

Measuring Appropriability in Research and Development with Item Response Models

Matthew Johnson
Dept. of Statistics
Carnegie Mellon University
Pittsburgh, PA 15213
masjohns@stat.cmu.edu

Wesley Cohen
Dept. of Social and Decision Sciences
Carnegie Mellon University
Pittsburgh, PA 15213
wc02@andrew.cmu.edu

Brian Junker
Dept. of Statistics
Carnegie Mellon University
Pittsburgh, PA 15213
brian@stat.cmu.edu

February 16, 1999

1 Introduction

The extent to which firms are able to capture, or appropriate, the profits created by their innovations is thought to be a key determinant of the amount of Research and Development (R&D) they do. Firms lose the incentive to innovate if they are unable to capture the returns on their innovations, and may perform R&D at a less than socially optimal level. For this reason, for example, the National Institute of Standards and Technology's (NIST) Advanced Technology Program (ATP) invests in such innovations, and by cost-sharing research to foster new, innovative technologies, benefits the U.S. economy (NIST 1998).

Uncompensated or undercompensated information flowing from innovators to rivals, called spillovers, are a primary reason for the inability of a firm to appropriate the returns due to its own innovation. Information spills over from one company to another through several channels, including word-of-mouth, publications, common suppliers, joint ventures, customers, and patents. The profits captured by the innovating firm may be reduced when a greater amount of information spills over from that firm to a competing firm.

Economists call the mechanisms that firms use to protect the profits created by their innovations appropriability mechanisms (Cohen, 1996). The most studied appropriability mechanism, patents, are intended to lessen the effects of spillovers by giving the patenting firm monopoly rights to their innovation. However, patents seem to be effective only in a small number of industries (Mansfield 1986; Levin, Klovorick, Nelson, and Winter, 1987). Other appropriability mechanisms include: lead time (being first to the market with the new product or process), other legal mechanisms such as design registration and copyrights, complementary sales and service, or manufacturing facilities, and secrecy. These mechanisms are also not expected to have homogeneous success across industries.

In this paper we use the results of the 1994 Carnegie Mellon Survey of Industrial R&D in the U.S. Manufacturing Sector (CMS; Cohen, 1996) to model and examine the effectiveness of six appropriability mechanism on product and process innovations in 72 U.S. industries. The survey questions, described further in Section 2, address subjects such as information flows, time until competitors imitate innovations, the amount of R&D focused on product versus process innovations, and the effectiveness of various appropriability mechanisms. R&D unit directors from 1489 units responded to the survey, and self-reported the industries they do work in. Our goal is to determine how well the various appropriability mechanisms work in protecting firms' profits; and to determine how the effectiveness of these appropriability mechanisms varies across industries, and varies according to the focus of the R&D unit: product or process innovations.

In Section 3 we develop a mixed effects generalized linear model (Stiratelli, Laird, and Ware 1984) to help characterize the responses to the CMS survey questions on appropriability mechanisms. This class of statistical models contains item response theory (IRT; van der Linden and Hambleton, 1996) models which are commonly used in educational testing. We adapt a fixed effects generalized linear model to model within respondent cross question dependence. This model is similar to the Partial Credit Model (PCM; Masters 1982) from IRT. Our model allows us to accommodate both the hierarchical structure of the data (respondents contained in industries), and

the polytomous nature of the responses.

Knowledge of between industry differences in the ability to appropriate R&D profits is important to governmental agencies such as the ATP in deciding which companies to help with the cost of innovation. The existence of industry effects in our model estimation allows the ATP to search for industries unable to capture the profits due to their innovations, and help them so R&D does not fall below socially optimal levels. In Section 4, using our statistical model, we investigate three hypotheses. We investigate whether industries differ in their ability to appropriate returns created by innovations. We investigate the effectiveness of each mechanism relative to the others. We also examine whether response behavior is a function of the focus, product versus process innovations, of the R&D.

2 The 1994 Carnegie Mellon Survey of Industrial R&D in the U.S. Manufacturing Sector

2.1 Overview

Our analysis of appropriability mechanisms used in the manufacturing sector relies on a subset of questions from The 1994 Carnegie Mellon Survey of Industrial R&D in the U.S. Manufacturing Sector (CMS; Cohen, 1996). The CMS builds on an earlier survey of appropriability and technological opportunity conditions in the American manufacturing sector by Levin, et al. (1987), and contains 63 multi-part questions.

The CMS was sent to 3240 R&D unit directors in the manufacturing sector. The sampling frame was built mostly from the Directory of American Research & Technology. The sample was then stratified into 74 industry groups such as pharmaceuticals, semiconductor, computer, steel, or automobile industries. Each stratum was sampled systematically with random starting points. The sampling procedure used sampling weights to oversample from smaller industries and Fortune 500 companies, and undersample from the larger industries. The sampling weights ranged from 1.0, for industries with fewer than 31 cases and Fortune 500 companies, and 0.24 for the largest industry, Measuring and Controlling Devices.

There has been a long debate on the relevance of sampling weights on statistical modeling in surveys such as the CMS where sampling weights have been employed. When likelihood-based approaches to inference are used, there is evidence suggesting that using the sampling weights is “at best irrelevant” (Fienberg, 1989). Our analysis consists of likelihood-based methods, and we therefore did not incorporate the sampling weights into the model.

Survey responses from 1489 units, or 46% of those sent out, were returned. Each respondent was classified into one of 77 industries according to what they identified as their focus industry from the CMS.

A followup survey of those that did not respond to the CMS was also conducted. Analysis of the non-response survey found that many of the non-respondents were in fact not in the manufacturing

sector, and should not have been in the sampling frame. Removing R&D units that are not in the manufacturing sector from the sample increases the response rate to 54%. Further analysis of the followup survey is ongoing.

Our analysis is based on four questions from the CMS, questions 1, 32, 33, and 46. For the convenience of the reader we list the four questions and describe the scoring of the responses.

Question 1: Please list the main industry or industries to which your R&D unit’s activities apply. If you list more than one, circle the one that is the principal focus of your R&D effort. We will refer to this industry as your *focus industry*.

If the respondent listed more than one, the circled industry was the industry the respondent was classified in. Questions 32 and 33 assess the effectiveness of six appropriability mechanisms in protecting product and process innovations.

Questions 32 and 33: During the last three years for what percent of your *innovations* were each of the following effective in protecting your firm’s competitive advantage for those innovations?

	0-10%	10-40%	41-60%	61-90%	91-100%
a. Secrecy	1	2	3	4	5
b. Patent Protection	1	2	3	4	5
c. Other Legal Mechanisms	1	2	3	4	5
d. Being First to Market	1	2	3	4	5
e. Complementary sales/service	1	2	3	4	5
f. Complementary manufacturing	1	2	3	4	5

Table 1: Appropriability mechanism effectiveness item response choices. The items were asked once for product innovations (Question 32), and once for process innovations (Question 33), producing 12 survey responses.

The questions were asked once for *product innovations* (Question 32), and once for *process innovations* (Question 33). The percentages tied to each score level were intended to help respondents use the full five-level Likert scale for each mechanism. Since it was not clear whether the respondents would interpret the percentages as intended (percent of innovations for which each mechanism protected the firm’s competitive advantage), as opposed for example to a subjective level of effectiveness for that firm’s product or process innovations generally, the responses were coded and analyzed using the five ordinal levels 1, 2, 3, 4, 5, instead of the percentages. In Section 2.2 we also explore the extent to which respondents did use the percentage category labels as intended.

In the remainder of the paper the responses to each of these questions will be referred to as the effectiveness of the corresponding appropriability mechanism on products or processes.

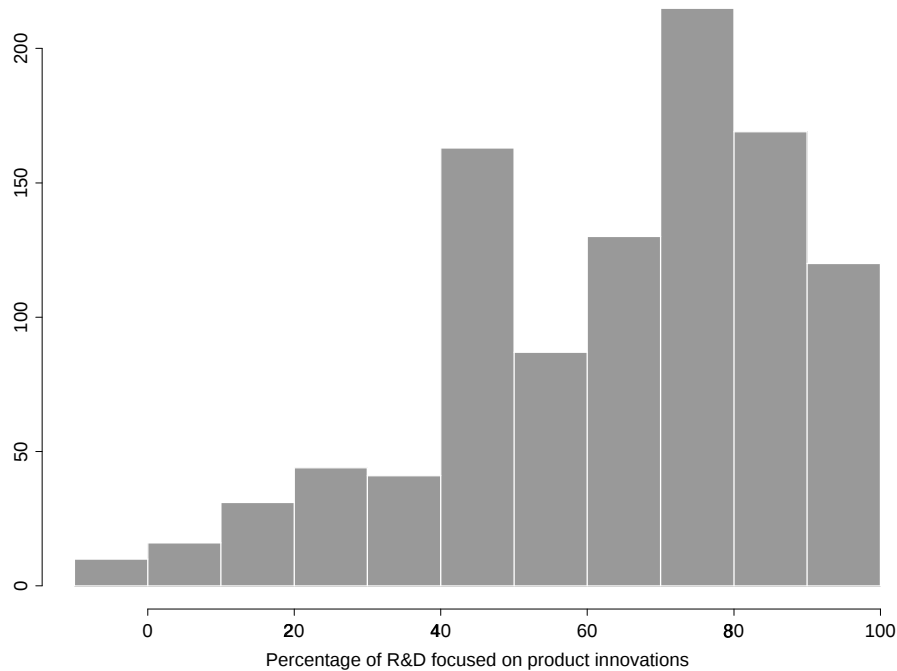


Figure 1: Histogram of the distribution of responses to question 46(b), the percentage of R&D focusing on new or improved products.

Question 46: Approximately what percentage of your R&D effort focuses on:

- a. New or improved processes _____%
- b. New or improved products _____%
- c. Other (specify)_____ %

The response to question 46(b) was used as a covariate to help explain differences in effectiveness of appropriability mechanisms for products versus processes in our formal analyses. The distribution of responses to question 46(b) appears in Figure 1.

The responses to these four questions were used for this analysis. After removing the respondents who left any one of the four questions blank, and those who come from industries with fewer than three observations there remain 1026 responses, classified into 72 industries.

2.2 Exploratory Analysis

2.2.1 Exploration of the percentage category weights

Recall from Section 2.1 that respondents were asked in CMS questions 32 and 33 to respond “for what percent of your innovations were each of the following effective in protecting your firm’s

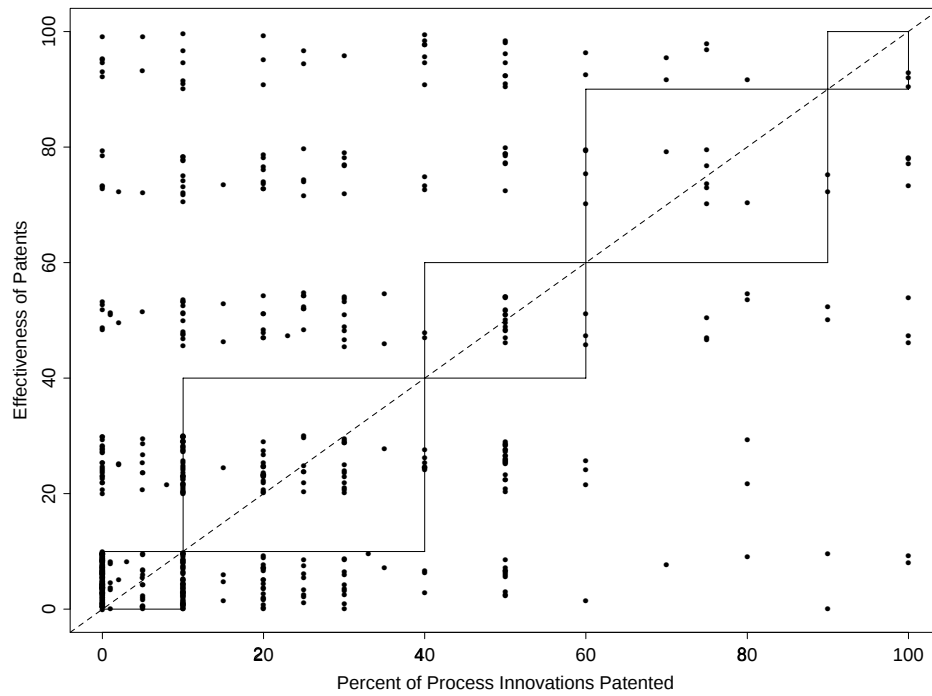
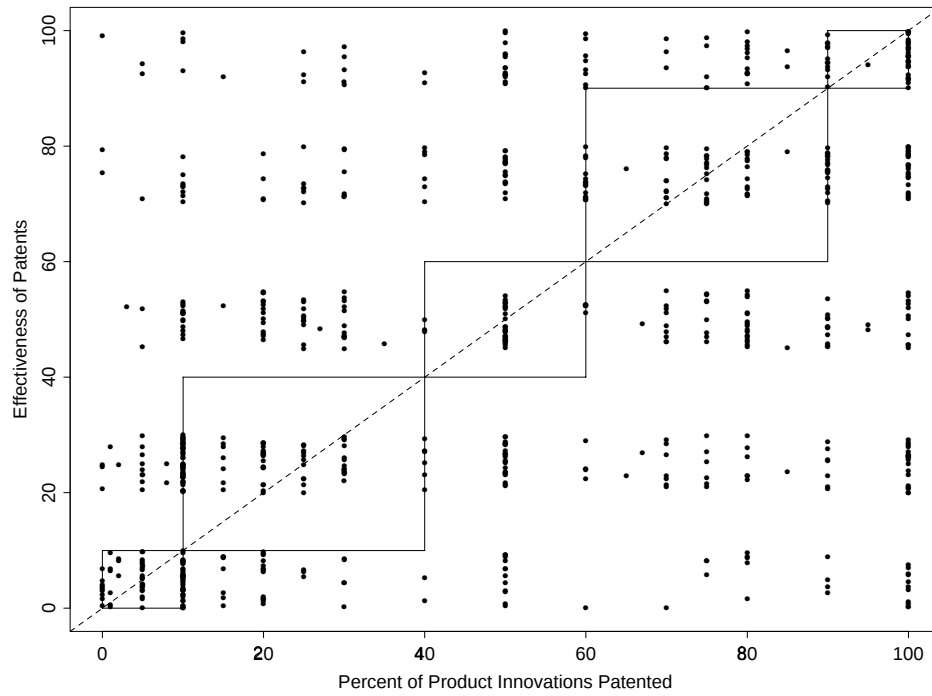


Figure 2: The comparison of percent of innovations patented to the percent of innovations on which patents were effective.

competitive advantage for those innovations?” for each of the six mechanisms, once for product innovations, and again for process innovations. It turns out that for one of these mechanisms, patents, respondents were also asked in CMS question 36 to write in the “approximate percent” of process and product innovations the firm applied patents for. Comparing the wording of questions 32 and 33 with the wording of question 36 on the CMS questionnaire, we see that if respondents are (a) calibrating their responses to agree with the percentage range category labels in Table 1, and (b) responding to these patent effectiveness questions as if they cover all possible uses of patents, then the responses to questions 32 and 33 (for what percent of process [resp. product] innovations were patents effective) should never be greater than the corresponding response to question 36 (for what percent of process [resp. product] innovations were patents applied for).

In Figure 2 we have plotted for each responding R&D unit their response to question 36 (percent applied for) on the horizontal axis and their responses to question 32 (33) (percent effective) on the vertical axis. The squares along the 45°-line represent the percentage intervals used in question 32 and 33, within which it is impossible to tell whether the reported percent effective (questions 32 and 33) is higher or lower than the reported percent applied for (question 36). Outside of these squares, if respondents are calibrating their answers to agree with the percentage intervals in Table 1, and assuming the question covers all possible uses of patents, we would expect to see no points above the 45° line. Clearly, this expectation fails. Either the respondents are not calibrating their answers to the given percentage intervals, or they are using patents for purposes other than protecting the competitive advantage for the patented innovations and only consider patents used for this purpose when responding to the patent effectiveness questions.

A plausible, less literal reading of question 32 (33) is along the lines of “for what percent of the innovations on which [eg. patents] were used to protect competitive advantage were they effective in protecting your firm’s competitive advantage for those innovations?”, i.e. effectiveness conditional on use, rather than effectiveness as an outcome subordinate to total use. The response behavior in Figure 2 would be consistent with this “conditional” reading of question 32 (33). Compounding the difficulty of interpreting the result of Figure 2 is the fact that, unlike the other appropriability mechanisms, patents are often used for reasons other than for protecting competitive advantage—for example, firms may patent innovations to reward or measure R&D personnel performance, or for use in broad cross-licensing negotiations (i.e. patent swapping). If respondents tacitly excluded these other uses of patents before answering question 32 (33), they could be calibrating their answers to the percentage intervals labeling each response category in Table 1, and yet percentages for question 32 (33) could exceed those of question 36.

Thus, the results of Figure 2 cast some doubt on the assumption that respondents are both calibrating their responses to the percentage intervals labeling each response category as well as considering all uses of patents. However these two sources of error cannot be separated in this study, and the previous paragraph gives several reasons that the second error is at least as plausible as the first. As a matter of caution, and to simplify the statistical modeling problem, we proceed in the rest of the paper to ignore the percentage interval labels and model the responses to each question as discrete ordinal Likert scale responses. In future work it would be interesting to further test and compare the purely ordinal approach we take, versus the partially censored continuous-response approach that is suggested by the percentage interval category weights.

2.2.2 Graphical Analysis

Graphical and exploratory analyses reveal apparent heterogeneity in the response behaviors in different industries. For example in Figures 3, 4, and 5, histograms of the counts of R&D units responding in each response category to each of the appropriability mechanism questions have been plotted, for each of the 72 industries. Here both the number of respondents from each industry and the response behavior of the respondents in that industry can be examined.

In Figure 3 we notice that the number of respondents from each industrie range from three in panels 7, 10, 29, 34, 40, to 62 in Medical Instruments (panel 70). Also in these figures we notice that secrecy appears to work equally well for process and product innovations, in contrast to what we see in patents, where respondents score the effectiveness of patents higher on product innovations than they do for process innovations. The opposite appears to be the case for lead time, in Figure 5, where lead time appears much more effective for process innovations than for product innovations. Finally, we can examine the varying effectiveness of the mechanisms across industries, and within industries. For example in the plots for the product patents questions, in Figure 4, it appears that Medical Instrument companies have a tendency to score in higher categories than the average industry. If we compare to Communications Equipment (panel 60), it appears that a respondent taken from the medical instrument industry is likely to score higher than a respondent from the communications industry.

2.2.3 Factor Analyses

To examine the within-subject between-question dependence, an exploratory factor analysis was done on the 12 appropriability mechanism items. The factor analysis was done using two methods. The first method assumes that the item responses are continuous with values in the range $(1, \dots, 5)$, while the second method uses the fact that the responses are ordinal.

Assuming the 12 responses from the appropriability questions were continuous, a factor analysis (Mardia, Kent, and Bibby, 1979) was done using the Splus (Statistical Science Inc., 1993) function `factanal`. Letting $Y_{i,j}$ be the item response to question j by respondent i , θ_i be the unobservable factor, λ_j be the loading matrix, η be the mean vector, and ε be the residual effect we have:

$$Y_{i,j} = \lambda_j' \theta_i + \eta_j + \varepsilon \quad (1)$$

$$\theta_i \sim \text{MVN}(0, \Sigma)$$

The percent of variance in Y explained by θ increases only from 49% to 52% when the number of factors increases from three to four; while increasing from two to three factors increases the amount of variance explained by 11%; see Figure 6. For this reason, we will assume the presence of three latent factors. The varimax rotated loadings appear in Table 3. The italic entries in Table 3 correspond to loadings which are greater than 0.4. Notice that the same mechanisms for process and product innovations load highly on the same factor, and that patents and other legal

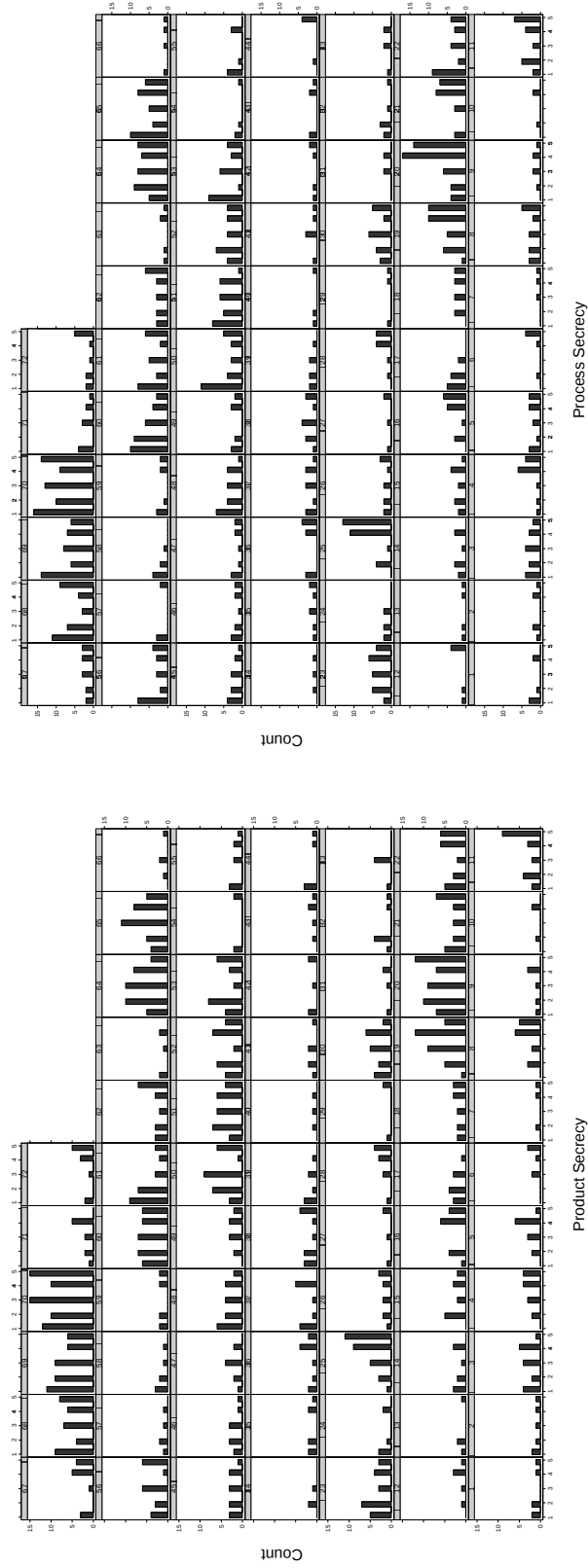


Figure 3: Secrecy mechanism effectiveness on products and processes.

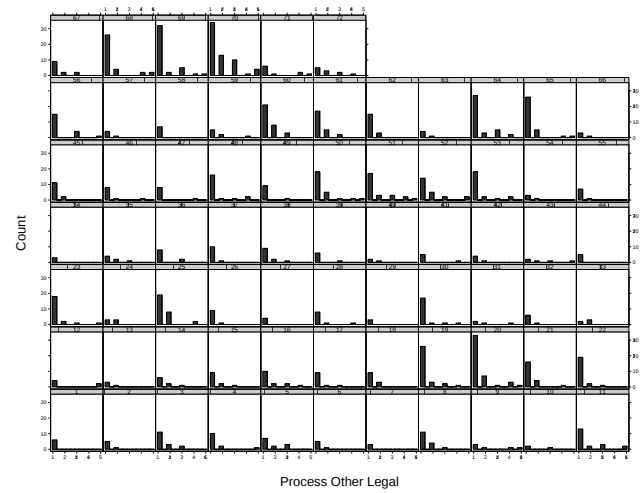
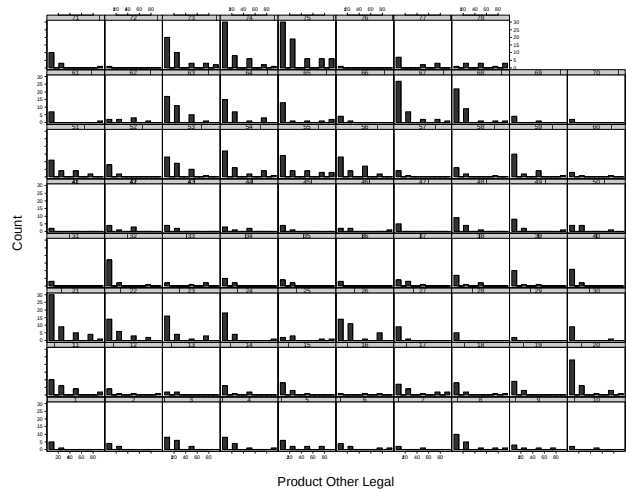
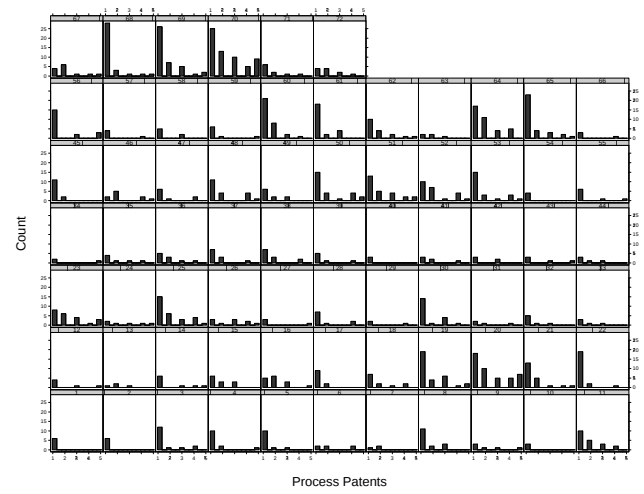
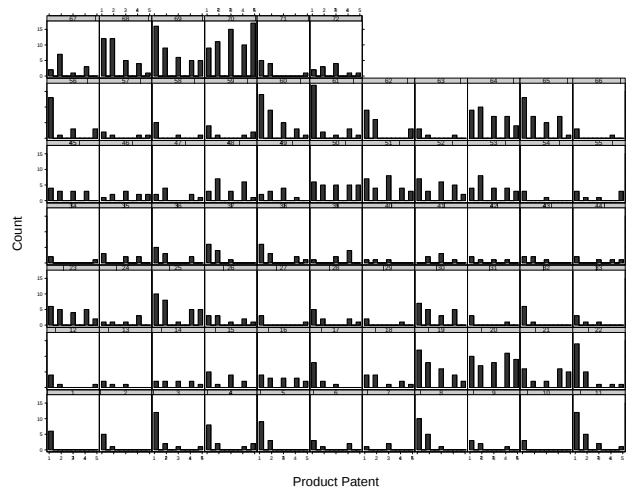


Figure 4: Legal mechanisms effectiveness on products and processes.

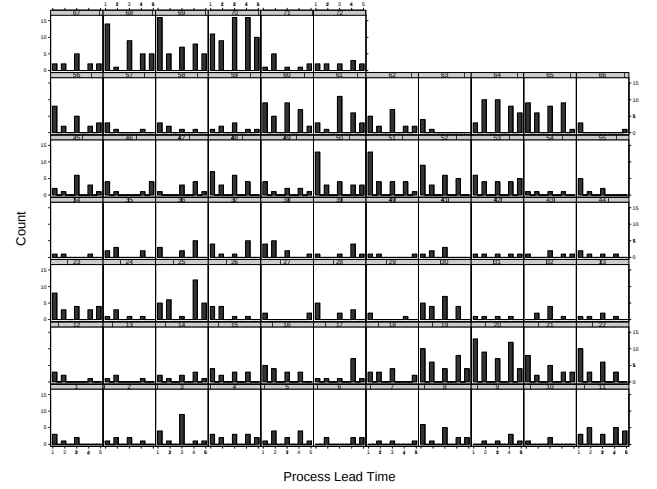
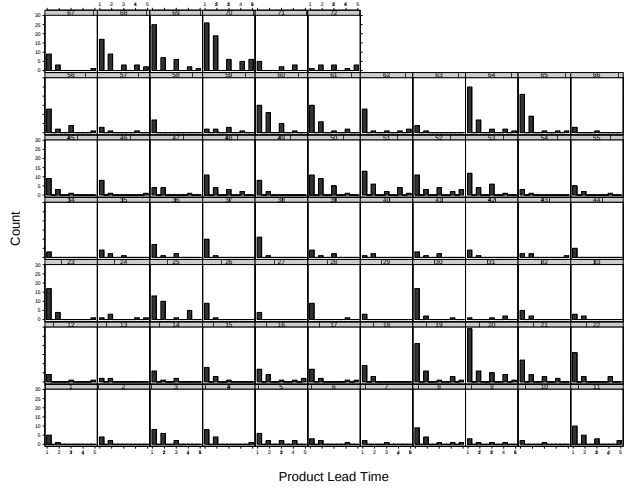
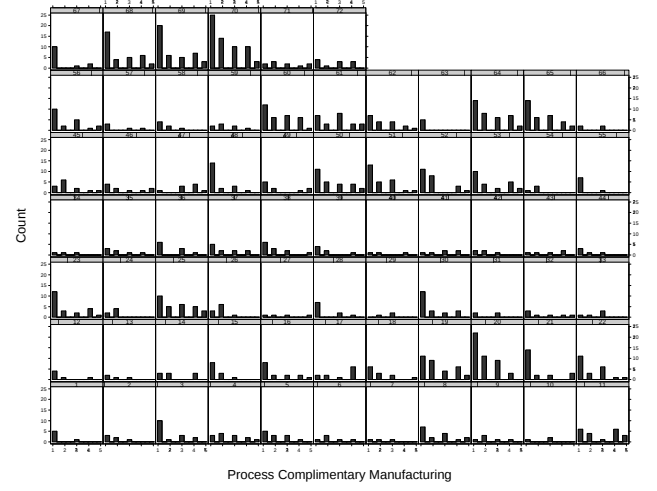
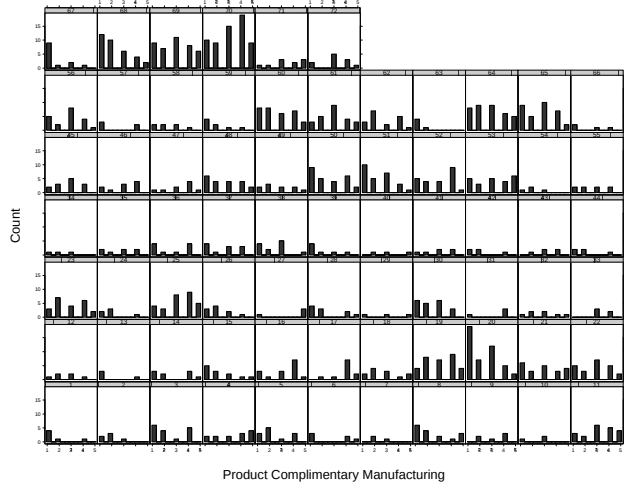
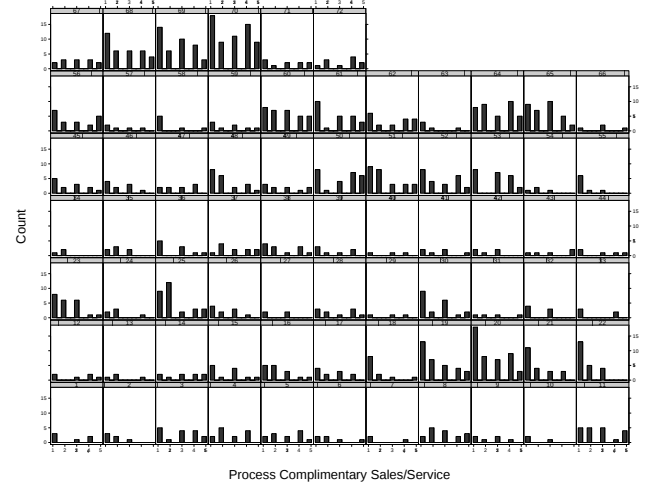
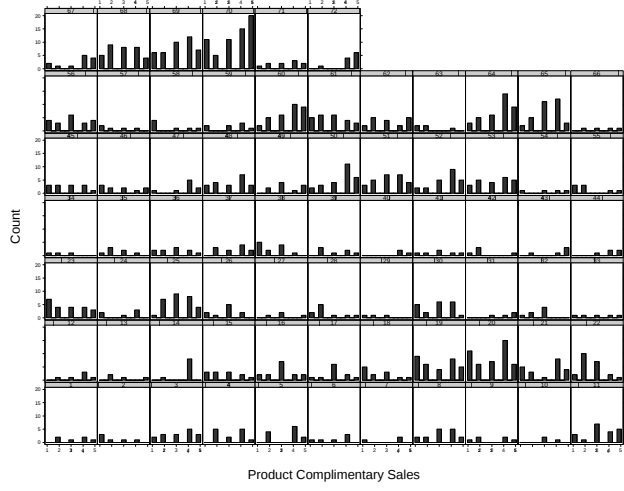


Figure 5: Complementary and Lead Time mechanisms on products and processes.

mechanisms load high on the same factor, and that complementary and lead time mechanisms load high on the first factor. This standard factor analysis suggests three factors, one corresponding to complementary mechanisms, one corresponding to legal mechanisms, and a third corresponding to secrecy.

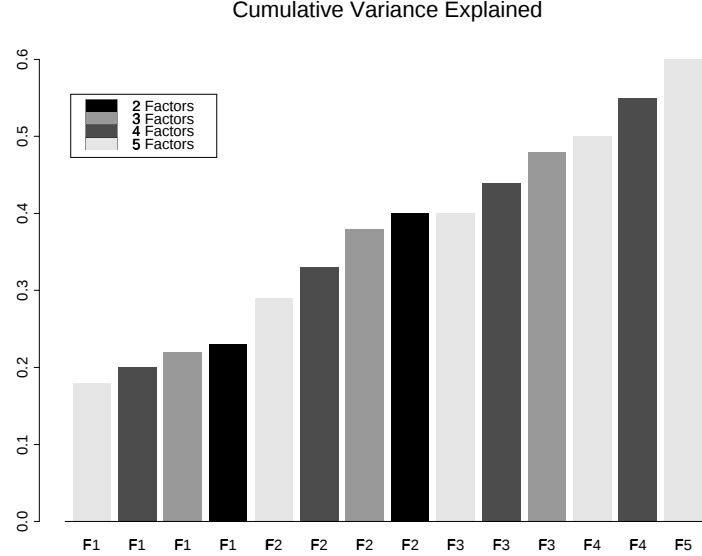


Figure 6: Cumulative variance explained by 2, 3, 4, and 5 factor factor analyses of the 12 appropriability mechanism items.

	Factor 1	Factor 2	Factor 3
SS loadings	2.7447197	1.9103095	1.2514021
Proportion Var	0.2287266	0.1591925	0.1042835
Cumulative Var	0.2287266	0.3879191	0.4922026

Table 2: Variance explained by three factor factor analysis of questions 32 and 33.

Standard factor analysis ignores the discrete nature of the data. Full-information factor analysis (Muraki and Carlson 1995) fits a factor analytic model to discrete ordinal data. Estimation of the full-information factor analysis model was carried out using Muraki’s software package POLYFACT. Using this software we find similar loadings to those found using the Splus function `factanal`. The varimax rotated loadings are in Table 4, with loadings greater than 0.2 italicized.

The goal of these preliminary factor analyses of the 12 questions on the effectiveness of appropriability mechanisms was to examine the dependence structure of the questions, as a guide in our more detailed modeling in Section 3 below. Both the continuous-data and discrete/ordinal-data factor analyses indicate that there are three fairly well-defined groups of appropriability mechanisms: secrecy mechanisms; legal mechanisms; and complementary capabilities and lead time mechanisms.

Question Description	Factor 1	Factor 2	Factor 3
32a–Product Secrecy	0.128	0.047	0.777
32b–Product Patents	-0.047	0.656	0.079
32c–Product Other Legal	0.223	0.608	0.116
32d–Product Lead Time	0.423	0.167	0.264
32e–Product Complementary Sales	0.704	0.026	0.032
32f–Product Complementary Manf.	0.752	0.033	0.072
33a–Process Secrecy	0.119	0.172	0.690
33b–Process Patents	0.051	0.729	0.077
33c–Process Other Legal	0.275	0.648	0.073
33d–Process Lead Time	0.514	0.258	0.225
33e–Process Complementary Sales	0.759	0.142	0.066
33f–Process Complementary Manf.	0.710	0.101	0.103

Table 3: Varimax rotated loadings using the S-plus function `factanal`.

Two aspects of the factor loadings in Table 3, from the continuous-data factor analysis, and Table 4, from the discrete-data factor analysis, deserve special mention. First, lead time mechanisms (questions 32d and 33d) do not load well on any factor in either three-factor model. To explore whether lead time mechanisms should be considered on a separate dimension, we reviewed the four-factor models. However, the fourth factor in both models is essentially a bipolar factor contrasting product and process questions, unrelated to lead time. Re-examining Figure 5, it seems that the reason for the weak lead time loadings is that there is little variability in the responses to these two survey items. However, because of its conceptual relation to other complementary capabilities we retained the lead-time items in our detailed analyses in Sections 3 and 4 below.

Second, nearly all of the discrete-data factor loadings in Table 4 are smaller in magnitude than the corresponding continuous-data loadings in Table 3. Differences in discrete-data and continuous-data factor models are not unusual—indeed the effect of estimating a continuous-data factor model on data that should be analyzed with a discrete-data factor model is that the continuous-data model usually contains an extra factor (Hambleton and Rovinelli 1986). The fact that, despite the attenuated loadings, both models in this case indicate three factors suggests a fairly strong three-factor structure that is not particularly sensitive to whether we model the observed data as continuous or discrete.

Nevertheless, the first factor seems less well-defined in the discrete-data factor model. In particular, other legal mechanisms to protect process innovations (question 33c) looks as though it has a marginally strong loading on Factor 1 (complementary capabilities and lead times) as well as a clear loading on Factor 2 (patents and other legal mechanisms). With weaker loadings overall in the discrete-data model, it is not surprising that one or more factors would be more poorly defined. Given the overall strength of the three-factor model we have identified, and our *a-priori* belief that legal mechanisms should be more closely related to patents than to complementary sales and manufacturing capabilities, we have kept the legal mechanisms with patents in Factor 2, and

kept complementary capabilities and lead time mechanisms together in Factor 1, in our detailed analyses below.

In Table 5 we summarize the results from these two exploratory factor analyses. We also indicate some indexing notation that will be useful in developing our general model in the next section: i indexes the R&D unit, j indexes the item, or appropriability mechanism, and m indexes the factor, or mechanism group.

Question	Factor 1	Factor 2	Factor 3
32a–Product Secrecy	0.112	0.044	0.678
32b–Product Patents	-0.048	0.565	0.070
32c–Product Other Legal	0.177	0.428	0.109
32d–Product Lead Time	0.234	0.100	0.151
32e–Product Complementary Sales	0.564	0.006	0.042
32f–Product Complementary Manf.	0.668	0.028	0.066
33a–Process Secrecy	0.092	0.162	0.534
33b–Process Patents	0.049	0.709	0.084
33c–Process Other Legal	0.290	0.567	0.072
33d–Process Lead Time	0.314	0.161	0.127
33e–Process Complementary Sales	0.679	0.116	0.060
33f–Process Complementary Manf.	0.569	0.102	0.071

Table 4: Varimax rotated loadings of the appropriability mechanisms items, using POLYFACT (Maraki, 1997).

3 Hierarchical model for ranking industries

In this section we describe the model, fitting, and hypothesis testing procedure we used for the main part of this study. The reader uninterested in theoretical details may skip to Section 4 after browsing through Section 3.1.

3.1 Model Description

R&D units’ responses on the effectiveness of the twelve appropriability mechanisms are expected to be affected by the industry in which the respondent does research; see Figures 3, 4, and 5. Patents for example are believed to have greater success in chemical industries. Also noticeable in Figures 3, 4, and 5 is that responses vary according to the mechanism of interest. The responses are also believed to be a function of the proportion of R&D effort focused on developing new products versus new processes. One would not expect a company that performs only limited research in process innovations to rate the effectiveness of the mechanisms on processes in the same way as

	Secrecy	Legal	Complementary Capabilities & Lead Time
Product	Secrecy $j = 1$	Patents $j = 2$ Other Legal $j = 3$	Comp. Man. $j = 4$ Comp. Sales/Service $j = 5$ Lead Time $j = 6$
Process	Secrecy $j = 7$	Patents $j = 8$ Other Legal $j = 9$	Comp. Man $j = 10$ Comp. Sales/Service $j = 11$ Lead Time $j = 12$
m_j Mechanism Group	1	2	3
Q_m Total	2	4	6

Table 5: Summary of factor analyses of the appropriability mechanism effectiveness items.

it would the mechanisms on product innovations. These considerations motivate our statistical model.

Let $Y_{i,j}$ be the response of respondent i to the j^{th} survey question ($i = 1, \dots, N$; $j = 1, \dots, Q$), and let x_i be the proportion of R&D effort company i focuses on product innovations. Also define $p_{jk} = \Pr\{\text{Response in category } k \text{ on } j^{th} \text{ question}\}$, $k = 1, \dots, K$. We will model the effects of these variables on the survey responses by regressing the log odds ratio of adjacent response categories on these variables:

$$\log \left\{ \frac{p_{j(k+1)}}{p_{jk}} \right\} = \mu_{\rho_i} + \beta_j(x_i - \bar{x}) - \alpha_j - \gamma_{j,k} \quad (2)$$

for $k = 1, \dots, K - 1$; $j = 1, \dots, Q$, where $\rho_i = 1, \dots, 72$ is the focus industry R&D unit i performs R&D in, and $\bar{x}$ is the average proportion of R&D effort on product innovations across all industries and R&D units. In the data considered here, $N=1026$, $Q=12$, and $K=5$.

Because we have set the constraint $\sum_{k=1}^{K-1} \gamma_{j,k} = 0$ for all j , α_j represents the point at which a respondent who satisfies $\mu_{\rho_i} + \beta_j(x_i - \bar{x}) > (=, <) \alpha_j$ has a greater (equal, lesser) chance of responding in category K versus category 1; see Figure 7. Similarly, a respondent satisfying $\mu_{\rho_i} + \beta_j(x_i - \bar{x}) > (=, <) \alpha_j + \gamma_{j,k}$ will have a greater (equal, lesser) chance of responding in category $k + 1$ versus k .

While both standard factor analysis and full-information factor analysis done with POLYFACT capture the within respondent across mechanism dependence, neither allows us to use the hierarchical or clustered structure of the data as we have done Equation (2). Also, Equation (2) may not explain enough of the within respondent across mechanism dependence through the variable $x_i - \bar{x}$.

To capture the within respondent across question dependence structure, and hierarchical struc-

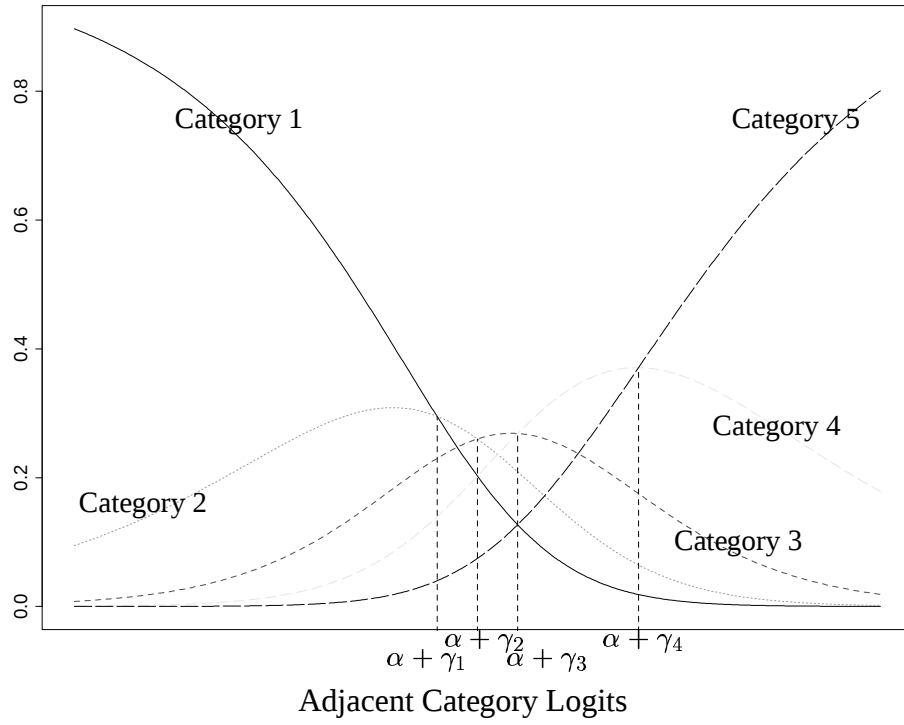


Figure 7: An example of item response probabilities (vertical axis) as they are affected by industry mean and proportion of R&D effort focused on product innovations (horizontal axis). $\alpha + \gamma_k$ is the point at which response $k + 1$ becomes more likely than response k . For each survey item the γ_k sum to zero.

ture, we now adapt the model in Equation (2) to a model that is similar to an item response theory model used in educational testing, known as the Partial Credit Model (PCM; Masters, 1982): We add a random intercept, θ_i , for each respondent in Equation (2). Using the information learned from the earlier factor analyses and to simplify analysis, the questions will be split into three mechanism groups (secrecy, legal, and complementary capabilities and lead time), and analyzed separately.

Our expanded model is:

$$\begin{aligned}
 & \bullet \text{ For secrecy mechanisms } (j = 1, 7), \\
 & \log \left\{ \frac{P_{j,k+1}(\theta_{i,1})}{P_{j,k}(\theta_{i,1})} \right\} = \theta_{i,1} - \alpha_j - \gamma_{j,k} + \beta_j(x_i - \bar{x}) \\
 & \bullet \text{ For legal mechanisms } (j = 2, 3, 8, 9), \\
 & \log \left\{ \frac{P_{j,k+1}(\theta_{i,2})}{P_{j,k}(\theta_{i,2})} \right\} = \theta_{i,2} - \alpha_j - \gamma_{j,k} + \beta_j(x_i - \bar{x}) \\
 & \bullet \text{ For complementary and lead time mechanisms } (j = 4, 5, 6, 10, 11, 12), \\
 & \log \left\{ \frac{P_{j,k+1}(\theta_{i,3})}{P_{j,k}(\theta_{i,3})} \right\} = \theta_{i,3} - \alpha_j - \gamma_{j,k} + \beta_j(x_i - \bar{x})
 \end{aligned} \tag{3}$$

Here θ_{i,m_j} is the random effect for R&D unit i , and mechanism group m . If $K=2$, this is a linear logistic reparameterization of the Rasch Model (van der Linden and Hambleton, 1996). To model industry effects (respondent i is in industry ρ_i) we will allow the mean of the latent distribution to depend on industry ($\theta_{i,m_j} \sim N(\mu_{\rho_i,m_j}, \sigma_{m_j}^2)$). Analogous to our discussion of model 2 the R&D unit satisfying $\theta_{i,m_j} + \beta_j(x_i - \bar{x}) > (=, <) \alpha_j$ has a greater (equal, lesser) chance of responding in category K versus category 1 on mechanism j , and similarly for the other category relationships; see again Figure 7.

Let us also define

$$\log \left\{ \frac{P_{j,k+1}(\theta_{i,m_j})}{P_{j,k}(\theta_{i,m_j})} \right\} = Z_{j,k}(\theta_{i,m_j})$$

$i = 1, 2, \dots, N$; $j = 1, 2, \dots, Q$; $k = 1, \dots, K - 1$ For the model in Equation (3) we have the following likelihood (to simplify notation we will let ϕ denote the vector of all model parameters).

$$L(\mathbf{y}; \phi) = \prod_{i=1}^N \prod_{j=1}^Q \frac{\exp(\sum_{k=1}^{y_{i,j}} Z_{j,k}(\theta_{i,m_j}))}{\exp(\sum_{l=1}^K \sum_{k=1}^l Z_{j,k}(\theta_{i,m_j}))} \tag{4}$$

For the CMS we have $N=1026$ responses from R&D unit directors on $Q=12$ questions from the CMS. Each unit is classified into one of 72 industries. Referring back to Section 2.1 we see that for each question each R&D unit responds to one of $K=5$ categories.

3.2 Model Fitting

To estimate parameters in various versions of the model in Equation (3) we used two different approaches.

The marginal maximum likelihood method implemented in CONQUEST (Wu, Adams, and Wilson, 1997) uses an E-M algorithm that treats the θ_{i,m_j} as censored or missing data, and allows relatively fast fits even for the full 3-dimensional model specified in 3, but the constraint $\beta_j \equiv 0$ is necessary. We used CONQUEST to confirm the factor structure discovered in our exploratory analyses in Section 2.2 and to check that correlations between $\theta_{i,1}$, $\theta_{i,2}$, $\theta_{i,3}$ were not large.

It was difficult to fit the full latent regression of innovation effort, estimating β_j as free parameters, in CONQUEST. Also, certain model comparisons detailed below were not easy to make using the asymptotic likelihood ratio chi-squared framework. Therefore a final model fitting was done in a fully Bayesian formulation of the model in Equation (3), using Markov chain Monte Carlo methods with the program BUGS (Bayesian Inference Using Gibbs Sampling; Spiegelhalter, Thomas, Best, and Gilks, 1996). Separate one-dimensional models of the form 3 were fitted for each appropriability index $\theta_{i,1}$, $\theta_{i,2}$, $\theta_{i,3}$ because BUGS slows considerably when the number of parameters grows within a model.

Appendix A contains the CONQUEST and BUGS model files for fitting Equation (3).

3.3 Model Checks

In order to check the validity of model (3), we choose a test statistic $T(y)$ and calculate the p-value for some fixed value of the model parameter ϕ . If the p-value falls below some predetermined threshold we determine the model does not fit the observed data sufficiently well.

For the PCM we have used here, two statistics often used are the unweighted, “outfit” and weighted, “infit” mean square statistics (Masters, 1997) because they focus attention on diagnosing misfit of particular items. We will use the outfit statistic for checking our model.

$$\text{Outfit for Question } j: \quad T_j(y|\phi) = \sum_{i=1}^N \frac{(y_{i,j} - E_{i,j})^2}{NW_{i,j}}, \quad (5)$$

where $y_{i,j}$ is respondent i ’s response to question j , $E_{i,j}$ is the expected value of $Y_{i,j}$ conditional on the parameter vector

$$\phi = (\theta_{1,1}, \dots, \theta_{N,1}, \theta_{1,2}, \dots, \theta_{N,2}, \theta_{1,3}, \dots, \theta_{N,3}, \alpha_1, \dots, \alpha_Q, \beta_1, \dots, \beta_Q, \gamma_{11}, \dots, \gamma_{QK}),$$

and $W_{i,j}$ is the variance of $Y_{i,j}$ also conditional on ϕ .

The outfit statistic $T_j(y|\phi)$ will also be conditional on the nuisance parameter vector ϕ through the expected value and the variance. If we can calculate the posterior distribution of our model parameter vector ϕ , then we may find the *posterior predictive p-value* (Gelman, Meng, and Stern 1996): the posterior predictive p-value is the expected value of the classical p-value over the posterior distribution of the parameter vector given the model H and the observed data y .

Mechanism	on Products	Processes
Secrecy	0.200	0.076
Patents	0.000	0.838
Other legal	0.289	0.993
Lead time	0.000	0.018
Comp. sales	0.371	0.998
Comp. manf.	0.974	0.745

Table 6: Posterior predictive p-values found using the outfit statistic for each survey question. Full model 3 was fit separately within each latent appropriability index as indicated by the groupings. See also Tables 1, and 5 for additional details.

We can approximate the posterior predictive p-value will be approximated by comparing the observed values of the test statistic $T_i(y|\phi)$ to values of the test statistic for data simulated from the model at all values of the MCMC sample produced by BUGS. Let $\mathbf{y}_1^*, \dots, \mathbf{y}_M^*$ be data simulated from the model with corresponding parameters $\phi_1, \dots, \phi_M$ taken from the Markov Chain generated by BUGS. Then we may estimate the posterior predictive p-value for the fit of the i^{th} item as

$$p \approx \frac{\#\{s : T_i(\mathbf{y}|\phi_s) < T_i(\mathbf{y}_s^*|\phi_s); s = 1, \dots, M\}}{M}.$$

If this value is small (less than 0.05, say), then there is reason for concern about the fit of our model to that particular question.

In Figures 8 and 9 we have plotted the simulated versus observed values of the outfit statistic for the twelve mechanism questions so the models can be visually checked. Plotted points below the diagonal line contribute to the posterior predictive p-value. For numerical examination of the model we look to Table 6 where there are the computed p-values found using the posterior predictive procedure.

Examining the plots and the calculated p-values, the responses to the survey questions concerning product patents are problematic. It is not clear why the model does not fit this mechanism very well. The other two problematic survey questions are the two concerning lead time for both process and product innovations. The fact that the model does not work well with these survey questions is not a complete surprise. Examining again the rotated factor loadings in Section 2.2.3 produced using POLYFACT, we find that the lead time mechanisms do not load nearly as high as the other four mechanisms in this group, and hence it is possible that lead time should not be included in this mechanism group.

3.4 Hypothesis Testing

To test for industry effects, mechanism effects, and the effect of the amount of R&D focused on product versus process innovations, we have a nested hypothesis testing problem. Each of the three

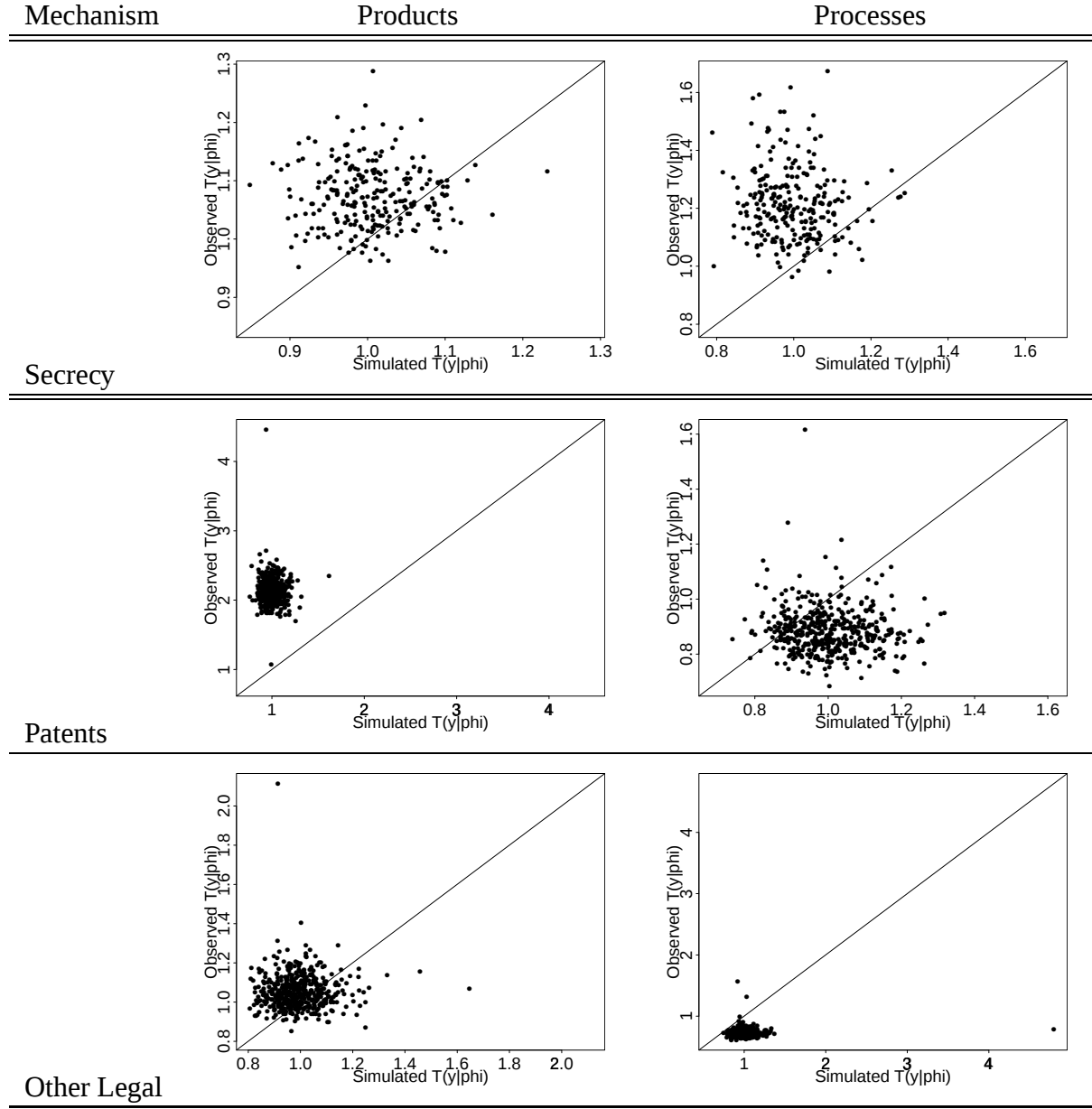


Figure 8: Posterior predictive plots using the outfit statistic for each survey question (p-values is proportion of scatter plot below the 45° line). Full model in Equation (3) was fit separately within each mechanism group. These plots are for Secrecy and Legal mechanism survey questions.

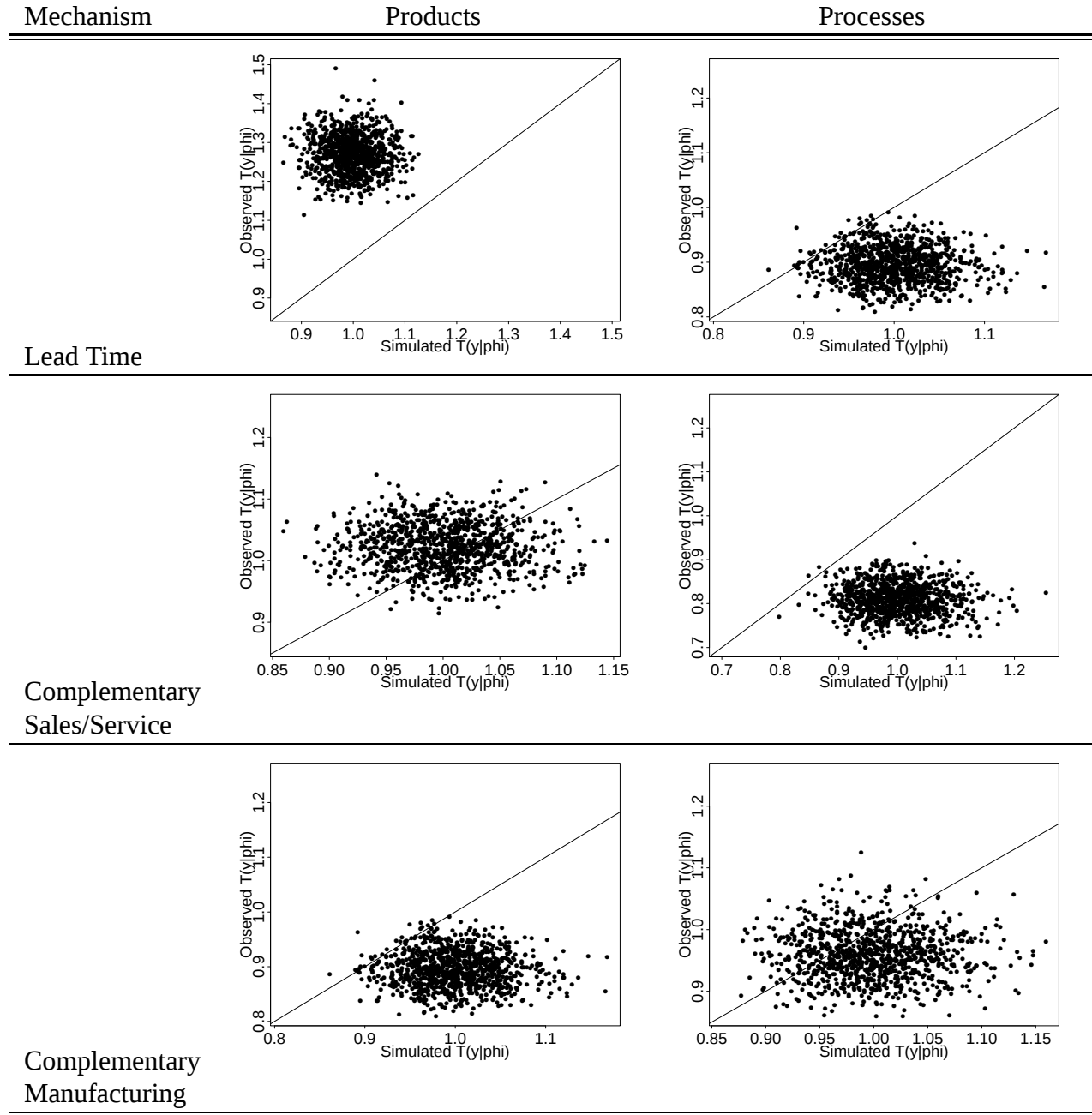


Figure 9: Posterior predictive plots using the outfit statistic for each survey question (p-value is proportion of scatter plot below the 45° line. Full model in Equation (3) was fit separately within each mechanism group. These plots are for Complementary and Lead Time mechanisms.

effects will be tested for each of the three mechanism groups (secrecy, legal, and complementary capabilities and lead time). Recall from Equation (3) that $\mu_{\rho,m}$ is the mean of the random effect of a firm performing primarily in industry ρ for mechanism group m ; α_j is the mechanism parameter for mechanism j , where there are Q_m mechanisms in group m ; and β_j is the latent slope parameter for mechanism j . Our three hypotheses for each mechanism group are then (see Equation (2) and Equation (3) for definitions of parameters):

1. Mechanism effects: $H_1 : \alpha_1 = \dots = \alpha_{Q_m} = 0$ v. $A_2 : \alpha_i \neq \alpha_j$ for some $i \neq j$
2. Product effort effects: $H_2 : \beta_1 = \dots = \beta_{Q_m} = 0$ v. $A_3 : \beta_i \neq 0$ for some i .
3. Industry effects: $H_3 : \mu_1 = \dots = \mu_{72} = 0$ v. $A_1 : \mu_i \neq \mu_j$ for some $i \neq j$

A common procedure for testing nested hypotheses like these is to use likelihood ratio tests (LRT). For the PCM, it has been found that the LRT follows an asymptotic χ^2 distribution (Wu, Adams, and Wilson 1997). However, we cannot be sure that our sample size is large enough to compare our LRT to a χ^2 distribution for our hypothesis tests. If we instead use estimates of the posterior distributions generated by BUGS, we can use Bayes factors (BF; Kass and Raftery, 1995), a Bayesian analogue of LRT whose distribution we may estimate directly. If we assume, for each t , that the prior probabilities for models H_t and A_t are equal, we have a formula for the BF similar to that of the common LRT. The difference is that we now integrate $L(\mathbf{y}; \boldsymbol{\phi})$ over the posterior distribution of the parameter vector $\boldsymbol{\phi}$ instead of taking the supremum over the null hypothesis space, or the supremum over union of the null hypothesis and alternative.

Using the Markov chain simulations from BUGS, we approximate $\Pr\{data|A_t\}$ and $\Pr\{data|H_t\}$ with the harmonic mean (Kass and Raftery, 1995) of the likelihood for models A_t and H_t . To test the hypotheses we then compare our estimated BF to the guidelines for evidence against model H_t described by Kass and Raftery (1995). The results of these tests will be described in Sections 4.1 and 4.2 below.

BF	$2\ln\{\text{BF}\}$	Evidence against the null hypothesis
1-3	0-2	Not worth more than a bare mention
3-20	2-6	Positive
20-150	6-10	Strong
>150	>10	Very Strong

Table 7: Cutoffs for interpreting Bayes Factors, (Kass and Raftery, 1995).

4 Results

4.1 Comparing the effectiveness of the mechanisms

Recall from Section 3.1 that the question parameter α_j is the threshold for the random effect θ_{i,m_j} , above which respondent i is more likely to answer in the highest category than they are in the lowest category. It can also be interpreted as the average of the random effect thresholds at which scoring in category k is more likely than scoring in $k - 1$. Thus α_j is a measure of mechanism ineffectiveness: the higher α_j , the less likely a R&D unit will respond in a high category on item j .

Using the Bayes Factor procedure discussed in 3.4 we test for the presence of differences in the effectiveness of the mechanisms. The null hypothesis, H_1 , is that the twelve appropriability mechanisms are equally ineffective. We test this hypothesis against the alternative that at least one of the mechanisms is either more or less effective. These hypotheses were tested separately within each mechanism group. Questions concerning the effectiveness of legal mechanisms, and complementary/lead time mechanisms exhibit strong evidence that there exists differences in these mechanisms ($2\ln\{\text{BF}\}=516$ (legal), 484 (comp./lead time)). However there is only slight evidence of this in secrecy mechanisms ($2\ln\{\text{BF}\}=3.8$).

Table 8 summarizes the posterior medians and 95% equal-tailed intervals for mechanism ineffectiveness. It appears that R&D units use secrecy as an appropriability mechanism more effectively than the other two groups (legal, complementary and lead time), and legal mechanisms are the least effective. It is also noteworthy that the effectiveness of secrecy on product innovations does not greatly differ from the effectiveness of secrecy on process innovations. This was highlighted above using Bayes factors where we found only slight evidence of differences, $2\ln\{\text{BF}\}=3.8$, in the effectiveness of secrecy on product v. process innovations. The fact that the secrecy estimates are close to zero indicates that an average respondent from an average industry is equally likely to respond in the highest or the lowest effectiveness category; for the other two mechanism groups the average R&D unit is more likely to score in the lowest response category than it is to score in the highest response category. In all cases but lead time, these appropriability mechanisms appear to be more effective with product innovations than they are with process innovations. This result may have arisen because the majority of innovations done are product innovations.

The point estimates in Table 8 can also be used to rank the six different mechanisms on their effectiveness to capture profits created by innovative efforts: secrecy is considered by the responding R&D units to be the most effective of the appropriability mechanisms, followed by complementary manufacturing, lead time, complementary sales, and patents; and other legal mechanisms are found to be the least effective.

To graphically understand the estimates in Table 8, as well as the item-step parameters, γ_{jk} in Equation (2) and (3) (numerical summaries not given), we turn to Figures 10, 11, and 12 which display industry means and category thresholds for each of the three mechanism groups respectively. By fixing a value on the vertical axis we may examine both the ineffectiveness parameters

Mechanism	Product	Process
Secrecy	-0.05 (-.21,.09)	0.01 (-.14,.15)
Patents	1.01 (.86,1.16)	1.68 (1.52,1.86)
Other Legal	1.87 (1.65,2.06)	2.29 (2.06,2.51)
Being First to the Market	0.21 (.01,.44)	0.30 (.10,.66)
Complementary Sales/Service	0.56 (.22,.73)	0.76 (.57,.98)
Complementary Manufacturing	0.18 (-.01,.36)	0.01 (-.20,.21)

Table 8: Posterior medians of mechanism ineffectiveness measures, α_j (95% equal-tailed intervals in parentheses).

and the item-step parameters at a fixed level of product innovation effort. The lines in the graph are the thresholds at which it becomes most likely to score above the category indicated. All the points discussed earlier in this section are also noticeable in these plots. In particular there is a noticeable shift of the thresholds in the legal mechanisms, indicating this mechanism group is considered the least effective by the respondents.

Figures 10, 11, and 12 also help us to understand how the survey respondents of the survey are using the five categories. In particular it appears that the difference between threshold three and threshold two is considerably smaller than the differences between the other adjacent thresholds. This indicates that respondents are not using category two as much as the other categories.

4.2 The effect of the proportion of R&D focused on product innovations

The model in Equation (3) also allows us to examine the effect of the proportion of R&D focused on product innovations on R&D units' response behavior through the term $\beta_j(x_i - \bar{x})$. All other things being equal, a positive value of β_j means that mechanism j is considered to be more effective by R&D units focusing an above average proportion of R&D on product innovations, and the mechanism is considered less effective by firms focusing less than average of innovations on products.

Again using the Bayes Factor procedure discussed in 3.4 we test for the presence of this effect in our data. We test the null hypothesis, H_2 , that response behavior is not affected by the proportion of R&D one focuses on product innovations versus the alternative, A_2 that the response behavior is affected by this. Recall that the three mechanism groups are being fit separately and we therefore test for this effect separately for each of the three groups. In all three mechanism groups the point

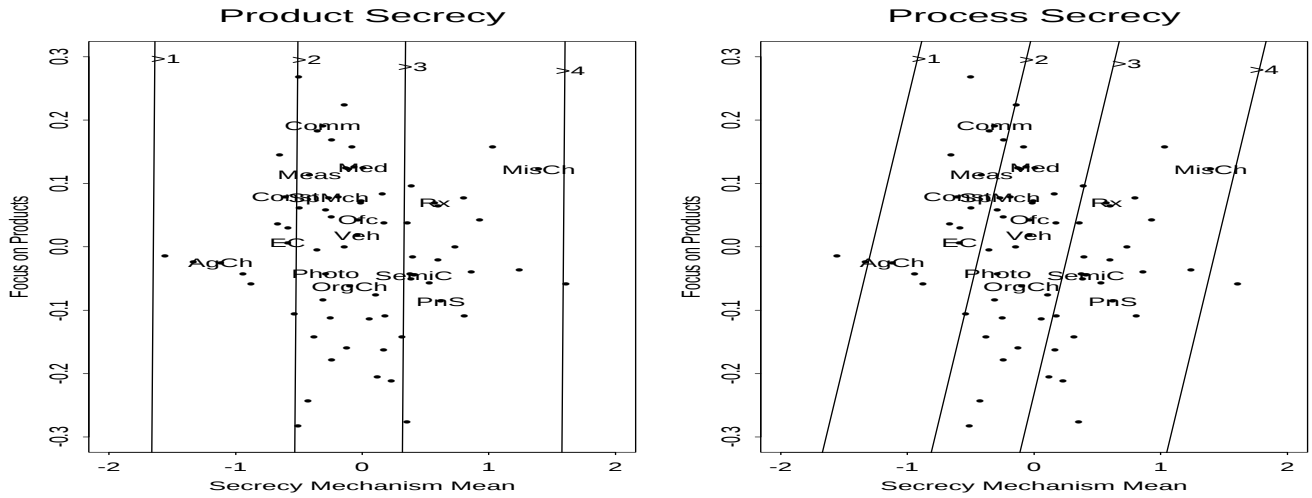


Figure 10: Contour plots of the item response surfaces, and industry means for secrecy mechanism survey questions. Top 15 industries (see Table 10) indicated on the plots.

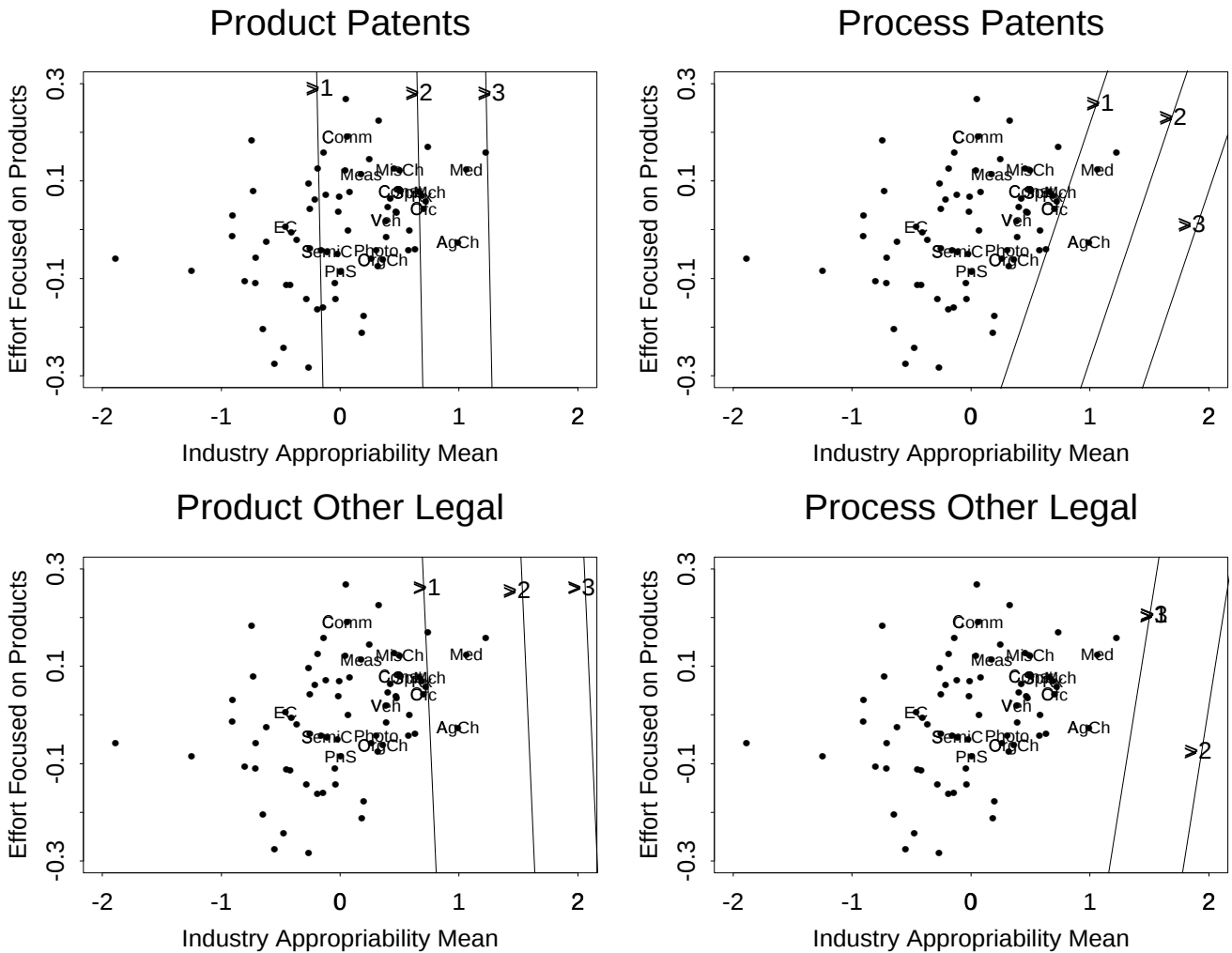


Figure 11: Contour plots of the item response surfaces, and industry means for legal mechanism survey questions. Top 15 industries (see Table 10) indicated on the plots.

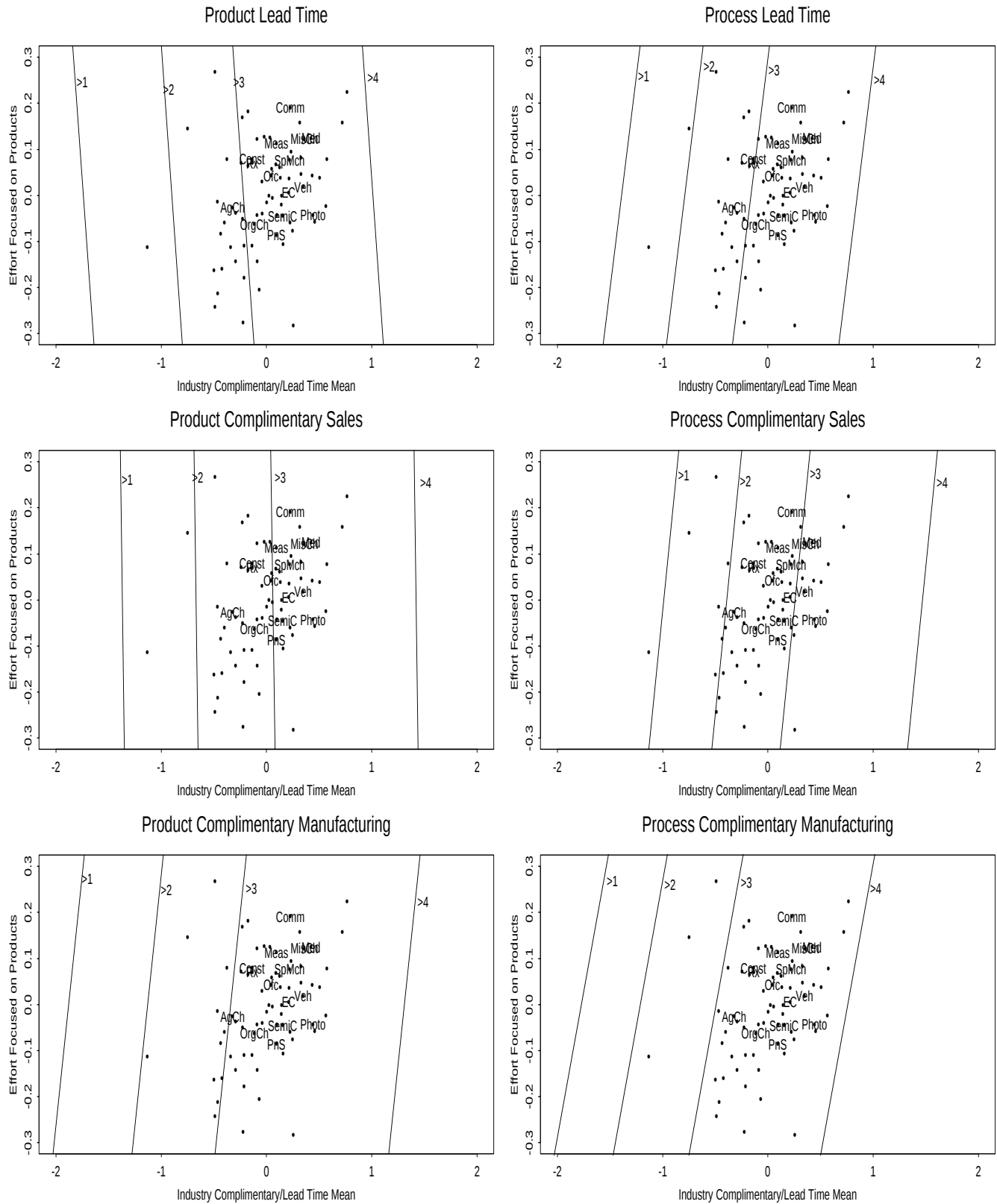


Figure 12: Contour plots of the item response surfaces, and industry means for complementary/lead time mechanism survey questions. Top 15 industries (see Table 10) indicated on the plots.

Mechanism	Product	Process
Secrecy	-.04 (-.64,.46)	-1.21 (-1.75,-.62)
Patents	0.08 (-.44,.57)	-1.38 (-1.88,-.86)
Other Legal	0.05 (-.40,.71)	-0.66 (-1.21,-.08)
Being First to the Market	0.32 (-.12,.66)	-0.53 (-.90,-.17)
Complementary Sales/Service	0.05 (-.35,.52)	-.43 (-.81,.06)
Complementary Manufacturing	-0.47 (-.84,-.01)	-0.77 (-1.20,-.40)

Table 9: Posterior median of the effect of focus on product innovations versus process innovations, β_j (95% equal-tailed intervals in parentheses).

estimates for $2\ln\{\text{BF}\}$ indicate very strong evidence against the hypothesis that there is no effect of R&D focus: $2\ln\{\text{BF}\} = 21$ (secrecy), 91 (legal mechanisms), 97 (complementary/lead time). However if we calculate estimate 95% intervals for these Monte Carlo estimates of the Bayes Factors we have little evidence that $2\ln\{\text{BF}\}$ is above zero: $[-93,96]$ secrecy, $[13,159]$ legal, $[-13,111]$ (complementary/lead time). These highly volatile estimates are due to the instability of the harmonic mean used in approximation of the Bayes factor (this is a drawback of the computational method, not necessarily a weakness in the signal provided by the data).

The volatility of the Monte Carlo estimates of the Bayes factors also occurs in the calculation of Bayes factors for testing mechanism ineffectiveness but does not affect our conclusions for legal and complementary/lead time mechanisms because we have such strong evidence against the null hypothesis.

Table 9 summarizes the posterior median and 95% equal-tailed intervals for β_j , the effect of R&D focus for each of the twelve survey mechanism questions. The proportion of R&D focused on product innovations does not appear to have any effect on how R&D units are responding to questions concerning product innovations, with the exception of the negative effect this proportion has on the effectiveness of complementary manufacturing.

This is in contrast to the survey questions concerning process innovations. For process innovations, respondents reporting an above average proportion of R&D focused on product innovations score the mechanisms less effective. This is true for all process mechanisms with the exception of the effectiveness of lead time which is not affected by the proportion.

To understand the differences of the effect of focus on product innovations we again turn to Figures 10, 11, and 12. Now we will fix a point on the horizontal axis and examine what happens as the proportion of R&D focused on product innovations vary. For example, if we examine the

effect on responses to the effectiveness of secrecy on process innovations we find a large effect. If we examine a firm with an average appropriability random effect ($\theta_{i,m_j}=0$), we see that this firm may be most likely to respond above the third category if it performs 20% to 30% more of its innovations on processes than the average. However, if an individual with the same random effect performs an average proportion of R&D on process innovations, the firm will most likely score above the second category, not the third.

If we contrast the plots for process innovations to those for product innovations we see that the contour lines for process innovations are steeper. Also noting that in most cases that the contour line for product innovation mechanisms are nearly vertical. This is consistent with our observations in Table 9 that there is little or no effect of the amount of R&D focused on product innovations on the reported effectiveness of the mechanisms on product innovations.

4.3 The effect of industries and ranking the industries

Let us turn now to the effect of the focus industry the responding R&D unit performs its research and development in. To test for the existence of this effect we use the estimated Bayes factors discussed in 3.4. We test the null hypothesis, H_3 , that there is no difference across industries in the ability to appropriate profits created by R&D to the alternative A_3 , that at least one industry responds either in higher or lower categories on the appropriability mechanisms survey questions with respect to the others. Again, the three mechanism groups are analyzed separately. We find that the estimated Bayes factors for all three mechanism groups indicate significant differences in the ability of industries to appropriate using these mechanisms ($2\ln\{BF\}=125$ for secrecy, 70 for legal mechanisms, 30 for complementary/lead time mechanisms). However, similarly to what we encountered when testing for the effect of differences in R&D focus, we find our Monte Carlo estimates of the Bayes Factors to be highly volatile, with the following intervals for secrecy, legal, comp./lead time respectively: [25,111], [-22,161], and [-30,70]. This suggests we can only be confident that industries differ in their use of secrecy to capture returns created by their innovations (again, our uncertainty here is due to weakness of the computational method, not necessarily lack of signal in the data).

Finally, we rank the industries according to their ability to use the three mechanism groups. In doing so we will be able to find industries in which there is an increased chance of market failure, because R&D units focusing innovations in low-ranking industries may find it unprofitable to perform R&D, increasing the chances of such market failures. We define the rank r_ρ of industry ρ on mechanism group m by:

$$r_\rho^{(m)} = \sum_{i=1}^{72} I_{\{\mu_{\rho,m} \geq \mu_{i,m}\}} \quad (6)$$

where $\mu_{\rho,m}$ is the latent mean of industry ρ 's appropriability in mechanism group m , $I_{\{\mu_{\rho,m} \geq \mu_{i,m}\}} = 1$ if $\mu_{\rho,m} \geq \mu_{i,m}$ and 0 else. The posterior distributions of the industry rankings of the 15 "top industries" as defined in Cohen 1996 (see Table 10) are summarized in Figure 13. The industry abbreviation is placed at the distribution median, and the line for each industry extends from the

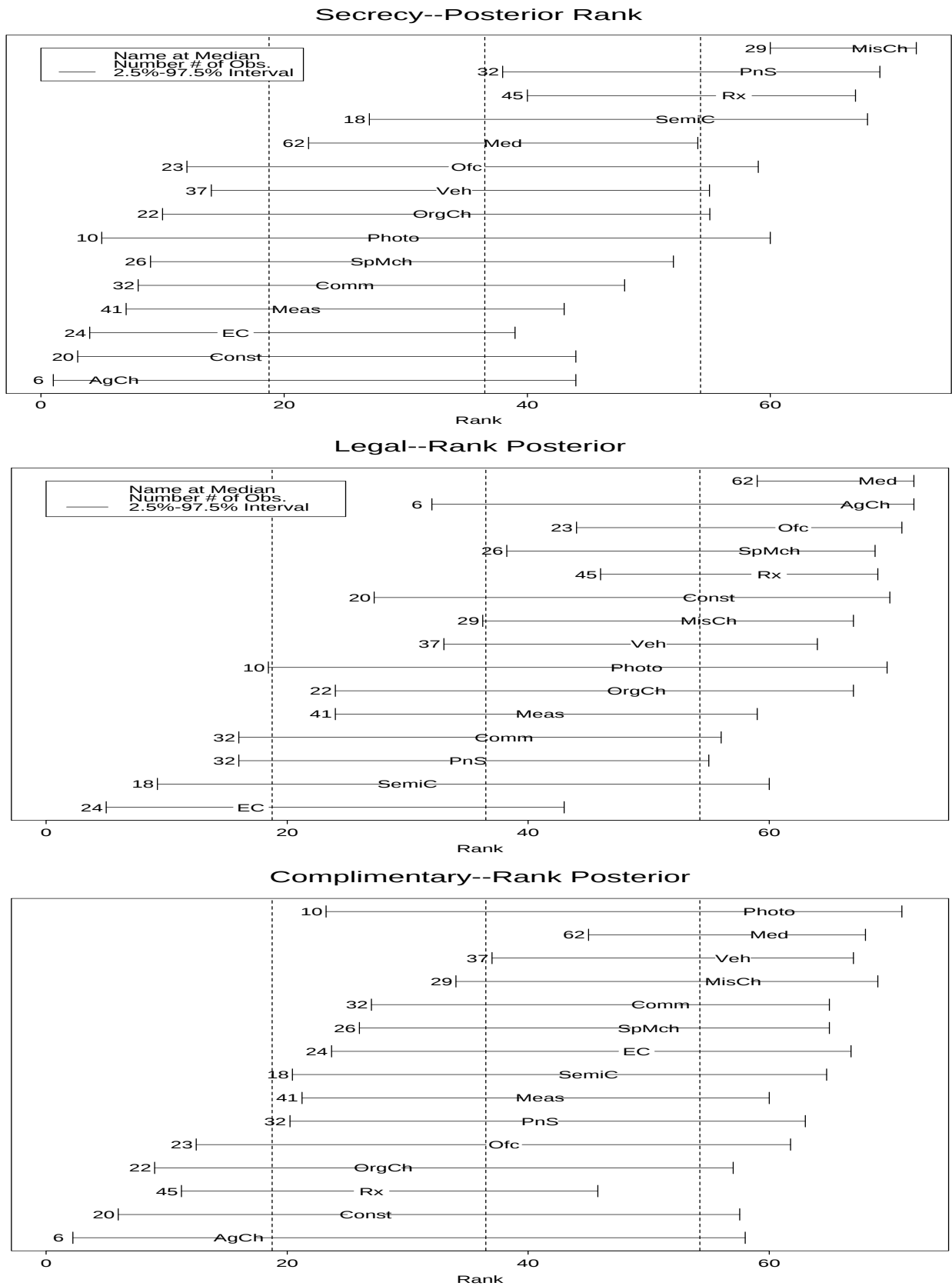


Figure 13: Posterior distributions of rank for the top 15 industries.

Abbreviation	Industry
MisCh	Miscellaneous Chemicals
PnS	Plastics and Synthetics
Rx	Pharmaceuticals
SemiC	Semi-Conductors
Med	Medical Instruments
Ofc	Office Equipment and Computers
Veh	Motor Vehicles and Supplies
OrgCh	Organic Chemicals
Photo	Photography Equipment and Supplies
SpMch	Special Machinery
Comm	Communications
Meas	Measuring Devices
EC	Electronic Components
Const	Construction Equipment
AgCh	Agricultural Chemicals

Table 10: Top 15 R&D industries.

2.5% quantile to the 97.5% quantile. The number of responding R&D units in each industry is listed to the left of that industry’s “ranking interval.”

Examining the three plots in Figure 13 we see that the data do not suggest a great deal of separation between the industries. We can however compare the industries to median rank of all industries, 36.5. Doing this we notice that for secrecy mechanisms Plastics and Synthetics (PnS), and Pharmaceutical (Rx) companies fall above the median rank, and Miscellaneous Chemicals (MisCh) 95% interval lies entirely above the 75% quantile of rank. This shows that companies in these industries rank better than at least half of the other industries in the ability to capture their profits created by innovation using the secrecy mechanisms. Similarly we find that Office Supplies and Computers (Ofc), and Special Machinery (SpMch) use legal mechanisms effectively; Motor Vehicle (Veh) companies use complementary capabilities to their advantage; and Medical Instrument (Med) companies use both legal mechanisms and complementary capabilities effectively.

5 Conclusions

We have developed a statistical model similar to the Partial Credit Model popularized in educational testing to model the responses of individual firms to questions from the Carnegie Mellon Survey of Industrial Research and Development in the Manufacturing Sector on the effectiveness of six different mechanisms used to protect the profits created by their innovations. We modeled responses to these survey questions using (a) the industry that the responding R&D unit is primarily interested in performing R&D in, (b) the effectiveness of the mechanisms being questioned,

and (c) the proportion of innovations that the responding firm has focused on product innovations. In addition to variation characterized by these covariates, we incorporated a random effect in the model to characterize the within-unit across-mechanism response dependences.

With this model, we examined differences in the ability of the six mechanisms to aid in capturing profits created by an R&D unit's innovations; differences in how the amount of R&D focused on product innovations affects responses; and differences in responses across focus industries.

Using Markov chain Monte Carlo estimates of the Bayes factors to test for these effects we found strong evidence for all effects discussed with the possible exception of differences in the effectiveness of secrecy for product versus process innovations. These estimates are however very unstable as indicated by very wide Monte Carlo 95% intervals. By examining these intervals we only find strong evidence for the existence of industry effects in secrecy mechanisms, differences in the effectiveness of the legal mechanisms, and complementary/lead time mechanisms, and for an effect of the focus of a firms R&D on the effectiveness of legal mechanisms.

We explored these effects in more detail by examining the posterior distributions of the model parameters and by examining these estimates graphically. Differences of the effectiveness of each mechanism, and the effect of innovation focus were examined using graphs similar to threshold plots used in educational testing. Industry effects were more closely examined by looking at graphical summaries of the posterior distributions of the three appropriability rankings of the 15 top R&D industries. The industries Miscellaneous Chemicals, Plastics and Synthetics, and Pharmaceutical are all ranked above the median in their ability to use secrecy effectively. Medical, Office and Computers, Special Machinery and Pharmaceutical are industries ranked above the median in their ability to use legal mechanisms. Medical, and Motor Vehicles were found rank above the median in their ability to use complimentary and lead time mechanisms.

From our analysis it appears that if governmental agencies are to aid R&D units, that they must take into account multiple factors. Our analysis has shown that there exist at least three factors correlated to an R&D units responses: the mechanism, the proportion of R&D focused on product innovations and the R&D unit's primary industry. The ability to understand how responses differ according to the industry the respondent focus innovations in, whether the question is about process or product innovations, and how the responses differ according to the focus on product v. process innovations may allow governmental agencies such as the ATP to decide which companies or industries are in need of more help.

A Fitting in CONQUEST

To fit the simplified model

$$\log \left\{ \frac{P_{j,k+1}(\theta)}{P_{j,k}(\theta)} \right\} = \theta + \mu_p - \beta_j - \gamma_{j,k}$$

in CONQUEST we have created the following command file `rnd.cmd`.

```
datafile rnd.dat;
set update=yes,warnings=no;
format indust 2-4 response 11-34 (a2);
model item+indust+item*step;
score (1,2,3,4,5) (0,1,2,3,4) ( ) ( ) !items(1,7);
score (1,2,3,4,5) ( ) (0,1,2,3,4) ( ) !items(2,3,8,9);
score (1,2,3,4,5) ( ) ( ) (0,1,2,3,4) !items(4,5,6,10,11,12);
import init_par << par.init;
import init_cov << cov.init;
export par >> par.exp;
export cov >> cov.exp;
estimate !converge=.0001;
show >> rnd.shw;
quit;
```

Responses from R&D units to the CMS are coded from one to five, however we wish to model these in CONQUEST as though they were from zero to four. Using Conquest it is possible to separate the questions into three different groups and fit responses to a three-dimensional latent vector. This is achieved by the three score statements. We then run this using the command line interface syntax `submit rnd.cmd`;

B Fitting in BUGS

Recall that we were unable to fit the entire three-dimensional model using bugs. For this reason we model each set of responses separately using bugs. Here we demonstrate how to model the four legal mechanism questions.

```
model LEGAL;

const    N=1026,           #define the constant for number of responses
         Q=4,              #define the number of questions for this group
         K=5,              #define the number of categories for the quests
         I=72;             #define the number of industries
```

```

var
    mu[I],          #means for each industry
    tau,            #precision
    gamma[Q,K-2],   #item*step parameters
    alpha[Q],       #item effect
    z[N,Q,K],       #working array
    theta[N],       #latent variables
    p[N,Q,K],       #probabilities
    beta[Q],        #slope parameter
    x[N],           #proportion of effort on products
    y[N,Q],         #responses to survey questions
    rho[N],         #industry numbers for each respondent
    d[N,Q];         #working array

data
    y in "pat.resp",      #responses in datafile pat.resp
    rho in "ind.data",    #industry numbers in ind.data
    x in "prd.data";      #proportions of patent R&D on products
                        #in prd.data

inits in "pat.in";       #initial values in pat.in

#####
#   Specifying the Likelihood   #
#####

for(i in 1:N){
    theta[i]~dnorm(mu[rho[i]],tau);
    for(j in 1:Q){
        y[i,j]~dcat(p[i,j,]);
        z[i,j,1]<-1;
        p[i,j,1]<-1/sum(z[i,j,]);
        d[i,j]<-theta[i]-alpha[j]+beta[j]*x[i];
        for(k in 2:4){
            log(z[i,j,k])<-(k-1)*d[i,j]-sum(gamma[j,1:(k-1)]);
            p[i,j,k]<-z[i,j,k]/sum(z[i,j,]);
        }
        log(z[i,j,K])<-(K-1)*d[i,j];
        p[i,j,K]<-z[i,j,K]/sum(z[i,j,]);
    }
}
}
#####
#   Specifying Priors           #

```

```
#####
```

```
for(j in 1:Q){
  for(k in 1:(K-2)){
    gamma[j,k]~dnorm(0,.5);
  }
}
for(j in 1:Q){
  alpha[j]~dnorm(0,.5);
}
for(i in 1:(I-1)){
  mu[i] ~ dnorm(0,.5);
}
mu[I]<- -sum(mu[1:(I-1)]);      #industry means constrained to zero
for(i in 1:Q){
  beta[i]~dnorm(0,.5);
}
tau ~ dgamma(1.45,.45);
}
```

The data in the file `pat.resp`, `ind.data`, and `prd.data` are modeled according to the likelihood described below the Specifying Likelihood header. All parameters are given $N(0, 2)$ priors (BUGS syntax is `dnorm(mean, precision)`), with the exception of the latent variable, $N(\mu_{\rho_i}, 1/\tau)$, and the latent precision, $\Gamma(1.45, .45)$.

References

- Cohen, W.M., (1996), "Within and Cross Industry Spillovers and Appropriability, A Proposal to the Advanced Technology Program."
- Fienberg, S.E., (1989), "Modeling Considerations: Discussion from a Modeling Perspective," *Panel Surveys*, New York: John Wiley & Sons.
- Gelman A., X.L. Meng, , H. Stern, (1996), "Posterior Predictive Assessment of Model Fitness Via Realized Discrepancies," *Statistica Sinica*, **6**, 733-807.
- Gilks, W.R., S. Richardson, and D. Spiegelhalter, (1995), *Practical Markov Chain Monte Carlo*. New York: Chapman & Hall.
- Hambleton, R.K., and R.J. Rovinelli, (1986), "Assessing the dimensionality of a set of test items," *Applied Psychological Measurement*, **10**, 287-302.
- Kass, R.E. and A.E. Raftery, (1995), "Bayes Factors," *Journal of the American Statistical Association*, **90**, 773–795.
- Levin, R.C., A.K. Klovorick, R.R. Nelson, and S.G. Winter, (1987), "Appropriating the Returns from Industrial R&D," *Brookings Papers on Economic Activity*, 783–820.
- Mansfield, E., (1986), "Patents and Innovation: An Empirical Study," *Management Science*, **32**, 173-181.
- Mardia, K.V., J.T. Kent, and J.M. Bibby, (1979), *Multivariate Analysis*, 255–280.
- Masters, G.N., (1982), "A Rasch Model for Partial Credit Scoring," *Psychometrika*, **47**, 149–174.
- Muraki, E., and J.E. Carlson, (1995), "Full-Information Factor Analysis for Polytomous Item Responses," *Applied Psychological Measurement*, **19**, 73–90.
- National Institute of Standards and Technology, (1998), "Advanced Technology Website," <http://www.atp.nist.gov>
- Spiegelhalter, D.J., A. Thomas, N.G. Best, and W.R. Gilks, (1995), *BUGS: Bayesian Inference Using Gibbs Sampling Version 0.50*, MRC Biostatistics Unit, Cambridge.
- Statistical Sciences, Inc. (1993), *S-PLUS Reference Manual, Version 3.2*, Seattle: StatSci, a division of MathSoft, Inc.
- Stiratelli, R., N. Laird, and J.H. Ware, (1984), "Random effects models for serial observations with binary responses," *Biometrics*, **40**, 961-971.
- van der Linden, W.J., and R.K. Hambleton, (1996) *Handbook of Modern Item Response Theory*, New York: Springer-Verlag.
- Wu, M.L., R.J. Adams, and M.R. Wilson, (1997), *ConQuest, Generalized Item Response Modeling Software*, ACER.