
D. Additional Analyses of Rank-Ordered Choice Data

This appendix presents additional analyses of the ranking sequence used in the survey. Section D.1 presents separate estimation results for the rank-ordered probit model for each dataset collected – ANES and FFRISP – and presents an LRT to consider whether or not the data can be pooled without significantly affecting parameter estimation. Section D.2 discusses the econometrics of several rank-ordered choice models and considers how the different model assumptions might affect WTP estimation. Section D.3 estimates a series of rank-ordered probit models and shows how the fully specified model of Chapter 8 represents the best fit for the rank-ordered choice data.

D.1 Does Pooling the Data Affect Parameter Estimation?

Table D.1 summarizes the covariates used to estimate the rank-ordered probit model in Chapter 8 over the two datasets, ANES and FFRISP.^{1, 2} None of the covariates are significantly different between the two datasets.³

1. Most discrete variables were defined as what respondents stated. For example, “married_own” equals 1 if the respondent answered “married” to the marriage question and “own a home” to the home-ownership question. Respondents who refused to answer either of those questions were left with a value of zero for the “married_own” dummy variable. The same approach was used for all other variables included in the model, with the exception of the education variable, “educ,” which was imputed to the median value of 10 (some college) for nonrespondents, and income, which is described in footnote 2. The final dataset for model estimation has 3,183 observations, which represents all respondents that completed at least part of the ranking sequence.

2. When respondents provided bracketed responses to the income question, we placed their income at the mid-point of their bracket for analysis. To assign respondents in the top income group, “Greater than \$175,000,” we referred to the 2009 ACS, which reports that of households with incomes greater than \$175,000, the median household is in the \$200,000 to \$249,999 range. We placed income for these respondents at the mid-point of this range: \$225,000.

A total of 237 respondents did not respond to the income question. For these observations, we used a hotdeck procedure, which randomly drew income responses from other observations in the dataset that were similar in terms of work status and educational attainment. Specifically, the procedure stratified over a work/not work binary variable and a binary variable for whether or not the respondent had at least an associate’s degree. This procedure allowed us to keep all observations in the dataset, while preserving sample variation.

We also applied the hotdeck procedure to observations that provided wide, open-ended brackets, such as “Greater than \$50,000.” A total of 31 respondents provided such open-ended brackets. For these cases, the hotdeck procedure randomly drew responses out of the set of observations that had income in the same open-ended bracket.

Table D.1. Comparison of covariates by dataset

Variable	Mean	Standard error	95% confidence interval	
Income				
ANES	64,677.740	1,344.307	62,041.980	67,313.510
FFRISP	63,001.340	1,785.135	59,501.240	66,501.430
Education				
ANES	10.165	0.055	10.058	10.273
FFRISP	10.068	0.079	9.914	10.223
Married_own				
ANES	0.534	0.014	0.505	0.562
FFRISP	0.510	0.019	0.473	0.547
Strong environmentalist				
ANES	0.166	0.010	0.146	0.186
FFRISP	0.184	0.015	0.155	0.212
Very strong environmentalist				
ANES	0.030	0.004	0.022	0.037
FFRISP	0.030	0.006	0.019	0.042
Def_visit				
ANES	0.301	0.013	0.275	0.327
FFRISP	0.283	0.017	0.249	0.316
Times				
ANES	3.161	0.240	2.689	3.632
FFRISP	3.338	0.423	2.509	4.166

3. Mean estimates for “married” are significantly different between the two datasets. Mean estimates for “own_home” are not significantly different. To facilitate pooling, the two variables were combined for modeling purposes. A LRT on combining these two variables finds that it does not significantly affect estimation of the rank-ordered probit model. $LRT = 3.41$, which is less than the 0.05 critical value of a $\chi^2_{(2)}$, which is 5.99.

Table D.2 presents the results of the rank-ordered probit model for the two separate datasets.⁴ Many covariates are significant at the 90% level with both datasets, including cost, the No-Fishing Zones Program dummy variable, Education X Fish, Strong environmentalist X Fish, Strong environmentalist X Ship, and Def_visit X Ship. Several variables, however, are significant with the ANES dataset but not with FFRISP dataset, including Income X Fish, Married_own X Fish, Married_own X Ship, Times X Fish, Times X Ship, and Def_visit X Fish. It is likely that the larger sample size of the ANES dataset is driving this difference in obtaining significance. The ANES dataset has 2,289 observations compared to 894 in the FFRISP dataset. The FFRISP dataset has two covariates – the Reef Repair Program dummy variable (“ship”) and Very strong environmentalist X Fish - that obtain significant t-statistics that the ANES dataset does not obtain.

Table D.2. Rank-ordered probit model estimation results

Variable	ANES	FFRISP
Cost	−0.002*** (−4.257)	−0.003*** (−4.584)
Fish	0.216** (2.475)	0.356*** (2.878)
Ship	0.029 (0.405)	0.204** (2.020)
Variables with the no-fishing zones program		
Income X Fish	0.001* (1.712)	0.002 (1.597)
Education X Fish	0.053** (2.176)	0.050* (1.829)
Married_own X Fish	−0.184** (−2.497)	−0.169 (−1.582)
Strong environmentalist X Fish	0.762*** (5.591)	0.542*** (3.614)
Very strong environmentalist X Fish	0.348 (1.607)	0.693* (1.846)
Def_visit X Fish	0.402*** (3.218)	0.169 (1.040)
Times X Fish	0.008*** (2.923)	0.004 (0.999)

4. In this appendix, estimation results for the variance-covariance matrices are presented in their transformed states. See Appendix E, footnote 5, for a discussion of how this transformation was made.

Table D.2. Rank-ordered probit model estimation results (cont.)

Variable	ANES	FFRISP
Variables with the reef repair program		
Income X Ship	0.000 (0.380)	−0.001 (−1.293)
Education X Ship	−0.013 (−0.640)	−0.026 (−1.091)
Married_own X Ship	−0.205*** (−2.841)	−0.068 (−0.727)
Strong environmentalist X Ship	0.621*** (5.523)	0.284** (2.318)
Very strong environmentalist X Ship	0.243 (1.335)	0.409 (1.425)
Def_visit X Ship	0.313** (2.436)	0.353** (2.058)
Times X Ship	0.007*** (2.736)	0.005 (1.314)
Insigma3	−0.119** (−2.087)	−0.144* (−1.791)
Insigma4	0.518*** (9.838)	0.507*** (7.369)
atanhr3_2	0.868*** (10.475)	0.838*** (7.452)
atanhr4_2	1.315*** (14.715)	1.269*** (10.203)
atanhr4_3	1.188*** (12.835)	1.057*** (8.966)
loglikelihood	−6,181.390	−2,440.570
t-stats are reported below the coefficient estimates in parentheses.		
*** Indicates significance at the 99% confidence level.		
** Indicates significance at the 95% confidence level.		
* Indicates significance at the 90% confidence level.		

To test the hypothesis that the two datasets do not differ significantly in terms of how they fit the rank-ordered probit model, an LRT can be conducted. The LRT compares the log-likelihood of the pooled model (found in Chapter 8) with the sum of the log-likelihoods from the two separate models. It is calculated as $-2[\log-L_{\text{Pooled}} - (\log-L_{\text{ANES}} + \log-L_{\text{FFRISP}})]$. Under the null hypothesis, it is distributed $\chi^2_{(22)}$. For these models, the calculated LRT is 32.60, which is less than the 0.05 critical value of a $\chi^2_{(22)}$, which is 33.92. We cannot reject the hypothesis that the two datasets do not differ in terms of the rank-ordered probit model. In other words, the two datasets do not differ significantly under this model. This result allows us to conclude that pooling the datasets to estimate the rank-ordered probit model, and ultimately WTP, is warranted. All subsequent analyses in this appendix and in Chapters 7 and 8 are therefore conducted on the pooled dataset.

D.2 Alternative Choice Models

The rank-ordered probit model is used in Chapter 8 to model the responses to the ranking sequence - Q10, Q13, and Q15 - and then to estimate WTP. There are, however, several different econometric models that have been presented in the literature as possible approaches to fitting rank-ordered data. This section discusses these different models. Equation D.1 gives the probability of observing a full ranking of alternatives k, l, m , and n for individual i :

$$\begin{aligned} P_i &= P_i(U_{ik} > U_{il} > U_{im} > U_{in}) \\ &= P(e_{il} - e_{ik} < V_{ik} - V_{il}, e_{im} - e_{il} < V_{il} - V_{im}, e_{in} - e_{im} < V_{im} - V_{in}) \end{aligned} \quad (\text{D.1})$$

The assumption made about the underlying joint distribution of the error terms – the e 's – is what determines which of the econometric models will be employed to fit the data. Appendix E assumes the errors are distributed normal, which leads to the rank-ordered probit model, the model used to estimate WTP in Chapter 8. If, however, the assumption was that the error terms were distributed extreme value – a common assumption made in the literature – then the error differences would be distributed logistic. This makes the probability that individual i selects alternative k as the first alternative in the ranking sequence:

$$P_i = P_i(U_{ik} > U_{il}, U_{ik} > U_{im}, U_{ik} > U_{in}) = \frac{e^{V_{ik}}}{e^{V_{ik}} + e^{V_{il}} + e^{V_{im}} + e^{V_{in}}} \quad (\text{D.2})$$

This is the conditional logit model, which can be estimated based on the responses to Q10, ignoring the follow-up rankings.

The probability that individual i selects the full ranking of k , l , m , and n is:

$$P_i = \frac{e^{V_{ik}}}{e^{V_{ik}} + e^{V_{il}} + e^{V_{im}} + e^{V_{in}}} \times \frac{e^{V_{il}}}{e^{V_{il}} + e^{V_{im}} + e^{V_{in}}} \times \frac{e^{V_{im}}}{e^{V_{im}} + e^{V_{in}}}, \quad (\text{D.3})$$

which is the rank-ordered logit model (Hausman and Ruud, 1987). This model is sometimes referred to as “exploded logit” because of a counter-intuitive quirk: the conditional probabilities – the second and third terms in Equation D.3 – are identical to the unconditional probabilities, meaning that the data could just as well be set up as a sequence of three separate choices made by three separate individuals (Train, 2009). Essentially, a rank-ordered logit model does not accumulate statistical information about an individual as it fits that individual’s sequence of choices. The implications of this issue are discussed in more detail in Section D.3.

There are other well-known issues associated with estimating conditional or rank-ordered logit models. The first is the independence of irrelevant alternatives (IIA). IIA results from the assumption that errors across a choice set are independent or uncorrelated. If an alternative were added (or removed) from the choice set, the logit model would predict that the original (or remaining) alternatives would be chosen in the same proportions as before the addition (removal). This can lead to some counter-intuitive results, especially if one or more of the alternatives is a close substitute.

A second issue with the logit specification is that errors are assumed to be identically distributed, meaning that they must have equal variances. In this study, the fourth alternative – the Full Program - is the combination of alternatives two and three (the separate No-Fishing Zones and Reef Repair programs). One might expect the error variance for this fourth alternative to be larger than for the other alternatives; in other words, there might be heteroskedasticity across the choice set. The standard logit specification cannot accommodate this.

The rank-ordered probit model, on the other hand, does accommodate these issues. The joint normal assumption of probit allows error terms across the choice set to be correlated and to have different variances. Allowing for correlated error terms gives the rank-ordered probit model a way of keeping track of the sequence of choices made by the same individual: the rankings are not treated as separate choices. This is why rank-ordered probit is chosen as the most appropriate model for the type of data being examined in this study. There is, however, another model in the literature that overcomes the issues discussed above: the “mixed logit” or “random coefficients” model (Hensher and Greene, 2003; Train, 2009).

The idea behind mixed logit is that the coefficients – the β ’s – might vary across the population. In other words, there might be heterogeneity in preferences for the programs. The β ’s are generally assumed to be normally or log-normally distributed, but other distributions are possible

as well. Assuming a normal distribution, the probability of observing individual i 's ranking (k , l , m , or n) becomes:

$$P_i = \int \frac{e^{V_{ik}}}{e^{V_{ik}} + e^{V_{il}} + e^{V_{im}} + e^{V_{in}}} \times \frac{e^{V_{il}}}{e^{V_{il}} + e^{V_{im}} + e^{V_{in}}} \times \frac{e^{V_{im}}}{e^{V_{im}} + e^{V_{in}}} g(b|q) dq \quad (D.4)$$

Because each individual is assumed to have a unique coefficient vector, β , the mixed logit model overcomes the problem of choices by the same individual being treated as separate choices made by separate individuals. The estimated coefficients link the sequence of choices an individual makes. Further, because of the more complex functional form, the mixed logit model overcomes the IIA problem. Mixed logit is also very flexible; in addition to random coefficients, the researcher can generate systematic errors by adding random, alternative-specific constants with means restricted to zero (Olsen, 2009). For example, to accommodate the potential higher variance for the fourth alternative, discussed above, the researcher could add a parameter to V_{i4} with zero mean and positive variance.

In preliminary examinations of the choice models, we found that a mixed logit model with correlated alternative-specific random parameters and parameterized, alternative-specific error terms predicts very similar WTP results to the fully specified rank-ordered probit model presented in Chapter 8. The Team chose to estimate the rank-ordered probit model because the model directly incorporates potential correlation and heteroskedasticity in its basic set-up. The significance of correlation and heteroskedasticity can be directly tested by looking at the t-statistics or by comparing nested models via the LRT. Section D.3 presents these tests.

D.3 Model Estimation Results

In this section, the rank-ordered probit model is presented under a range of assumptions. We present the basic rank-ordered probit model under the assumption of homoskedasticity (constant error variance across alternatives) and zero correlation in error terms across alternatives. We also present the model under the assumptions of heteroskedasticity and correlation, respectively. LRTs are conducted to determine whether each of these assumptions is warranted. The LRT on the full model of Chapter 8 – that includes both assumptions of heteroskedasticity and correlation – is also presented and the conclusion is drawn that the full model is warranted for this application.

The second column in Table D.3 presents the estimation results for the basic (or restricted) rank-ordered probit model, with the restrictions being a constant error variance (homoskedasticity) and zero correlation among error terms across alternatives. The WTP estimates based on this model are presented in Table D.4.

Table D.3. Rank-ordered probit model estimation results (N = 3,183)

Variable	Basic model	Model with heteroskedasticity	Model with correlation	Final model
Cost	−0.002*** (−4.536)	−0.002*** (−5.699)	−0.002*** (−6.102)	−0.002*** (−5.437)
Fish	0.416*** (5.508)	0.423*** (5.757)	0.204*** (4.701)	0.245*** (2.845)
Ship	0.163*** (2.590)	0.283*** (4.607)	0.052 (1.411)	0.071 (1.044)
Variables with the no-fishing zones program				
Income X Fish	0.002** (2.445)	0.002** (2.329)	0.001** (2.537)	0.002** (2.279)
Education X Fish	0.080*** (3.790)	0.075*** (3.520)	0.040*** (2.871)	0.049*** (2.652)
Married_own X Fish	−0.210*** (−2.947)	−0.212*** (−3.035)	−0.139*** (−3.079)	−0.179*** (−3.031)
Strong environmentalist X Fish	0.797*** (7.252)	0.815*** (7.427)	0.523*** (7.155)	0.691*** (6.697)
Very strong environmentalist X Fish	0.524** (2.230)	0.521** (2.358)	0.332** (2.274)	0.440** (2.365)
Def_visit X Fish	0.369*** (3.212)	0.392*** (3.404)	0.244*** (3.316)	0.333*** (3.399)
Times X Fish	0.009*** (3.160)	0.009*** (3.081)	0.005*** (2.904)	0.007*** (2.762)
Variables with the reef repair program				
Income X Ship	−0.001 (−0.754)	0.000 (−0.501)	0.000 (−0.284)	0.000 (−0.210)
Education X Ship	−0.023 (−1.318)	−0.008 (−0.465)	−0.017 (−1.415)	−0.020 (−1.313)
Married_own X Ship	−0.178** (−2.538)	−0.181*** (−2.691)	−0.124*** (−2.873)	−0.167*** (−2.938)
Strong environmentalist X Ship	0.574*** (6.800)	0.573*** (6.671)	0.382*** (6.655)	0.512*** (6.048)
Very strong environmentalist X Ship	0.326* (1.868)	0.296* (1.765)	0.217* (1.917)	0.294* (1.945)

Table D.3. Rank-ordered probit model estimation results (N = 3,183) (cont.)

Variable	Basic model	Model with heteroskedasticity	Model with correlation	Final model
Def_visit X Ship	0.371*** (3.222)	0.356*** (3.147)	0.244*** (3.263)	0.333*** (3.237)
Times X Ship	0.008*** (2.907)	0.008*** (2.982)	0.005*** (2.878)	0.006*** (2.810)
lnsigmaP1		-1.028*** (-7.240)		
lnsigmaP2		0.319*** (7.191)		
atanhrP1			1.118*** (19.610)	
atanhrP2			1.028*** (20.784)	
atanhrP3			0.885*** (18.046)	
lnsigma3				-0.124*** (-2.633)
lnsigma4				0.512*** (12.614)
atanhr3_2				0.862*** (11.592)
atanhr4_2				1.304*** (15.113)
atanhr4_3				1.159*** (14.014)
Loglikelihood	-9,489.950	-9,170.380	-9,034.300	-8,638.250

t-stats are reported below the coefficient estimates in parentheses.

*** Indicates significance at the 99% confidence level.

** Indicates significance at the 95% confidence level.

* Indicates significance at the 90% confidence level.

Table D.4. WTP results based on the four models

	WTP	Standard error	95% confidence interval	
Basic model				
No-fishing zones program	\$292.57	\$46.80	\$200.85	\$384.29
Reef repair program	\$91.18	\$14.17	\$63.40	\$118.96
Full program	\$383.75	\$55.21	\$275.53	\$491.96
Model with heteroskedasticity				
No-fishing zones program	\$272.52	\$34.92	\$204.09	\$340.96
Reef repair program	\$138.03	\$17.42	\$103.88	\$172.17
Full program	\$410.55	\$50.31	\$311.94	\$509.15
Model with correlation				
No-fishing zones program	\$215.29	\$24.01	\$168.24	\$262.34
Reef repair program	\$54.13	\$11.65	\$31.29	\$76.97
Full program	\$269.42	\$28.78	\$213.02	\$325.82
Final model				
No-fishing zones program	\$224.81	\$32.19	\$161.72	\$287.89
Reef repair program	\$62.82	\$21.73	\$20.23	\$105.40
Full program	\$287.62	\$48.04	\$193.46	\$381.78

The third column in Table D.3 presents the estimation results for the rank-ordered probit model that allows for heteroskedasticity across alternatives. The estimable variance terms for the Reef Repair Program and the Full Program (σ_3 and σ_4 , respectively) are both significantly different from the assumed base variance of 1. Specifically, the variance term for the Reef Repair Program is significantly less than 1, implying a smaller error variance for this program, and the variance term for the Full Program is significantly greater than 1, confirming our prior hypothesis about this parameter. A LRT confirms that our prior expectation that relaxing the homoskedasticity assumption significantly improves model estimation. The calculated LRT is 639.14, which is well above the 0.05 critical value of a $\chi^2_{(2)}$, which is 5.99.

Table D.3, column 4, presents the results of the rank-ordered probit model under the assumptions of homoskedasticity (constant variance) but potentially non-zero correlations across alternatives. The estimable correlation terms (see Appendix E for details on estimability of the variance-covariance matrix under rank-ordered probit) are all significantly different from zero, indicating that estimation of this model is improved by relaxing the zero-correlation restriction. The LRT comparing this model to the basic model with no correlation is 911.3, which is well above the

0.05 critical value of a $\chi^2_{(3)}$, which is 7.82, confirming our prior expectation that there is correlation across alternatives.

To test whether both heteroskedasticity and correlation significantly affect model estimation, we conducted LRTs to compare the final model of Chapter 8. These calculated LRTs are:

- } The LRT comparing the basic rank-ordered probit to the final model of Chapter 8 is 1,703.4, which is well above the 0.05 critical value of a $\chi^2_{(5)}$, which is 11.07.
- } The LRT comparing the heteroskedastic model to the final model of Chapter 8 is 1,064.26, which is well above the 0.05 critical value of a $\chi^2_{(3)}$.
- } The LRT comparing the model with correlation to the final model of Chapter 8 is 792.10, which is well above the 0.05 critical value of a $\chi^2_{(2)}$.

These tests confirm our prior expectation that the fully specified rank-ordered probit model that allows for both correlation and heteroskedasticity in the error terms across alternatives is the best model to fit the ranked data collected in this study.

References

- Hausman, J.A. and P.A. Ruud. 1987. Specifying and testing econometric models for rank-ordered data. *Journal of Econometrics* 34:83- 104.
- Hensher, D.A. and W.H. Greene. 2003. The mixed logit model: The state of practice. *Transportation* 30:133- 176.
- Olsen, S.B. 2009. Choosing between internet and mail survey modes for choice experiment surveys considering non-market goods. *Environmental and Resource Economics* 44:591- 610.
- Train, K.E. 2009. *Discrete Choice Methods with Simulation*. 2nd ed. Cambridge University Press, New York.