
Inference for Network Regression Models with Community Structure

Mengjie Pan¹ Tyler H. McCormick² Bailey K. Fosdick³

Abstract

Network regression models, where the outcome comprises the valued edge in a network and the predictors are actor or dyad-level covariates, are used extensively in the social and biological sciences. Valid inference relies on accurately modeling the residual dependencies among the relations. Frequently homogeneity assumptions are placed on the errors which are commonly incorrect and ignore critical, natural clustering of the actors. In this work, we present a novel regression modeling framework that models the errors as resulting from a community-based dependence structure and exploits the subsequent exchangeability properties of the error distribution to obtain parsimonious standard errors for regression parameters.

1. Introduction

Researchers are often interested in how relations between pairs of actors are related to observable covariates, such as demographic, sociological and geographic factors. For example, Ward & Hoff (2007) examine political and institutional effects on international trade and find that the domestic political framework of the exporter and importer are important factors of the trade; Aker (2010) explores the impact of mobile phones on the price difference of grain between a pair of markets and find that the introduction of mobile phone service explains a reduction in grain price dispersion; Fafchamps & Gubert (2007) explore the role of geographic proximity on risk sharing among agriculture workers in the Philippines and find that intra-village mutual insurance links are largely determined by social and geographical proximity, potentially since personal/geographical closeness facilitates enforcement.

^{*}Equal contribution ¹Facebook, Seattle, Washington, USA ²Department of Statistics and Department of Sociology, University of Washington, Seattle, Washington, USA ³Department of Statistics, Colorado State University, Fort Collins, Colorado, USA. Correspondence to: Mengjie Pan <panmj0221@gmail.com>, Tyler H. McCormick <tylermc@uw.edu>, Bailey Fosdick <bailley.fosdick@colostate.edu>.

In this work, we focus on a case where continuous relations between pairs of actors are modeled as a linear function of observable covariates. Continuous, pairwise relations can be represented as a network with directed, weighted edges. We assume a set of observed covariates for the actors and dyads (ordered actor pairs) and wish to study the association between the relational response and covariates. Efficient inference for the effect of the covariates on the relations requires accurate modeling of the dependence between the regression errors. Our main contribution is a novel non-parametric block-exchangeability assumption on the covariance structure of the error vector suitable for when there is excess hidden block variation in the network beyond that accounted for by the covariates. When the underlying error structure of data satisfies the model assumptions, we show that our inference procedures have correct confidence interval coverage. We present algorithms to both estimate the latent block structure and estimate the corresponding standard errors of the regression coefficients. Our goal in the paper is to provide a new approach to model and estimate the dependence between residual relations, which bridges the gap between the existing non-parametric estimators.

Let n be the observed number of individuals, y_{ij} be the directed relational response from actor i to actor j , and $\mathbf{X}_{ij} = [1, X_{1,ij}, \dots, X_{(p-1),ij}]^T$ be a $(p \times 1)$ vector of covariates. We assume there is no relation from an actor i to itself. The regression model can be expressed

$$y_{ij} = \beta^T \mathbf{X}_{ij} + \xi_{ij}, \quad i, j \in \{1, \dots, n\}, i \neq j. \quad (1)$$

We model the error vector $\Xi = [\xi_{21}, \xi_{31}, \dots, \xi_{n1}, \dots, \xi_{1n}, \dots, \xi_{(n-1)n}]^T \sim N(\mathbf{0}, \Omega)$, where $\Omega = \text{Var}(\Xi)$ is a $n(n-1)$ by $n(n-1)$ symmetric matrix. For example, if we are interested in how geographical and demographic factors affect number of mobile calls between actors, then y_{ij} is the number of mobile calls from actor i to actor j and X_{ij} may include the actors' geographical distances and their mobile plans. Making inference on β then provides insights into how a change in geographical distance or mobile plans is associated with a change in the number of mobile calls.

In order to get accurate estimation of the standard error of β and thus a confidence interval with the correct coverage, we need to pose assumptions on the error structure that is satisfied by the data. The challenge in modeling $\Omega = \text{Var}(\Xi)$ is that ξ_{ij} and ξ_{kl} are likely correlated whenever the

relation pairs share a member, i.e. $\{i, j\} \cap \{k, l\} \neq \emptyset$ (Kenny et al. (2006)). The residuals represent variation in relational observations not accounted by observable covariates, and two residuals which both involve actor A may be affected by actor A's individual effects. For example, if the residual of number of mobile calls from actor A to actor B is negative, we may expect number of mobile calls from actor A to actor C is likely also less than expected under the mean model, because actor A does not use mobile calls a lot. Another example is the case of reciprocal relations (Miller & Kenny (1986)). The residuals of number of mobile calls from actor A to actor B and from actor B to actor A are likely correlated, because they involve the same pair of actors and this dependence is often not fully captured by covariates.

One set of approaches to model the covariance structure Ω is to impose parametric distributional assumptions on the error vector or model the error covariance structure directly (Hoff (2005), Ward & Hoff (2007), Hoff et al. (2011), Hoff (2015)). While these approaches produce interpretable representations of underlying residual structure, they always assume the error structure is consistent with an underlying parametric model.

Another set of approaches to model the covariance structure, Ω , is using non-parametric methods. However, existing approaches either make no distributional assumptions and estimate $\mathcal{O}(n^3)$ parameters (see dyadic clustering estimator in Fafchamps & Gubert (2007)), or assume exchangeability of the error vector Ξ and estimate five parameters (Marrs et al. (2017)). The former approach results in a standard error estimator for β that is extremely flexible yet extremely variable, whereas the latter approach assumes all actors are identically distributed and results in a relatively restricted estimator. The former approach is appealing when a researcher does not have any information on the error structure and wants to allow for heterogeneity, while the later approach is appealing when a researcher is more confident the errors are exchangeable and thus can enjoy the simplicity of the error structure and a fixed, small number of covariance parameters. Nevertheless, there are likely cases where a researcher has some information about the error structure, but the errors are not exchangeable. This calls for an approach that bridges the gap between these two existing methods.

We propose an alternative block-exchangeable standard error estimator that assumes that actors have block memberships and actors within the same block are exchangeable (i.e. have relations that are identically distributed). Heterogeneity based on unobserved variables are quite common in networks, and relational observations between actors in the same block may have different patterns than relations between actors in different blocks. The stochastic block model (Holland et al. (1983), Snijders & Nowicki (1997)) and degree-corrected stochastic block model (Karrer & Newman

(2011)) have been proposed to model connectivity between actors based on latent block memberships and actor degree heterogeneities. Spectral clustering algorithms (Rohe et al. (2011), Qin & Rohe (2013)) have also been proposed to estimate the hidden block membership for these models. By imposing the exchangeability assumption on the error vector conditioned on block membership of the actors, we take into account possible block structure in the network residuals and allow for heterogeneity between blocks. Specifically, we propose an algorithm that estimates the covariance matrix $\hat{\Omega}$ given the block memberships, as well as a second algorithm to estimate block memberships using spectral clustering. We present theoretical results proving the block-exchangeable estimator outperforms the exchangeable estimator when the errors are block-exchangeable. Critically, if the distribution of the covariates is dependent on block membership, we see a larger difference in standard errors from the block-exchangeable estimator compared to those from the exchangeable estimator.

2. Previous Methodology

In a linear regression model of form (1), there are a number of ways to model Ω . Fafchamps & Gubert (2007) propose a maximally flexible model for Ω subject to the single condition that $\text{Cov}(\xi_{ij}\xi_{kl}) = 0$ if dyads (i, j) and (k, l) do not share a member, i.e. $\{i, j\} \cap \{k, l\} = \emptyset$. No additional structure is placed on the $\mathcal{O}(n^3)$ remaining covariance terms. This method is known as dyadic clustering, denoted here 'DC', and we let Ω_{DC} denote the covariance matrix under the Fafchamps & Gubert (2007) assumption. Fafchamps & Gubert (2007) propose a simple way to estimate the elements in Ω_{DC} : $\widehat{\text{Cov}}(\xi_{ij}, \xi_{kl}) = r_{ij}r_{kl}$, where r_{ij} and r_{kl} are the residuals of the corresponding relations. While the DC estimator is extremely flexible, the estimator $\hat{\Omega}_{DC}$ contains $\mathcal{O}(n^3)$ parameters and each element is estimated by a single product of residuals. This makes the estimator highly variable, which consequently leads to highly variable β standard errors estimates.

In order to ease the computational burden and decrease the variance of dyadic clustering estimator, Marrs et al. (2017) propose an exchangeability assumption on the error vector and a simple moment-based estimator for the covariance parameters resulting in $\hat{\Omega}_E$. The errors in a relational data model are jointly exchangeable if the probability distribution of the error vector is invariant under simultaneous permutation of the rows and columns. Li & Loken (2002) argue that data generated under the variance component model, which assumes that the observation can be decomposed additively into multiple actor-level components, and the Social Relation Model (Warner et al. (1979), Cockerham & Weir (1977)) satisfy this exchangeability assumption. Under exchangeability and the assumption that the covariance

between relations involving non-overlapping dyads is zero, [Marrs et al. \(2017\)](#) shows that there are five non-zero parameters in Ω (see Figure 1), notably one variance σ^2 and four covariances $\{\sigma_A^2, \sigma_B^2, \sigma_C^2, \sigma_D^2\}$. They estimate these five parameters by averages of the corresponding residual products, greatly reducing the variance of the estimator $\hat{\Omega}_E$ compared to $\hat{\Omega}_{DC}$.

While the number of parameters is significantly reduced under the exchangeability assumption, this assumption may be violated in practice in many scientific settings. For example, when a network has block structure (i.e. community structure), such that actors in different blocks have different behavior patterns, this needs to be accounted for. For instance, conditioned on the covariates, relations in one block may have larger variation than those among actors in another block, thus violating the exchangeability assumption where a single variance is shared among all relations. In the mobile calls example, variance of phone calls among employed actors and that among unemployed actors may likely be different, meaning that without information on actor employment status, residual heterogeneity is likely present. This motivates us to consider a block-exchangeability assumption on Ω .

3. Block-exchangeability

With the dyadic cluster estimator making a single assumption but yielding too many parameters and the exchangeable estimator making strong assumptions, we propose a block-exchangeability assumption that compromises between imposing assumptions on error vector and model complexity. In a network of B latent blocks, let g_i denote the block assignment of actor i : $g_i \in \{1, \dots, B\}$. We propose the following definition of **block-exchangeability** as conditional exchangeability ([Lindley et al. \(1981\)](#)) of Ξ given g :

Definition 3.1. The errors in a relational data model are jointly block-exchangeable if $P(\Xi)$, the probability distribution of the error vector, is invariant under permutation of the rows and columns within each block:

$$P(\Xi) = P(\prod(\Xi)) \text{ such that } g_i = g_{\pi(i)} \text{ and } g_j = g_{\pi(j)},$$

where $\prod(\Xi) = \{\xi_{\pi(i)\pi(j)}\}$ is the residual matrix with its rows and columns reordered according to permutation operator π .

A different exchangeable block assumption in the regression settings is discussed in [McCullagh \(2005\)](#), where the distribution of observations is invariant under permutations that preserve the block-to-block relationship structure, i.e. permutations π such that $B(i, j) = B(\pi(i), \pi(j)) \forall i, j$, where $B(i, j) = 1$ if $g_i = g_j$, and $B(i, j) = 0$ otherwise. There are two key differences between this assumption and that we propose. One is that block-exchangeability

in [McCullagh \(2005\)](#) is on the observations, whereas we propose block-exchangeability on the errors. The other is that the permutation in [McCullagh \(2005\)](#) only requires that $B(i, j) = B(\pi(i), \pi(j)) = 0$, meaning observations that are in different blocks remain in different blocks after permutation.

Under our block-exchangeability assumption and conditional on block membership, the covariance between two arbitrary errors ξ_{ij} and ξ_{kl} takes one of the following six values depending on the block memberships $\{g_i, g_j, g_k, g_l\}$ and relationships among the indices $\{i, j, k, l\}$:

$$\begin{aligned} \text{Var}(\xi_{ij}) &= \sigma_{(g_i, g_j)}^2 & \text{Cov}(\xi_{ij}, \xi_{il}) &= \phi_{B(g_i, \{g_j, g_l\})} \\ \text{Cov}(\xi_{ij}, \xi_{ji}) &= \phi_{A\{g_i, g_j\}} & \text{Cov}(\xi_{ij}, \xi_{kj}) &= \phi_{C(g_j, \{g_i, g_k\})} \\ \text{Cov}(\xi_{ij}, \xi_{kl}) &= 0 & \text{Cov}(\xi_{ij}, \xi_{ki}) &= \phi_{D(g_i, g_j, g_k)} \end{aligned}$$

where $\{\}$ denotes unordered set and $()$ denotes ordered set. Note that this notation is an expansion of that introduced in [Marrs et al. \(2017\)](#) such that all non-zero parameters are now indexed by node block memberships.

	y_{BA}	y_{CA}	y_{DA}	y_{AB}	y_{CB}	y_{DB}	y_{AC}	y_{BC}	y_{DC}	y_{AD}	y_{BD}	y_{CD}
y_{BA}	#	&	&	#	#	#	\$	&		\$	&	
y_{CA}	&	&	#	#	#		&	^	-	&		+
y_{DA}	&	#	&	#		#	&		+	&	^	-
y_{AB}	#	#	#	#	&	&	&	\$		&	\$	
y_{CB}	#	#		&	&	#	^	&	-		&	+
y_{DB}	#		#	&	#	&		&	+	^	&	-
y_{AC}	\$	&	&	&	^		+	+	-	-		+
y_{BC}	&	^		\$	&	&	+	+	-		-	+
y_{DC}		-	+		-	+	-	-	-	+	+	+
y_{AD}	\$	&	&	&		^	-		+	+	+	-
y_{BD}	&		^	\$	&	&		-	+	+	+	-
y_{CD}		+	-		+	-	+	+	+	-	-	+

Figure 1. Visualization of covariance matrix for a network of four actors. Under the block-exchangeability assumption, where A and B are in one block and C and D are in another block, entries shaded with the same color and symbol share the same parameter value. Conversely, under the exchangeability assumption, entries with the same color share the same value.

Figure 1 shows a visualization of Ω_B for a simple network of four actors $\{A, B, C, D\}$, where actors A and B are in Block 1 and actors C and D are in Block 2. Under both exchangeability and block-exchangeability assumption, the blank entries indicate a covariance value of zero between non-overlapping dyads (y_{ij}, y_{kl}) where $\{i, j\} \cap \{k, l\} = \emptyset$. Under the block-exchangeability assumption, each color denotes a dyad configuration and conditioned on the color, each symbol denotes a parameter indexed by the actor block memberships. Thus entries with the same color and symbol share the same parameter value. For example, $\text{Var}(\xi_{CA}) = \text{Var}(\xi_{DA}) = \text{Var}(\xi_{CB}) = \text{Var}(\xi_{DB}) = \sigma_{(2,1)}^2$, as denoted

Covariance term	Number of parameters
$\text{Var}(\xi_{ij}) = \sigma^2$	B^2
$\text{Cov}(\xi_{ij}, \xi_{ji}) = \phi_A$	$B(B+1)/2$
$\text{Cov}(\xi_{ij}, \xi_{il}) = \phi_B$	$B^2(B+1)/2$
$\text{Cov}(\xi_{ij}, \xi_{kj}) = \phi_C$	$B^2(B+1)/2$
$\text{Cov}(\xi_{ij}, \xi_{ki}) = \phi_D$	B^3

Table 1. Number of parameters of each covariance type under the block-exchangeability assumption, where B is the number of actor blocks. Note that if there are fewer than three actors in a given block, some of these parameters will not appear in Ω_B .

by the blue & in Figure 1, because the sender is in Block 2 and the receiver is in Block 1. On the contrary, under the exchangeability assumption of Marrs et al. (2017), entries with the same color share the same value. Therefore, Ω_B has more parameters than Ω_E , while maintaining the same places for zero-valued entries.

Figure 2 shows the configurations of relation pairs under the block-exchangeability assumption. Each circle contains dyad configurations of the same type under exchangeability in Ω_E . However, under block exchangeability, there is variability with each configuration based on actor block memberships and these variations are shown within each circle. For example, the top left circle shows the variance parameters under block-exchangeability: $\sigma_{(1,1)}^2, \sigma_{(1,2)}^2, \sigma_{(2,1)}^2$, and $\sigma_{(2,2)}^2$ corresponding to every ordered pair of blocks. In contrast, under the exchangeability assumption, all variance terms share the same parameter value σ^2 . The top right circle shows four block-exchangeability parameters for the configuration of relations pairs of the form (y_{ij}, y_{kj}) . In the top left corner of this circle, the common receiver actor B is in Block 1, sender A is in Block 1 and sender C is in Block 2 and therefore $\text{Cov}(\xi_{AB}, \xi_{CB}) = \phi_{B,(1,\{1,2\})}$. Because we only have two actors in each block, the case when $\text{Cov}(\xi_{ij}, \xi_{kj}) = \phi_{B,(1,\{1,1\})}$ is not shown in Figure 2 since it would require three actors i, j , and k in Block 1. In general the number of block-exchangeable parameters belonging to each configuration type depends on the number of blocks B (see Table 1).

As shown in Table 1, the number of parameters in Ω_B is on the order of $\mathcal{O}(B^3)$. This is substantially greater than the number of parameters under the exchangeability assumption, which is five regardless of network size, yet significantly smaller than the number of parameters for dyadic clustering, which is on the order of $\mathcal{O}(n^3)$. The block-exchangeability assumption balances between imposing assumptions on error vector and model complexity, in an attempt to model the covariance matrix with a reasonable number of parameters while keeping the assumptions feasible for real world applications.

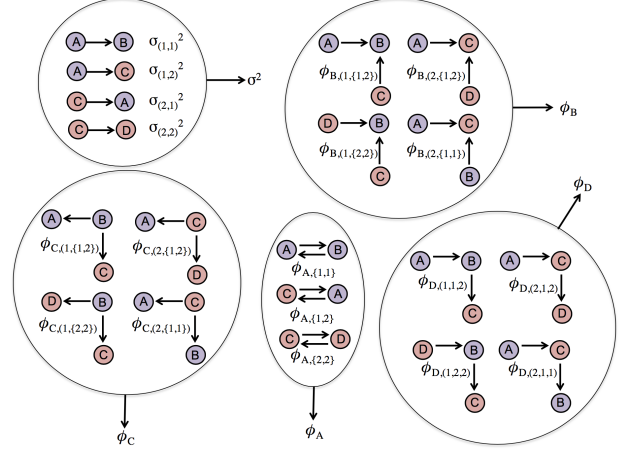


Figure 2. Configurations of directed relation pairs under the block-exchangeability assumption in a simple network of four actors $\{A, B, C, D\}$, where A and B are in one block (indicated by purple color) and C and D are in the other (indicated by light coral color). Each circle represents one dyad configuration, and the parameters within the circle correspond to those under the block-exchangeability assumption. The circles denote which block-exchangeable parameters share the same value under the exchangeability assumption and the associated exchangeable parameter is positioned at the end of the arrow outside the circle.

4. Network Regression with Block-exchangeable Errors

Assuming the errors are block-exchangeable and there are B blocks, we now present algorithms that produce standard error estimates for the coefficients β in a linear regression model (1). Let X be the design matrix, $\hat{\beta} = (X^T X)^{-1} X^T y$ denote the ordinary least squares estimate of β and $r_{ij} = y_{ij} - X \hat{\beta}$ denote the residual for observation y_{ij} .

4.1. Known Blocks

Given block memberships, the estimate of each block-exchangeable parameter is formed by the empirical average of the products of the residual pairs of the same block-dyad configuration type. To formally describe the estimator, let $[B] = \{1, \dots, B\}$ and $[n] = \{1, \dots, n\}$. Let M index the five dyad configurations $M \in \{\sigma^2, \phi_A, \phi_B, \phi_C, \phi_D\}$, and Q_M denote the set of block pairs/triplets for dyad configuration M given $[B]$. Thus $Q_{\sigma^2} = \{(u, v) : u, v \in [B]\}$, and $Q_{\phi_B} = \{(u, \{v, w\}) : u, v, w \in [B]\}$. We explicitly define all other sets Q_M in the supplementary material. Furthermore, let $\Phi_{M,q}$, where $q \in Q_M$, denote the set of ordered relation pairs that have the configuration M and block specification q . Thus $\Phi_{\sigma^2, (u,v)} = \{[(i, j), (i, j)] : i, j \in [n], i \neq j, g_i = u, g_j = v\}$. All other sets $\Phi_{\phi_A, \{u,v\}}, \Phi_{\phi_B, (u, \{v, w\})}$,

Algorithm 1 Known block estimation of Ω_B

Input: residuals $\{r_{ij} : i \neq j\}$, number of blocks B , block memberships $\{g_i\}$

Output: $\hat{\Omega}_B$

1. For each configuration type M and block combination q , calculate the set of residual products associated with $\Phi_{M,q}$:

$$\mathbf{R}_{M,q} = \{r_{jk}r_{mn} : [(j,k), (m,n)] \in \Phi_{M,q}\}$$

3. Estimate the parameters in Ω_B using the empirical average of the corresponding residual products:

$$\hat{\theta}_{M,q} = \frac{\sum_{t:t \in \mathbf{R}_{M,q}} t}{|\mathbf{R}_{M,q}|}$$

where $\theta_{M,q}$ is the block-exchangeable parameter corresponding to M and q .

$\Phi_{\phi_C, (u, \{v, w\})}$ and $\Phi_{\phi_D, (u, v, w)}$ are explicitly defined in the supplementary material. Algorithm 1 formally describes estimation on Ω_B in this known block setting.

4.2. Unknown Blocks

When block memberships are unknown, we propose spectral clustering to estimate them by constructing a similarity matrix from the regression residuals (see Algorithm 2). For each actor i and each dyad configuration M , we extract all pairs of relation residuals that involve actor i as the overlapping actor in the given configuration. Let $\Phi_{M,i}$ denote the set of relation pairs that involve a specific actor i in configuration type $M \in \{\sigma^2, \phi_A, \phi_B, \phi_C, \phi_D\}$. For example, $\Phi_{\sigma^2, i} = \{[(i, j), (i, j)] : j \in [n], i \neq j\} \cup \{[(j, i), (j, i)] : j \in [n], i \neq j\}$ and $\Phi_{\phi_B, i} = \{[(i, j), (i, k)] : j, k \in [n], i \neq j \neq k\}$. Complete definitions of all other sets are provided in the supplementary material. We compute the Kolmogorov-Smirnov statistic between the distribution of residual products that involve actor i and the distribution that involve actor j for each configuration type M and combine these to create a similarity measure between actors i and j . Unnormalized spectral clustering is then performed on the resulting similarity matrix to obtain block membership estimates (Von Luxburg (2007)).

When block memberships are known, we apply Algorithm 1 to obtain $\hat{\Omega}_B$. When block membership are unknown, we apply Algorithm 2 to estimate the block memberships $\{\hat{g}_i\}$ and then apply Algorithm 1 using $\{\hat{g}_i\}$ as an input to obtain $\hat{\Omega}_B$. The value K in step 4 of Algorithm 2 is a tuning parameter and is used to construct a K-nearest neighbor weighted adjacency matrix for input to the spectral clustering. Maier et al. (2007) prove that choosing $K = c_1 n - c_2 \log(n) + c_3$, where $c_1, c_2 \geq 0$ and c_3 are all constants, provides an opti-

Algorithm 2 Block membership estimation

Input: residuals $\{r_{ij} : i \neq j\}$, number of blocks B , K for nearest neighbor graph

Output: estimated block memberships $\{\hat{g}_i\}$

1. For each actor i and configuration $M \in \{\sigma^2, \phi_A, \phi_B, \phi_C, \phi_D\}$, calculate the set of residual products for $\Phi_{M,i}$:

$$\mathbf{R}_{M,i} = \{r_{ab}r_{cd} : [(a,b), (c,d)] \in \Phi_{M,i}\}$$

2. Let $F_{i,M}$ be the empirical distribution function for $\mathbf{R}_{M,i}$. For each pair of actors i and j and for each M , calculate the Kolmogorov-Smirnov statistic

$$KS_{i,j,M} = \sup_x |F_{i,M}(x) - F_{j,M}(x)|.$$

3. For each pair of actors i and j , define

$$s_{ij} = 1 - \left(\sum_{M \in \{\sigma^2, \phi_A, \phi_B, \phi_C, \phi_D\}} KS_{i,j,M} \right) / 5.$$

4. Let $W = (w_{ij})_{i,j=1,\dots,n}$ denote the weighted adjacency matrix, where

$$w_{ij} = w_{ji} = \begin{cases} s_{ij}, & \text{if } i \in KNN(j) \text{ or } j \in KNN(i) \\ 0, & \text{otherwise} \end{cases}$$

where KNN denotes K-nearest neighbor.

5. Perform unnormalized spectral clustering on weighted graph W to get estimated blocks $\hat{g}_i$, where $\hat{g}_i \in [B] \forall i$.

mal choice of K . We found that for our simulation setting $K = 0.2n$ worked well. When block memberships are known, computation of $\hat{\Omega}_B$ is quite inexpensive because the algorithm simply extracts all dyad pairs with the same covariance and averages the residual products (e.g. 5 seconds when $n = 80$, 20 seconds when $n = 160$ on standard machine). When the block memberships are unknown, Step 1 and 2 of Algorithm 2 may be expensive if the network size is large. In these cases, we suggest a modification of Step 2. Instead of letting $F_{i,M}$ be the empirical distribution function for $\mathbf{R}_{M,i}$, we modify $F_{i,M}$ to be the empirical distribution function for quantiles of $\mathbf{R}_{M,i}$. This reduces the size of the set $\mathbf{R}_{M,i}$, which decreases the storage cost as well as the computational cost of computing Kolmogorov-Smirnov statistic.

5. Theoretical Analysis of Estimator

If the block-exchangeability assumption is appropriate, then our method provides accurate estimation of the regression coefficient standard errors, and confidence intervals constructed with such standard errors have the correct coverage. This is why an accurate estimation of standard errors is

important in inference on the coefficients.

Given $\hat{\Omega}_B$ from Algorithm 1, the sandwich covariance estimator can be used to estimate the standard error of the ordinary least squares estimate $\hat{\beta}$:

$$\hat{V}(\hat{\beta}) = (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T \hat{\Omega}_B \mathbf{X} (\mathbf{X}^T \mathbf{X})^{-1}. \quad (2)$$

Observe that entries in $\hat{V}(\hat{\beta})$ are entries in $\hat{\Omega}_B$ weighted by functions of $\mathbf{X}$. It is possible that even with the (incorrect) exchangeable covariance structure in Ω , we still obtain accurate standard error estimation of $\hat{\beta}$ because the difference $\hat{\Omega}_B - \hat{\Omega}_E$ averages out over $\mathbf{X}$. Here we quantify the difference between the standard error estimator of $\hat{\beta}$ under the assumption of exchangeability and block-exchangeability, as a function of covariates $\mathbf{X}$, block assignments $\{g_i\}$, and the true Ω .

Consider a simple linear regression model with only one covariate:

$$y_{ij} = \beta_0 + \beta_1 X_{ij} + \xi_{ij}, \quad (3)$$

where y_{ij} is the observed relation, X_{ij} is a scalar covariate, ξ_{ij} is the error term, β_0 is the intercept, and β_1 is the covariate coefficient. In addition, assume there is a two block structure in the network, with block sizes n_1 and n_2 , respectively, where $n_1 + n_2 = n$. Under the assumption that the error vector is block-exchangeable, we show that the difference in $\hat{V}(\hat{\beta})$ with the exchangeable estimator $\hat{\Omega}_E$, denoted $\hat{V}_E(\hat{\beta})$, and that with the block-exchangeable estimator $\hat{\Omega}_B$, denoted $\hat{V}_B(\hat{\beta})$, converges in probability to a matrix that depends on the distribution of the covariate, block assignments, and parameters in Ω_B .

THEOREM 5.1. Assume (a) the error vector satisfies the block-exchangeability assumption, with two blocks of sizes n_1 and n_2 , (b) $\mathbf{X}$ is a full rank $(n(n-1) \times 2)$ matrix, (c) covariates $\{X_{ij}\}$ are independent and identically distributed, (d) the fourth moment of the errors and covariates are bounded, (e) errors Ξ and $\mathbf{X}$ are independent, and (f) the number of blocks B is $\mathcal{O}(1)$. As $n_1 \rightarrow \infty, n_2 \rightarrow \infty$, and $n_1/n_2 \rightarrow \alpha$, where α is a constant such that $0 < \alpha < \infty$,

$$n \left(\hat{V}_B(\hat{\beta}) - \hat{V}_E(\hat{\beta}) \right) \xrightarrow{p} c(\mathbf{X}). \quad (4)$$

where $c(\mathbf{X})$ is a weighted linear combination of the differences between the true block exchangeable parameters and the corresponding exchangeable parameters (when the block exchangeable parameters are appropriately averaged within configuration type) and convergence is pointwise. Furthermore, when X_{ij} is independent of g_i and g_j , $c(\mathbf{X}) = \mathbf{0}$ and thus the estimators are asymptotically equivalent.

Proof of this theorem is provided in the supplementary materials. The corresponding exchangeable parameter σ^2 under block-exchangeability is a weighted average of

$\sigma_{(1,1)}^2, \sigma_{(1,2)}^2, \sigma_{(2,1)}^2$ and $\sigma_{(2,2)}^2$. Note we can interpret σ^2 as the common variance term if the error vector is in fact exchangeable. We use the same logic for the other four configurations, and recognize that the magnitude and sign of the difference in standard errors using block-exchangeable estimator and exchangeable estimator are determined by a sum of weighted differences of all five configurations. Therefore, whether the exchangeable estimator has over- or under- coverage depends on parameters in Ω_B , $\{g_i\}$, and $\{X_{ij}\}$.

The second part of the theorem notes that even if the differences $\sigma_{(1,1)}^2 - \sigma^2, \sigma_{(1,2)}^2 - \sigma^2, \sigma_{(2,1)}^2 - \sigma^2, \sigma_{(2,2)}^2 - \sigma^2$ are nonzero, as long as the covariate X_{ij} is independent of block memberships g_i and g_j , on average the difference will disappear after adjusted by weights. This is a critical insight, because we see that in order for the block-exchangeable estimator to have lower bias than exchangeable estimator, we need (1) the error vector satisfies block exchangeability but not exchangeability, and (2) the distribution of X_{ij} is correlated with g_i and g_j .

6. Simulations

To evaluate the performance of our proposed block-exchangeable error model, we generate data from a modified latent space model (Hoff, 2005), which satisfies the requirements for block exchangeability. We consider a simple regression model with one covariate, as in (3) where both coefficients equal 1. We consider three settings for the relationship between the covariate and block structure, and three types of covariates. Figure 3 shows the coverage of 95% confidence intervals for β_1 for all nine simulation settings. The first column represents the cases where the covariate X_{ij} is uncorrelated with block membership, the second column represents the case where relations with high variance in X_{ij} are correlated with low variance errors ξ_{ij} , and the third column represents the case where relations with high variance in X_{ij} also have high variance errors ξ_{ij} . The rows represent different covariates: the first row is a binary indicator of actors sharing an attribute $X_{ij,1} = \mathbb{1}_{[X_i=X_j]}$, the second row represents the absolute difference between an actor attribute $X_{ij,2} = |X_i - X_j|$, and the third row represents a pairwise covariate with block structure $X_{ij,3} \sim N(0, a_{g_i, g_j}^2)$. We generated 1000 errors for each of 500 simulations of the covariates and block memberships, and considered networks of size 20, 40, 80, and 160. We consider four estimators of Ω that are then plugged into the sandwich estimator (2) to obtain a confidence interval for β_1 . The red box shows the coverage using the block-exchangeable estimator conditioned on the true block membership (Algorithm 1), the blue box shows the coverage using the block-exchangeable estimator with the estimated block membership (Algorithms 1 and 2), the yellow box shows the coverage using exchange-

able estimator, and the purple box shows the coverage using the dyad clustering estimator. For each boxplot, the middle line indicates the median coverage, the top and bottom boundaries indicate the 90% and 10% percentiles, and the top and bottom whiskers indicate the 97.5% and 2.5% percentiles.

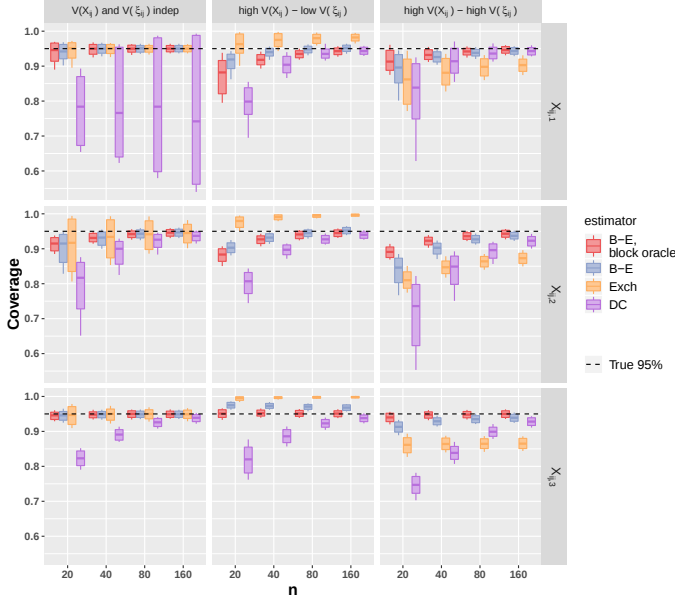


Figure 3. Coverage of 95% confidence interval for β_1 under three error settings and three covariate types for the block-exchangeable standard error estimator conditioned on the true block memberships (oracle), block-exchangeable standard error estimator with estimated blocks, exchangeable standard error estimator, and dyadic clustering standard error estimator.

The block-exchangeable estimator performs similarly to the exchangeable estimator when the covariate is uncorrelated the errors, while the block-exchangeable estimator substantially outperforms the exchangeable estimator when the covariate is correlated the errors. This is consistent with our theoretical results in Section 5. When high variance in a relation’s covariate X_{ij} is associated with low variance in the error ξ_{ij} , we observe that the exchangeable estimator is conservative, and the bias in coverage probability increases with increasing network size. On the contrary, the coverage bias of block-exchangeable estimator decreases with increasing network size. When high variance in a relation’s covariate X_{ij} is associated with high variance in the error ξ_{ij} , the exchangeable estimator is anti-conservative, and its performance improves little with increasing network size. Most notably, at $n = 160$, the exchangeable estimator’s coverage is worse than the dyadic clustering estimator, which is evidence that estimators with strict assumptions perform worse than distribution-free estimators when the

assumptions are violated. In addition, we observe that the differences between the oracle block-estimator using true block memberships and the block-estimator using estimated block membership decreases with increasing network size, suggesting that our block estimation gets better with increasing network size.

7. Air Traffic Data

We demonstrate our method on data representing passenger volume between US airports (Bureau of Transportation Statistics, 2016). The data consist of origin, destination, and number of passengers by month for $n = 573$ airports for all months of 2016. The number of passenger seats is a right-tailed skewed distribution, so we use the values $y'_{ij} = \log(y_{ij} + 1)$ as the relational observations for regression model in (1). For covariates, we calculated the great circle distance between two airports using their longitudes and latitudes. Additionally, we identified the county of the municipality of each airport, and found the total GDP of that county from of Economic Analysis (2015) and average payroll of an employed person from Bureau (2015). We standardized the distance, GDP, and average payroll measures before using them as covariates in the model.

An additional complication in this data is that, for most airports, there is no direct traffic between them. Using ordinary least squares on only the positive observations results in an inconsistent estimator of β (Wooldridge, 2001) (Chapter 16.3). Therefore, instead, we estimated both the regression coefficients β and covariance parameters using a maximum pseudo-likelihood approach (Arnold & Strauss, 1991; Besag, 1975; Strauss & Ikeda, 1990). We use techniques similar to Fieuws & Verbeke (2006) and Solomon & Weissfeld (2017) for longitudinal observations on the same individual, but modified the approach to account for network structure.

A required input to the pseudo-likelihood estimation procedure is known or estimated block memberships. A preliminary estimate of $\hat{\beta}_E$ was obtained assuming exchangeable errors. These coefficient estimates were then used to compute residuals $r_{ij} = y_{ij} - \hat{\beta}_E X_{ij}$, $\forall y_{ij} > 0$, and Algorithm 2 was performed on the residuals for positive observations to obtain block membership estimates. Given the block memberships, the covariate effects β and the block-exchangeable covariance parameters could be simultaneously estimated using the pseudo-likelihood framework.

To numerically optimize the pseudo-likelihood, we used `optim` in R, with `method="L-BFGS-B"`. We do not set bounds on β , but did place a lower bound of $1e^{-2}$ for all variance parameters and a bound of $[-0.9, 0.9]$ for all correlation parameters. We used the eigengap method, which locates a large gap between two subsequent eigenvalues, to

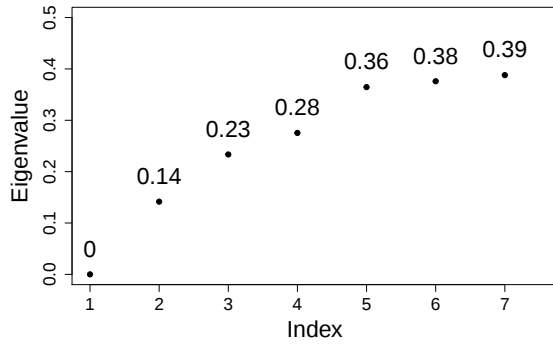


Figure 4. Smallest eigenvalues of the Laplacian matrix in increasing order

choose the number of blocks B . Figure 4 shows the smallest seven eigenvalues in increasing order. The gap between λ_2 and λ_3 is larger than the gap between λ_3 and λ_4 , suggesting that $B = 2$ is a reasonable choice. Figure 4 also shows that the gap between λ_4 and λ_5 is large, suggest $B = 4$ may also be appropriate. When running the spectral clustering algorithm with $B = 3$ and $B = 4$, the smallest block size contained just two airports. Therefore, we proceeded with fitting a block-exchangeable covariance estimator with two blocks, which resulted in one block estimated to have 49 airports, and the other having 524 airports. Full details are provided in the supplementary materials.

Table 2 shows 95% confidence intervals of coefficients using exchangeable estimator and block-exchangeable estimator. We see that the distance between airports is negatively associated with number of passenger seats, while GDP and average payroll of both departure and arrival airports' counties are positively associated with traffic between airports. Compared to the exchangeable estimator, the block-exchangeable estimator returns large effects of economic factors on airport traffic.

8. Discussion

In this paper, we propose a novel block-exchangeable estimator to estimate the standard errors of regression coefficients, assuming block-exchangeability. Our proposed estimator bridges the gap between the existing dyadic clustering estimator, where no distributional assumptions are made, and the exchangeable estimator, where the joint distribution of errors are assumed to be exchangeable. Through theory and simulations, we have shown that when latent block memberships are correlated with the generative process of the covariates, our block-exchangeable estimator outperforms the exchangeable estimator by having less bias

	intercept	distance	GDP _i
Exch	(-29.00, -28.94)	(-7.27, -7.21)	(1.89, 1.95)
B-E	(-26.25, -26.18)	(-6.99, -6.92)	(3.38, 3.46)

	GDP _j	payroll _i	payroll _j
Exch	(1.89, 1.95)	(1.42, 1.48)	(1.41, 1.47)
B-E	(3.44, 3.51)	(2.99, 3.06)	(2.96, 3.04)

Table 2. 95% confidence intervals of coefficients using exchangeable estimator and block-exchangeable estimator. Distance represents the standardized great circle distance between two airports, GDP_i and GDP_j denote standardized GDP for departure and arrival airport, respectively, and payroll_i and payroll_j denote standardized average payroll for departure and arrival airport, respectively.

of coverage, and outperforms dyadic clustering estimator by having less variance.

Because there may not exist a link between every pair of actors in real network data, we extend our estimation algorithms to a case where we assume relational observations are zero left censored. In contrast to the method of moments approach we propose for uncensored data, we use a maximum pseudo-likelihood approach to estimate both the regression coefficients and covariance parameters simultaneously. Although maximum pseudo-likelihood estimates are less preferable to maximum likelihood estimates, using the likelihood directly is not computationally feasible in this censored data setting.

Although we focus our discussion on the impact of block dependence on inference for regression coefficients, possibly equally as interesting, is how the covariance structure, and inferred block structure, is impacted by the inclusion of covariates. In many settings—namely where a researcher is conducting experiments on graphs or wants to make causal claims—the role of covariates is often paramount. As an example, if a researcher can identify covariates that induce very strong residual block structure, these blocks may suffice for units for randomized in a causal inference study.

There are a few limitations of our work, and we discuss them here. We consider linear regression and continuous relational observations on a fully connected network, and assume actors are sampled randomly. A future direction for this work includes extending it to respondent-driven samples. Extending this approach to the generalized linear model framework is unfortunately nontrivial due to the coupling of the relation mean and variance in non-Gaussian link functions. Additionally, if the block sizes are unbalanced, the variance of the estimated parameters associated with the smallest block is presumably largest. Comparing the performances of different estimators at various levels of unbalanced block size is a direction for future study. Finally,

in the case of unknown block memberships, Algorithm 2 attempts to identify memberships based on similarities between the distribution of actor residual products. Computing these similarity scores is computationally intensive and in our examples, required a matter of hours using a standard laptop with codes written in R and not optimized for efficiency.

Acknowledgements

We thank the anonymous reviewers that provided feedback on our work. This work was partially supported by NSF awards IOS-1856229 and DMS-1737673, as well as the National Institute Of Mental Health of the National Institutes of Health under Award Number DP2MH122405.

References

- Aker, J. C. Information from markets near and far: Mobile phones and agricultural markets in Niger. *American Economic Journal: Applied Economics*, 2(3):46–59, 2010.
- Arnold, B. C. and Strauss, D. Pseudolikelihood estimation: Some examples. *Sankhyā: The Indian Journal of Statistics, Series B*, pp. 233–243, 1991.
- Besag, J. Statistical analysis of non-lattice data. *Journal of the Royal Statistical Society: Series D (The Statistician)*, 24(3):179–195, 1975.
- Bureau, U. S. C. 2015 SUSB annual data tables by establishment industry, 2015. data retrieved from <https://www.census.gov/data/tables/2015/econ/susb/2015-susb-annual.html>.
- Bureau of Transportation Statistics. On-flight market passengers, 2016. data retrieved from https://www.transtats.bts.gov/DL_SelectFields.asp?Table_ID=292.
- Cockerham, C. C. and Weir, B. S. Quadratic analyses of reciprocal crosses. *Biometrics*, pp. 187–203, 1977.
- Fafchamps, M. and Gubert, F. The formation of risk sharing networks. *Journal of Development Economics*, 83(2): 326–350, 2007.
- Fieuws, S. and Verbeke, G. Pairwise fitting of mixed models for the joint modeling of multivariate longitudinal profiles. *Biometrics*, 62(2):424–431, 2006.
- Hoff, P. D. Bilinear mixed-effects models for dyadic data. *Journal of the American Statistical Association*, 100(469): 286–295, 2005.
- Hoff, P. D. Multilinear tensor regression for longitudinal relational data. *The Annals of Applied Statistics*, 9(3): 1169, 2015.
- Hoff, P. D. et al. Separable covariance arrays via the tucker product, with applications to multivariate relational data. *Bayesian Analysis*, 6(2):179–196, 2011.
- Holland, P. W., Laskey, K. B., and Leinhardt, S. Stochastic blockmodels: First steps. *Social Networks*, 5(2):109–137, 1983.
- Karrer, B. and Newman, M. E. Stochastic blockmodels and community structure in networks. *Physical Review E*, 83(1):016107, 2011.
- Kenny, D. A., Kashy, D. A., and Cook, W. L. *Dyadic data analysis*. Guilford press, 2006.
- Li, H. and Loken, E. A unified theory of statistical analysis and inference for variance component models for dyadic data. *Statistica Sinica*, pp. 519–535, 2002.
- Lindley, D. V., Novick, M. R., et al. The role of exchangeability in inference. *The Annals of Statistics*, 9(1):45–58, 1981.
- Maier, M., Hein, M., and Von Luxburg, U. Cluster identification in nearest-neighbor graphs. In *International Conference on Algorithmic Learning Theory*, pp. 196–210. Springer, 2007.
- Marrs, F. W., Fosdick, B. K., and McCormick, T. H. Standard errors for regression on relational data with exchangeable errors. *arXiv preprint arXiv:1701.05530*, 2017.
- McCullagh, P. Exchangeability and regression models. 2005.
- Miller, L. C. and Kenny, D. A. Reciprocity of self-disclosure at the individual and dyadic levels: A social relations analysis. *Journal of Personality and Social Psychology*, 50(4):713, 1986.
- of Economic Analysis, B. Gdp by county, 2015. data retrieved from <https://www.bea.gov/data/gdp/gdp-county>.
- Qin, T. and Rohe, K. Regularized spectral clustering under the degree-corrected stochastic blockmodel. In *Advances in Neural Information Processing Systems*, pp. 3120–3128, 2013.
- Rohe, K., Chatterjee, S., Yu, B., et al. Spectral clustering and the high-dimensional stochastic blockmodel. *The Annals of Statistics*, 39(4):1878–1915, 2011.
- Snijders, T. A. and Nowicki, K. Estimation and prediction for stochastic blockmodels for graphs with latent block structure. *Journal of Classification*, 14(1):75–100, 1997.

- Solomon, G. and Weissfeld, L. Pseudo maximum likelihood approach for the analysis of multivariate left-censored longitudinal data. *Statistics in Medicine*, 36(1):81–91, 2017.
- Strauss, D. and Ikeda, M. Pseudolikelihood estimation for social networks. *Journal of the American Statistical Association*, 85(409):204–212, 1990.
- Von Luxburg, U. A tutorial on spectral clustering. *Statistics and computing*, 17(4):395–416, 2007.
- Ward, M. D. and Hoff, P. D. Persistent patterns of international commerce. *Journal of Peace Research*, 44(2): 157–175, 2007.
- Warner, R. M., Kenny, D. A., and Stoto, M. A new round robin analysis of variance for social interaction data. *Journal of Personality and Social Psychology*, 37(10):1742, 1979.
- Wooldridge, J. M. Econometric analysis of cross section and panel data. 2001.