

Ordinal Regression

Earlier ("Analysis of Ordinal Contingency Tables"), we considered ordinal variables in contingency tables. Such models do not generally assume designation of an explanatory (independent) and response (dependent) variable, but they also are limited in inclusion of covariates and more complex models that involve interactions between continuous and categorical predictors and so on. In discussing regression models thus far, the focus has been on binary response variables. But both logistic and probit regression models, however, can be applied to ordinal response variables with more than two ordered categories, such as response options of "never," "sometimes," and "a lot," which do not necessarily have equal distance between the values.¹ The application of these models is typically to response variables with 3 or 4 rank-ordered categories, and when there are 5 or more categories, moderate sample size, and fairly symmetrically distributed variables, there will be minimal loss of power if ordinary least squares is used instead of ordinal regression (Kromrey & Rendina-Gobioff, 2002; Taylor, West, & Aiken, 2006).²

For outcomes that can be considered ordinal, it is generally better to use all of the ordinal values rather than collapsing into fewer categories or dichotomizing variables, even with a sparse number of responses in some categories. Collapsing categories has been shown to reduce statistical power (Ananth & Kleinbaum 1997; Manor, Mathews, & Power, 2000) and increase Type I error rates (Murad, Fleischman, Sadetzki, Geyer, & Freedman, 2003).

Ordered Logit

Ordered logit models are logistic regressions that model the change among the several ordered values as a function of each unit increase in the predictor. With three or more ordinal responses, there are several potential forms of the logistic regression model. By far, the most common is the *cumulative logit* model, which can be conceptualized in terms of the threshold model with underlying an Y^* latent (or unobserved) continuous distribution. Instead of two observed categories and only a single threshold, there are $J - 1$ thresholds for J ordinal observed values. The logistic model can be stated similarly to the binary logistic model, except the logit is a cumulative logit.

$$\text{logit}[P(Y \leq j | X)] = \alpha_j + \beta X$$

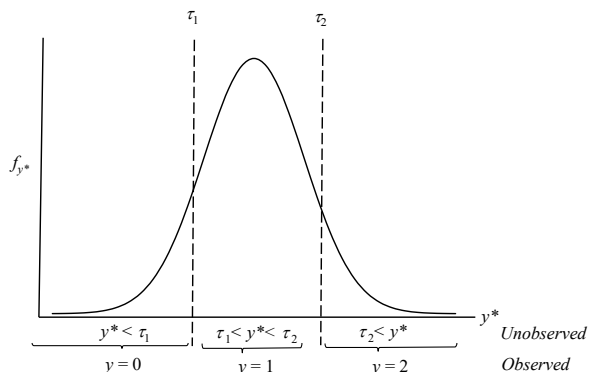
The log of the probability that Y has a value greater than the lower values given X is modeled. It is assumed that the same effect occurs for each level comparison of the ordered responses, so that the increase or decrease in odds for each unit increase in X is the same for the increment from $\ln[P(Y \leq 1)]$ to $\ln[P(Y \leq 2)]$ as from $\ln[P(Y \leq 2)]$ to $\ln[P(Y \leq 3)]$. In other words, "slopes" for predicting the logit are parallel over all of the ordered categories of the response. The curves for the probabilities will be S -curves that have the same shape and increment at the same rate in the center of the curve. *Cumulative odds ratios*, e^{β} , can also be defined in a parallel fashion to the binary case and interpreted as the odds increase from one category on the response to the next for each increment in X , where the odds are assumed to be equal across all response categories (sometimes referred to as the proportional odds model; McCullagh, 1980). A score test of the proportionality assumption is available in some software programs.

A novel aspect of the ordered logit model is that there are multiple ($J - 1$) intercepts. Each intercept is an estimate of the threshold, τ , from the Y^* distribution.³ The threshold falls between any two intercepts, $\tau_{j-1} < Y^* < \tau_j$, so that there will always be one fewer intercepts than response categories.

¹ As a reminder, we are only concerned with special treatment of binary and ordinal dependent variables, because ordinary least squares (linear) regression has assumptions about the conditional distribution (residuals). There is no need for any special treatment of binary and ordinal independent variables in linear regression (or otherwise).

² See also the "Levels of Measurement and Choosing the Correct Statistical Test" handout for my univariate statistics course for more detail and references.

³ The Azen and Walker text does not use a symbol for the threshold, but the Long (1997) reading uses τ .



The logistic transformation can be used to estimate the predicted probability in each category.

$$P(Y \leq j) = \frac{e^{\alpha_j + \beta X}}{1 + e^{\alpha_j + \beta X}} - \frac{e^{\alpha_{j-1} + \beta X}}{1 + e^{\alpha_{j-1} + \beta X}}$$

In SPSS, SAS, and R, ordinal logit analysis can be obtained through several different procedures. SPSS does not provide odds ratios using the ordinal regression procedure, but odds ratios can be obtained by exponentiation of the coefficients.⁴ In the case of an ordinal outcome with three or more categories, the odds ratio for the logit model represents the odds of the higher category as compared to all lower categories combined. In other words, it is a cumulative odds ratio representing the increased likelihood to the next highest category relative to the lower categories for each unit increase in the predictor. The Wald, likelihood ratio, and score tests are generally printed for parameter and model fit, and pseudo- R^2 values can provide some approximation of the variance accounted for in the outcome by the predictors.

Two other ordered logistic models are less frequently used. The *continuation ratio logit model* compares a response to all other response categories above it. This approach usually involves separate logistic models. The *adjacent category logit model* compares a response category to the next response category above it.

Ordered Probit

An alternative approach to regression when the response is ordinal is probit regression (see the “Generalized Linear Models” handout). The ordinal probit model has a probit link and standard normal error distribution. The threshold conceptualization is useful for the probit model as well and similar proportionality assumptions apply. Because probit models involve a normal distribution for Y^* , the thresholds are standardized score values, with most values occurring between approximately -3 and +3. As with a binary outcome, the logit and probit analysis will nearly always lead to the same conclusions (Long, 1997). The unstandardized coefficients are the change in the normal distribution value (change in the z -score) for each unit change in X . The unstandardized coefficients represent the change in the predicted z -score for every unit change in the predictor. Standardized coefficients are commonly available in software programs (or, if not, the predictors can be standardized), offering familiar interpretation and a convenient method of gauging magnitude of effect.

Software Examples

Below is an example from Karen Seccombe's project focusing on healthcare among welfare recipients in Oregon. The outcome for this model is a response to a question about how often the respondent cut meal sizes because of affordability, an indicator of food insecurity. Responses to two questions were coded into a single ordinal variable with three values, 0 = never or rarely, 1 = some months but not every month, and 2 = almost every month.

⁴ Note that with the ordinal regression procedure in SPSS and R using the logit link function, the threshold is -1 times the constant obtained in the logistic regression, so you will see opposite signed constant values in SPSS and R compared with SAS.

Case Processing Summary

		N	Marginal Percentage
cutmeal how often cut meal size	0 never or rarely	424	77.7%
	1 some months but not every month	56	10.3%
	2 almost every month	66	12.1%
Valid		546	100.0%
Missing		96	
Total		642	

Ordinal Logit Model in SPSS

```
plum cutmeal with mosmed depress1 educat marital
/link = logit
/print= parameter.
```

Model Fitting Information

Model	-2 Log Likelihood	Chi-Square	df	Sig.
Intercept Only	572.929			
Final	543.454	29.475	4	.000

Link function: Logit.

Parameter Estimates

	Estimate	Std. Error	Wald	df	Sig.	95% Confidence Interval	
						Lower Bound	Upper Bound
Threshold [cutmeal = 0]	1.884	.351	28.864	1	.000	1.197	2.571
[cutmeal = 1]	2.659	.363	53.548	1	.000	1.947	3.371
Location mosmed	.012	.024	.230	1	.631	-.036	.059
depress1	.201	.039	26.158	1	.000	.124	.278
educat	-.035	.115	.091	1	.762	-.260	.190
marital	.463	.235	3.887	1	.049	.003	.922

Link function: Logit.

Ordered Logit Model in R

Note: The `polr` function requires the outcome be a factor, and does not like categorical predictors. So, I converted predictors that were nonnumeric to numeric [I use `lessR` command below, but base R can be used too, e.g., `d$mosmed <- as.numeric(d$mosmed)`]. Missing data are also problematic with `polr`, so I used the following listwise deletion routine to remove cases with missing data on any of the variables in the model (using `lessR` code).

```
#make variables numeric for listwise deletion (different var types=double and char - won't work)
> library(lessR)
> d <- Transform(cutmeal = (as.numeric(cutmeal)))
> d <- Transform(mosmed = (as.numeric(mosmed)))
> d <- Transform(depress1 = (as.numeric(depress1)))
> d <- Transform(educat = (as.numeric(educat)))
> d <- Transform(marital = (as.numeric(marital)))

> library(lessR)
#listwise deletion to match n from regression (needed to make sure nested test has same n)
> d <- Subset(cutmeal!='NA' & mosmed!='NA' & depress1!='NA' & educat!='NA' & marital!='NA')
#always double check variable type changes and listwise deletion using str(d) and descriptive analysis

> #polr requires response to be a factor, so transform
> d$cutmeal <- factor(d$cutmeal)

> library(MASS)
> model <- polr(cutmeal ~ mosmed + depress1 + educat + marital, data=d, contrasts=NULL, method=c("logistic"))
> summary(model, digits = 3)
```

Re-fitting to get Hessian

```
Call:
polr(formula = cutmeal ~ mosmed + depress1 + educat + marital,
      data = d, contrasts = NULL, method = c("logistic"))
```

Coefficients:
value Std. Error t value

```
mosmed      0.0115      0.0239      0.483
depress1    0.2009      0.0391      5.137
educat      -0.0347      0.1153     -0.301
marital     0.4626      0.2365      1.956
```

Intercepts:

```
Value Std. Error t value
0|1 1.884 0.351 5.361
1|2 2.659 0.364 7.306
```

Residual Deviance: 718.9374

AIC: 730.9374

```
> #use AER.coefest and coefci for tests and confidence intervals
> library("AER")
> coefest(model)
```

Re-fitting to get Hessian

z test of coefficients:

```
Estimate Std. Error z value Pr(>|z|)
mosmed 0.011520 0.023850 0.4830 0.62910
depress1 0.200902 0.039107 5.1372 0.0000002788310690
educat -0.034684 0.115283 -0.3009 0.76352
marital 0.462583 0.236459 1.9563 0.05043
0|1 1.883932 0.351405 5.3611 0.0000000826971867
1|2 2.658861 0.363930 7.3060 0.0000000000002753
```

```
> coefci(model)
```

Re-fitting to get Hessian

```
2.5 % 97.5 %
mosmed -0.035330933 0.05837016
depress1 0.124081506 0.27772335
educat -0.261142777 0.19177462
marital -0.001909381 0.92707519
```

```
> #obtain odds ratios
> exp(cbind(OR=coef(model), confint(model)))
Waiting for profiling to be done...
```

Re-fitting to get Hessian

```
OR 2.5 % 97.5 %
mosmed 1.0115862 0.9656504 1.060473
depress1 1.2225055 1.1323385 1.320390
educat 0.9659105 0.7681237 1.208022
marital 1.5881708 0.9927728 2.513934
```

Probit Model in SPSS

Probit models in SPSS can be specified in several different ways. I use the PLUM procedure, but the user can use the *Ordinal* procedure (specifying probit link) or the *Probit* procedure through the menus. The *Probit* procedure requires specification of a variable with the count of total observed, so it is a less convenient approach. SPSS now has a *Generalized Linear Models* option through the menus in which ordinal logistic, probit models, Poisson, and negative binomial models can be tested.

```
plum cutmeal with mosmed depress1 educat marital
/link = probit
/print= parameter summary.
```

Model Fitting Information

Model	-2 Log Likelihood	Chi-Square	df	Sig.
Intercept Only	572.929			
Final	541.558	31.371	4	.000

Link function: Probit

Pseudo R-Square

Cox and Snell	.056
Nagelkerke	.075
McFadden	.042

Link function: Probit

Parameter Estimates

		Estimate	Std. Error	Wald	df	Sig.	95% Confidence Interval	
							Lower Bound	Upper Bound
Threshold	[cutmeal = 0]	1.150	.198	33.700	1	.000	.761	1.538
	[cutmeal = 1]	1.584	.203	61.063	1	.000	1.187	1.982
Location	mosmed	.006	.014	.218	1	.640	-.020	.033
	depress1	.121	.023	27.759	1	.000	.076	.166
	educat	-.013	.065	.040	1	.841	-.141	.115
	marital	.260	.136	3.664	1	.056	-.006	.525

Link function: Probit

As noted in the previous handout, standardized coefficients could be obtained in SPSS by prestandardizing the predictor variables using the same N (e.g., using `DESCRIPTIVE VARS=mosmed(zmosmed)`) and ignoring the significance tests in the output.

Probit Model in R

(Same precautions regarding missing data and nonnumeric variables apply to probit models).

```
> library(MASS)
> model <- polr(cutmeal ~ mosmed + depress1 + educat + marital, data=d, contrasts=NULL, method=c("probit"))
> summary(model, digits = 3)
```

Re-fitting to get Hessian

Call:

```
polr(formula = cutmeal ~ mosmed + depress1 + educat + marital,
     data = mydata, contrasts = NULL, method = c("probit"))
```

Coefficients:

	Value	Std. Error	t value
mosmed	0.00667	0.0136	0.488
depress1	0.12133	0.0230	5.281
educat	-0.01374	0.0649	-0.212
marital	0.26103	0.1359	1.921

Intercepts:

	Value	Std. Error	t value
0 1	1.152	0.198	5.829
1 2	1.587	0.203	7.833

Residual Deviance: 717.29

AIC: 729.29

```
> #use AER coeftest and coefci for tests and confidence intervals
> coeftest(model1)
```

Re-fitting to get Hessian

z test of coefficients:

	Estimate	Std. Error	z value	Pr(> z)
mosmed	0.0063709	0.0136609	0.4664	0.64096
depress1	0.1210496	0.0229848	5.2665	0.000000139036058455
educat	-0.0130156	0.0648567	-0.2007	0.84095
marital	0.2596249	0.1359329	1.9099	0.05614
0 1	1.1495693	0.1976394	5.8165	0.000000006009265638
1 2	1.5843137	0.2025726	7.8210	0.000000000000005242

```
> coefci(model1)
```

Re-fitting to get Hessian

	2.5 %	97.5 %
mosmed	-0.020464160	0.03320602
depress1	0.075899142	0.16620014
educat	-0.140417971	0.11438673
marital	-0.007397164	0.52664703

```
#nested LR comparison assumes listwise deletion used so that N is the same for both nested models
> model0 <-polr(cutmeal ~ 1,data=d,contrasts=NULL,method=c("probit"))
> summary(model,digits = 3)
> model1 <-polr(cutmeal ~ mosmed + depress1 + educat + marital,data=d,contrasts=NULL,method=c("probit"))
> summary(model,digits = 3)
```

```
#requests likelihood ratio (G-squared) comparing the deviances from the two models
> anova(model0,model1,test="Chisq")
```

Likelihood ratio tests of ordinal regression models

Response: cutmeal	Model	Resid. df	Resid. Dev	Test	Df	LR stat.	Pr(Chi)
1	1	544	748.4121				
2	mosmed + depress1 + educat + marital	540	717.0415	1 vs 2	4	31.37067	0.000002572061

```
> #use AER coeftest and coefci for tests and confidence intervals
> coeftest(model1)
```

Re-fitting to get Hessian

z test of coefficients:

	Estimate	Std. Error	z value	Pr(> z)
mosmed	0.0063709	0.0136609	0.4664	0.64096
depress1	0.1210496	0.0229848	5.2665	0.000000139036058455
educat	-0.0130156	0.0648567	-0.2007	0.84095
marital	0.2596249	0.1359329	1.9099	0.05614
0 1	1.1495693	0.1976394	5.8165	0.000000006009265638
1 2	1.5843137	0.2025726	7.8210	0.000000000000005242

```
> coefci(model1)
```

Re-fitting to get Hessian

	2.5 %	97.5 %
mosmed	-0.020464160	0.03320602
depress1	0.075899142	0.16620014
educat	-0.140417971	0.11438673
marital	-0.007397164	0.52664703

```
#obtaining the psuedo-R-sq values with modEVA package requires use of glm not polr
> model3=glm(cutmeal ~ mosmed + depress1 + educat + marital,data=d,family=binomial(link="probit"))
> summary(model3)
```

```
library(modEVA)
RsqGLM(model=model3) #model on right side of equal sign is name of my model above
```

NOTE: Tjur R-squared applies only to binomial GLMs

```
$`CoxSnell`
[1] 0.04747709
```

```
$Nagelkerke
[1] 0.07255096
```

```
$McFadden
[1] 0.04578147
```

```
$Tjur
[1] NA
```

```
$sqPearson
[1] 0.05009599
```

```
#can get standardized coefficients with reghelper
```

```
> library(reghelper)
> beta(model3, x = TRUE, y = FALSE)
```

Sample Write-Up I report only on the ordinal logistic. The probit write-up would be the same except there is no OR and the standardized coefficients would hopefully be reported. I computed the OR by using e^B)

An ordered logit model was estimated to investigate whether months on medical insurance, depression, education, and marital status predict how often meals were cut (“never,” “some months,” “almost every month”). Together, the predictors accounted for a significant amount of variance in the outcome, likelihood ratio $\chi^2(4) = 31.371$, $p < .001$. Only depression, $B = .201$, $SE = .039$, $OR = 1.22$, $p < .001$, and marital status, $B = .463$, $SE = .235$, $OR = 1.59$, $p = .049$, significantly independently predicted the

frequency of cutting meals. Each point increase on the depression scale was associated with about 22% increase in the frequency of cutting meals compared to the lower frequency categories. Married individuals were approximately 50% more likely to have an increase in the frequency of cutting meals compared to the lower categories. Overall the model accounted for approximately 4% of the variance in the outcome, McFadden's pseudo- $R^2 = .042$.

References and Further Reading

- Ananth, C. V., and Kleinbaum, D. G. (1997), "Regression Models for Ordinal Responses: A Review of Methods and Applications," *International Journal of Epidemiology*, 26, 1323-1333.
- Kromrey, J. D., & Rendina-Gobioff, G. (2002). An empirical comparison of regression analysis strategies with discrete ordinal variables. *Multiple linear regression viewpoints*, 28(2), 30-43.
- Long, J.S. (1997). *Regression models for categorical and limited dependent variables*. Thousand Oaks, CA: Sage.
- Manor, O., Matthews, S., & Power, C. (2000). Dichotomous or Categorical Response? Analysing Self-Rated Health and Lifetime Social Class," *International Journal of Epidemiology*, 29, 149-157
- McCullagh, P. (1980). Regression models for ordinal data. *Journal of the royal statistical society. Series B (Methodological)*, 109-142.
- Taylor, A. B., West, S. G., & Aiken, L. S. (2006). Loss of power in logistic, ordinal logistic, and probit regression when an outcome variable is coarsely categorized. *Educational and psychological measurement*, 66(2), 228-239.
- Murad, H., Fleischman, A., Sadetzki, S., Geyer, O., & Freedman, L. S. (2003). Small samples and ordered logistic regression: does it help to collapse categories of outcome?. *The American Statistician*, 57(3), 155-160.
- Taylor, A. B., West, S. G., & Aiken, L. S. (2006). Loss of power in logistic, ordinal logistic, and probit regression when an outcome variable is coarsely categorized. *Educational and psychological measurement*, 66(2), 228-239.