

Lecture 20something – Rank likelihood

Alexander Volfovsky

November 29, 2018

Outline

Ordered outcomes in Bayesian inference

Rank likelihood

Likelihoods for Fixed Rank Nomination Networks with Applications to Friendship Networks from Add Health

Ordered outcomes

- ▶ Frequently encountered in health and social sciences.

Ordered outcomes

- ▶ Frequently encountered in health and social sciences.
- ▶ Example: educational levels (“high school” ,
“college” , “masters” , “med school”)

Ordered outcomes

- ▶ Frequently encountered in health and social sciences.
- ▶ Example: educational levels (“high school” ,
“college” , “masters” , “med school”)
- ▶ What's special about “order” ?

Ordered outcomes

- ▶ Frequently encountered in health and social sciences.
- ▶ Example: educational levels (“high school” ,
“college” , “masters” , “med school”)
- ▶ What’s special about “order” ?
- ▶ There is no immediate numerical scale for the data.

Ordered outcomes

- ▶ Frequently encountered in health and social sciences.
- ▶ Example: educational levels (“high school”, “college”, “masters”, “med school”)
- ▶ What’s special about “order”?
- ▶ There is no immediate numerical scale for the data.
- ▶ Coding 1 = “high school,” 2 “college,” and so on...

Ordered outcomes

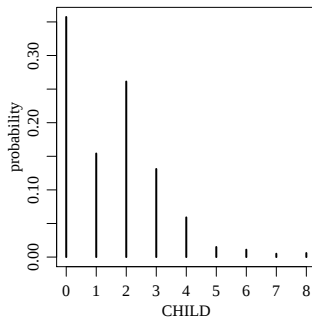
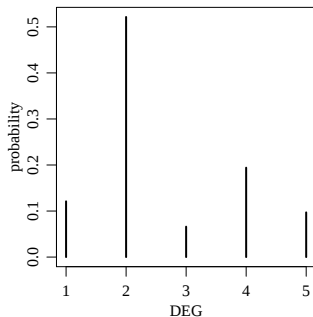
- ▶ Frequently encountered in health and social sciences.
- ▶ Example: educational levels (“high school” , “college” , “masters” , “med school”)
- ▶ What’s special about “order” ?
- ▶ There is no immediate numerical scale for the data.
- ▶ Coding 1 = “high school,” 2 “college,” and so on...
- ▶ College is not twice as much as high school.

Ordered outcomes

- ▶ Frequently encountered in health and social sciences.
- ▶ Example: educational levels (“high school” , “college” , “masters” , “med school”)
- ▶ What's special about “order” ?
- ▶ There is no immediate numerical scale for the data.
- ▶ Coding 1 = “high school,” 2 “college,” and so on...
- ▶ College is not twice as much as high school.
- ▶ We call these type of outcomes “ordinal non-numeric” .

Ordered outcomes

Data on education level and number of children from the 1994 General Social Survey.



Probit regression

Chapter 12.1.1 in Hoff

- ▶ In the previous example, it is not possible to use a numeric scale to model education level.

Probit regression

Chapter 12.1.1 in Hoff

- ▶ In the previous example, it is not possible to use a numeric scale to model education level.
- ▶ Imagine there is a numeric random variable “education” which is binned into the “education level” categories.

Probit regression

Chapter 12.1.1 in Hoff

- ▶ In the previous example, it is not possible to use a numeric scale to model education level.
- ▶ Imagine there is a numeric random variable “education” which is binned into the “education level” categories.
- ▶ If we know how to get from “education level” to “education” then we can just model that in a numeric way.

Probit regression

Chapter 12.1.1 in Hoff

- ▶ In the previous example, it is not possible to use a numeric scale to model education level.
- ▶ Imagine there is a numeric random variable “education” which is binned into the “education level” categories.
- ▶ If we know how to get from “education level” to “education” then we can just model that in a numeric way.
- ▶ If we can’t go directly we can treat the “education” variable as latent:

$$\begin{aligned}\epsilon_1, \dots, \epsilon_n &\stackrel{\text{iid}}{\sim} \text{normal}(0, 1) \\ \text{education}_i &= \beta^t x_i + \epsilon_i \\ \text{education_level}_i &= g(\text{education}_i)\end{aligned}$$

Probit regression

Chapter 12.1.1 in Hoff

$$\epsilon_1, \dots, \epsilon_n \stackrel{\text{iid}}{\sim} \text{normal}(0, 1)$$

$$\text{education}_i = \beta^t \mathbf{x}_i + \epsilon_i$$

$$\text{education_level}_i = g(\text{education}_i)$$

Peculiarities of setup:

- How do you specify g ?

Probit regression

Chapter 12.1.1 in Hoff

$$\epsilon_1, \dots, \epsilon_n \stackrel{\text{iid}}{\sim} \text{normal}(0, 1)$$

$$\text{education}_i = \beta^t \mathbf{x}_i + \epsilon_i$$

$$\text{education_level}_i = g(\text{education}_i)$$

Peculiarities of setup:

- How do you specify g ?

Any non-decreasing function that maps between the domain of the latent variable and the observed variable.

Probit regression

Chapter 12.1.1 in Hoff

$$\epsilon_1, \dots, \epsilon_n \stackrel{\text{iid}}{\sim} \text{normal}(0, 1)$$

$$\text{education}_i = \beta^t \mathbf{x}_i + \epsilon_i$$

$$\text{education_level}_i = g(\text{education}_i)$$

Peculiarities of setup:

- ▶ How do you specify g ?

Any non-decreasing function that maps between the domain of the latent variable and the observed variable.

- ▶ Variance of errors is fixed.

Probit regression

Chapter 12.1.1 in Hoff

$$\epsilon_1, \dots, \epsilon_n \stackrel{\text{iid}}{\sim} \text{normal}(0, 1)$$

$$\text{education}_i = \beta^t \mathbf{x}_i + \epsilon_i$$

$$\text{education_level}_i = g(\text{education}_i)$$

Peculiarities of setup:

- ▶ How do you specify g ?

Any non-decreasing function that maps between the domain of the latent variable and the observed variable.

- ▶ Variance of errors is fixed.

The scale can be defined by g .

Probit regression

Chapter 12.1.1 in Hoff

$$\epsilon_1, \dots, \epsilon_n \stackrel{\text{iid}}{\sim} \text{normal}(0, 1)$$

$$\text{education}_i = \beta^t \mathbf{x}_i + \epsilon_i$$

$$\text{education_level}_i = g(\text{education}_i)$$

Peculiarities of setup:

- ▶ How do you specify g ?

Any non-decreasing function that maps between the domain of the latent variable and the observed variable.

- ▶ Variance of errors is fixed.

The scale can be defined by g .

- ▶ What can be in β ?

Probit regression

Chapter 12.1.1 in Hoff

$$\epsilon_1, \dots, \epsilon_n \stackrel{\text{iid}}{\sim} \text{normal}(0, 1)$$

$$\text{education}_i = \beta^t \mathbf{x}_i + \epsilon_i$$

$$\text{education_level}_i = g(\text{education}_i)$$

Peculiarities of setup:

- ▶ How do you specify g ?

Any non-decreasing function that maps between the domain of the latent variable and the observed variable.

- ▶ Variance of errors is fixed.

The scale can be defined by g .

- ▶ What can be in β ?

The location can be defined by g .

Probit regression

Chapter 12.1.1 in Hoff

When the observed variable has K categories, g can be defined as:

$$\begin{aligned}y = g(z) &= 1 \text{ if } -\infty = g_0 < z < g_1 \\&= 2 \text{ if } g_1 < z < g_2 \\&\vdots \\&= K \text{ if } g_{K-1} < z < g_K = \infty\end{aligned}$$

The values $\{g_1, \dots, g_{K-1}\}$ are “thresholds” - when z moves past them, the category in y changes.

Probit regression

Chapter 12.1.1 in Hoff

When the observed variable has K categories, g can be defined as:

$$\begin{aligned}y = g(z) &= 1 \text{ if } -\infty = g_0 < z < g_1 \\&= 2 \text{ if } g_1 < z < g_2 \\&\vdots \\&= K \text{ if } g_{K-1} < z < g_K = \infty\end{aligned}$$

The values $\{g_1, \dots, g_{K-1}\}$ are “thresholds” - when z moves past them, the category in y changes.

- The unknown parameters in the model are the regression coefficients and the thresholds.

Probit regression

Chapter 12.1.1 in Hoff

When the observed variable has K categories, g can be defined as:

$$\begin{aligned}y = g(z) &= 1 \text{ if } -\infty = g_0 < z < g_1 \\&= 2 \text{ if } g_1 < z < g_2 \\&\vdots \\&= K \text{ if } g_{K-1} < z < g_K = \infty\end{aligned}$$

The values $\{g_1, \dots, g_{K-1}\}$ are “thresholds” - when z moves past them, the category in y changes.

- ▶ The unknown parameters in the model are the regression coefficients and the thresholds.
- ▶ Bayesian approach: normal priors.

Probit regression

Chapter 12.1.1 in Hoff

When the observed variable has K categories, g can be defined as:

$$\begin{aligned}y = g(z) &= 1 \text{ if } -\infty = g_0 < z < g_1 \\&= 2 \text{ if } g_1 < z < g_2 \\&\vdots \\&= K \text{ if } g_{K-1} < z < g_K = \infty\end{aligned}$$

The values $\{g_1, \dots, g_{K-1}\}$ are “thresholds” - when z moves past them, the category in y changes.

- ▶ The unknown parameters in the model are the regression coefficients and the thresholds.
- ▶ Bayesian approach: normal priors.
- ▶ Get joint posterior via a Gibbs sampler.

Probit regression

Full model

Model:

$$\epsilon_1, \dots, \epsilon_n \stackrel{\text{iid}}{\sim} \text{normal}(0, 1)$$

$$\text{education}_i = \beta^t \mathbf{x}_i + \epsilon_i$$

$$\text{education_level}_i = g(\text{education}_i)$$

Probit regression

Full model

Model:

$$\epsilon_1, \dots, \epsilon_n \stackrel{\text{iid}}{\sim} \text{normal}(0, 1)$$

$$\text{education}_i = \beta^t \mathbf{x}_i + \epsilon_i$$

$$\text{education_level}_i = g(\text{education}_i)$$

Priors:

$$\beta \sim \text{normal}(0, n(X^t X)^{-1})$$

$$\{g_1, \dots, g_{K-1}\} \sim p(g)(= \text{normal}(\mu, \Sigma))$$

Probit regression

Full conditionals: β

- ▶ Just like with ordinary regression: $p(\beta|y, z, g) \propto p(\beta)p(z|\beta)$.
- ▶ No dependence on g .
- ▶ Normal-normal conjugacy, so posterior is also normal:

$$\text{var}(\beta|z) = \frac{n}{n+1}(X^t X)^{-1}$$
$$\text{E}(\beta|z) = \frac{n}{n+1}(X^t X)^{-1}X^t z.$$

Probit regression

Full conditionals: Z

- ▶ The sampling distribution for the Z_i is $\text{normal}(\beta^t x_i, 1)$.
- ▶ The posterior of the Z_i is simply that constrained to be in the “correct” subinterval.
- ▶ If $Y_i = y_i$ then Z_i must be in the interval (g_{y_i-1}, g_{y_i}) .
- ▶ The posterior is a constrained normal.

Probit regression

Full conditionals: g

- ▶ Recall the prior on g is multivariate normal with mean vector μ and diagonal covariance matrix Σ .
- ▶ Conditional on all the Y and Z , the element g_k must lie between all the z_i for which $y_i = k$ and the z_i for which $y_i = k + 1$.
- ▶ With a normal prior, this is another constrained normal.

Sampling from a constrained normal

- ▶ We are interested in sampling g from a normal distribution with mean μ , variance σ constrained to the interval $(a, b) \in \bar{\mathbb{R}}$.
- ▶ Easiest way to do this by sampling a uniform random variable and doing an inverse CDF transform:

$$u \sim \text{Unif}\left(\Phi\left(\frac{a - \mu}{\sigma}\right), \Phi\left(\frac{b - \mu}{\sigma}\right)\right)$$
$$g = \mu + \sigma \Phi^{-1}(u)$$

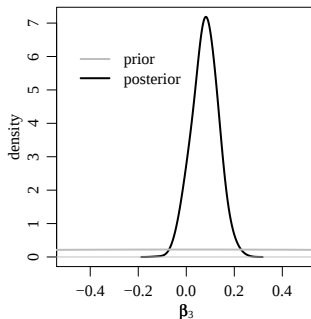
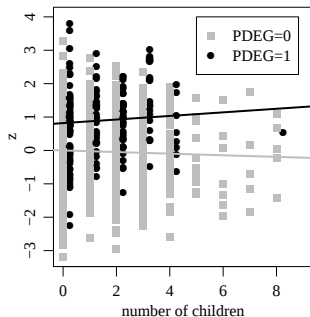
- ▶ R code is given by:

```
u <- runif(1, pnorm((a-mu)/sigma), pnorm((b-mu)/sigma))  
g = mu + sigma*qnorm(u)
```

Example

GSS data

Interested in modeling education level based on number of children, parental education level and an interaction.



Outline

Ordered outcomes in Bayesian inference

Rank likelihood

Likelihoods for Fixed Rank Nomination Networks with Applications to Friendship Networks from Add Health

Rank likelihood

- ▶ For data $z = (z_1, \dots, z_n)$ that comes from a distribution parametrized by β , the probability $p(z|\beta)$ as a function of β is known as the likelihood.

Rank likelihood

- ▶ For data $z = (z_1, \dots, z_n)$ that comes from a distribution parametrized by β , the probability $p(z|\beta)$ as a function of β is known as the likelihood.
- ▶ If we only know the ranks r of the data z then $p(r|\beta)$ is a marginal likelihood known as the rank likelihood.

Rank likelihood

- ▶ For data $z = (z_1, \dots, z_n)$ that comes from a distribution parametrized by β , the probability $p(z|\beta)$ as a function of β is known as the likelihood.
- ▶ If we only know the ranks r of the data z then $p(r|\beta)$ is a marginal likelihood known as the rank likelihood.
- ▶ It is marginal because it is given by the integral $\int p(z|\beta) dz$ where the integration is done over the region $\{z_{r_1} < z_{r_2} < \dots < z_{r_n}\}$.

Rank likelihood

- ▶ For data $z = (z_1, \dots, z_n)$ that comes from a distribution parametrized by β , the probability $p(z|\beta)$ as a function of β is known as the likelihood.
- ▶ If we only know the ranks r of the data z then $p(r|\beta)$ is a marginal likelihood known as the rank likelihood.
- ▶ It is marginal because it is given by the integral $\int p(z|\beta) dz$ where the integration is done over the region $\{z_{r_1} < z_{r_2} < \dots < z_{r_n}\}$.
- ▶ Detailed outline given by Pettitt (1982).

Rank likelihood

- ▶ For data $z = (z_1, \dots, z_n)$ that comes from a distribution parametrized by β , the probability $p(z|\beta)$ as a function of β is known as the likelihood.
- ▶ If we only know the ranks r of the data z then $p(r|\beta)$ is a marginal likelihood known as the rank likelihood.
- ▶ It is marginal because it is given by the integral $\int p(z|\beta) dz$ where the integration is done over the region $\{z_{r_1} < z_{r_2} < \dots < z_{r_n}\}$.
- ▶ Detailed outline given by Pettitt (1982).
- ▶ For continuous z , the information in the rank likelihood is the same as in the ranks of the data (hence the name).

Rank likelihood

- ▶ For data $z = (z_1, \dots, z_n)$ that comes from a distribution parametrized by β , the probability $p(z|\beta)$ as a function of β is known as the likelihood.
- ▶ If we only know the ranks r of the data z then $p(r|\beta)$ is a marginal likelihood known as the rank likelihood.
- ▶ It is marginal because it is given by the integral $\int p(z|\beta) dz$ where the integration is done over the region $\{z_{r_1} < z_{r_2} < \dots < z_{r_n}\}$.
- ▶ Detailed outline given by Pettitt (1982).
- ▶ For continuous z , the information in the rank likelihood is the same as in the ranks of the data (hence the name).
- ▶ For discrete data, there is less information due to the possibility of ties.

Rank likelihood (Pettitt, 1982)

- ▶ Let $z_1, \dots, z_n$ be independent where $z_i \sim \text{normal}(x_i\beta, 1)$.

Rank likelihood (Pettitt, 1982)

- ▶ Let $z_1, \dots, z_n$ be independent where $z_i \sim \text{normal}(x_i\beta, 1)$.
- ▶ We are interested in making inference about β when we only observe the ranks $r_1, \dots, r_n$ of the data.

Rank likelihood (Pettitt, 1982)

- ▶ Let $z_1, \dots, z_n$ be independent where $z_i \sim \text{normal}(x_i\beta, 1)$.
- ▶ We are interested in making inference about β when we only observe the ranks $r_1, \dots, r_n$ of the data.
- ▶ The marginal likelihood of the ranks is given by $f(r|\beta) = \Pr(z_{\alpha(1)} < \dots < z_{\alpha(n)}|\beta)$ where $\alpha(i) = j$ if z_j is the i th smallest.

Rank likelihood (Pettitt, 1982)

- ▶ Let $z_1, \dots, z_n$ be independent where $z_i \sim \text{normal}(x_i\beta, 1)$.
- ▶ We are interested in making inference about β when we only observe the ranks $r_1, \dots, r_n$ of the data.
- ▶ The marginal likelihood of the ranks is given by $f(r|\beta) = \Pr(z_{\alpha(1)} < \dots < z_{\alpha(n)}|\beta)$ where $\alpha(i) = j$ if z_j is the i th smallest.
- ▶ Pettitt constructs approximations for the integral over the order of the z s

$$f(r|\beta) = \text{const} \int \exp(-(z - X\beta)^t(z - X\beta)) dz.$$

Rank likelihood (Pettitt, 1982)

- ▶ Let $z_1, \dots, z_n$ be independent where $z_i \sim \text{normal}(x_i\beta, 1)$.
- ▶ We are interested in making inference about β when we only observe the ranks $r_1, \dots, r_n$ of the data.
- ▶ The marginal likelihood of the ranks is given by $f(r|\beta) = \Pr(z_{\alpha(1)} < \dots < z_{\alpha(n)}|\beta)$ where $\alpha(i) = j$ if z_j is the i th smallest.
- ▶ Pettitt constructs approximations for the integral over the order of the z s

$$f(r|\beta) = \text{const} \int \exp(-(z - X\beta)^t(z - X\beta)) dz.$$

- ▶ Inference via a prior on β and using the exact or approximate marginal likelihoods. (Monahan and Boos, 1992)

Rank likelihood (Pettitt, 1982)

- ▶ Let $z_1, \dots, z_n$ be independent where $z_i \sim \text{normal}(x_i\beta, 1)$.
- ▶ We are interested in making inference about β when we only observe the ranks $r_1, \dots, r_n$ of the data.
- ▶ The marginal likelihood of the ranks is given by $f(r|\beta) = \Pr(z_{\alpha(1)} < \dots < z_{\alpha(n)}|\beta)$ where $\alpha(i) = j$ if z_j is the i th smallest.
- ▶ Pettitt constructs approximations for the integral over the order of the z s

$$f(r|\beta) = \text{const} \int \exp(-(z - X\beta)^t(z - X\beta)) dz.$$

- ▶ Inference via a prior on β and using the exact or approximate marginal likelihoods. (Monahan and Boos, 1992)
- ▶ Theoretical guarantees (Bickel and Ritov, 1997, Hoff, Niu and Wellner, 2014).

Rank likelihood

Details

- ▶ If we observe $y_1 > y_2$ then we know that $g(Z_1) > g(Z_2)$ and so $Z_1 > Z_2$.

Rank likelihood

Details

- ▶ If we observe $y_1 > y_2$ then we know that $g(Z_1) > g(Z_2)$ and so $Z_1 > Z_2$.
- ▶ Observing y tells us that the Z_i must lie in

$$R(y) = \{z \in \mathbb{R}^n : z_{i_1} < z_{i_2} \text{ if } y_{i_1} < y_{i_2}\}.$$

Rank likelihood

Details

- ▶ If we observe $y_1 > y_2$ then we know that $g(Z_1) > g(Z_2)$ and so $Z_1 > Z_2$.
- ▶ Observing y tells us that the Z_i must lie in

$$R(y) = \{z \in \mathbb{R}^n : z_{i_1} < z_{i_2} \text{ if } y_{i_1} < y_{i_2}\}.$$

- ▶ Posterior inference for β does not change.

Rank likelihood

Details

- ▶ If we observe $y_1 > y_2$ then we know that $g(Z_1) > g(Z_2)$ and so $Z_1 > Z_2$.
- ▶ Observing y tells us that the Z_i must lie in

$$R(y) = \{z \in \mathbb{R}^n : z_{i_1} < z_{i_2} \text{ if } y_{i_1} < y_{i_2}\}.$$

- ▶ Posterior inference for β does not change.
- ▶ Posterior inference for Z_i does change since we must update each of the Z_i with respect to the other ones (there are no more cutoffs).

Rank likelihood

Details

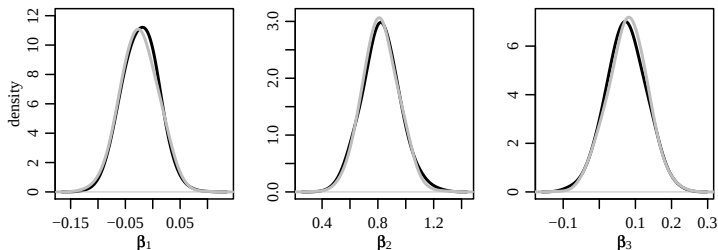
- ▶ If we observe $y_1 > y_2$ then we know that $g(Z_1) > g(Z_2)$ and so $Z_1 > Z_2$.
- ▶ Observing y tells us that the Z_i must lie in

$$R(y) = \{z \in \mathbb{R}^n : z_{i_1} < z_{i_2} \text{ if } y_{i_1} < y_{i_2}\}.$$

- ▶ Posterior inference for β does not change.
- ▶ Posterior inference for Z_i does change since we must update each of the Z_i with respect to the other ones (there are no more cutoffs).
- ▶ Posterior of Z_i conditional on the ordering and the value of the other Z_j and on β is a constrained normal with boundary given by $\max\{z_j : y_j < y_i\}$ and $\min\{z_j : y_i < y_j\}$.

GSS Example

Interested in modeling education level based on number of children, parental education level and an interaction.



Rank likelihood

Applications

- ▶ Semiparametric regression.
- ▶ Ordinal regression.
- ▶ Semiparametric copula estimation (Section 12.2).
- ▶ Network analysis.

Outline

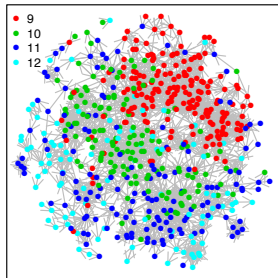
Ordered outcomes in Bayesian inference

Rank likelihood

Likelihoods for Fixed Rank Nomination Networks with Applications to Friendship Networks from Add Health

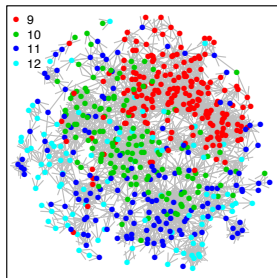
Social network data

- Datasets: PROSPER, NSCR, AddHealth



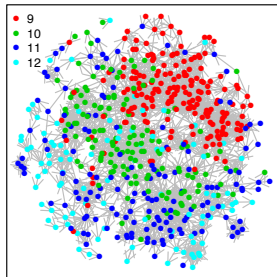
Social network data

- Datasets: PROSPER, NSCR, AddHealth
- Relate network characteristics to individual-level behavior



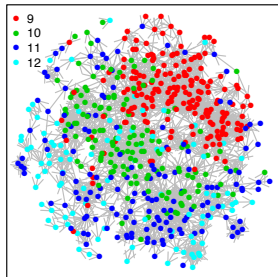
Social network data

- ▶ Datasets: PROSPER, NSCR, AddHealth
- ▶ Relate network characteristics to individual-level behavior
- ▶ Literature: ERGM, latent variable models



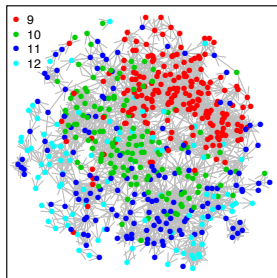
Social network data

- ▶ Datasets: PROSPER, NSCR, AddHealth
- ▶ Relate network characteristics to individual-level behavior
- ▶ Literature: ERGM, latent variable models
- ▶ Assumptions:
 - ▶ Data is fully observed
 - ▶ The support is the set of all sociomatrices



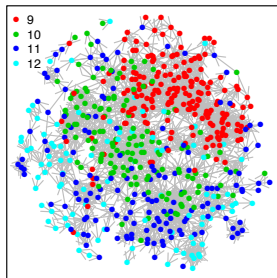
Social network data

- ▶ Datasets: PROSPER, NSCR, AddHealth
- ▶ Relate network characteristics to individual-level behavior
- ▶ Literature: ERGM, latent variable models
- ▶ Assumptions:
 - ▶ Data is fully observed
 - ▶ The support is the set of all sociomatrices
- ▶ In practice:
 - ▶ Ranked data
 - ▶ Censored observations



Social network data

- ▶ Datasets: PROSPER, NSCR, AddHealth
- ▶ Relate network characteristics to individual-level behavior
- ▶ Literature: ERGM, latent variable models
- ▶ Assumptions:
 - ▶ Data is fully observed
 - ▶ The support is the set of all sociomatrices
- ▶ In practice:
 - ▶ Ranked data
 - ▶ Censored observations



A type of likelihood that accommodates the ranked and censored nature of data from **Fixed Rank Nomination (FRN)** surveys and allows for estimation of regression effects.

FRN Outline

- ▶ Fixed rank nominations likelihood
- ▶ Rank based likelihood
- ▶ Binary likelihood
- ▶ Simulations
- ▶ AddHealth example

Notation

- $Y = \{y_{ij} : i \neq j\}$ is a sociomatrix of ordinal relationships

$$Y = \begin{pmatrix} - & y_{12} & \cdots & y_{1n} \\ y_{21} & - & & \\ \vdots & & - & \\ y_{n1} & & & - \end{pmatrix}$$

Notation

- ▶ $Y = \{y_{ij} : i \neq j\}$ is a sociomatrix of ordinal relationships
 $y_{ij} > y_{ik}$ denotes person i preferring person j to person k

$$Y = \begin{pmatrix} - & y_{12} & \cdots & y_{1n} \\ y_{21} & - & & \\ \vdots & & - & \\ y_{n1} & & & - \end{pmatrix}$$

Notation

- ▶ $Y = \{y_{ij} : i \neq j\}$ is a sociomatrix of ordinal relationships

$y_{ij} > y_{ik}$ denotes person i preferring person j to person k

$$Y = \begin{pmatrix} - & y_{12} & \cdots & y_{1n} \\ y_{21} & - & & \\ \vdots & & - & \\ y_{n1} & & & - \end{pmatrix}$$

- ▶ In FRN we observe a sociomatrix $S = \{s_{ij} : i \neq j\}$

Notation

- ▶ $Y = \{y_{ij} : i \neq j\}$ is a sociomatrix of ordinal relationships

$y_{ij} > y_{ik}$ denotes person i preferring person j to person k

$$Y = \begin{pmatrix} - & y_{12} & \cdots & y_{1n} \\ y_{21} & - & & \\ \vdots & & - & \\ y_{n1} & & & - \end{pmatrix}$$

- ▶ In FRN we observe a sociomatrix $S = \{s_{ij} : i \neq j\}$

$s_{ij} = 0$ if j is not nominated by i

$s_{ij} > s_{ik}$ if i scores j more highly than k

Notation

- ▶ $Y = \{y_{ij} : i \neq j\}$ is a sociomatrix of ordinal relationships

$y_{ij} > y_{ik}$ denotes person i preferring person j to person k

$$Y = \begin{pmatrix} - & y_{12} & \cdots & y_{1n} \\ y_{21} & - & & \\ \vdots & & - & \\ y_{n1} & & & - \end{pmatrix}$$

- ▶ In FRN we observe a sociomatrix $S = \{s_{ij} : i \neq j\}$

$s_{ij} = 0$ if j is not nominated by i

$s_{ij} > s_{ik}$ if i scores j more highly than k

Observed outdegree $d_i = \sum_{j \neq i} 1(s_{ij} > 0)$ satisfied $d_i \leq m$

Notation

- ▶ $Y = \{y_{ij} : i \neq j\}$ is a sociomatrix of ordinal relationships

$y_{ij} > y_{ik}$ denotes person i preferring person j to person k

$$Y = \begin{pmatrix} - & y_{12} & \cdots & y_{1n} \\ y_{21} & - & & \\ \vdots & & - & \\ y_{n1} & & & - \end{pmatrix}$$

- ▶ In FRN we observe a sociomatrix $S = \{s_{ij} : i \neq j\}$

$s_{ij} = 0$ if j is not nominated by i

$s_{ij} > s_{ik}$ if i scores j more highly than k

Observed outdegree $d_i = \sum_{j \neq i} 1(s_{ij} > 0)$ satisfied $d_i \leq m$

- ▶ For each likelihood, define the set relations between s_{ij} and y_{ij}

Notation

- ▶ $Y = \{y_{ij} : i \neq j\}$ is a sociomatrix of ordinal relationships

$y_{ij} > y_{ik}$ denotes person i preferring person j to person k

$$Y = \begin{pmatrix} - & y_{12} & \cdots & y_{1n} \\ y_{21} & - & & \\ \vdots & & - & \\ y_{n1} & & & - \end{pmatrix}$$

- ▶ In FRN we observe a sociomatrix $S = \{s_{ij} : i \neq j\}$

$s_{ij} = 0$ if j is not nominated by i

$s_{ij} > s_{ik}$ if i scores j more highly than k

Observed outdegree $d_i = \sum_{j \neq i} 1(s_{ij} > 0)$ satisfied $d_i \leq m$

- ▶ For each likelihood, define the set relations between s_{ij} and y_{ij}
- ▶ Statistical model $\{p(Y|\theta) : \theta \in \Theta\}$ assists in analysis

Model - standard approach (SRM)

- ▶ The social relations model (SRM) was introduced by Warner, Kenny and Stoto (1979).
- ▶ Probit model with row and column effects:

$$\begin{aligned}Y_{ij} &= \mathbf{1}_{Z_{ij} > 0} \\Z_{ij} &= \beta^t \mathbf{X}_{ij} + a_i + b_j + \epsilon_{ij} \\(a_i \ b_i) &\stackrel{\text{iid}}{\sim} \text{normal}(0, \Sigma_{ab}) \\ \text{cor}(\epsilon_{ij}, \epsilon_{ji}) &= \rho\end{aligned}$$

- ▶ Note that $\text{cov}(Z_{ij}, Z_{kl}) = 0$ unless Z_{ij} and Z_{kl} are in the same column, same row, or are reciprocal.

Model - extended approach

Multiplicative effects

$$\begin{aligned}Y_{ij} &= \mathbf{1}_{Z_{ij} > 0} \\Z_{ij} &= \beta^t X_{ij} + a_i + b_j + \mathbf{u}_i^t \mathbf{v}_j \epsilon_{ij} \\(a_i, b_i) &\stackrel{\text{iid}}{\sim} \text{normal}(0, \Sigma_{ab}) \\(u_i, v_i) &\sim \text{normal}(0, \Sigma_{uv}) \\\text{cor}(\epsilon_{ij}, \epsilon_{ji}) &= \rho\end{aligned}$$

Set relations

$$\begin{aligned}s_{ij} > s_{ik} &\Rightarrow y_{ij} > y_{ik} \\ s_{ij} = 0 \text{ and } d_i < m &\Rightarrow y_{ij} \leq 0 \\ s_{ij} > 0 &\Rightarrow y_{ij} > 0 \\ s_{ij} = 0 &\Rightarrow y_{ij} < 0\end{aligned}$$

Set relations

$$\begin{aligned}s_{ij} > s_{ik} &\Rightarrow y_{ij} > y_{ik} \\ s_{ij} = 0 \text{ and } d_i < m &\Rightarrow y_{ij} \leq 0 \\ s_{ij} > 0 &\Rightarrow y_{ij} > 0 \\ s_{ij} = 0 &\Rightarrow y_{ij} < 0\end{aligned}$$

s_i	—	—	—	—	—	—	—	—	—	—
y_i	—	—	—	—	—	—	—	—	—	—
	1	2	3	4	5	6	7	8	9	10

Set relations

$$\begin{aligned}s_{ij} > s_{ik} &\Rightarrow y_{ij} > y_{ik} \\ s_{ij} = 0 \text{ and } d_i < m &\Rightarrow y_{ij} \leq 0 \\ s_{ij} > 0 &\Rightarrow y_{ij} > 0 \\ s_{ij} = 0 &\Rightarrow y_{ij} < 0\end{aligned}$$

s_i	<u>4</u>	<u>3</u>	<u>2</u>	<u>1</u>	<u>0</u>	<u>0</u>	<u>0</u>	<u>0</u>	<u>0</u>	<u>0</u>
y_i	<u>1</u>	<u>2</u>	<u>3</u>	<u>4</u>	<u>5</u>	<u>6</u>	<u>7</u>	<u>8</u>	<u>9</u>	<u>10</u>

Set relations

$$\begin{aligned}
 s_{ij} > s_{ik} &\Rightarrow y_{ij} > y_{ik} \\
 s_{ij} = 0 \text{ and } d_i < m &\Rightarrow y_{ij} \leq 0 \\
 s_{ij} > 0 &\Rightarrow y_{ij} > 0 \\
 s_{ij} = 0 &\Rightarrow y_{ij} < 0
 \end{aligned}$$

s_i	<u>4</u>	<u>3</u>	<u>2</u>	<u>1</u>	<u>0</u>	<u>0</u>	<u>0</u>	<u>0</u>	<u>0</u>	<u>0</u>
y_i	<u>y_{i1}</u>	<u>y_{i2}</u>	<u>y_{i3}</u>	<u>y_{i4}</u>	<u>$0>$</u>	<u>$0>$</u>	<u>$0>$</u>	<u>$0>$</u>	<u>$0>$</u>	<u>$0>$</u>
	1	2	3	4	5	6	7	8	9	10

Set relations

$$\begin{aligned}s_{ij} > s_{ik} &\Rightarrow y_{ij} > y_{ik} \\ s_{ij} = 0 \text{ and } d_i < m &\Rightarrow y_{ij} \leq 0 \\ s_{ij} > 0 &\Rightarrow y_{ij} > 0 \\ s_{ij} = 0 &\Rightarrow y_{ij} < 0\end{aligned}$$

s_i	—	—	—	—	—	—	—	—	—	—
y_i	—	—	—	—	—	—	—	—	—	—
	1	2	3	4	5	6	7	8	9	10

Set relations

$$\begin{aligned}s_{ij} > s_{ik} &\Rightarrow y_{ij} > y_{ik} \\ s_{ij} = 0 \text{ and } d_i < m &\Rightarrow y_{ij} \leq 0 \\ s_{ij} > 0 &\Rightarrow y_{ij} > 0 \\ s_{ij} = 0 &\Rightarrow y_{ij} < 0\end{aligned}$$

s_i	<u>5</u>	<u>4</u>	<u>3</u>	<u>2</u>	<u>1</u>	<u>0</u>	<u>0</u>	<u>0</u>	<u>0</u>	<u>0</u>
y_i	<u>1</u>	<u>2</u>	<u>3</u>	<u>4</u>	<u>5</u>	<u>6</u>	<u>7</u>	<u>8</u>	<u>9</u>	<u>10</u>

Set relations

$$\begin{aligned}s_{ij} > s_{ik} &\Rightarrow y_{ij} > y_{ik} \\ s_{ij} = 0 \text{ and } d_i < m &\Rightarrow y_{ij} \leq 0 \\ s_{ij} > 0 &\Rightarrow y_{ij} > 0 \\ s_{ij} = 0 &\Rightarrow y_{ij} < 0\end{aligned}$$

s_i	<u>5</u>	<u>4</u>	<u>3</u>	<u>2</u>	<u>1</u>	<u>0</u>	<u>0</u>	<u>0</u>	<u>0</u>	<u>0</u>
y_i	<u>y_{i1}</u>	<u>y_{i2}</u>	<u>y_{i3}</u>	<u>y_{i4}</u>	<u>y_{i5}</u>	<u>?</u>	<u>?</u>	<u>?</u>	<u>?</u>	<u>?</u>
	1	2	3	4	5	6	7	8	9	10

Set relations

$$\begin{aligned}s_{ij} > s_{ik} &\Rightarrow y_{ij} > y_{ik} \\ s_{ij} = 0 \text{ and } d_i < m &\Rightarrow y_{ij} \leq 0 \\ s_{ij} > 0 &\Rightarrow y_{ij} > 0 \\ s_{ij} = 0 &\Rightarrow y_{ij} < 0\end{aligned}$$

Set relations

$$\begin{aligned}s_{ij} > s_{ik} &\Rightarrow y_{ij} > y_{ik} \\ s_{ij} = 0 \text{ and } d_i < m &\Rightarrow y_{ij} \leq 0 \\ s_{ij} > 0 &\Rightarrow y_{ij} > 0 \\ s_{ij} = 0 &\Rightarrow y_{ij} < 0\end{aligned}$$

- Define sets $F(S)$, $R(S)$ and $B(S)$

Set relations

$$\begin{aligned}s_{ij} > s_{ik} &\Rightarrow y_{ij} > y_{ik} \\ s_{ij} = 0 \text{ and } d_i < m &\Rightarrow y_{ij} \leq 0 \\ s_{ij} > 0 &\Rightarrow y_{ij} > 0 \\ s_{ij} = 0 &\Rightarrow y_{ij} < 0\end{aligned}$$

- ▶ Define sets $F(S)$, $R(S)$ and $B(S)$
- ▶ Define a model $\{p(Y|\theta) : \theta \in \Theta\}$

Set relations

$$\begin{aligned}s_{ij} > s_{ik} &\Rightarrow y_{ij} > y_{ik} \\ s_{ij} = 0 \text{ and } d_i < m &\Rightarrow y_{ij} \leq 0 \\ s_{ij} > 0 &\Rightarrow y_{ij} > 0 \\ s_{ij} = 0 &\Rightarrow y_{ij} < 0\end{aligned}$$

- ▶ Define sets $F(S)$, $R(S)$ and $B(S)$
- ▶ Define a model $\{p(Y|\theta) : \theta \in \Theta\}$
- ▶ Base inference on θ on a likelihood $\int dP(Y|\theta)$

$$\left. \begin{array}{l}
 s_{ij} > s_{ik} \Rightarrow y_{ij} > y_{ik} \\
 s_{ij} = 0 \text{ and } d_i < n \Rightarrow y_{ij} \leq 0 \\
 s_{ij} > 0 \Rightarrow y_{ij} > 0 \\
 s_{ij} = 0 \Rightarrow y_{ij} < 0
 \end{array} \right\} F(S)$$

$$\left. \begin{array}{ll}
 s_{ij} > s_{ik} & \Rightarrow y_{ij} > y_{ik} \\
 s_{ij} = 0 \text{ and } d_i < n & \Rightarrow y_{ij} \leq 0 \\
 s_{ij} > 0 & \Rightarrow y_{ij} > 0 \\
 s_{ij} = 0 & \Rightarrow y_{ij} < 0
 \end{array} \right\} F(S)$$

- Captures censoring in the data

$$\left. \begin{array}{l} s_{ij} > s_{ik} \Rightarrow y_{ij} > y_{ik} \\ s_{ij} = 0 \text{ and } d_i < n \Rightarrow y_{ij} \leq 0 \\ s_{ij} > 0 \Rightarrow y_{ij} > 0 \\ s_{ij} = 0 \Rightarrow y_{ij} < 0 \end{array} \right\} F(S)$$

- ▶ Captures censoring in the data
- ▶ Differentiates between different ranks

Rank

$$\begin{aligned} & s_{ij} > s_{ik} \Rightarrow y_{ij} > y_{ik} \} R(S) \\ s_{ij} = 0 \text{ and } d_j < n & \Rightarrow y_{ij} \leq 0 \\ s_{ij} > 0 & \Rightarrow y_{ij} > 0 \\ s_{ij} = 0 & \Rightarrow y_{ij} < 0 \end{aligned}$$

Rank

$$\begin{aligned} & s_{ij} > s_{ik} \Rightarrow y_{ij} > y_{ik} \} R(S) \\ & s_{ij} = 0 \text{ and } d_j < n \Rightarrow y_{ij} \leq 0 \\ & s_{ij} > 0 \Rightarrow y_{ij} > 0 \\ & s_{ij} = 0 \Rightarrow y_{ij} < 0 \end{aligned}$$

- Valid but not fully informative: $F(S) \subsetneq R(S)$

Rank

$$\begin{aligned} s_{ij} > s_{ik} &\Rightarrow y_{ij} > y_{ik} \} R(S) \\ s_{ij} = 0 \text{ and } d_j < n &\Rightarrow y_{ij} \leq 0 \\ s_{ij} > 0 &\Rightarrow y_{ij} > 0 \\ s_{ij} = 0 &\Rightarrow y_{ij} < 0 \end{aligned}$$

- ▶ Valid but not fully informative: $F(S) \subsetneq R(S)$
- ▶ Variants of this likelihood are used for semiparametric regression modeling and copula estimation

Rank

$$\begin{aligned} s_{ij} > s_{ik} &\Rightarrow y_{ij} > y_{ik} \} R(S) \\ s_{ij} = 0 \text{ and } d_j < n &\Rightarrow y_{ij} \leq 0 \\ s_{ij} > 0 &\Rightarrow y_{ij} > 0 \\ s_{ij} = 0 &\Rightarrow y_{ij} < 0 \end{aligned}$$

- ▶ Valid but not fully informative: $F(S) \subsetneq R(S)$
- ▶ Variants of this likelihood are used for semiparametric regression modeling and copula estimation
- ▶ Cannot estimate “sender” specific effects

Binary

$$\begin{array}{lcl} s_{ij} > s_{ik} & \Rightarrow & y_{ij} > y_{ik} \\ s_{ij} = 0 \text{ and } d_i < n & \Rightarrow & y_{ij} \leq 0 \\ \left. \begin{array}{l} s_{ij} > 0 \Rightarrow y_{ij} > 0 \\ s_{ij} = 0 \Rightarrow y_{ij} < 0 \end{array} \right\} & & B(S) \end{array}$$

Binary

$$\begin{array}{lcl} s_{ij} > s_{ik} & \Rightarrow & y_{ij} > y_{ik} \\ s_{ij} = 0 \text{ and } d_i < n & \Rightarrow & y_{ij} \leq 0 \\ \left. \begin{array}{l} s_{ij} > 0 \Rightarrow y_{ij} > 0 \\ s_{ij} = 0 \Rightarrow y_{ij} < 0 \end{array} \right\} & & B(S) \end{array}$$

- ▶ Neither fully informative nor valid!

Binary

$$\begin{array}{lcl} s_{ij} > s_{ik} & \Rightarrow & y_{ij} > y_{ik} \\ s_{ij} = 0 \text{ and } d_j < n & \Rightarrow & y_{ij} \leq 0 \\ \left. \begin{array}{l} s_{ij} > 0 \Rightarrow y_{ij} > 0 \\ s_{ij} = 0 \Rightarrow y_{ij} < 0 \end{array} \right\} & & B(S) \end{array}$$

- ▶ Neither fully informative nor valid!
- ▶ Discards information on the ranks

Binary

$$\begin{array}{lcl} s_{ij} > s_{ik} & \Rightarrow & y_{ij} > y_{ik} \\ s_{ij} = 0 \text{ and } d_i < n & \Rightarrow & y_{ij} \leq 0 \\ \left. \begin{array}{l} s_{ij} > 0 \Rightarrow y_{ij} > 0 \\ s_{ij} = 0 \Rightarrow y_{ij} < 0 \end{array} \right\} & & B(S) \end{array}$$

- ▶ Neither fully informative nor valid!
- ▶ Discards information on the ranks
- ▶ Ignores the censoring on the outdegrees

Binary

$$\begin{array}{lcl} s_{ij} > s_{ik} & \Rightarrow & y_{ij} > y_{ik} \\ s_{ij} = 0 \text{ and } d_i < n & \Rightarrow & y_{ij} \leq 0 \\ \left. \begin{array}{l} s_{ij} > 0 \Rightarrow y_{ij} > 0 \\ s_{ij} = 0 \Rightarrow y_{ij} < 0 \end{array} \right\} & & B(S) \end{array}$$

- ▶ Neither fully informative nor valid!
- ▶ Discards information on the ranks
- ▶ Ignores the censoring on the outdegrees
- ▶ In particular: $F(S) \not\subset B(S)$

Bayesian estimation

- The FRN, rank and binary likelihoods can be expressed as:

$$L_F(\theta : S) = \Pr(Y \in F(S) | \theta) = \int_{F(S)} dP(Y | \theta)$$

$$L_R(\theta : S) = \Pr(Y \in R(S) | \theta) = \int_{R(S)} dP(Y | \theta)$$

$$L_B(\theta : S) = \Pr(Y \in B(S) | \theta) = \int_{B(S)} dP(Y | \theta)$$

Bayesian estimation

- The FRN, rank and binary likelihoods can be expressed as:

$$L_F(\theta : S) = \Pr(Y \in \textcolor{red}{F}(S) | \theta) = \int_{\textcolor{red}{F}(S)} dP(Y | \theta)$$

$$L_R(\theta : S) = \Pr(Y \in \textcolor{blue}{R}(S) | \theta) = \int_{\textcolor{blue}{R}(S)} dP(Y | \theta)$$

$$L_B(\theta : S) = \Pr(Y \in \textcolor{green}{B}(S) | \theta) = \int_{\textcolor{green}{B}(S)} dP(Y | \theta)$$

- Generally intractable

Bayesian estimation

- ▶ The FRN, rank and binary likelihoods can be expressed as:

$$L_F(\theta : S) = \Pr(Y \in F(S) | \theta) = \int_{F(S)} dP(Y | \theta)$$

$$L_R(\theta : S) = \Pr(Y \in R(S) | \theta) = \int_{R(S)} dP(Y | \theta)$$

$$L_B(\theta : S) = \Pr(Y \in B(S) | \theta) = \int_{B(S)} dP(Y | \theta)$$

- ▶ Generally intractable
- ▶ Inference for θ can proceed using a MCMC approximation:

Bayesian estimation

- ▶ The FRN, rank and binary likelihoods can be expressed as:

$$L_F(\theta : S) = \Pr(Y \in \textcolor{red}{F}(S) | \theta) = \int_{\textcolor{red}{F}(S)} dP(Y | \theta)$$

$$L_R(\theta : S) = \Pr(Y \in \textcolor{blue}{R}(S) | \theta) = \int_{\textcolor{blue}{R}(S)} dP(Y | \theta)$$

$$L_B(\theta : S) = \Pr(Y \in \textcolor{green}{B}(S) | \theta) = \int_{\textcolor{green}{B}(S)} dP(Y | \theta)$$

- ▶ Generally intractable
- ▶ Inference for θ can proceed using a MCMC approximation:
 - ▶ Given the observed ranks S and a prior distribution $p(\theta)$

Bayesian estimation

- ▶ The FRN, rank and binary likelihoods can be expressed as:

$$L_F(\theta : S) = \Pr(Y \in F(S) | \theta) = \int_{F(S)} dP(Y | \theta)$$

$$L_R(\theta : S) = \Pr(Y \in R(S) | \theta) = \int_{R(S)} dP(Y | \theta)$$

$$L_B(\theta : S) = \Pr(Y \in B(S) | \theta) = \int_{B(S)} dP(Y | \theta)$$

- ▶ Generally intractable
- ▶ Inference for θ can proceed using a MCMC approximation:
 - ▶ Given the observed ranks S and a prior distribution $p(\theta)$
 - ▶ Generate a Markov chain whose stationary distribution is that of (θ, Y) given $Y \in F(S)$, $R(S)$ or $B(S)$

Bayesian estimation

- ▶ The FRN, rank and binary likelihoods can be expressed as:

$$L_F(\theta : S) = \Pr(Y \in F(S) | \theta) = \int_{F(S)} dP(Y | \theta)$$

$$L_R(\theta : S) = \Pr(Y \in R(S) | \theta) = \int_{R(S)} dP(Y | \theta)$$

$$L_B(\theta : S) = \Pr(Y \in B(S) | \theta) = \int_{B(S)} dP(Y | \theta)$$

- ▶ Generally intractable
- ▶ Inference for θ can proceed using a MCMC approximation:
 - ▶ Given the observed ranks S and a prior distribution $p(\theta)$
 - ▶ Generate a Markov chain whose stationary distribution is that of (θ, Y) given $Y \in F(S)$, $R(S)$ or $B(S)$
 - ▶ Gibbs: simulate $\theta \sim p(\theta | Y)$ and then given $Y \in F(S)$, $R(S)$ or $B(S)$, for $i \neq j$ simulate $y_{ij} \sim p(y_{ij} | \theta, Y_{-ij})$

Bayesian estimation

- ▶ The FRN, rank and binary likelihoods can be expressed as:

$$L_F(\theta : S) = \Pr(Y \in F(S) | \theta) = \int_{F(S)} dP(Y | \theta)$$

$$L_R(\theta : S) = \Pr(Y \in R(S) | \theta) = \int_{R(S)} dP(Y | \theta)$$

$$L_B(\theta : S) = \Pr(Y \in B(S) | \theta) = \int_{B(S)} dP(Y | \theta)$$

- ▶ Generally intractable
- ▶ Inference for θ can proceed using a MCMC approximation:
 - ▶ Given the observed ranks S and a prior distribution $p(\theta)$
 - ▶ Generate a Markov chain whose stationary distribution is that of (θ, Y) given $Y \in F(S)$, $R(S)$ or $B(S)$
 - ▶ Gibbs: simulate $\theta \sim p(\theta | Y)$ and then given $Y \in F(S)$, $R(S)$ or $B(S)$, for $i \neq j$ simulate $y_{ij} \sim p(y_{ij} | \theta, y_{ji})$

Bayesian estimation

- ▶ The FRN, rank and binary likelihoods can be expressed as:

$$L_F(\theta : S) = \Pr(Y \in F(S) | \theta) = \int_{F(S)} dP(Y | \theta)$$

$$L_R(\theta : S) = \Pr(Y \in R(S) | \theta) = \int_{R(S)} dP(Y | \theta)$$

$$L_B(\theta : S) = \Pr(Y \in B(S) | \theta) = \int_{B(S)} dP(Y | \theta)$$

- ▶ Generally intractable
- ▶ Inference for θ can proceed using a MCMC approximation:
 - ▶ Given the observed ranks S and a prior distribution $p(\theta)$
 - ▶ Generate a Markov chain whose stationary distribution is that of (θ, Y) given $Y \in F(S)$, $R(S)$ or $B(S)$
 - ▶ Gibbs: simulate $\theta \sim p(\theta | Y)$ and then given $Y \in F(S)$, $R(S)$ or $B(S)$, for $i \neq j$ simulate $y_{ij} \sim p(y_{ij} | \theta, y_{ji})$
- ▶ Allows for imputation of missing s_{ij}

Bayesian Estimation for FRN

Model: $Y \sim p(Y|\theta)$, $\theta \in \Theta$

Bayesian Estimation for FRN

Model: $Y \sim p(Y|\theta)$, $\theta \in \Theta$

Data: $Y \in F(S)$

Bayesian Estimation for FRN

Model: $Y \sim p(Y|\theta)$, $\theta \in \Theta$

Data: $Y \in F(S)$

Estimation: Given $p(\theta)$, $p(\theta|Y \in F(S))$ can be approximated by a Gibbs sampler:

Bayesian Estimation for FRN

Model: $Y \sim p(Y|\theta)$, $\theta \in \Theta$

Data: $Y \in F(S)$

Estimation: Given $p(\theta)$, $p(\theta|Y \in F(S))$ can be approximated by a Gibbs sampler:

- ▶ Simulate $\theta \sim p(\theta|Y)$.

Bayesian Estimation for FRN

Model: $Y \sim p(Y|\theta)$, $\theta \in \Theta$

Data: $Y \in F(S)$

Estimation: Given $p(\theta)$, $p(\theta|Y \in F(S))$ can be approximated by a Gibbs sampler:

- ▶ Simulate $\theta \sim p(\theta|Y)$.
- ▶ Simulate $y_{ij} \sim p(y_{ij}|\theta, Y_{-ij}, Y \in F(S))$:

Bayesian Estimation for FRN

Model: $Y \sim p(Y|\theta)$, $\theta \in \Theta$

Data: $Y \in F(S)$

Estimation: Given $p(\theta)$, $p(\theta|Y \in F(S))$ can be approximated by a Gibbs sampler:

- ▶ Simulate $\theta \sim p(\theta|Y)$.
- ▶ Simulate $y_{ij} \sim p(y_{ij}|\theta, Y_{-ij}, Y \in F(S))$:
 1. $s_{ij} > 0$: $y_{ij} \sim p(y_{ij}|\theta, Y_{-ij})1_{y_{ij} \in (a,b)}$ where $a = \max(y_{ik} : s_{ik} < s_{ij})$ and $b = \min(y_{ik} : s_{ik} > s_{ij})$.

Bayesian Estimation for FRN

Model: $Y \sim p(Y|\theta)$, $\theta \in \Theta$

Data: $Y \in F(S)$

Estimation: Given $p(\theta)$, $p(\theta|Y \in F(S))$ can be approximated by a Gibbs sampler:

- ▶ Simulate $\theta \sim p(\theta|Y)$.
- ▶ Simulate $y_{ij} \sim p(y_{ij}|\theta, Y_{-ij}, Y \in F(S))$:
 1. $s_{ij} > 0$: $y_{ij} \sim p(y_{ij}|\theta, Y_{-ij})1_{y_{ij} \in (a,b)}$ where $a = \max(y_{ik} : s_{ik} < s_{ij})$ and $b = \min(y_{ik} : s_{ik} > s_{ij})$.
 2. $s_{ij} = 0$ and $d_i < m$: $y_{ij} \sim p(y_{ij}|Y_{-ij}, \theta)1_{y_{ij} \leq 0}$.

Bayesian Estimation for FRN

Model: $Y \sim p(Y|\theta)$, $\theta \in \Theta$

Data: $Y \in F(S)$

Estimation: Given $p(\theta)$, $p(\theta|Y \in F(S))$ can be approximated by a Gibbs sampler:

- ▶ Simulate $\theta \sim p(\theta|Y)$.
- ▶ Simulate $y_{ij} \sim p(y_{ij}|\theta, Y_{-ij}, Y \in F(S))$:
 1. $s_{ij} > 0$: $y_{ij} \sim p(y_{ij}|\theta, Y_{-ij})1_{y_{ij} \in (a,b)}$ where $a = \max(y_{ik} : s_{ik} < s_{ij})$ and $b = \min(y_{ik} : s_{ik} > s_{ij})$.
 2. $s_{ij} = 0$ and $d_i < m$: $y_{ij} \sim p(y_{ij}|Y_{-ij}, \theta)1_{y_{ij} \leq 0}$.
 3. $s_{ij} = 0$ and $d_i = m$: $y_{ij} \sim p(y_{ij}|Y_{-ij}, \theta)1_{y_{ij} \leq \min(y_{ik} : s_{ik} > 0)}$

Simulations

Letting θ before represent the parameters in a regression model, we generated Y from the following Social Relations Model:

$$\begin{aligned}y_{ij} &= \beta^t x_{ij} + a_i + b_j + \epsilon_{ij} \\ \begin{pmatrix} a_i \\ b_i \end{pmatrix} &\stackrel{\text{iid}}{\sim} \text{normal}(0, \Sigma_{ab}) \\ \begin{pmatrix} \epsilon_{ij} \\ \epsilon_{ji} \end{pmatrix} &\stackrel{\text{iid}}{\sim} \text{normal}\left(0, \sigma^2 \begin{pmatrix} 1 & \rho \\ \rho & 1 \end{pmatrix}\right)\end{aligned}$$

Simulations

Letting θ before represent the parameters in a regression model, we generated Y from the following Social Relations Model:

$$\begin{aligned}y_{ij} &= \beta^t x_{ij} + a_i + b_j + \epsilon_{ij} \\ \begin{pmatrix} a_i \\ b_i \end{pmatrix} &\stackrel{\text{iid}}{\sim} \text{normal}(0, \Sigma_{ab}) \\ \begin{pmatrix} \epsilon_{ij} \\ \epsilon_{ji} \end{pmatrix} &\stackrel{\text{iid}}{\sim} \text{normal}\left(0, \sigma^2 \begin{pmatrix} 1 & \rho \\ \rho & 1 \end{pmatrix}\right)\end{aligned}$$

Mean model: $\beta^t x_{ij} = \beta_0 + \beta_r x_{ir} + \beta_c x_{jc} + \beta_{d_1} x_{ij1} + \beta_{d_2} x_{ij2}$

Simulations

Letting θ before represent the parameters in a regression model, we generated Y from the following Social Relations Model:

$$y_{ij} = \beta^t x_{ij} + a_i + b_j + \epsilon_{ij}$$
$$\begin{pmatrix} a_i \\ b_i \end{pmatrix} \stackrel{\text{iid}}{\sim} \text{normal}(0, \Sigma_{ab})$$
$$\begin{pmatrix} \epsilon_{ij} \\ \epsilon_{ji} \end{pmatrix} \stackrel{\text{iid}}{\sim} \text{normal}\left(0, \sigma^2 \begin{pmatrix} 1 & \rho \\ \rho & 1 \end{pmatrix}\right)$$

Mean model: $\beta^t x_{ij} = \beta_0 + \beta_r x_{ir} + \beta_c x_{jc} + \beta_{d_1} x_{ij1} + \beta_{d_2} x_{ij2}$

- ▶ x_{ir}, x_{jc} : individual level variables
- ▶ x_{ij1} : pair specific variable
- ▶ x_{ij2} : co-membership in a group

Simulations

Letting θ before represent the parameters in a regression model, we generated Y from the following Social Relations Model:

$$\begin{aligned} y_{ij} &= \beta^t x_{ij} + a_i + b_j + \epsilon_{ij} \\ \begin{pmatrix} a_i \\ b_i \end{pmatrix} &\stackrel{\text{iid}}{\sim} \text{normal}(0, \Sigma_{ab}) \\ \begin{pmatrix} \epsilon_{ij} \\ \epsilon_{ji} \end{pmatrix} &\stackrel{\text{iid}}{\sim} \text{normal}\left(0, \sigma^2 \begin{pmatrix} 1 & \rho \\ \rho & 1 \end{pmatrix}\right) \end{aligned}$$

Mean model: $\beta^t x_{ij} = \beta_0 + \beta_r x_{ir} + \beta_c x_{jc} + \beta_{d1} x_{ij1} + \beta_{d2} x_{ij2}$

- ▶ x_{ir}, x_{jc} : individual level variables
- ▶ x_{ij1} : pair specific variable
- ▶ x_{ij2} : co-membership in a group

$$\begin{aligned} \beta_r = \beta_c = \beta_{d1} = \beta_{d2} = 1 \text{ and } \beta_0 = -3.26 \\ x_{ir}, x_{ic}, x_{ij1} \stackrel{\text{iid}}{\sim} N(0, 1) \quad x_{ij2} = z_i z_j / .42 \text{ for } z_i \stackrel{\text{iid}}{\sim} \text{binary}(1/2) \end{aligned}$$

Simulations

Letting θ before represent the parameters in a regression model, we generated Y from the following Social Relations Model:

$$\begin{aligned} y_{ij} &= \beta^t x_{ij} + a_i + b_j + \epsilon_{ij} \\ \begin{pmatrix} a_i \\ b_i \end{pmatrix} &\stackrel{\text{iid}}{\sim} \text{normal}(0, \Sigma_{ab}) \\ \begin{pmatrix} \epsilon_{ij} \\ \epsilon_{ji} \end{pmatrix} &\stackrel{\text{iid}}{\sim} \text{normal}\left(0, \sigma^2 \begin{pmatrix} 1 & \rho \\ \rho & 1 \end{pmatrix}\right) \end{aligned}$$

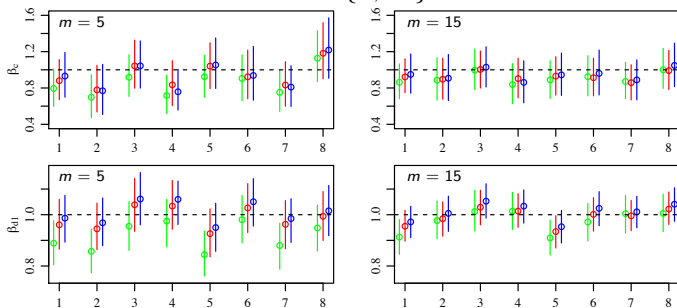
Mean model: $\beta^t x_{ij} = \beta_0 + \beta_r x_{ir} + \beta_c x_{jc} + \beta_{d1} x_{ij1} + \beta_{d2} x_{ij2}$

- ▶ x_{ir}, x_{jc} : individual level variables
- ▶ x_{ij1} : pair specific variable
- ▶ x_{ij2} : co-membership in a group

$$\begin{aligned} \beta_r = \beta_c = \beta_{d1} = \beta_{d2} = 1 \text{ and } \beta_0 = -3.26 \\ x_{ir}, x_{ic}, x_{ij1} \stackrel{\text{iid}}{\sim} N(0, 1) \quad x_{ij2} = z_i z_j / .42 \text{ for } z_i \stackrel{\text{iid}}{\sim} \text{binary}(1/2) \end{aligned}$$

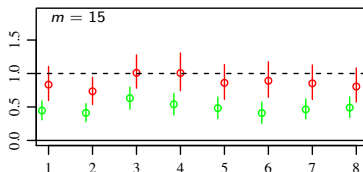
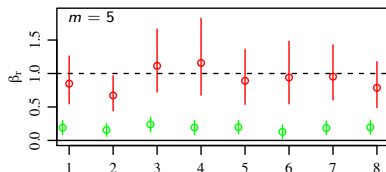
Simulations - Censoring

8 simulations for each $m \in \{5, 15\}$ with 100 nodes each



Confidence intervals under the three different likelihood for column and an iid dyadic variable. The groups of three CIs are based on binary, FRN and rank likelihoods from left to right.

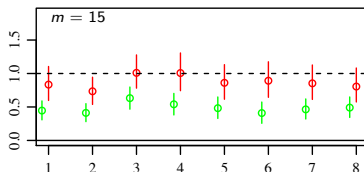
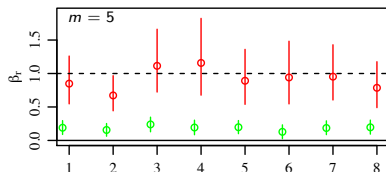
Simulations - Censoring



- Rank likelihood cannot estimate row effects

$$Y \in R(S) \iff Y + c1^t \in R(S) \forall c \in \mathbb{R}^m$$

Simulations - Censoring

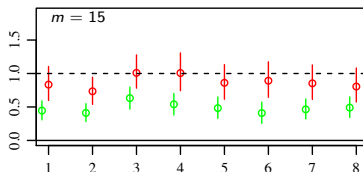
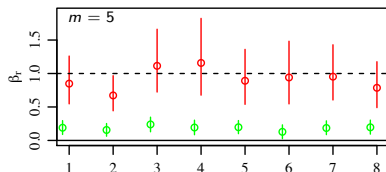


- Rank likelihood cannot estimate row effects

$$Y \in R(S) \iff Y + c1^t \in R(S) \forall c \in \mathbb{R}^m$$

- Binary likelihood poorly estimates row effects

Simulations - Censoring



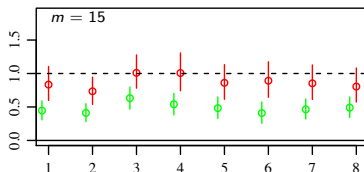
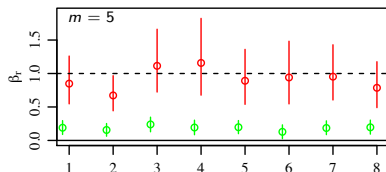
- Rank likelihood cannot estimate row effects

$$Y \in R(S) \iff Y + c1^t \in R(S) \forall c \in \mathbb{R}^m$$

- Binary likelihood poorly estimates row effects

Large amount of censoring

Simulations - Censoring



- Rank likelihood cannot estimate row effects

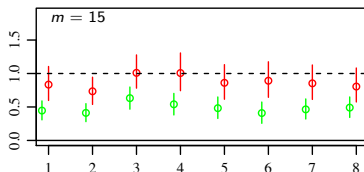
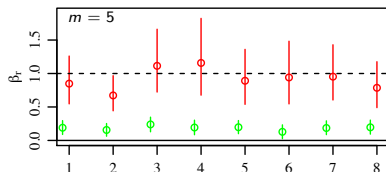
$$Y \in R(S) \iff Y + c1^t \in R(S) \forall c \in \mathbb{R}^m$$

- Binary likelihood poorly estimates row effects

Large amount of censoring

$\Rightarrow$ Heterogeneity of censored outdegrees is low

Simulations - Censoring



- Rank likelihood cannot estimate row effects

$$Y \in R(S) \iff Y + c1^t \in R(S) \forall c \in \mathbb{R}^m$$

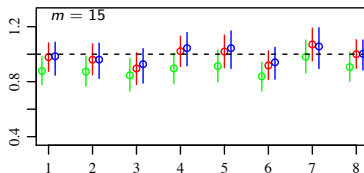
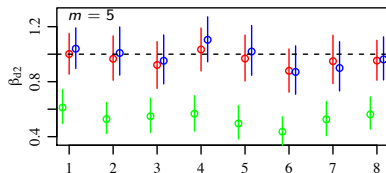
- Binary likelihood poorly estimates row effects

Large amount of censoring

$\Rightarrow$ Heterogeneity of censored outdegrees is low

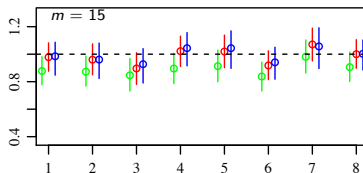
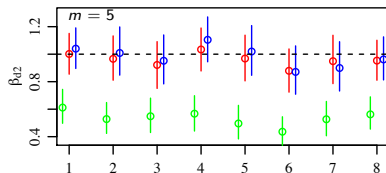
$\Rightarrow$ Regression coefficients estimated too low

Simulations - Censoring



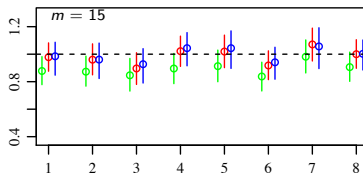
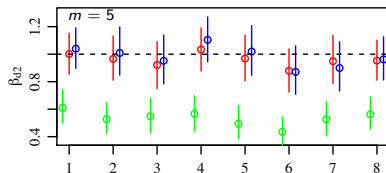
Recall: $x_{ij2} \propto z_i z_j$, an indicator of comembership to a group

Simulations - Censoring



Recall: $x_{ij2} \propto z_i z_j$, an indicator of comembership to a group
Ignore the censoring

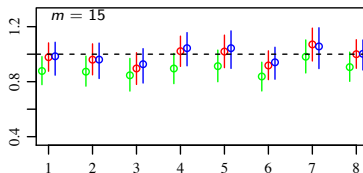
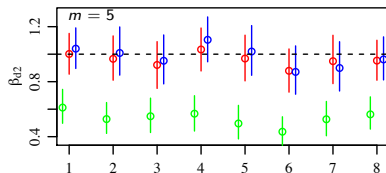
Simulations - Censoring



Recall: $x_{ij2} \propto z_i z_j$, an indicator of comembership to a group
Ignore the censoring

⇒ Binary likelihood underestimates row variability

Simulations - Censoring



Recall: $x_{ij2} \propto z_i z_j$, an indicator of comembership to a group
Ignore the censoring

⇒ Binary likelihood underestimates row variability

⇒ Underestimate the variability in x_{ij2}

Simulations - information in the ranks

Simulations - information in the ranks

Let $C(S)$ be the set of values for which the following is true:

$$s_{ij} > 0 \Rightarrow y_{ij} > 0$$

$$s_{ij} = 0 \text{ and } d_i < n \Rightarrow y_{ij} \leq 0$$

$$\min \{y_{ij} : s_{ij} > 0\} \geq \max \{y_{ij} : s_{ij} = 0\}$$

Simulations - information in the ranks

Let $C(S)$ be the set of values for which the following is true:

$$s_{ij} > 0 \Rightarrow y_{ij} > 0$$

$$s_{ij} = 0 \text{ and } d_i < n \Rightarrow y_{ij} \leq 0$$

$$\min \{y_{ij} : s_{ij} > 0\} \geq \max \{y_{ij} : s_{ij} = 0\}$$

We refer to $L_C(\theta : S) = \Pr(Y \in C(S) | \theta)$ as the censored binary likelihood.

Simulations - information in the ranks

Let $C(S)$ be the set of values for which the following is true:

$$s_{ij} > 0 \Rightarrow y_{ij} > 0$$

$$s_{ij} = 0 \text{ and } d_i < n \Rightarrow y_{ij} \leq 0$$

$$\min \{y_{ij} : s_{ij} > 0\} \geq \max \{y_{ij} : s_{ij} = 0\}$$

We refer to $L_C(\theta : S) = \Pr(Y \in C(S) | \theta)$ as the censored binary likelihood.

- Recognizes censoring but ignores information in the ranks

Simulations - information in the ranks

Let $C(S)$ be the set of values for which the following is true:

$$\begin{aligned}s_{ij} > 0 &\Rightarrow y_{ij} > 0 \\ s_{ij} = 0 \text{ and } d_i < n &\Rightarrow y_{ij} \leq 0 \\ \min \{y_{ij} : s_{ij} > 0\} &\geq \max \{y_{ij} : s_{ij} = 0\}\end{aligned}$$

We refer to $L_C(\theta : S) = \Pr(Y \in C(S) | \theta)$ as the censored binary likelihood.

- ▶ Recognizes censoring but ignores information in the ranks
- ▶ Performs similarly to FRN in the previous study

Simulations - information in the ranks

Let $C(S)$ be the set of values for which the following is true:

$$\begin{aligned}s_{ij} > 0 &\Rightarrow y_{ij} > 0 \\ s_{ij} = 0 \text{ and } d_i < n &\Rightarrow y_{ij} \leq 0 \\ \min \{y_{ij} : s_{ij} > 0\} &\geq \max \{y_{ij} : s_{ij} = 0\}\end{aligned}$$

We refer to $L_C(\theta : S) = \Pr(Y \in C(S) | \theta)$ as the censored binary likelihood.

- ▶ Recognizes censoring but ignores information in the ranks
- ▶ Performs similarly to FRN in the previous study
- ▶ Less precise than FRN when m is big

Simulations - information in the ranks

Let $C(S)$ be the set of values for which the following is true:

$$s_{ij} > 0 \Rightarrow y_{ij} > 0$$

$$s_{ij} = 0 \text{ and } d_i < n \Rightarrow y_{ij} \leq 0$$

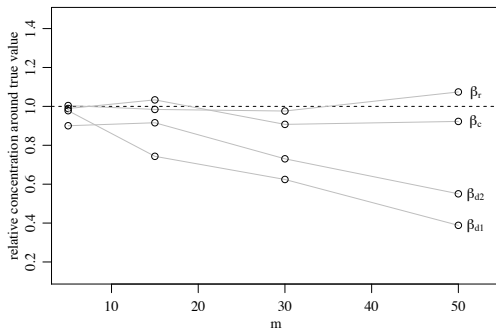
$$\min \{y_{ij} : s_{ij} > 0\} \geq \max \{y_{ij} : s_{ij} = 0\}$$

We refer to $L_C(\theta : S) = \Pr(Y \in C(S) | \theta)$ as the censored binary likelihood.

- ▶ Recognizes censoring but ignores information in the ranks
- ▶ Performs similarly to FRN in the previous study
- ▶ Less precise than FRN when m is big
- ▶ When $m \ll n$, most of the information found by considering ranked/unranked individuals as groups rather than the relative ordering of the ranked individuals.

Simulations - information in the ranks

Same setup as before, but average uncensored outdegree is m

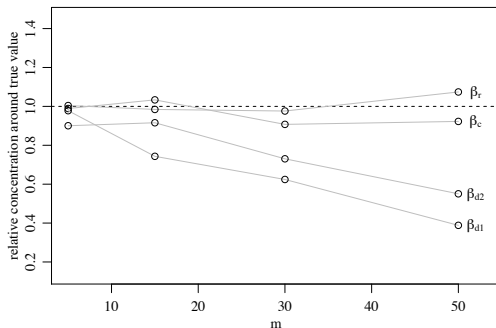


Relative concentration around true value of each parameter:

Measured by $E \left[(\beta - 1)^2 | F(S) \right] / E \left[(\beta - 1)^2 | C(S) \right]$ for each β

Simulations - information in the ranks

Same setup as before, but average uncensored outdegree is m

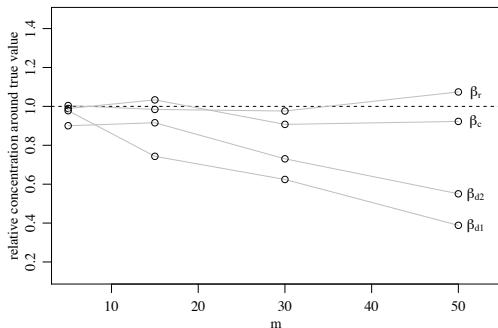


Relative concentration around true value of each parameter:

Measured by $E \left[(\beta - 1)^2 | F(S) \right] / E \left[(\beta - 1)^2 | C(S) \right]$ for each β

Simulations - information in the ranks

Same setup as before, but average uncensored outdegree is m

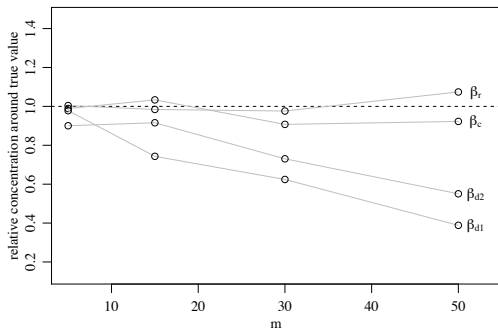


Relative concentration around true value of each parameter:

Measured by $E \left[(\beta - 1)^2 | F(S) \right] / E \left[(\beta - 1)^2 | C(S) \right]$ for each β

Simulations - information in the ranks

Same setup as before, but average uncensored outdegree is m

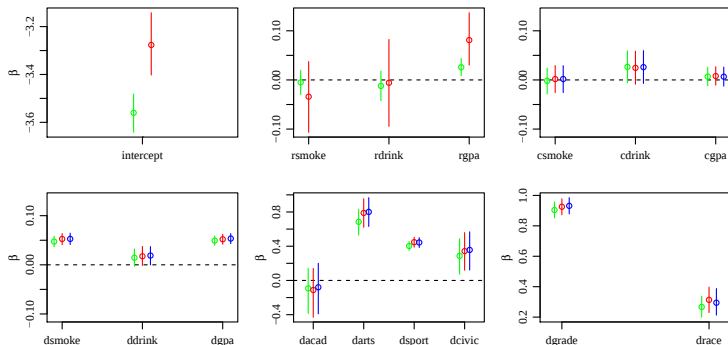


Relative concentration around true value of each parameter:

Measured by $E \left[(\beta - 1)^2 | F(S) \right] / E \left[(\beta - 1)^2 | C(S) \right]$ for each β

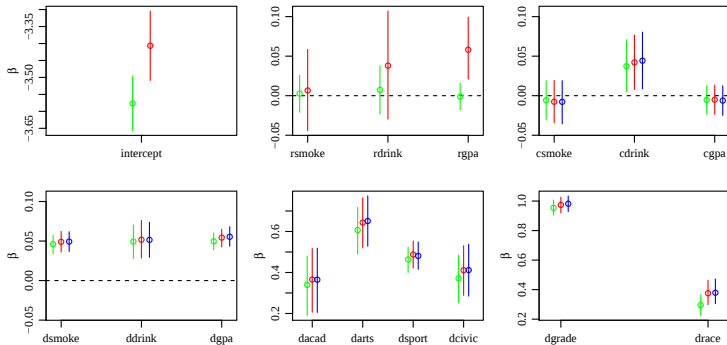
- ▶ When $m \ll n$, most of the information found by considering ranked/unranked individuals as groups rather than the relative ordering of the ranked individuals.

AddHealth Data - Results



- ▶ 622 males were asked to rank up to 5 male friends
- ▶ Fit a mean model with row, column and dyadic effects for smoking, drinking and gpa as well as dyadic effects for comembership in activities and grade, and a similarity-in-race measure.

AddHealth Data - Results (Females)



Results across schools

Likelihood	intercept	row	column	mean-zero dyadic	other dyadic
binomial	0.89 , 1.68	2.22 , 2.95	1.02 , 1.03	1.06 , 1.06	1.20 , 1.09
rank	NA , NA	NA , NA	1.05 , 0.98	0.99 , 0.99	1.06 , 0.98

Results across schools

Likelihood	intercept	row	column	mean-zero dyadic	other dyadic
binomial	0.89 , 1.68	2.22 , 2.95	1.02 , 1.03	1.06 , 1.06	1.20 , 1.09
rank	NA , NA	NA , NA	1.05 , 0.98	0.99 , 0.99	1.06 , 0.98

FRN compared to Binary and Rank likelihoods:

- ▶ First: Average relative magnitudes of parameter estimates
- ▶ Second: Average relative CI widths

Summary: Under binary:

- ▶ $\hat{\beta}_0$ too negative, standard errors too small.
- ▶ $|\hat{\beta}_r|$ too small, standard errors too small.

FRN conclusion

- ▶ Binary likelihood is likely to underestimate the effects of regressors with variation among the nominators of relations.
- ▶ Accounting for censoring is an important first step
- ▶ The FRN likelihood can be used in conjunction with latent variable models to capture network features such as transitivity, clustering or stochastic equivalence.