we have already discussed by taking first derivative. Similarly in this case also P value is expressed with this. We are taking dP/dX . beta 1 will separate due to dX and F prime or F, you can write down a F prime or F, it is up to you, how you are writing to indicate the first derivative.

The effect of change in X on P depends on the value of X. In practice we usually evaluate the marginal effect at the sample average of $\bar{X}$ that is marginal effect conditioned upon the sample mean of each of the variable. When X is binary it is not clear what does the sample average mean? When the X is binary it is very difficult to interpret.

So, the marginal effect then measures the probability difference between 1 and 0. So, that is basically probability of 1 given X equal to 1 minus probability of 1 when it is 0. So, 0 and 1, if you compare the probability difference that is basically called the marginal effect. The marginal value we are going to get in this case is nothing but, if it is categorical it is not the average that matters rather the probability of occurrence in both the case. when you subtract it is basically marginal effect.

(Refer Slide Time: 16:18)

CAUTION FOR COMPARING MODELS

- ❑ Though the models (LPM, logit and probit) do the same thing (estimating probabilities), one has to be careful in interpreting the coefficients estimated.
- ❑ The coefficients are not directly comparable.
- ❑ But there are certain ways to compare the models. By multiplying coefficients by certain value:

$$\begin{aligned}\beta_{logit} &= 4 \beta_{LPM} \\ \beta_{LPM} &= 0.25 \beta_{logit} \\ \beta_{logit} &= 1.6 \beta_{probit}\end{aligned}$$
$$\begin{aligned}\beta_{probit} &= 0.625 \beta_{logit} \\ \beta_{probit} &= 2.5 \beta_{LPM}\end{aligned}$$


































































































































































































































































































































































































































(Refer Slide Time: 18:09)

EXAMPLE: PROBIT REGRESSION MODEL

- ❑ We are using the same dataset (BCM_practice.dta) and same variables (check lecture on LPM for variable description) as in LPM model.
- ❑ running the same model (i.e. what factors explain the choices of women being a necessity entrepreneur) with probit regression approach.



16

We can also compare our coefficient by taking all the coefficients together. So, we have a practice dataset that is BCM practice we are going to provide you. We are going to check all three models together and running those commands for the same dependent variable that is women being necessity entrepreneur with Probit regression approach as the first. Then we will compare.

(Refer Slide Time: 18:49)

- ❑ use `"BCM_practice.dta", clear`

```
probit enterprise_type enterprise_age age_square i.sector  
i.mix_activity i.prblm_facd i.assistance_rcvd i.trad_sec  
i.southstates i.location_entprise i.social_category  
i.bank_account i.growth_status i.activity_group
```

- ❑ The basic `probit` commands report **coefficient estimates** and the **underlying standard errors**. These coefficients are the **index coefficients** and what we can only say is the **direction of the effect** and **partial effects** on the Probit index/score. They do not correspond to the average partial effects.



17

Then we will discuss their coefficient estimation, their standard errors then direction of the coefficient and the partial effects. We will discuss some of those things right now, alright.

(Refer Slide Time: 19:28)

The screenshot shows the Stata 15.1 interface with the command window displaying the following results:

```

. probit enterprise_type enterprise_age age_square i.sector i.mix_activity i.pshin_fend i.assistance_cvd i.trad_sec i.southstates
      > i.location_enterprise i.social_category i.bank_account i.growth_status i.activity_group

Iteration 0: log likelihood = -3203.1881
Iteration 1: log likelihood = -2138.8495
Iteration 2: log likelihood = -2105.7635
Iteration 3: log likelihood = -2105.709
Iteration 4: log likelihood = -2105.709

Probit regression      Number of obs   =    6,492
                      LR chi2(128)      =   2194.96
                      Prob > chi2       =    0.0000
                      Pseudo R2        =    0.3426

Log likelihood = -2105.709

      enterprise_type      Coef.   Std. Err.      z    P>|z|     [95% Conf. Interval]
-----+-----
enterprise_age      0.000388   0.000012   3.08   0.000   -0.000083   0.000476
age_square      -0.001495   0.000210   -0.81   0.420   -0.0009817   0.0024827
sector
  urban      -1.348293   0.462944   -2.90   0.003   -2.236051   -0.460534
mix_activity
  No      0.094659   1.185742   0.08   0.933   -2.130604   2.321422
  2.pshin_fend      1235644   0.403587   3.11   0.002   0.568234   2503090
  2.assistance_cvd      9817468   2135514   0.38   0.702   -3367864   5003299
  trad_sec      1833633   6935709   1.97   0.049   -5005471   347359
  nontraditional      0775424   5009774   1.52   0.128   -9223515   1774744
location_enterprise
  outside      -1.331322   0.489371   -2.71   0.008   -2.284481   -0.378163
social_category
  SC      4478568   6971108   0.64   0.520   -777823   6581905
  OBC      248546   9702071   0.03   0.980   -1727278   409804
  others      -0.224247   0.711503   -0.31   0.752   -1.428788   1.160253
bank_account
  enterprise_name      -1.071529   1.098005   -0.97   0.330   -3.24889   1.105832
  both      -1.248042   2240594   -0.00   0.999   -4.68112   3.185036
  noaccounts      3994478   0.487349   8.21   0.000   0.3994478   0.487349
  
```

The screenshot shows the Stata 15.1 interface with the command window displaying the following results:

```

. probit enterprise_type enterprise_age age_square i.sector i.mix_activity i.pshin_fend i.assistance_cvd i.trad_sec i.southstates
      > i.location_enterprise i.social_category i.bank_account i.growth_status i.activity_group

Iteration 0: log likelihood = -3203.1881
Iteration 1: log likelihood = -2138.8495
Iteration 2: log likelihood = -2105.7635
Iteration 3: log likelihood = -2105.709
Iteration 4: log likelihood = -2105.709

Probit regression      Number of obs   =    6,492
                      LR chi2(128)      =   2194.96
                      Prob > chi2       =    0.0000
                      Pseudo R2        =    0.3426

Log likelihood = -2105.709

      enterprise_type      Coef.   Std. Err.      z    P>|z|     [95% Conf. Interval]
-----+-----
enterprise_age      0.000388   0.000012   3.08   0.000   -0.000083   0.000476
age_square      -0.001495   0.000210   -0.81   0.420   -0.0009817   0.0024827
sector
  urban      -1.348293   0.462944   -2.90   0.003   -2.236051   -0.460534
mix_activity
  No      0.094659   1.185742   0.08   0.933   -2.130604   2.321422
  2.pshin_fend      1235644   0.403587   3.11   0.002   0.568234   2503090
  2.assistance_cvd      9817468   2135514   0.38   0.702   -3367864   5003299
  trad_sec      1833633   6935709   1.97   0.049   -5005471   347359
  nontraditional      0775424   5009774   1.52   0.128   -9223515   1774744
location_enterprise
  outside      -1.331322   0.489371   -2.71   0.008   -2.284481   -0.378163
social_category
  SC      4478568   6971108   0.64   0.520   -777823   6581905
  OBC      248546   9702071   0.03   0.980   -1727278   409804
  others      -0.224247   0.711503   -0.31   0.752   -1.428788   1.160253
bank_account
  enterprise_name      -1.071529   1.098005   -0.97   0.330   -3.24889   1.105832
  both      -1.248042   2240594   -0.00   0.999   -4.68112   3.185036
  noaccounts      3994478   0.487349   8.21   0.000   0.3994478   0.487349
  
```

So, this is where we are going to operate. It is here. So, we will apply the same data. Then we go to the estimation on the same data. We are simply deriving; we have derived the Probit result. So, Probit result is here. first we need to check this value, the Chi square value and its significance level. These are first requirement. Then we come to the interpretation but it only gives the coefficient. But marginal effect is more important.

(Refer Slide Time: 19:42)

□ use “BCM_practice.dta”, clear

```
probit enterprise_type enterprise_age age_square i.sector  
i.mix_activity i.prblm_facd i.assistance_rcvd i.trad_sec  
i.southstates i.location_entrprise i.social_category  
i.bank_account i.growth_status i.activity_group
```

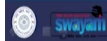
□ The basic **probit** commands report **coefficient estimates** and the **underlying standard errors**. These coefficients are the **index coefficients** and what we can only say is the **direction of the effect** and **partial effects** on the Probit index/score. They do not correspond to the average partial effects.

<pre>probit enterprise_type enterprise_age age_square i.sector i.mix_activity i.prblm_facd > i.bank_account i.growth_status i.activity_group Iteration 0: log likelihood = -3203.1881 Iteration 1: log likelihood = -2138.8491 Iteration 2: log likelihood = -2105.7631 Iteration 3: log likelihood = -2105.709 Iteration 4: log likelihood = -2105.709 Probit regression Number of obs = 6,492 LR chi2(19) = 2194.96 Prob > chi2 = 0.0000 Pseudo R2 = 0.3426 Log likelihood = -2105.709</pre>									
enterprise_type	Coef.	Std. Err.	z	P> z	[95% Conf. Interval]				
enterprise_age	.0006358	.0002012	0.08	0.938	-.0154383	.0167099			
age_square	-.0001695	.0002103	-0.81	0.420	-.0005981	.0002427			
sector									
urban	-.1348293	.0452946	-2.98	0.003	-.2236001	-.0460534			
mix_activity									
no	.0090409	.1187742	0.08	0.933	-.2130404	.2321422			
2.prblm_facd	.1535666	.0493997	3.11	0.002	.0568234	.2503098			
3.assistance_rcvd	.0817668	.2138514	0.38	0.702	-.3267864	.5003199			
trad_sec									
nontraditional	.1839633	.0935709	1.97	0.049	.0005677	.367359			
2.southstates	.0775624	.0509774	1.52	0.128	-.0223515	.1774764			
location_entrprise									
outsideHR	-.1331122	.0483077	-27.56	0.000	-1.425803	-1.236441			
social_category									
SC	.4678568	.0971108	4.82	0.000	.277523	.6581905			
OBC	.268566	.0720871	3.73	0.000	.127278	.409854			
others	-.0234267	.0711503	-0.33	0.742	-.1628788	.1160253			
bank_account									
enterprise_name	-.1.071529	.1109005	-9.66	0.000	-1.28889	-.8541678			
both	-.1.248062	.2240096	-5.57	0.000	-1.687112	-.809011			
noaccounts	.3994678	.0457349	8.73	0.000	.3098291	.4891066			
growth_status									
1	.2560845	.0465302	5.50	0.000	.1648863	.3472821			
3	.2788949	.0802176	3.48	0.001	.1216713	.4361184			
activity_group									
trade	.1914512	.1080314	1.77	0.076	-.0202864	.4031888			
services	-.5512999	.0530573	-10.39	0.000	-.6552903	-.4473095			
_cons	1.004728	.2601325	3.87	0.000	.4908781	1.516579			

We have to go to the slide. So, coefficients are not important. We have to find out the partial effects. That is more important. This is what the result we derived.

(Refer Slide Time: 19:58)

- ❑ The model took 4 iterations to converge. The log likelihood is -2105.709. the log likelihood is used to compare nested models.
- ❑ The likelihood ratio (LR) chi2 is 2194.96 and it is highly statistically significant. That means, it fits significantly better than a model with no predictors.
- ❑ **Interpretation of the coefficients:** The estimate of coefficients in the probit model CANNOT be interpreted as the change in the probability that $Y = 1$ associated with a unit change in X 's.



19

<pre> probit enterprise_type enterprise_age age_square i.sector i.min_activity i.prbin_fact > i.bank_account i.growth_status i.activity_group Iteration 0: log likelihood = -3203.1801 Iteration 1: log likelihood = -2138.8491 Iteration 2: log likelihood = -2105.7631 Iteration 3: log likelihood = -2105.709 Iteration 4: log likelihood = -2105.709 Probit regression Number of obs = 6,492 LR chi2(19) = 2194.96 Prob > chi2 = 0.0000 Pseudo R2 = 0.3426 Log likelihood = -2105.709 </pre>									
enterprise_type	Coef.	Std. Err.	z	P> z	[95% Conf. Interval]				
enterprise_age	.0006358	.0002012	0.08	0.938	-.0154383	.0167099			
age_square	-.0001695	.0002103	-0.81	0.420	-.0005917	.0002427			
sector									
urban	-.1348293	.0452946	-2.98	0.003	-.2236051	-.0460534			
min_activity									
no	.0090409	.1185742	0.08	0.933	-.2130404	.2321422			
2 public_fact	.1535646	.0493997	3.11	0.002	.0568234	.2503098			
2 assistance_cvd	.0817649	.2135514	0.38	0.702	-.3247864	.5003199			
trad_sec									
nontraditional	.1839633	.0935709	1.97	0.049	.0005677	.367359			
2 southstates	.0775624	.0509774	1.52	0.128	-.0223515	.1774764			

location_enterprise									
outsideHR	-1.331122	.0483077	-27.56	0.000	-1.425803	-1.236441			
social_category									
SC	.4678568	.0971108	4.82	0.000	.277523	.6581905			
OBC	.268566	.0720871	3.73	0.000	.127278	.409854			
others	-.0234267	.0711503	-0.33	0.742	-.1628788	.1160253			
bank_account									
enterprise_name	-1.071529	.1109005	-9.66	0.000	-1.28889	-.8541678			
both	-1.248062	.2240096	-5.57	0.000	-1.687112	-.809011			
noaccounts	.3994678	.0457349	8.73	0.000	.3098291	.4891066			
growth_status									
1	.2560845	.0465302	5.50	0.000	.1648863	.3472821			
3	.2788949	.0802176	3.48	0.001	.1216713	.4361184			
activity_group									
trade	.1914512	.1080314	1.77	0.076	-.0202864	.4031888			
services	-.5512999	.0530573	-10.39	0.000	-.6552903	-.4473095			
_cons	1.004728	.2601325	3.87	0.000	.4908781	1.516579			



18

The model took 4 iteration at this moment, 4 iteration to converge. The log likelihood is of minus 2105. It is given here on the log likelihood on the very beginning. The log likelihood is used to compare some nested models if you have. The likelihood ratio that is Chi square test is of 2194 and highly significant that I have already discussed. this fits significantly better then a model with no predictor. This is what it is interpreted.

Let us interpret the coefficient. The estimate of the coefficient of the Probit model cannot be interpreted as the change in the probability that Y equal to 1, associated with the unit change in X. That is not going to be interpreted because of the categorical values.

(Refer Slide Time: 21:01)

- ❑ We can only check for the sign (whether positively or negatively related) and significance of the coefficient.
- ❑ In analysing binary choice models the parameter of interest are not the index coefficients, rather the **marginal/ partial effects**.
- ❑ The command for this task is **margins**.
- ❑ First run the probit model then:
margins, dydx(*) atmeans

<pre> probit enterprise_type enterprise_age age_square i.sector i.mix_activity i.prbin_fact > i.bank_account i.growth_status i.activity_group Iteration 0: log likelihood = -3203.1801 Iteration 1: log likelihood = -2138.8491 Iteration 2: log likelihood = -2105.7631 Iteration 3: log likelihood = -2105.709 Iteration 4: log likelihood = -2105.709 Probit regression Number of obs = 6,492 LR chi2(19) = 2194.96 Prob > chi2 = 0.0000 Pseudo R2 = 0.3426 Log likelihood = -2105.709 </pre>									
enterprise_type	Coef.	Std. Err.	z	P> z	[95% Conf. Interval]				
enterprise_age	.0006358	.0002012	0.08	0.938	-.0154383	.0167099			
age_square	-.0001695	.0002103	-0.01	0.420	-.0005917	.0002427			
sector									
urban	-.1348293	.0452946	-2.98	0.003	-.2236001	-.0460534			
mix_activity									
no	.0090409	.1135742	0.08	0.933	-.2130404	.2321422			
2 public_fund	.1335646	.0493997	3.11	0.002	.0360234	.2303098			
2 assistance_cvd	.0817648	.1135514	0.38	0.702	-.3367864	.5003199			
trade_rec									
nontraditional	.1839633	.0935709	1.97	0.049	.0005677	.367359			
2 southstates	.0775624	.0509774	1.52	0.128	-.0223515	.1774764			

location_enterprise									
outsideHR	-1.331122	.0483077	-27.56	0.000	-1.425803	-1.236441			
social_category									
SC	.4678568	.0971108	4.82	0.000	.277523	.6581905			
OBC	.248566	.0720871	3.73	0.000	.127278	.409854			
others	-.0234267	.0711503	-0.33	0.742	-.1628788	.1160253			
bank_account									
enterprise_name	-1.071529	.1109005	-9.66	0.000	-1.28889	-.8541678			
both	-1.248042	.2240096	-5.57	0.000	-1.687112	-.809011			
noaccounts	.3994678	.0457349	8.73	0.000	.3098291	.4891066			
growth_status									
1	.2560845	.0465302	5.50	0.000	.1648869	.3472821			
3	.2788949	.0802176	3.48	0.001	.1216713	.4361184			
activity_group									
trade	.1914512	.1080314	1.77	0.076	-.0202864	.4031888			
services	-.5512999	.0530573	-10.39	0.000	-.6552903	-.4473095			
_cons	1.004728	.2601325	3.87	0.000	.4968781	1.516579			

We have to take the marginal values. And we can only check for the sign whether positively linked or negatively linked. I think coefficient has already been mentioned; positive, negative in the coefficient table in the very first column.

Then coming to analyzing binary choice model, the parameter of interest are not the index coefficient rather the marginal or partial effects. The command for this task we are already dealt in Logit. Here we can also go by margin with the same command dY/dX asterisk atmeans. We will get the result like this.

(Refer Slide Time: 21:43)

The change in the probability for one instant change in age is .01 percentage points (pp), in age square .003 pp. none of the effects here are significant

	Delta method				
	dy/dx	Std. Err.	z	P> z	[95% Conf. Interval]
enterprise_age	.0001268	.001435	0.08	0.938	-.0030777 .0033313
age_square	-.0000338	.0000419	-0.81	0.420	-.000116 .000484
sector					
urban	-.0248226	.0089951	-2.98	0.003	-.0444328 -.0091924
mix_activity					
No	.0019121	.0228812	0.08	0.933	-.0429343 .0467585
2.phile.fact	.0314243	.0104741	3.02	0.003	.0110955 .052133
2.assistance_xred	.0170781	.0466495	0.37	0.714	-.0743133 .1085095
read_soc					
nontraditional	-.0349377	.0168478	-2.07	0.039	-.0683666 -.0015089
2.nontraditional	.0156778	.0104457	1.50	0.133	-.0047853 .036151
location_enterprise					
outsideRR	-.3291332	.0130459	-25.19	0.000	-.35476 -.3035424
social_category					
SC	.0870969	.0179164	4.86	0.000	.0519815 .1222123
OBC	.0540457	.0163193	3.44	0.001	.0240806 .0840509
others	-.005709	.0172486	-0.33	0.741	-.0395155 .0280976
bank_account					
enterprise_name	-.367435	.0432488	-8.50	0.000	-.4522011 -.2826688
bank	-.4375417	.0866483	-5.04	0.000	-.6088893 -.2672142
microfinance	.0744109	.0088675	8.42	0.000	.059331 .0894909
growth_status					
2	.0536295	.0100559	5.33	0.000	.0339203 .0733387
3	.0574641	.0152043	3.79	0.000	.0278643 .0874639
activity_group					
trade	.0305668	.0162246	1.88	0.060	-.0023329 .0633664
services	-.1356072	.0141709	-9.57	0.000	-.1633816 -.1078328

Being from urban area decreases the probability of necessity entrepreneur by 2.7 percentage points and is highly statistically significant. Other values are interpreted in the same way.

Note: dy/dx for factor levels is the discrete change from the base level.

The result here is interpreted like the earlier way we interpreted. Here also the same way we will discuss. Like for age this is now 0.1, 0.01 percent points. Age square it is 0.003. Similarly for urban and rural comparison 2.7 times for urban it is 2.7, 0.027 is there. So, in percentage it is 2.7 percentage points lower probability of being necessity entrepreneur. So, you can go through further details and let us go further.

(Refer Slide Time: 22:14)

❑ The choice of logit and probit model is on the researcher. But there exist one difference between logit and probit model that is, when the residual follows the **logistic distribution** we apply **logit model** and when the residual follows the **standard normal distribution** we apply **probit model**.

❑ We run the probit model and then predict the values of residuals:

`predict r`

To check the distribution of residual, we use kernel density function:

`kdensity r, normal`

The choice of Logit and Probit model is on the researcher but there exists one difference between Logit and Probit model, Logit model, when the residual follows the logistic distribution you apply Logit model. When the residual follows a standard normal distribution we should apply Probit model.

We run the Probit model and then predict the values of residuals. We will go by the first, we have already done the Probit regression. We will predict the residuals. Then we will go by kernel density function, kernel density of its normal function. Then everything will be fine.

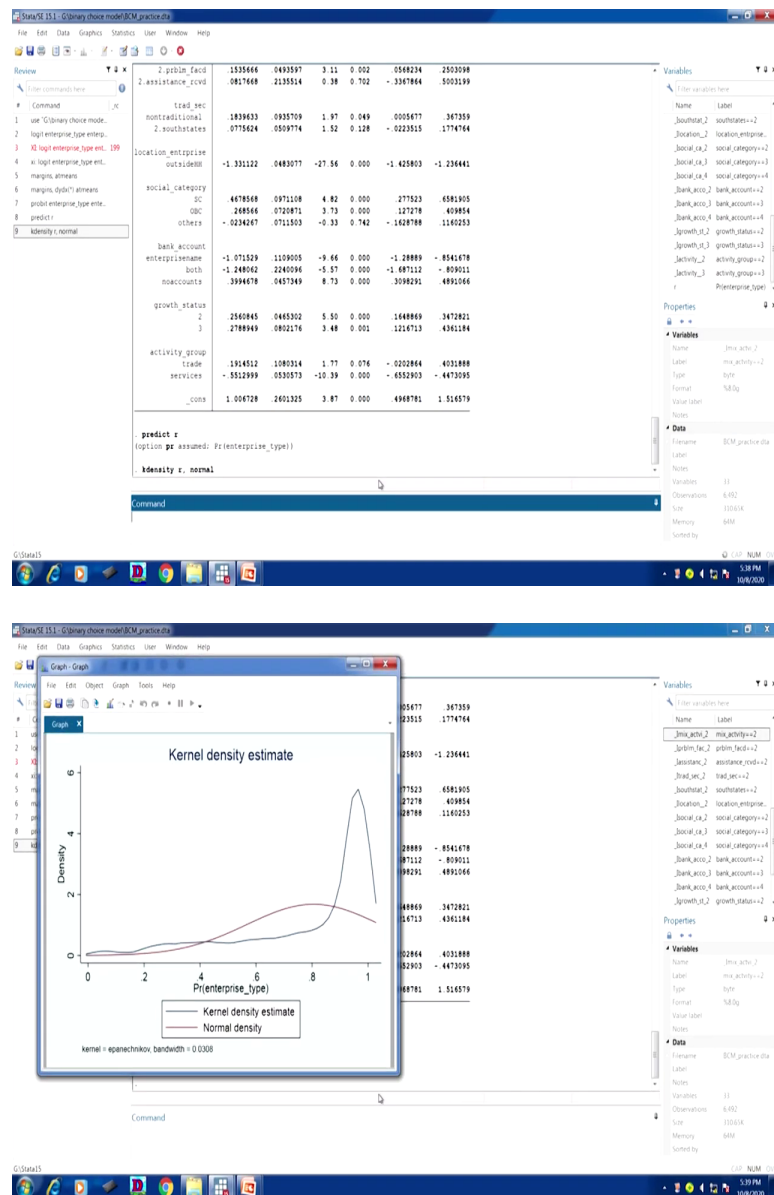
(Refer Slide Time: 23:03)

The screenshot shows the Stata 15.1 interface with the following components:

- Review Panel:** Displays the command `probit` and the list of variables included in the model: `2. assistanc_rcd`, `2. southstates`, `location_enterprise`, `outsidekm`, `social_category`, `bank_account`, `enterprisename`, `both`, `noaccounts`, `growth_status`, `activity_group`, `trade`, `services`, and `_cons`.
- Variables Panel:** Lists the variables used in the model, including `enterprise_age`, `enterprise_age^2`, `location_enterprise`, `location_enterprise^2`, `location_enterprise^3`, `location_enterprise^4`, `location_enterprise^5`, `location_enterprise^6`, `location_enterprise^7`, `location_enterprise^8`, `location_enterprise^9`, `location_enterprise^10`, `location_enterprise^11`, `location_enterprise^12`, `location_enterprise^13`, `location_enterprise^14`, `location_enterprise^15`, `location_enterprise^16`, `location_enterprise^17`, `location_enterprise^18`, `location_enterprise^19`, `location_enterprise^20`, `location_enterprise^21`, `location_enterprise^22`, `location_enterprise^23`, `location_enterprise^24`, `location_enterprise^25`, `location_enterprise^26`, `location_enterprise^27`, `location_enterprise^28`, `location_enterprise^29`, `location_enterprise^30`, `location_enterprise^31`, `location_enterprise^32`, `location_enterprise^33`, `location_enterprise^34`, `location_enterprise^35`, `location_enterprise^36`, `location_enterprise^37`, `location_enterprise^38`, `location_enterprise^39`, `location_enterprise^40`, `location_enterprise^41`, `location_enterprise^42`, `location_enterprise^43`, `location_enterprise^44`, `location_enterprise^45`, `location_enterprise^46`, `location_enterprise^47`, `location_enterprise^48`, `location_enterprise^49`, `location_enterprise^50`, `location_enterprise^51`, `location_enterprise^52`, `location_enterprise^53`, `location_enterprise^54`, `location_enterprise^55`, `location_enterprise^56`, `location_enterprise^57`, `location_enterprise^58`, `location_enterprise^59`, `location_enterprise^60`, `location_enterprise^61`, `location_enterprise^62`, `location_enterprise^63`, `location_enterprise^64`, `location_enterprise^65`, `location_enterprise^66`, `location_enterprise^67`, `location_enterprise^68`, `location_enterprise^69`, `location_enterprise^70`, `location_enterprise^71`, `location_enterprise^72`, `location_enterprise^73`, `location_enterprise^74`, `location_enterprise^75`, `location_enterprise^76`, `location_enterprise^77`, `location_enterprise^78`, `location_enterprise^79`, `location_enterprise^80`, `location_enterprise^81`, `location_enterprise^82`, `location_enterprise^83`, `location_enterprise^84`, `location_enterprise^85`, `location_enterprise^86`, `location_enterprise^87`, `location_enterprise^88`, `location_enterprise^89`, `location_enterprise^90`, `location_enterprise^91`, `location_enterprise^92`, `location_enterprise^93`, `location_enterprise^94`, `location_enterprise^95`, `location_enterprise^96`, `location_enterprise^97`, `location_enterprise^98`, `location_enterprise^99`, `location_enterprise^100`.
- Properties Panel:** Shows the variable `enterprise_age` with a label of `enterprise age` and a type of `float`.
- Command Window:** Displays the command `. predict r` and the output `(option pr assumed: Pr(enterprise_type))`.

So, we will go by predict r that is our residuals, have been predicted. We will go by Kernel density. So, the k density, r , a value predicted, the predicted values of r have already come.

(Refer Slide Time: 23:32)



We are giving k density of r. we are going to give a normality test of it. It gives a diagram. It is now estimating or deriving. Still it is running. Look at the diagram, the graph of the predicted values that is highlighted in blue color. The red one is the normal density function, the standard normal density function. The estimated one is in blue or black color.

Lookat very clearly. The graph seems asymmetric. It is not symmetric. It is nowhere closing to the normal distribution. The normal density function is not overlapping with this. Or nowhere near to the normal density function. So, that is the reason why we are not accepting these

distributions to be normal. So, in this case it is suggested that we should apply Logit model. Probit is not the best fit. Now, this is what we said already.

(Refer Slide Time: 24:40)

COMPARISON OF MODELS AND PARAMETERS

```

quietly regress enterprise_type enterprise_age age_square i.sector
i.mix_actvity i.prblm_facd i.assistance_rcvd i.trad_sec i.southstates
i.location_entrpriase i.social_category i.bank_account i.growth_status
i.activity_group
estimates store blpm

quietly logit enterprise_type enterprise_age age_square i.sector
i.mix_actvity i.prblm_facd i.assistance_rcvd i.trad_sec i.southstates
i.location_entrpriase i.social_category i.bank_account i.growth_status
i.activity_group
estimates store blolit
  
```

The screenshot shows the Stata 15.1 interface with the following components:

- Command Window:** Contains the commands:


```

. use "C:\practise\choice model\BCM_practice.dta"
. quietly regress enterprise_type enterprise_age age_square i.sector i.mix_actvity i.prblm_facd i.assistance_rcvd i.trad_sec i.southstates i.location_entrpriase i.social_category i.bank_account i.growth_status i.activity_group
. estimates store blpm
. quietly logit enterprise_type enterprise_age age_square i.sector i.mix_actvity i.prblm_facd i.assistance_rcvd i.trad_sec i.southstates i.location_entrpriase i.social_category i.bank_account i.growth_status i.activity_group
. estimates store blolit
      
```
- Results Window:** Displays the output of the regress command, showing coefficients for various variables like bank_account, enterprise_type, enterprise_age, age_square, i.sector, i.mix_actvity, i.prblm_facd, i.assistance_rcvd, i.trad_sec, i.southstates, i.location_entrpriase, i.social_category, i.bank_account, i.growth_status, and i.activity_group.
- Variables Window:** Lists the variables used in the analysis, including bank_account, enterprise_type, enterprise_age, age_square, i.sector, i.mix_actvity, i.prblm_facd, i.assistance_rcvd, i.trad_sec, i.southstates, i.location_entrpriase, i.social_category, i.bank_account, i.growth_status, and i.activity_group.

We are going to compare all the models together and to decide what is the best fit and how we should go for it. Alright, so we have three. Like if you go by this command quietly regress with the LPM model, for the LPM model, quietly regress if you do it, what we do? We can quickly do that. Let me show you. We can quickly copy this command and show you how it works. All those... So, this basically Logit and LPM we are running together.

Alright, so both the results, here you can see the new estimated values has already been derived for LPM and for Logit. So, we are going to show the Probit as well. We are also comparing all those three together. In this command we are going to compare as well.

(Refer Slide Time: 25:57)

The screenshot shows the Stata 15.1 interface with the command window displaying the results of three models: Probit, Logit, and LPM. The command used is `estimates table bign blpm blgit bprobit, t stata(11)`. The results are presented in a table with columns for the variable, the Probit coefficient, the Logit coefficient, and the LPM coefficient. The command window also shows the command `estimates table bign blpm blgit, t stata(11)`.

Variable	bign	blgit	bprobit
enterprise_g	.00048059	0.43	
age_square	-.0000449	-1.06	
sector			
urban		.03787455	-3.37
mix_activity			
no	.00046958	0.04	
probit_faced			.02612111

Alright, first command is for Probit. You can see from this here Logit blpm is for LPM. Then second one is for Logit then third one is for Probit.

(Refer Slide Time: 26:18)

The screenshot shows the Stata 15.1 interface with the command window displaying the command `estimates table bign blpm blgit bprobit, t stata(11)`. The command window also shows the command `estimates table bign blpm blgit, t stata(11)`.

Variable	LPM	Logit	Probit
assistance_d	2.89	0.7933183	0.8176476
trad_sec	2.05	0.21	0.38
southstates	1.97	0.4388563	0.396333
location_entrprise	1.81	0.1471811	0.7756243
social_cat_y	4.82	-2.4053173	-1.331122
bank_account	4.82	-24.55	-27.56
growth_status	4.82	0.8284259	0.4785676
activity_group	4.82	0.44767679	0.6856801
others	4.82	0.44	0.73
estimates store bprobit	4.82	-0.0562422	-0.2242474
estimates table blpm	4.82	-0.67	-0.33
estimates table bprobit	4.82	-1.8317239	-1.0715288
estimates table bprobit	4.82	-9.36	-9.66
estimates table bprobit	4.82	-2.27440123	-1.2481617
estimates table bprobit	4.82	-5.19	-5.57
estimates table bprobit	4.82	-72448104	399446782
estimates table bprobit	4.82	8.71	8.73

So, after that we are comparing all these three together. This is for LPM, linear probability model then Logit then Probit. from all those models we have got a comparison clearly. First column is for LPM than second and third order Logit and Probit respectively. Each of the coefficients can be compared and how they differ each other. That some of them we have already discussed in the previous equation, that 4 times related, 2.6 times related that you can compare and find out. Alright so this is what we have done.

(Refer Slide Time: 26:49)

```
quietly probit enterprise_type enterprise_age age_square
i.sector i.mix_actvity i.prblm_facd i.assistance_rcvd i.trad_sec
i.southstates i.location_entrprise i.social_category
i.bank_account i.growth_status i.activity_group
```

```
estimates store bprobit
```

```
estimates table blpm blogit bprobit, tstats(N ll)
```

T value and other stats like number of observation and log likelihood is specified.

```
estimates table blpm blogit bprobit, star(.01 .05 .10)
```

Star option is specified to get significance along with coefficients.

- ❑ `quietly` command is used to suppress terminal output.
- ❑ `Estimates store` command stores the estimation result of different models.
- ❑ `Estimates table` command organizes estimation results from one or more models in a single formatted table.
- ❑ Similarly we can also compare marginal effects of three models.

Here I just wanted to mention that T value and other stats like number of observation and log likelihood is specified very clearly. So, here first of all we specify with T and, T and its log likelihood values. But you can also specify its significance level. If you specify the significance level the significance with star mark is also going to come. The star mark against the coefficient will be there. Once the star mark is there you can compare which coefficient which variable is significant in which particular model and it will be easy for you to decide alright.

So, that is the reason why `quietly` command is used to suppress terminal output. And these `estimates store` basically gives the estimation result of different models together. And `estimates table` particularly. `store` keep it in the model and `table` keeps in all those models together. Alright, so then basically in single format table.

(Refer Slide Time: 28:13)

Marginal effects:

```
quietly regress enterprise_type enterprise_age age_square  
i.sector i.mix_activity i.prblm_facd i.assistance_rcvd i.trad_sec  
i.southstates i.location_entprise i.social_category  
i.bank_account i.growth_status i.activity_group
```

```
quietly margins, dydx(*) atmeans
```

This command gives similar values as coefficient of ols model.

```
estimates store marginlpm
```



27

Similarly we can also compare marginal effects of these three models. Similarly we can type quietly regress with i, quietly regress that is basically the LPM model. Marginal effect is with the same regression model itself whereas after the regression only in case of LPM we have to quietly margin. Margin and dy/dx star atmeans that will give you the marginal effect. But usually that is not necessary in case of continuous data. And LPM since it considers to be continuous series, though the dependent variable is binary that is why we are running marginal effect. Then similar command estimates store margin LPM it will save.

(Refer Slide Time: 29:01)

```

quietly logit enterprise_type enterprise_age age_square i.sector
i.mix_activity i.prblm_facd i.assistance_rcvd i.trad_sec
i.southstates i.location_entprise i.social_category
i.bank_account i.growth_status i.activity_group

quietly margins, dydx(*) atmeans
estimates store marginlogit

quietly probit enterprise_type enterprise_age age_square
i.sector i.mix_activity i.prblm_facd i.assistance_rcvd i.trad_sec
i.southstates i.location_entprise i.social_category
i.bank_account i.growth_status i.activity_group

```

```

quietly margins, dydx(*) atmeans
estimates store marginprobit

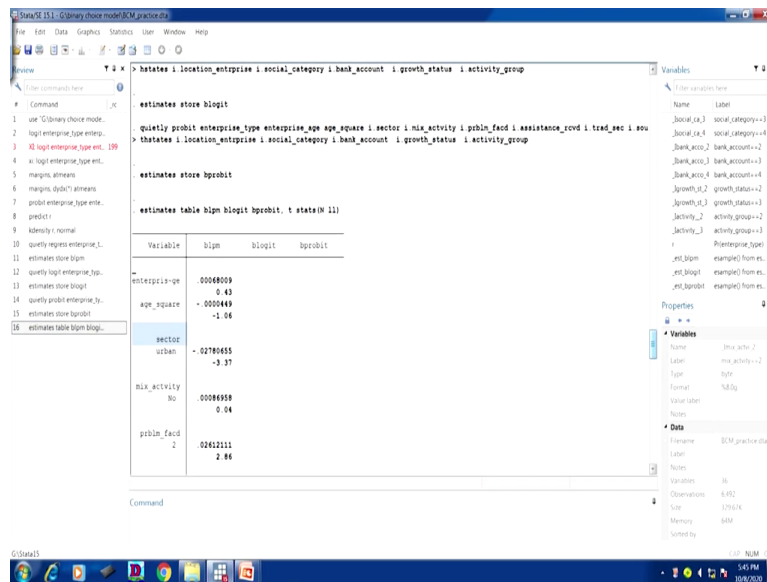
estimates table marginlpm marginlogit marginprobit, star(.01
.05 .10)

```

From the result we can compare the coefficients.

Then we go for quietly logit, then with all those commands for Probit as well. Probit after running Logit then we discuss about the marginal. Then we also store it, store will save it. Similarly after Probit we will go for the next slide its marginal value atmeans. Then we store it. After storing all those things we can take all those three models together. When you take three models together will be able to decide which one is better.

(Refer Slide Time: 29:49)



Likewise we have just shown in 3 columns. Here, 3 columns if you are confused, here only coefficients we have shown. But in our other slides we have already mentioned that if you run by the marginal effect the marginal Coefficient will be visible. Alright, so I am not running several times because of paucity of time. So, with that we finished the discussion with LPM, Logit and Probit. And we decided which one is most fitted based on their coefficient, their significance level, based on the distribution function and so on.

Coming to Tobit regression model. The last model in this series is on Tobit which is also an extension of the Probit model. Originally developed by James Tobin as you might have heard as

Nobel Laureate in 1958. Often the dependent variable is constrained. Many cases it is censored. There is a limit.

Like those who are estimating some policy implementation of a particular like PDS scheme for example and that will be applicable for people with below poverty line. So, all the income earners are not going to be considered. Somewhere you have to censor. You have to define a benchmark level of income. To that only you are going to give. If you want to give more number of people in that caveat then you can change that censor limit. And accordingly we can run the regression.

(Refer Slide Time: 31:23)

INTRODUCTION

- ❑ An extension of the probit model.
- ❑ Originally developed by James Tobin (Nobel laureate) in 1958.
- ❑ Often the dependent variable is constrained (or censored).
- ❑ Takes on a positive value for some observations and zero for other observations.
- ❑ Represents non-continuous data as there is a large cluster of observations at zero.
- ❑ Such samples are called **censored samples**.



31

- ❑ Therefore, such models are also known as **censored regression model** or **limited dependent variable model**.
- ❑ Using OLS leads to biased estimates of the parameters.



32

So, usually it takes on a positive value for some observation and others as 0. If you have censored then censored values are considered to be 0. And these represent non-continuous data as there is a large cluster of observation at 0. Such symbols are also called censored samples. Basically those are considered to be 0 are called censored samples.

Therefore such models are also known as censored regression or also called limited dependent variable models. Using OLS , method leads to biased estimates because, in any case this has categorized into 0 and 1. These are not a continuous series all totally though the non-censored

variables or the values are of continuous, could be continuous and different than that of the existing Logit and Probit model.

(Refer Slide Time: 32:22)

An Example:

- ❑ The dataset containing information on the value added of any enterprise. In this case some enterprises have positive value added while others have zero value added or report negative. Those who have reported negative value added are considered value added equal to zero.
- ❑ Thus, in this case the enterprises are divided into 2 groups say n_1 and n_2 .
- ❑ n_1 whom we have information on the regressors (type of enterprise, location, age of the enterprise etc.) as well as regressand (value added.)



33

- ❑ n_2 about whom we have information only on the regressors but not on the regressand.
- ❑ For OLS estimates of the parameters obtained from the subset of n_1 will be biased as well as inconsistent.
- ❑ There is another type of sample distinguished from the censored sample called **truncated sample** in which information on the regressor is available only if the regressand is observed.



34

So, let us understand in terms of example. We have already referred the dataset of the enterprises context. The dataset containing information on the value added of any enterprise. In this case some enterprise have positive values added while others have 0 value added. Or they report negative.

So far as the value added of the enterprises is concerned, enterprises and their value added figures are available in the 73rd round. Those who have been reported negative value added are considered value added equal to 0 because no question of negative value added is valid. You can make it to 0. So, thus in this case enterprises are divided into two groups, say n_1 or n_2 . So, n_1 whom we have information on the regressors like type of enterprise, their location, age of the enterprise as well as regression that is value added we are discussing.

So, n_2 is all about whom we have information only on the regressor but not on the regressand because it is of either 0 values or negative values, or no information. So, for OLS estimates of the parameters obtained from the subset of n_1 will be biased as well as inconsistent. So, there is another type of sample distinguished from the censored sample called truncated sample and information on the regressor is available only if the regressand is observed. So, when only regressand is observed then in that case we get the truncated sample.

(Refer Slide Time: 34:01)

STATISTICAL REPRESENTATION OF THE MODEL

Consider a linear model:

$$y_i = x_i\beta + \varepsilon_i, \quad \varepsilon_i \sim N(0, \sigma^2)$$

First, let us find the **cumulative distribution function (CDF) of $(y_i | x_i)$** :

$$F(y_i | x_i) \equiv P(Y \leq y_i | x_i)$$

Upper case F typically denotes CDF. Y is stochastic variable. The above equation defines CDF that the probability of stochastic variable Y is smaller than or equal to y_i . Here, y_i is sort of realization. Putting value of stochastic variable from linear model:

So, in some mathematical interpretation we can simply identify that model is Y of x_i beta of epsilon, the error distribution is distributed with 0 and sigma square variance. So, let us find the cumulative distribution function of y_i given x_i where we are limiting the y with certain level of y_i given x_i .

So, here a limit is given correctly, less than or equal to that is going to be valid in our case. Or entire y_i is not going to be defined in this particular distribution. The uppercase that is uppercase of F typically denotes CDF function or the cumulative distribution function. Y is stochastic variable. The above equation defines CDF that the probability of stochastic variable Y is smaller than or equal to that y_i we are defining here. Here y_i is sort of realization. Putting value of stochastic variable from linear model.

(Refer Slide Time: 35:13)

$$F(y_i | x_i) \equiv P(x_i\beta + \varepsilon_i \leq y_i | x_i)$$

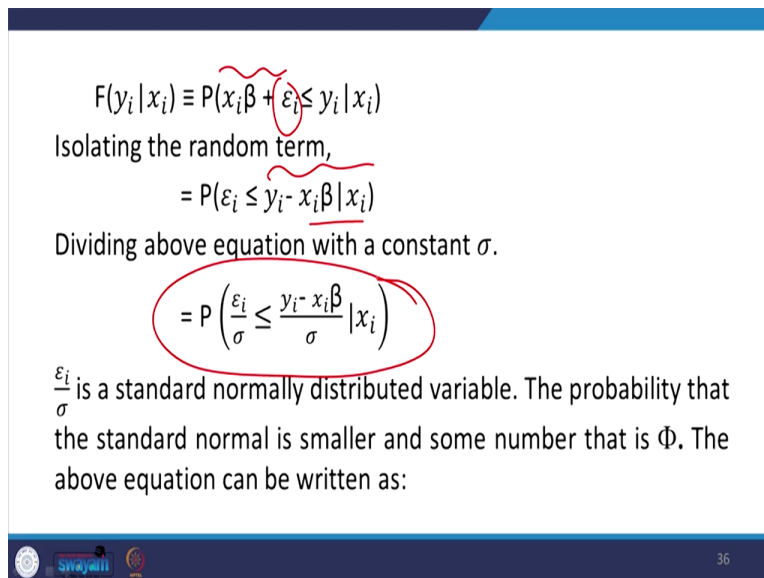
Isolating the random term,

$$= P(\varepsilon_i \leq y_i - x_i\beta | x_i)$$

Dividing above equation with a constant σ .

$$= P\left(\frac{\varepsilon_i}{\sigma} \leq \frac{y_i - x_i\beta}{\sigma} | x_i\right)$$

$\frac{\varepsilon_i}{\sigma}$ is a standard normally distributed variable. The probability that the standard normal is smaller and some number that is Φ . The above equation can be written as:



36

$$= \Phi\left(\frac{y_i - x_i\beta}{\sigma}\right)$$

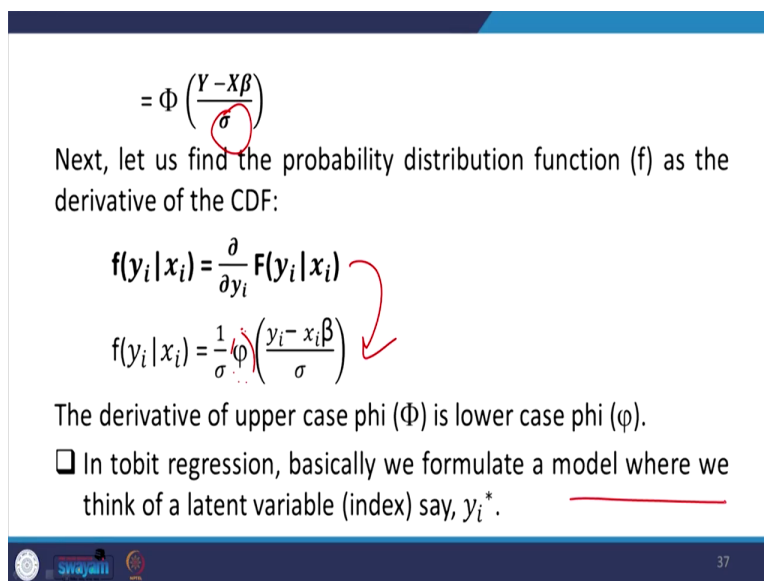
Next, let us find the probability distribution function (f) as the derivative of the CDF:

$$f(y_i | x_i) = \frac{\partial}{\partial y_i} F(y_i | x_i)$$

$$f(y_i | x_i) = \frac{1}{\sigma} \phi\left(\frac{y_i - x_i\beta}{\sigma}\right)$$

The derivative of upper case phi (Φ) is lower case phi (ϕ).

☐ In tobit regression, basically we formulate a model where we think of a latent variable (index) say, y_i^* .



37

Similarly after discussing this, Y can for the disaggregate to its beta coefficients and epsilon with the limit. Error can be defined. Error distribution is defined as y_i minus of this estimated

coefficients. The distribution where we are going to estimate given x_i can be defined with this probabilistic value.

This is the epsilon with respect to the standard deviation is a normal distributed variable. The probability that the standard normal is smaller or some number that is indicated with phi. The above equation can be written phi, and in terms of phi as a function of y_i minus x_i beta divided by sigma which we have defined just a minutes back. That is the probability distribution function.

(Refer Slide Time: 36:06)

$$= \Phi \left(\frac{y - x\beta}{\sigma} \right)$$

Next, let us find the probability distribution function (f) as the derivative of the CDF:

$$f(y_i | x_i) = \frac{\partial}{\partial y_i} F(y_i | x_i)$$

$$f(y_i | x_i) = \frac{1}{\sigma} \phi \left(\frac{y_i - x_i \beta}{\sigma} \right)$$

The derivative of upper case phi (Φ) is lower case phi (ϕ).

- In tobit regression, basically we formulate a model where we think of a latent variable (index) say, y_i^* .

- This y_i^* is censored above or below some value. That is, it is not observable for part of the population.
- For example, assume that y_i^* is enterprise's value added and for a randomly drawn enterprises the actual value of wealth recorded up to some threshold, say, ₹10,00,000, but above that level only the fact that value added was more than ₹10,00,000 is recorded.
- Similarly, negative values can also be censored.
- So, based on these assumptions we can split the probability density function as follows:

$$= \Phi \left(\frac{y - x\beta}{\sigma} \right)$$

Next, let us find the probability distribution function (f) as the derivative of the CDF:

$$f(y_i | x_i) = \frac{\partial}{\partial y_i} F(y_i | x_i)$$

$$f(y_i | x_i) = \frac{1}{\sigma} \phi \left(\frac{y_i - x_i \beta}{\sigma} \right)$$

The derivative of upper case phi (Φ) is lower case phi (ϕ).

□ In tobit regression, basically we formulate a model where we think of a latent variable (index) say, y_i^* .

By taking their derivative we get the marginal value. And the marginal value since sigma is there and these are constants. So, this is a small phi and indicate the marginal difference of the equation. This follows from here. And in Tobit regression basically we formulate model where we think of a latent variable, say that is y_i , latent variable on observed variables since probabilistic issues are attached to that variable. So, that is why we are mentioning as latent here. So, y_i star is censor above or below a certain value. That it is basically not observable for part of the population. And so part of the population since not observable we can censor it, alright.

In this example like y_i star is enterprise value added and for randomly drawn enterprises actual value of the wealth recorded up to some threshold let it be 10 lakhs but above that level only the fact that the value added was more than 10 lakh is recorded. Similarly negative values can also be censored if above value and below value or negative value can also be censored. Whichever we require for the necessity of the model we can censor accordingly. Based on the assumptions we can split the probability density function like the way we defined here, like this, alright. So, there are different limits. Like we can define 1 if it is less than certain limit, 2 if it is with another bracket, we can categorize their probability limit.

Let us come to the example to estimate it correctly. How you can estimate it? Somewhere we have to understand that there are some limits defined out of the total data we can censor at the below, at the top, or at the medium. So, you can give their coefficient accordingly.

(Refer Slide Time: 38:34)

EXAMPLE: TOBIT REGRESSION MODEL

- ❑ For understanding tobit model a sample of NSS 73rd round data (tobit_practice.dta) for the year 2015-16 have been used.
- ❑ Our data contains 3509 observations.
- ❑ Our objective is to explore the factors that explain the performance of female entrepreneurs.
- ❑ Our dependent variable is gross value added of a enterprise (GVA-in rs) that ranges from -20700 to 619622.
- ❑ assume that our data is censored so that we could not observe a negative GVA.

39

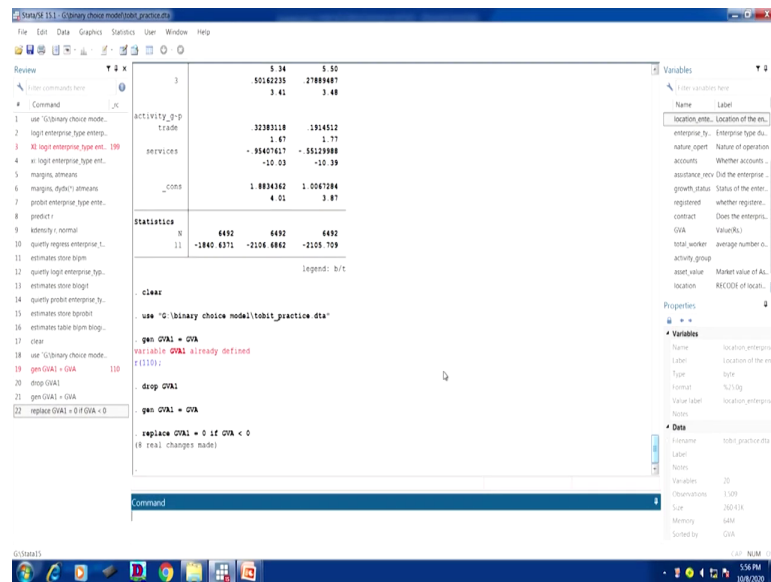
So, from the 73rd round of the data we have a practice dataset called tobit practice dataset. We are going to open right now. So, our data contains 3509 observations. Our objective here is to explore the factors which explain the performance of female entrepreneurs with a certain caveat. So, dependent variable here to be estimated with a censor value for its gross value added of the enterprise. Gross value measured is in rupees that ranges from minus 20700 to 619622. These are the distribution. Since minus values are there so it is a concern for us. We need to censor it.

(Refer Slide Time: 39:18)

- ❑ **Censor the GVA:**
 - ❑ If the true GVA is 0 or less, all we know is that GVA is equal to 0.
 - ❑ Let's generate a new variable:
`gen GVA1 = GVA`
 - ❑ Replacing any value that is zero or less with zero
`replace GVA1 = 0 if GVA < 0`
(8 real changes made)

It means 8 values of GVA were negative

40



So, how could be able to censor it? In the data, so let us open the data first. Then we can be able to discuss. So, first of all, just a minute, alright. So, we are going to open the data that is tobit practice. This has already been opened. We are going to do something very quickly, since GVA has already been defined we need to compare GVA 1 with the value if GVA has less than 0 or negative values. Let us redefine the GVA with a new variable called GVA1. If that is having GVA less than 0, alright.

So, it is here. So, we are defining a value here, then this. So, GVA1 we have already defined so we can drop that variable at this moment. That is there, alright. Once again we are going to run the same command. So, GVA1 is defined now. GVA1 is here. It is visible at the end, GVA1. And now we need to replace GVA1 is equal to 0, we want to censor it with 0 value if GVA1 is less than 0. So, very quickly we have got 8 real changes. That means 8 cases where the values are negative. those have been converted to 0, isn't it? So, let us move on and we wanted to find some other interesting interpretation.

(Refer Slide Time: 41:17)

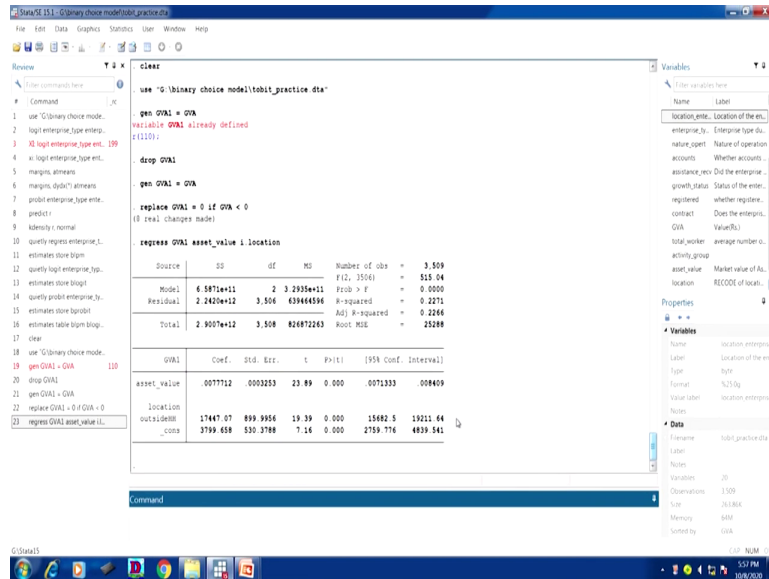
- ❑ Keep in mind that our dependent variable GVA1 is not dichotomous but continuous in nature.
 - ❑ Our variable of interest are asset value (continuous) and location of the enterprise (dichotomous).
 - ❑ Summarize and describe the variables of interest to check missing values and storage type of variable.
 - ❑ As dependent variable is continuous we can also run OLS.
- ```
regress GVA1 asset_value i.location
```



41

Keep in mind that our dependent variable that is GVA is not dichotomous but continuous in nature. So, that is basically 0 and above, alright. That is interesting to note but usually in Logit and Probit it has to be mandatorily dichotomous. So, our variable of interest here is asset value. That is continuous and also location and any other interesting variable you wanted to include you can include. So, we are going to test with this one as the dependent variable GVA1 right now. So, will simply go by ordinary least square method. Since it's a continuous variable so let us compare how Tobit and these two can be compared.

(Refer Slide Time: 42:11)



So, this is our ordinary regression coefficient that has been defined with respect to asset value, location of the enterprises, alright.

(Refer Slide Time: 42:22)

□ We can estimate a tobit model by instructing the software to censor the data:

Command: `tobit GVA1 asset_value i.location, ll`

Here, we have used already censored variable GVA1

Lower limit (ll)

```

tobit GVA1 asset_value i.location, ll
Defining starting values:
Grid node 0: log likelihood = -40451.031
Fitting full model:
Iteration 0: log likelihood = -40451.031
Iteration 1: log likelihood = -40451.013
Iteration 2: log likelihood = -40451.013

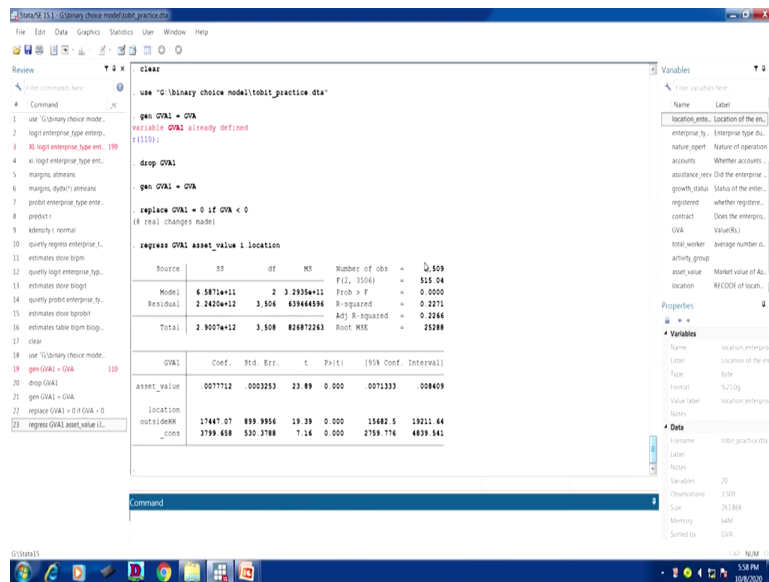
Tobit regression Number of obs = 3,509
 Uncensored = 3,499
 Limits: lower = 0 Left-censored = 10
 upper = +inf Right-censored = 0

 LR chi2(2) = 901.96
 Prob > chi2 = 0.0000
 Pseudo R2 = 0.0110

Log likelihood = -40451.013

```

|             | GVA1 | Coef.    | Std. Err. | t     | P> t  | [95% Conf. Interval] |
|-------------|------|----------|-----------|-------|-------|----------------------|
| asset_value |      | .0077753 | .0003258  | 23.87 | 0.000 | .0071366 .008414     |
| location    |      | 1.74493  | 901.5007  | 19.38 | 0.000 | 15701.78 19236.82    |
| _cons       |      | 3741.558 | 531.3524  | 7.04  | 0.000 | 2499.767 4783.35     |
| var( GVA1)  |      | 6.41e+08 | 1.53e+07  |       |       | 6.12e+08 6.72e+08    |



Let us go to the explanation through Tobit. So, we are going to define through Tobit now. We have defined already for ordinary least square. So, let us define through Tobit. So, Tobit is going to give little different result and that is interesting to note. And here the Tobit, our command is here, or at the end we have a lower limit, lower limit in general we have written. One specific lower limit we can also mention. So, when we have defined this ll, at the end ll, not 11, please read this is ll, lower limit. In general the lower limit we have defined. Now, this is the Tobit result and this is the OLS result.

(Refer Slide Time: 43:12)

☐ We can estimate a tobit model by instructing the software to censor the data:

Command:

tobit GVA1 asset\_value i.location, ll

Here, we have used already censored variable GVA1

Lower limit (ll)

```
tobit GVA1 asset_value i.location, ll
Refining starting values:
Grid node 0: log likelihood = -40451.031
Fitting full model:
Iteration 0: log likelihood = -40451.031
Iteration 1: log likelihood = -40451.013
Iteration 2: log likelihood = -40451.013

Tobit regression Number of obs = 3,509
 Unobserved = 3,499
 Limits: lower = 0 Left-censored = 10
 upper = *inf Right-censored = 0

 LR chi2(2) = 901.96
 Prob > chi2 = 0.0000
 Pseudo R2 = 0.0110

Log likelihood = -40451.013
```

|             | Coef.    | Std. Err. | t     | P> t  | [95% Conf. Interval] |
|-------------|----------|-----------|-------|-------|----------------------|
| GVA1        |          |           |       |       |                      |
| asset_value | .0077753 | .0003258  | 23.87 | 0.000 | .0071366 .008414     |
| location    |          |           |       |       |                      |
| outside09   | 17469.3  | 901.5007  | 19.38 | 0.000 | 15701.78 19236.82    |
| _cons       | 3741.558 | 531.3524  | 7.04  | 0.000 | 2699.767 4783.35     |
| vsize(GVA1) | 6.41e+08 | 1.53e+07  |       |       | 6.12e+08 6.72e+08    |

- Using the uncensored variable (GVA), we could instruct the software to censor it in the estimation by using the subcommand: `ll(0)`

Command:

`Tobit GVA asset_value i.location, ll(0)`

As we can see, both the results are same, there is a little variation in the coefficients values of OLS and tobit model

If have any upper limit (u). If both, ll() u() can be added

```
tobit GVA asset_value i.location, ll(0)

Refining starting values:
Grid node 0: log likelihood = -40451.031

Fitting full model:
Iteration 0: log likelihood = -40451.031
Iteration 1: log likelihood = -40451.013
Iteration 2: log likelihood = -40451.013

Tobit regression Number of obs = 3,509
 Uncensored = 3,499
 Limits: lower = 0 Left-censored = 10
 upper = -inf Right-censored = 0

 LR chi2(2) = 901.96
 Prob > chi2 = 0.0000
 Pseudo R2 = 0.0110

Log likelihood = -40451.013
```

|             | Coef.    | Std. Err. | z     | P> z  | [95% Conf. Interval] |
|-------------|----------|-----------|-------|-------|----------------------|
| asset_value | .0077753 | .0003258  | 23.87 | 0.000 | .0071366 .008414     |
| location    |          |           |       |       |                      |
| outsideR01  | 17469.3  | 901.5007  | 19.38 | 0.000 | 15701.78 19236.82    |
| _conc       | 3741.558 | 531.3526  | 7.04  | 0.000 | 2699.767 4783.35     |
| var(a_GVA)  | 6.41e+08 | 1.53e+07  |       |       | 6.12e+08 6.72e+08    |

So, what we will do? We will discuss. Here we have used already a censored variable that is GVA1. GVA1 is a censored variable. Using the uncensored variable GVA, we can go by the GVA also like this. So, GVA, uncensored variable so how it looks like? We can compare.

(Refer Slide Time: 43:39)

```

. tobit gva asset_value i.location, ll(0)

Refining starting values:
Grid node 0: log likelihood = -40451.031

Fitting full model:
Iteration 0: log likelihood = -40451.031
Iteration 1: log likelihood = -40451.013
Iteration 2: log likelihood = -40451.013

Tobit regression Number of obs = 3,509
 Unobserved = 3,499
 Limits: lower = 0 Left-censored = 10
 upper = +inf Right-censored = 0
 LR chi2(2) = 901.96
 Prob > chi2 = 0.0000
 Pseudo R2 = 0.0110

Log likelihood = -40451.013

+-----+-----+-----+-----+-----+-----+
| GVA | Coef. | Std. Err. | t | P>|t| | [95% Conf. Interval] |
+-----+-----+-----+-----+-----+-----+
| asset_value | .0077753 | .0003258 | 23.87 | 0.000 | .0071366 .008414 |
+-----+-----+-----+-----+-----+-----+
| location | 17469.3 | 901.5007 | 19.38 | 0.000 | 15701.78 19236.82 |
+-----+-----+-----+-----+-----+-----+
| outsideR | 3741.558 | 531.3526 | 7.04 | 0.000 | 2699.767 4783.35 |
+-----+-----+-----+-----+-----+-----+
| _cons | 6.41e+08 | 1.53e+07 | | | 6.12e+08 6.72e+08 |
+-----+-----+-----+-----+-----+-----+

```

❑ Using the uncensored variable (GVA), we could instruct the software to censor it in the estimation by using the subcommand: , ll(...)

Command:  
Tobit GVA asset\_value i.location, ll(0)

If have any upper limit (ul) if both, ll() ul() can be added

As we can see, both the results are same, there is a little variation in the coefficients values of OLS and tobit model

GVA's are uncensored variable. So, in this model it is not GVA1, this is only GVA where the negative values are also included. what is interesting to note here is that, please mark it, here we are mentioning in the command is lower limit. If you have any upper limit you can write down with at the end ul as the upper limit. If both are there, lower limit and upper limit with the particular value, here we are considering 0 as our lower limit because below that we have already censored.

So, upper limit if you have a particular value maybe 66000 or may be 50000 above we are not going to discuss that limit you can set. As you can see in the both the models that there is a little variation in the coefficient values of OLS and Tobit.

(Refer Slide Time: 44:34)

- ❑ Tobit regression coefficients are interpreted in the same manner as ols regression coefficients.
- ❑ For a one unit increase in asset value, there is .008 point increase in the predicted value of GVA.
- ❑ The marginal effects are just the same as from the regression model.

- ❑ Using the uncensored variable (GVA), we could instruct the software to censor it in the estimation by using the subcommand: , ll(...)

Command:  
Tobit GVA asset\_value i.location, ll(0)

As we can see, both the results are same, there is a little variation in the coefficients values of OLS and tobit model

If have any upper limit (ul). If both, ll() ul() can be added

```
tobit GVA asset_value i.location, ll(0)

Refining starting values:
Grid node 0: log likelihood = -40451.031

Fitting full model:
Iteration 0: log likelihood = -40451.031
Iteration 1: log likelihood = -40451.013
Iteration 2: log likelihood = -40451.013

Tobit regression Number of obs = 3,499
 Unobserved = 3,499
 Left-censored = 10
 Right-censored = 0
 Limits: lower = 0
 upper = +inf
 LD chi2(2) = 901.96
 Prob > chi2 = 0.0000
 Pseudo R2 = 0.0110
Log likelihood = -40451.013
```

|             | Coef.    | Std. Err. | z     | P> z  | [95% Conf. Interval] |
|-------------|----------|-----------|-------|-------|----------------------|
| asset_value | .007773  | .0003258  | 23.87 | 0.000 | .0071366 .008414     |
| location    |          |           |       |       |                      |
| outsideR08  | 17469.3  | 901.5007  | 19.38 | 0.000 | 15701.78 19236.82    |
| _con        | 3741.558 | 531.3524  | 7.04  | 0.000 | 2699.767 4783.35     |
| var(e_GVA)  | 6.41e+08 | 1.53e+07  |       |       | 6.12e+08 6.72e+08    |

Tobit regression coefficients are interpreted in the same manner as the OLS coefficients. For one unit increase in the asset value there is 8 percent, 0.008 percent increase in the predictive value of GVA. So, here it is the case. With the asset value case 0.00777 so roughly 0.008, alright. So, that is the interpretation.

Coming to this we wanted to say that, not a problem we are just finalizing. For a one unit increase in asset value it is of 0.008 point increase in the predicted value of GVA. The marginal effects are just the same as from the regression model. So, basically all other test you can do it the way we have already done it. So, rest of the details you can experiment and find out. If you have difficulties please come up with the doubts. We will be very happy to explain you in detail. So, all the qualitative dependent variable models, all the dummy variable models have already been completed. So, this is all for today in this lecture. From the next week onwards we are going to unfold discussion of panel regression using STATA. So, thank you very much.