

Reliable network inference from unreliable data: A tutorial on latent network modeling using STRAND

Daniel Redhead^{1*}, Richard McElreath¹, Cody T. Ross¹

¹ Department of Human Behaviour, Ecology and Culture

Max Planck Institute for Evolutionary Anthropology, Deutscher Platz 6, 04103 Leipzig, Germany.

*Corresponding author: Daniel Redhead, daniel_redhead@eva.mpg.de

ABSTRACT

Social network analysis provides an important framework for studying the causes, consequences, and structure of social ties. Standard self-report measures—e.g., as collected through the popular ‘name-generator’ method—however, do not provide an impartial representation of transfers, interactions, or social relationships. At best, they represent perceptions filtered through the cognitive biases of respondents. Individuals may, for example, report transfers that did not really occur, or forget to mention transfers that really did. The propensity to make such reporting inaccuracies is both an individual-level and item-level characteristic—variable across members of any given group. Past research has highlighted that many network-level properties are highly sensitive to such reporting inaccuracies. However, there remains a dearth of easily deployed statistical tools that account for such biases. To address this issue, we introduce a latent network model that allows us to jointly estimate parameters measuring both reporting biases and a latent, underlying social network. Building upon past research, we conduct several simulation experiments in which network data are subject to various reporting biases, and find that these reporting biases strongly impact our ability to accurately infer fundamental network properties. These impacts are not adequately addressed using standard approaches to network reconstruction (i.e., treating either the union or the intersection of double-sampled data as the true network), but are appropriately resolved through the use of our latent network models. To make implementation of our models easier for end-users, we provide a fully-documented R package, STRAND, and include a tutorial illustrating its functionality when applied to empirical food/money sharing data from a rural Colombian population.

Key words. latent variable models, measurement error, social network analysis, social relations.

1. Introduction

Social network analysis has become a popular framework for representing the social relationships that occur within diverse and complex social systems. From school classrooms across North America and the Netherlands (Dijkstra et al. 2015; Harris 2013), to rural villages in India (Power and Ready 2018), and small-scale subsistence communities in Latin America (Koster and Leckie 2014; Pisor et al. 2020; Redhead and von Rueden 2021), the use of social network analysis allows researchers to examine the relationships and interactions (i.e., ties or links) between individuals (i.e., nodes) within bounded or closed communities (i.e., networks). The rich structure of social relationships that characterizes humans—and many other socially-living species (Krause et al. 2009; Pinter-Wollman et al. 2014)—constitutes an important force in shaping individuals’ perceptions of, beliefs about, and positions within, their wider communities. The structure of these relationships can guide a plethora of important social (Lin 2002; Redhead and Power 2021), economic and financial (Banerjee et al. 2013; Jackson et al. 2017), epidemiological and health-related (Bansal et al. 2010; Holt-Lunstad et al. 2010; Smith and Christakis 2008; Ready et al. 2020), and evolutionary (Kurvers et al. 2014) outcomes.

The important roles that network structure and social position have for understanding many social and behavioral phenomena (Borgatti et al. 2009; Buyalskaya et al. 2021), and the biases that shape an individuals’ perceptions of their social relationships (e.g. Bernard et al. 1979; Killworth and Bernard 1976, 1979)—and thus influence conventional self-report measures of social networks—have frequently been noted across the social sciences. This work has inspired a great deal of effort to reduce such biases through research design and sampling (e.g. Marsden 2005). However, there remains a paucity of models and applications that correct for these biases (but see Butts 2003; Young et al. 2020).

Social network data are typically collected via elicitation of self-report nominations; respondents provide the names of other individuals with whom they have specific kinds of interactions or relationships (this self-report method is often referred to as ‘the name generator method’; Marsden 1990). Such data, however, cannot be considered impartial accounts of the true state of social ties. Rather, they reflect individuals’ subjective perceptions of their own social realities (Krackhardt 1987). Importantly, the perceptions of two individuals about a single relationship can conflict with one another. For example, one person may report giving a loan to a second

person, but that second person might not report having ever received a loan from the first person. In such cases, it is not immediately clear how researchers should code the resulting social network. Should a directed tie be coded as present if either respondent reports it? Or only if both do?

There has been a wealth of empirical evidence suggesting that non-trivial amounts of bias and measurement error characterize self-reported social network data (e.g., Bernard et al. 1979; Bell et al. 2007; Campbell and Lee 1991; Eagle and Proeschold-Bell 2015; Kossinets 2006; Marin 2004; Pustejovsky and Spillane 2009). However, most contemporary research still firmly relies on the use of self-report instruments to collect such data. It is also standard practice to use statistical models that do not control for reporting accuracy; the mainstream statistical tools used in the social sciences (e.g., exponential random graph models (Lusher et al. 2013), stochastic actor-oriented models (Snijders 2017), and social relations models (Kenny and La Voie 1984)) treat reported data as if it were ground-truth.

We begin our paper by reviewing the literature concerning the most likely sources of bias in data collected using the *name generator method*, with a specific focus on recall, frequency, attribute-related, and question order biases. Following this, we present a latent network modeling framework that extends tools introduced by Butts (2003) and Young et al. (2020) to assess and account for these biases. To allow end-users to easily implement our models, we have organized our simulation and analysis tools into a fully-documented R package, STRAND, available at: <https://github.com/ctross/STRAND>. We then conduct several simulation experiments to examine the accuracy and robustness of our modeling framework. Finally, we apply our latent network models to an empirical data-set, provide end-users with a set-by-step tutorial for using STRAND, and highlight the potential use of latent network models across the social sciences.

1.1. The (noisy) measurement of social networks

Social network analysis allows researchers to parse the potential mechanisms that influence social relationships at different levels—i.e., the individual level (e.g., popularity), dyadic level (e.g., attribute-similarity), or higher-order levels (e.g., the formation of transitive groups where ‘friends of friends become friends’). Given the promise of such an approach, network-based frameworks are now widely used for assessing the formation and maintenance of friendships (Ball and Newman 2013; Krackhardt and Kilduff 1999; Selfhout et al. 2010), examining theories of social support/cooperation (Koster and Leckie 2014; Lakey and Cohen 2000; Nolin 2012; Power 2017; von Rueden et al. 2019; Zhu et al. 2013) and/or animus (Gervais 2017; Pisor and Ross 2021), describing the dynamics of drug and alcohol use (Knecht et al. 2011; Ready et al. 2020), and measuring interaction rates, spatial proximity, and tolerance in human and non-human animals (Crofoot et al. 2011; Farine et al. 2016; Eagle et al. 2009; DeTroy et al. 2021).

It is common for social networks to be measured using free-list self-report nominations (i.e., in ‘name

generator’ designs; Marin and Hampton 2007), or self-reports of the relationship between every potential dyad in which the respondent could appear (i.e., in roster-based designs; Marsden 2005). While social relationships can be captured using several different measurement instruments (see Marsden 2005; Ross and Redhead 2021, for reviews), the most popular questionnaire-based approaches fall into two broad categories: event-recall questions and perceptual questions.

With event-recall questions, researchers attempt to record specific relationships based on time- or frequency-anchored recall of specific events. For example, to record food sharing partnerships, participants may be asked to: “Name the individuals who have given you at least 1kg of meat at any point in the last three months”. With perceptual questions, researchers attempt to characterize more qualitative or hypothetical relationships. For example, respondents might be asked to: “Name the individuals who you could go to for emotional support during a time of hardship”. This second category of question is generally used to measure broader, relatively qualitative and emotive relationships, the most popular of which is friendship (Furman 1996). Many view these different approaches as examining conceptually distinct entities. The first approach—often referred to as the ‘classical’ or ‘objectivist’ framework—attempts to evaluate relationships based on observable quantities. The latter approach—sometimes called the ‘cognitivist’ framework—assesses looser, cognitive constructs, and does not consider self-report networks as representations of anything other than the participants’ perceptions of network structure, which have no ‘true’ or correct form (Bernard et al. 1979; Krackhardt 1987; Marsden 1990).

Examining the overall accuracy of network measurements is an important and straightforward endeavor through an ‘objectivist’ lens—as one may consider a self-report data-set about transfers of tangible goods, for example, to be an approximation of the ‘true’, underlying network of transfers. The use of recall-based measurement instruments lends itself to such a perspective, as these kinds of questionnaires aim to roughly capture actual, observable interactions. That is, an individual may report having loaned money to another individual within their group, and that report can, in principle, be validated by actual observations of financial exchanges. Consequently, a fundamental aim of the objectivist approach is to reduce the influence of measurement error on resultant inferences (Bernard et al. 1979). In ideal cases, this goal can be accomplished at the level of research design—for example, by implementing data collection methods that make use of objective information, rather than self-reports (e.g., Koster et al. 2013)—in other cases, it can be accomplished statistically (Butts 2003; Young et al. 2020).

The idea of accounting for measurement error and variation in respondent accuracy is, however, complicated when considering questions about qualitative or counterfactual relationships as these constructs might not have quantifiable ‘true’ states. Thus, the relevance of possible ‘inaccuracies’ in such forms of data has been subject to enduring debate. Several classical papers have recommended that researchers focus on determining the

cognitive structures that guide reports of such relationships, rather than attempting to characterize the accuracy of responses (Freeman 1992; Freeman et al. 1987; Hammer 1980). Nonetheless, empirical network data on such relationships almost universally feature anomalous patterns of responses. For example, respondents on opposite ends of a dyad have a tendency to disagree about the simple presence or absence of a dyadic tie, as there is rarely ever perfect (or even moderate) reciprocity in friendship nominations (Ball and Newman 2013). If researchers have questions about the roles that friendship ties have on other dyadic characteristics, then the inability to define if a given tie exists remains a serious methodological problem—and one that can be addressed statistically without having to treat differential perceptions of social connections as ‘inaccurate’ in any philosophical sense.

By adopting a Bayesian latent network approach, we gain the ability to build a more realistic characterization of social relationships that balances the weight given to different reports of the same tie, based on the extent to which each person’s reports are generally in concordance or discordance with the reports of others in their own community. Our latent network approach combines estimation of standard social network parameters (e.g., stochastic block structure, in- and out-degree distributions, and dyad-level random effects) with estimation of parameters that measure individual-specific propensities to over-nominate ties that others in the community do not acknowledge as existing, or under-nominate ties that others in the community do acknowledge as existing. Because the latent network approach provides posterior estimates of tie probability for each dyad, un-ambiguous ties will have narrower posterior distributions, and thus contribute more information to downstream analyses, than ties which are ambiguous.

In what remains of this section, we overview the most salient biases that may impact the quality of self-report social network data. Although we will generally frame our examples using ‘objectivist’ language, we do this only for expositional clarity; our approach is useful for both event-recall and perceptual network questions, but fewer linguistic acrobatics are needed when referencing event-recall data.

1.1.1. Cognition, recall, and frequency biases

The factors that shape and constrain an individual’s memory and perceptions of social relations have been a central focus in the design of survey and interview methods across the social sciences. Research on social cognition has highlighted a diversity of ways in which individuals process social information, contrasting the differential effects of variation in perception and response accuracy (Bandura 2002; Fiske and Taylor 1991). Individual differences (Marin 2004), context variation (Casciaro 1998), and question content and framing (Kogovšek and Ferligoj 2005) can all cause individuals to respond in inconsistent ways about their social relationships.

Classical studies in social network analysis propose that humans use schemata and cognitive social structures to recall their relationships (Carley 1986; Fiske 1995;

Freeman 1992; Krackhardt 1987; Lewin 1951). These mental representations of social relationships produce a non-random pattern in free-response network data (see Brands 2013; Smith et al. 2020, for reviews). For example, there is evidence that individuals’ perceptions of their social networks are shaped by their own relative position within a social hierarchy (Walker 1976). Additionally, responses may be impacted by desires for structural balance—i.e., if individual i is friends with individual j , but is hostile with individual k , then i may be inclined to perceive individual j as also being hostile with individual k (Crockett 1982; De Soto 1960; Heider 1982). Social identities, or affiliations with specific social groups, have also been shown to pattern responses (Freeman 1992)—e.g., by causing individuals to structure their responses based on a given attribute, such as occupation or ethnicity.

Self-report social network data may also be shaped by other cognitive processes. For example, individuals are generally prone to forget a non-trivial proportion of their ties. There is also individual-level variation in the rate of forgetting events or relations as a function of the amount of time that has elapsed between a given interaction and the point of measurement, even across relatively short timescales (Brewer 2000). Alongside this, an individual’s memory and perception of a social relationship is likely structured by the overall frequency of interactions that they have with a given individual (Hammer 1984). In an effort to fully collect network data, researchers sometimes prompt respondents, asking if there is anyone else that they have forgotten to mention. This method may increase the number of ties within a reported network (i.e., network density), but the strength or nature of connections to individuals mentioned after prompting may be different than those mentioned without prompting (Marin 2004).

1.1.2. Attribute-related biases

A number of individual attributes and characteristics are associated with variation in perceptions of network structure (Casciaro 1998; Flynn et al. 2006). Such perceptual differences can cause self-report network data to diverge from a true underlying network of interest. Here, we focus on the biases in nomination caused by attributes related to social status, because these biases are most salient in the extant literature (e.g., Redhead and Power 2021). There are, however, likely to be a multitude of other individual-level attributes that may create analogous biases across a broad range of research settings.

Specific attributes may influence an individual’s ability to accurately recall their own social ties. Individuals in positions of high power or status (i.e., group members who are conferred disproportionate social influence, authority, and esteem relative to other group members: Redhead et al. 2018) often have less precise perceptions of social relationships than those lower in power or status (Freeman 1992; Keltner et al. 2003; Press et al. 1969). High status individuals may likely operate within dense local network neighbourhoods, and interact or have relationships with a large number of individuals. By having such a high number of ties, these individuals

may be prone to forget—and thus not report—a significant proportion of their true ties relative to individuals of lower status. For example, across a series of experiments, [Simpson et al. \(2011\)](#) found that undergraduate students who had been ‘primed’ with high power were less able to recall the true state of a hypothetical vignette network than those with low power. Similar findings have been shown in real-world groups and organizations, where individuals who occupy positions near the top of a social hierarchy ([Shakya et al. 2017](#)) or are more central within a network ([Grippa and Gloor 2009](#)) have more asymmetrical incoming and outgoing ties, and, more generally, are less accurate in their reports of social ties.

Alongside this, individual-level attributes may bias the reports of others about their social ties. That is, individuals may be more likely to be nominated as friends by others within the group, even if they are not really connected in the ‘true’ network. Empirical evidence suggests that individuals are also more likely to report interactions and relationships with community members high in social status and power than others with lower status or power ([Marineau et al. 2018](#)). Likewise, at a dyadic level, individuals may be more inclined to recall or report connections based on shared attributes (i.e., homophily), such as gender or ethnicity, which may in turn inflate the levels of homophily observed in self-reported network data ([Flynn et al. 2010](#)).

1.1.3. The effects of question order

The order in which a researcher asks social network questions within a given survey or interview can have striking effects on reported outcomes. Researchers typically ask multiple name generator questions during a given interview in order to measure social relationships across different network layers. There are two issues to consider. First, eliciting data across multiple network layers is time consuming and often induces participant fatigue, with participant responsiveness declining with each question asked ([Pustejovsky and Spillane 2009](#); [Tourangeau and Rasinski 1988](#); [Yousefi-Nooraie et al. 2019](#)). Second, the names elicited in earlier questions may prime responses in subsequent questions.

Regarding the first issue, participant fatigue can inflate the level of ‘false negative’ non-nominations in a reported network and may result in biased estimates of network-relevant statistics. A well-known example of this effect is the controversial ‘shrinking’ of friendship networks among United States citizens between 1985 and 2004, where there was a downward secular trend in the average number of reported confidants within a nationally-representative sample of participants in the General Social Survey ([McPherson et al. 2006](#)). Re-examinations of the data have since revealed that this finding was likely an artifact produced by participant fatigue due to survey design and interviewer effects ([Fischer 2009](#); [Paik and Sanchagrin 2013](#)).

Concerning the second issue, evidence suggests that there are contamination effects that result from earlier questions influencing nomination decisions in subsequent questions ([Tourangeau and Rasinski 1988](#); [Schwarz 1999](#)). These contamination effects can pro-

duce response patterns where participants erroneously duplicate names from a previous question (see [Ready and Power 2021](#)). Antithetically, participants can sometimes modify their nominations to reduce redundancy in naming (i.e., purposely failing to mention names from a previous question; [Pustejovsky and Spillane 2009](#)).

Question order effects may likely be compounded by research designs that double sample social network data. Double-sampled designs ask all participants to report both directions of a given tie, providing information that can help verify the accuracy of reported relationships ([Nolin 2008](#)). For example, within such a design, individuals may be asked who they have given advice to, and who they have received advice from. It seems reasonable to assume that there should be a high correlation between an individual’s report of who they give advice to and corresponding reports of who has received advice from them. However, this correlation may be considered somewhat spurious if, for instance, these double-sampled questions are asked in a fixed order, and more-so if they are asked in direct succession ([Ready and Power 2021](#)). Participants may be primed to mentally ‘copy and paste’ their responses, or even respond by directly stating “the same people as before” instead of listing individuals.

1.2. Incorporating measurement error

In recent years, applied mathematicians have developed a multitude of statistical frameworks for analyzing social network data. In the social and behavioural sciences, researchers have advanced methods that parse variation in node-level degree distributions, dyadic-level characteristics, and higher-order structures ([Van Duijn and Vermunt 2006](#)). These approaches have been built upon in psychology and psychometrics, with researchers incorporating item response theory ([De Ayala 2013](#)) into dyadic models with multiple indicators (e.g., the social relations model; [Gin et al. 2020](#)). However, these methods generally assume that the observed or reported data represent the ‘true’ structure of relationships, in spite of the fact that such measurements may be biased or internally inconsistent ([Ball and Newman 2013](#)). Such an approach is thus likely to generate inaccurate estimates of network structure, and cause misleading conclusions to be drawn about key phenomena ([Kossinets 2006](#); [Newman 2018](#)).

To resolve these issues, social scientists can draw upon recent work in statistical physics and complexity science, where several approaches have been developed to estimate underlying network structure from unreliable and noisy data ([Butts 2003](#); [Clauset et al. 2008](#); [Guimerà and Sales-Pardo 2009](#); [Peixoto 2018](#); [Young et al. 2020](#)). These approaches generally involve a hierarchical Bayesian framework, consisting of a joint model of the data generating process for a latent, unobserved network, and a set of parameters that measure and adjust for response reliability. The model that we introduce below builds upon the architecture introduced in these past approaches, but integrates a more realistic generative model of human social network data, and adds parameters for response biases that have not previously been accounted for.

2. A latent network modeling framework

Here, we introduce a latent network approach to estimate a set of parameters governing both true network connections and the reporting biases of respondents. In an attempt to balance clarity and insight, we keep the model of “true” network connections as simple as possible, while ensuring that it remains empirically realistic. The generative model that we use is the union of a Stochastic Block Model (SBM; [Peixoto 2019](#)) with discrete, non-overlapping blocks to generate gross-level sub-structure, and a Social Relations Model (SRM; [Kenny and La Voie 1984](#)) to characterize individual-level variation in in-degree and out-degree (i.e., the frequency of making and receiving ties), and capture reciprocal dyadic relationships. Any implementation of our approach would need to consider if these generative model parameters are sufficient for that given application, or whether further generalizations of our model are necessary.

2.1. A model of ‘true’ network connections

2.1.1. Block or community structure

Let the adjacency matrix, y , denote the true (unobserved) network; the elements $y_{[i,j]} \in \{0, 1\}$ reflect the presence or absence of directed ties (e.g., resource transfers) between pairs of individuals. The community of N individuals is assumed to be divided into K non-overlapping blocks, or sub-communities. We further assume that each individual belongs to only one block, that the block of type k has sample size N_k , and that the block of individual i is returned by the function $b(i)$. We let the elements of the square matrix, B , denote the log-odds of a tie between individuals in various blocks. For example, suppose there are $K = 3$ blocks. The B matrix is then:

$$B = \begin{bmatrix} \beta_{1 \rightarrow 1} & \beta_{1 \rightarrow 2} & \beta_{1 \rightarrow 3} \\ \beta_{2 \rightarrow 1} & \beta_{2 \rightarrow 2} & \beta_{2 \rightarrow 3} \\ \beta_{3 \rightarrow 1} & \beta_{3 \rightarrow 2} & \beta_{3 \rightarrow 3} \end{bmatrix} \quad (1)$$

where the parameters on the diagonal control the probabilities of within-group ties, and the parameters on the off-diagonals control the probabilities of between-group ties. For instance, if ties tend to flow in one direction—e.g., from individuals in group 1 to individuals in group 2—then the parameters in one triangle may be larger than in the opposite triangle—e.g., $\beta_{1 \rightarrow 2} > \beta_{2 \rightarrow 1}$. If desired, the parameters in B can be made a function of other parameters and covariate data, or be constrained in various ways (e.g., forcing $B = B^T$ if transfers between blocks are equally likely in both directions).

Assuming only a block-induced sub-structure, the generative model for $y_{[i,j]}$ may then be written as:

$$y_{[i,j]} \sim \text{Bernoulli}(\text{Logistic}(B_{[b(i),b(j)]})) \quad (2)$$

where the probability of a tie from individual i in block $b(i)$ to individual j in block $b(j)$ is controlled by the corresponding entry in the square matrix, $B[b(i), b(j)]$.

In typical cases, the diagonal elements of B , which control the frequency of ties within a block, will have higher prior weight than the off-diagonal elements. Priors

should also depend on sample size, N , so that the resultant network density approximates empirical networks. Basic priors could be:

$$\beta_{k \rightarrow k} \sim \text{Normal}(\text{Logit}(\frac{0.1}{\sqrt{N_k}}), 1.5) \quad (3)$$

$$\beta_{k \rightarrow \tilde{k}} \sim \text{Normal}(\text{Logit}(\frac{0.01}{0.5\sqrt{N_k} + 0.5\sqrt{N_{\tilde{k}}}}), 1.5) \quad (4)$$

where $k \rightarrow k$ indicates a diagonal element and $k \rightarrow \tilde{k}$ indicates an off-diagonal element.

Empirical accounts of real-world social networks, however, indicate that there is substantially more structure—e.g., individual-level variation in in- and out-degree, and dyad-level effects—that should be incorporated to produce realistic social networks ([Carrington et al. 2005](#); [Doreian and Conti 2012](#); [Snijders et al. 2006](#)). As such, we integrate the insights from the SRM approach ([Kenny and La Voie 1984](#)) and its extensions (e.g., [Koster and Leckie 2014](#); [Pisor et al. 2020](#)) into Eq. 2.

2.1.2. The social relations model

In the SRM, individuals have random effects that control their propensity to send outgoing, and receive incoming, ties. Similarly, specific dyads may be more or less likely than chance to share reciprocal ties. Thus, in the SRM, dyadic random effects are also included to capture this tendency. By integrating the SRM and the SBM, Eq. 2 may be replaced by Eq. 5:

$$y_{[i,j]} \sim \text{Bernoulli}(\text{Logistic}(\phi_{[i,j]})) \quad (5)$$

where:

$$\phi_{[i,j]} = B_{[b(i),b(j)]} + \lambda_{[i]} + \pi_{[j]} + \delta_{[i,j]} + \dots \quad (6)$$

Here, B is the SBM intercept matrix, λ is a vector of individual-specific sender/nominator effects governing out-degree, π is a vector of individual-specific receiver/target effects governing in-degree, δ is a matrix of dyadic effects governing dyadic reciprocity, and the ellipse signifies any linear model of coefficients and focal, recipient, or dyadic covariates. For example, if S is a person-specific measure, like social status, and Q is a dyad-specific measure, like a network of kinship ties, then the ellipse may be replaced with: $\kappa_{[1]}S_{[i]} + \kappa_{[2]}S_{[j]} + \kappa_{[3]}Q_{[i,j]}$, to give the effects of social status on both sending and receiving transfers and the effects of kinship on transfers in either direction.

To complete the SRM definition, we model the sender and receiver effects jointly using a multivariate normal distribution. This allows for generalized correlations at the individual level—e.g., if individuals who tend to support others are also more likely to be supported by others:

$$\begin{pmatrix} \lambda_{[i]} \\ \pi_{[i]} \end{pmatrix} \sim \text{MV Normal}\left(\begin{pmatrix} 0 \\ 0 \end{pmatrix}, \begin{pmatrix} \sigma_\lambda^2 & \sigma_\lambda \sigma_\pi \rho \\ \sigma_\lambda \sigma_\pi \rho & \sigma_\pi^2 \end{pmatrix}\right) \quad (7)$$

For computational reasons ([Stan Development Team 2021b](#); [Lewandowski et al. 2009](#)), it is better to implement Eq. 7 by defining:

$$\begin{pmatrix} \lambda_{[i]} \\ \pi_{[i]} \end{pmatrix} = \begin{pmatrix} \sigma_\lambda \\ \sigma_\pi \end{pmatrix} \circ \left(L * \begin{pmatrix} \hat{\lambda}_{[i]} \\ \hat{\pi}_{[i]} \end{pmatrix}\right) \quad (8)$$

where L is a Cholesky factor from the decomposition of the 2×2 correlation matrix with ρ on the off-diagonal, and $\lambda_{[i]} \sim \text{Normal}(0, 1)$ and $\hat{\pi}_{[i]} \sim \text{Normal}(0, 1)$ are unit-normal random effects. Weak priors may then be independently specified on the variance and correlation terms (Lewandowski et al. 2009):

$$\sigma_\lambda \sim \text{Exponential}(1.5) \quad (9)$$

$$\sigma_\pi \sim \text{Exponential}(1.5) \quad (10)$$

$$L \sim \text{LKJ Cholesky}(2.0) \quad (11)$$

We also use the above approach to define the dyad-level random effects:

$$\begin{pmatrix} \delta_{[i,j]} \\ \delta_{[j,i]} \end{pmatrix} = \begin{pmatrix} \sigma_\delta \\ \sigma_\delta \end{pmatrix} \circ \left(L_\delta * \begin{pmatrix} \hat{\delta}_{[i,j]} \\ \hat{\delta}_{[j,i]} \end{pmatrix} \right) \quad (12)$$

where $\hat{\delta}_{[i,j]} \sim \text{Normal}(0, 1)$ have unit-normal priors, and the variance and correlation terms have weak priors:

$$\sigma_\delta \sim \text{Exponential}(1.5) \quad (13)$$

$$L_\delta \sim \text{LKJ Cholesky}(2.0) \quad (14)$$

Under this model, ρ provides an indication of generalized reciprocity—i.e., whether those who give more (to anyone) also receive more (from anyone)—and ρ_δ provides a measure of dyadic reciprocity—i.e., whether the probability of focal i giving to alter j , increases with the probability that focal j gives to alter i .

2.1.3. The observation model

Following previous latent network modelling frameworks (e.g., Butts 2003; Young et al. 2020), the matrix of true ties, $y_{[i,j]}$, is assumed to generate a tensor of observable variables $x_{[i,j,q,t]}$, where q is an index for the specific type of question or observable variable and t is a time point. These x variables can have any arbitrary distribution and relation to the ties represented by $y_{[i,j]}$. Here, we consider two types of ties:

1. Survey data for which each i nominates a set of alters indexed by j . For example, in question/layer $q = 1$ at time point $t = T$, the variable $x_{[i,j,q=1,t=T]}$ might represent a binary nomination by individual i of a resource transfer from i to j , and, at the same point in time, in question/layer $q = 2$, the variable $x_{[i,j,q=2,t=T]}$ might represent a binary nomination by individual i of a resource transfer from j to i . Such self-report data, however, may be unreliable. Individuals i and j vary in the extent to which they recall true transfers and/or falsely report non-existent transfers. Likewise, even for a single focal respondent, the reliability may vary by the direction of nomination, $i \rightarrow j$ versus $j \rightarrow i$. Most critically, i 's report of j 's aid may disagree with j 's report of the same directed relationship.
2. Observable data on directed exchanges or interactions between individuals i and j , such as transfers of gifts or resources, which flow/occur at average rate $\gamma_{[q,t]}$. These data—possibly documented by focal follows or by GPS tagging—may be more reliable than self-reports, but constraints (e.g., resource constraints) may also make it impossible for focal

respondents to actually share with all social ties in a given time period, and data collection may be much more onerous.

For simplicity, we assume binary outcome data across all network layers, and model $x_{[i,j,q,t]}$ using a Bernoulli distribution:

$$x_{[i,j,q,t]} \sim \text{Bernoulli}(\psi_{[i,j,q]}) \quad (15)$$

$$\psi_{[i,j,q]} = \alpha_{[i,q]}(1 - y_{[i,j]}) + \beta_{[i,q]}y_{[i,j]} \quad (16)$$

Here, $\alpha_{[i,q]}$ gives the probability of individual i reporting a tie in network layer q when no such tie exists in the true network (i.e., when $y_{[i,j]} = 0$). We refer to this loosely as a ‘false positive rate’. Likewise, $\beta_{[i,q]}$ gives the probability of individual i reporting a tie in network layer q when such a tie does exist in the true network (i.e., when $y_{[i,j]} = 1$). We refer to this loosely as a ‘true-tie recall rate’. Because α and β vary by person and layer, our approach allows some individuals and network layers to provide more accurate information about the underlying true network than others.

The $\alpha_{[i,q]}$ and $\beta_{[i,q]}$ parameters can also be constructed as functions of covariates specific to individuals, blocks, or network types. For example, if the social status, $S_{[i]}$, of each individual influences reporting accuracy, then we might define the accuracy parameters as:

$$\text{logit}(\alpha_{[i,q]}) = \text{logit}(\mu_{\alpha_{[q]}}) + \sigma_{\alpha_{[q]}}\hat{\alpha}_{[i,q]} + \eta_{\alpha_{[q]}}S_{[i]} + \dots \quad (17)$$

$$\text{logit}(\beta_{[i,q]}) = \text{logit}(\mu_{\beta_{[q]}}) + \sigma_{\beta_{[q]}}\hat{\beta}_{[i,q]} + \eta_{\beta_{[q]}}S_{[i]} + \dots \quad (18)$$

where $\mu_{\alpha_{[q]}}$ and $\mu_{\beta_{[q]}}$ give the average false positive and true-tie recall rates when the contribution of individual-level effects is null, and $\sigma_{\alpha_{[q]}}$ and $\sigma_{\beta_{[q]}}$ scale the variation in false positive and true tie-recall rates resulting from individual-level unit-normal random effects. Alongside this, $\eta_{\alpha_{[q]}}$ and $\eta_{\beta_{[q]}}$ control the effects of some covariate—e.g., $S_{[i]}$ or social status—on false positive and true-tie recall rates in network layer q .

Higher order priors for this model might be:

$$\mu_{\alpha_{[q]}} \sim \text{Beta}(1.0, 12.0) \quad (19)$$

$$\sigma_{\alpha_{[q]}} \sim \text{Exponential}(1.0) \quad (20)$$

$$\eta_{\alpha_{[q]}} \sim \text{Normal}(0.0, 1.0) \quad (21)$$

$$\mu_{\beta_{[q]}} \sim \text{Beta}(12.0, 1.0) \quad (22)$$

$$\sigma_{\beta_{[q]}} \sim \text{Exponential}(1.0) \quad (23)$$

$$\eta_{\beta_{[q]}} \sim \text{Normal}(0.0, 1.0) \quad (24)$$

2.2. Question order effects

In addition to overall false positive and true-tie recall rates, practitioners often report that when double sampled social networks are collected using the name-generator method, some respondents have a strong tendency to report ties that are reciprocal, even when the other half of the dyad does not report the same bidirectional connection (Ready and Power 2021). In other words, for many individuals, the names nominated in

question $q = 1$ (e.g., *who did you provide financial aid to in the last 30 days?*) may be repeated in question $q = 2$ (e.g., *who gave you financial aid over the last 30 days?*) due to a kind of question-order effect. This effect may result from a cognitive bias to view social relationships as reciprocal, or may arise simply from priming. Whatever the cause, it can obscure many important network characteristics, especially reciprocity (see, [Ready and Power 2021](#), for an in-depth discussion).

To measure and account for this potential bias, we can modify Eq. 15 slightly to represent responses in the second name-generator question as a mixture of processes. In the first of two cases, if $x_{[i,j,q=1,t]} = 0$, then we model $x_{[i,j,q=2,t]}$ exactly as in Eq. 15. In other words, if a name was not mentioned in the first name-generator question, then it cannot be copied blindly into the second response. However, in the second case, if $x_{[i,j,q=1,t]} = 1$, and a name was mentioned in the first name-generator question, then, with probability $\theta_{[i]}$, we assume that the same name is erroneously copied into the second name-generator response, and with probability $\hat{\theta}_{[i]} = 1 - \theta_{[i]}$, we assume that the response is not erroneously copied but arises under the probability model given by $\psi_{[i,j,q=2,t]}$:

$$x_{[i,j,q=2,t]} \sim \begin{cases} \text{Bernoulli}(\psi_{[i,j,q=2,t]}), & \text{if } x_{[i,j,q=1,t]} = 0 \\ \theta_{[i]} \text{Bernoulli}(1.0) + \\ \hat{\theta}_{[i]} \text{Bernoulli}(\psi_{[i,j,q=2,t]}), & \text{if } x_{[i,j,q=1,t]} = 1 \end{cases} \quad (25)$$

As before, we assume that the extent of this question-order response bias is an individual-level characteristic:

$$\text{logit}(\theta_{[i]}) = \text{logit}(\mu_\theta) + \sigma_\theta \hat{\theta}_{[i]} + \eta_\theta S_{[i]} + \dots \quad (26)$$

where μ_θ controls the average rate of response duplication when the contribution of individual-level effects is null, σ_θ scales the variation in individual-level unit-normal random effects, and η_θ controls the effects of some covariate—e.g., $S_{[i]}$ or social status—on the response duplication rate.

Higher order priors may be specified as:

$$\mu_\theta \sim \text{Beta}(3.0, 12.0) \quad (27)$$

$$\sigma_\theta \sim \text{Exponential}(1.0) \quad (28)$$

$$\eta_\theta \sim \text{Normal}(0.0, 1.0) \quad (29)$$

2.3. Recency and frequency biases

Another source of methodological noise that researchers may wish to examine and account for relates to the recency of interactions, as well as their number. For example, imagine that network layers $q = 1$ and $q = 2$ are self-report nominations as described above, and that network layer $q = 3$ is an adjacency matrix derived from observations of transfers of material goods (such as food or money) from i to j , over a total of $t \in (1, \dots, T)$ time-points. Then, at time-point T , individual i may be more likely to recall a true tie to an alter j in the self-report question layer $q = 1$ if there were many (or quite recent) observed resource transfers from i to j over the period from $t = 1$ to $t = T$. Specifically, let us assume that the variables $x_{[i,j,q=3,t]}$ have been realized for $t \in (1, \dots, T)$, as:

$$x_{[i,j,q,t]} \sim \text{Bernoulli}(\psi_{[i,j,q,t]}) \quad (30)$$

$$\psi_{[i,j,q,t]} = \alpha_{[i,q]}(1 - y_{[i,j]}) + \beta_{[i,q]}y_{[i,j]}\gamma_{[q,t]} \quad (31)$$

and we are modelling the self-report data, $x_{[i,j,q=1,t=T]}$. We assume that the log-odds of a reported transfer in $x_{[i,j,q=1,t=T]}$ can be affected by the real observed transfers in $x_{[i,j,q=3,t]}$. We first assume that the increment in log-odds, ω , due to a single observed transfer follows an exponential decay function with time:

$$\omega_{[t]} = \zeta \exp(-\xi(T - t)) \quad (32)$$

where ζ gives the increment in log-odds to the true-tie recall rate when $t = T$ (i.e., when there was a very recent observed transfer), and ξ controls the rate at which the increment in log-odds declines with time.

Weak, positive constrained, priors can be specified on ζ and ξ :

$$\zeta \sim \text{Gamma}(3, 1) \quad (33)$$

$$\xi \sim \text{Exponential}(2) \quad (34)$$

Then, in the first layer, $q = 1$, for each directed dyad, we can sum over the relevant values in ω (i.e., the ones corresponding to $x_{[i,j,q=3,t]} = 1$):

$$\Omega_{[i,j,q=1]} = \sum_{t=1}^T \begin{cases} \omega_{[t]}, & \text{if } x_{[i,j,q=3,t]} = 1 \\ 0, & \text{otherwise} \end{cases} \quad (35)$$

Likewise, for the second network layer, $q = 2$, we start by defining a log-odds offset, but this time we need to use the values in layer $q = 3$ corresponding to transfers from j to i (i.e., where $x_{[j,i,q=3,t]} = 1$):

$$\Omega_{[i,j,q=2]} = \sum_{t=1}^T \begin{cases} \omega_{[t]}, & \text{if } x_{[j,i,q=3,t]} = 1 \\ 0, & \text{otherwise} \end{cases} \quad (36)$$

Finally, for $q \in \{1, 2\}$ and $t = T$, we integrate these offsets into the outcome models (e.g., Eqs. 15 or 25), by writing:

$$x_{[i,j,q,t]} \sim \text{Bernoulli}(\psi_{[i,j,q,t]}) \quad (37)$$

$$\psi_{[i,j,q,t]} = \alpha_{[i,q]}(1 - y_{[i,j]}) + \text{logistic}(\text{logit}(\beta_{[i,q]}) + \Omega_{[i,j,q]})y_{[i,j]} \quad (38)$$

2.4. Computing the posterior

To compute probabilities of observed variables, x , we marginalize over the unknown discrete variables, y . Specifically, the probability of $x_{[i,j,q,t]}$ is given by linking the model of the ‘true’ latent network and the observation model as follows:

$$\begin{aligned} \Pr(x_{[i,j,q,t]} | \Theta, \Phi) &= \Pr(y_{[i,j]} = 1 | \Phi) \Pr(x_{[i,j,q,t]} | \Theta, y_{[i,j]} = 1) \\ &\quad + \Pr(y_{[i,j]} = 0 | \Phi) \Pr(x_{[i,j,q,t]} | \Theta, y_{[i,j]} = 0) \end{aligned} \quad (39)$$

where Θ is the list of all relevant parameters in the observation model and Φ is the list of all relevant parameters in the model of the ‘true’ latent network.

Then, after sampling, we can recover the posterior probability of $y_{[i,j]} = 1$ given x using Bayes’ rule:

$$\Pr(y_{[i,j]} = 1 | x, \Theta, \Phi) = \frac{\Pr(y_{[i,j]} = 1 | \Phi) \Pr(x | \Theta, y_{[i,j]} = 1)}{\Pr(x | \Theta, \Phi)}$$

This can be calculated using MCMC samples of the parameters in Θ and Φ .

3. Model validation

To validate our approach, we first implement our model of network connections generatively to simulate a realistic ‘true’ network (see Figure 1a). Next, we simulate double-sampled self-reporting from this ‘true’ network, subject to individual-level reporting biases. Then, we apply standard methods to convert these double-sampled, self-report networks into single-layer networks (see Figures 1c and 1d). Finally, we fit our Bayesian latent network model to the reported data, and attempt to recover the ‘true’ underlying network (see Figure 1b). At this point, we evaluate the performance of each network reconstruction approach by assessing: 1) graph-level metrics (see sections 3.1 and 3.2), and 2) node-level metrics (see supplementary section 5.3). Each time, we contrast the inferred networks with the true network. Additionally, we evaluate the performance of our model by: 3) ensuring that generative parameter values can be recovered (see supplementary section 5.1), and 4) ensuring that node-level attributes in the inferred network are correlated with the node-level attributes in the generative model (see supplementary section 5.2).

3.1. Simulation experiments with the base model

Figure 2 displays the results of the base model. Data were simulated such that the only source of bias was in reporting rates (i.e., the recall rate of true ties, and the false positive rate). Each sub-figure illustrates a parameter sweep for a focal parameter, holding other parameters constant at moderate levels. For example, in frame 5a, we first simulate data under a given average false positive rate, μ_α (see Eq. 17). We then use our model to identify the latent network, and we calculate the relevant network-level metrics. Following this, we increment μ_α , holding all other simulation parameters constant, and repeat the process. Each sub-figure visualizes a sweep across a range of plausible parameter values. In frame 5b, for instance, we fix μ_α , and instead vary σ_α , which controls the amount of between-individual variation in false positive rate.

We compare the results of our model to the results obtained by taking the intersection (i.e., coding a tie as present only if both individuals within a dyad state there to be such a tie) or the union (i.e., coding a tie as present if at least one individual within a dyad states there to be such a tie) of nominations, as these are two of the most common techniques for processing double-sampled data in the social sciences (Ready and Power 2021). While there are a handful of cases where either the intersection or the union can accurately capture a specific network metric, one would need to know the reporting bias rates *a priori* in order to select the correct metric to use. Moreover, there are many locations in the parameter space where neither method recovers accurate network metrics. *In contrast, our latent network approach dynamically adapts, and almost universally recovers the true network metrics—even in cases where the reporting biases are set unrealistically high.*

3.2. Simulation experiments with the question-order model

Figure 3 shows the results of the base model with additional control for question-order effects. In this experiment, data were simulated such that biases in reporting rates (i.e., the recall rate of true ties, and the false positive rate) were still present, but a further name-duplication bias across questions was added (i.e., respondents were more likely to claim that they received food from a given alter in question 2, if they reported giving food to that same alter in question 1). As before, the latent network approach almost universally captures the true network metrics, even in cases where the reporting biases are set unrealistically high. In contrast, neither standard method allows for accurate inference about the true network when question-order effects are present. As discussed by Ready and Power (2021), the use of the union rule to merge double-sampled networks is especially problematic when question-order effects are present.

3.3. Simulation experiments of the “status effect”

The true in- or out-degree distribution for a node may often be linked to some key covariate, like social status, that may also affect reporting bias parameters. Any such covariate that influences true network structure, and affects the strength of reporting biases in the observation sub-model, has the potential to severely confound inference. In such cases, double-sampled self-report data may be insufficient to resolve the effects of covariates on either sub-component.

As an example, we can generate data in which some covariate, S —e.g., social status—has a strong effect on outgoing transfers: empirically, a high-status individual might provide help to a large number of people in a given community. However, by virtue of having such a large number of out-going ties to recall, these same high-status individuals might be less likely to recall all of their true outgoing ties during an interview. As such, we can simulate a response bias sub-model in which S has a strong negative effect on the recall rate of true outgoing ties. Finally, we can fit our model on the simulated data, and attempt to recover the effects of social status on both tie existence and recall rate.

Figure 4 plots the results of this analysis, using a range of parameter sweeps. As before, the data that define the true underlying network are held constant, but the parameters governing response accuracy vary. In all parameter sweeps, the effects of status on out-degree (positive) and on recall of true ties (negative) are held fixed. Across all parameter sweeps, we either underestimate or entirely fail to detect the true positive effect of status on out-degree. Likewise, we generally fail to detect the effect of status on recall of true ties. In sum, the status bias in the self-report network can substantially mask the effect of status on out-degree. In order to accurately detect the effects of status, we need to integrate information from other network layers, especially ones that are less susceptible to reporting biases.

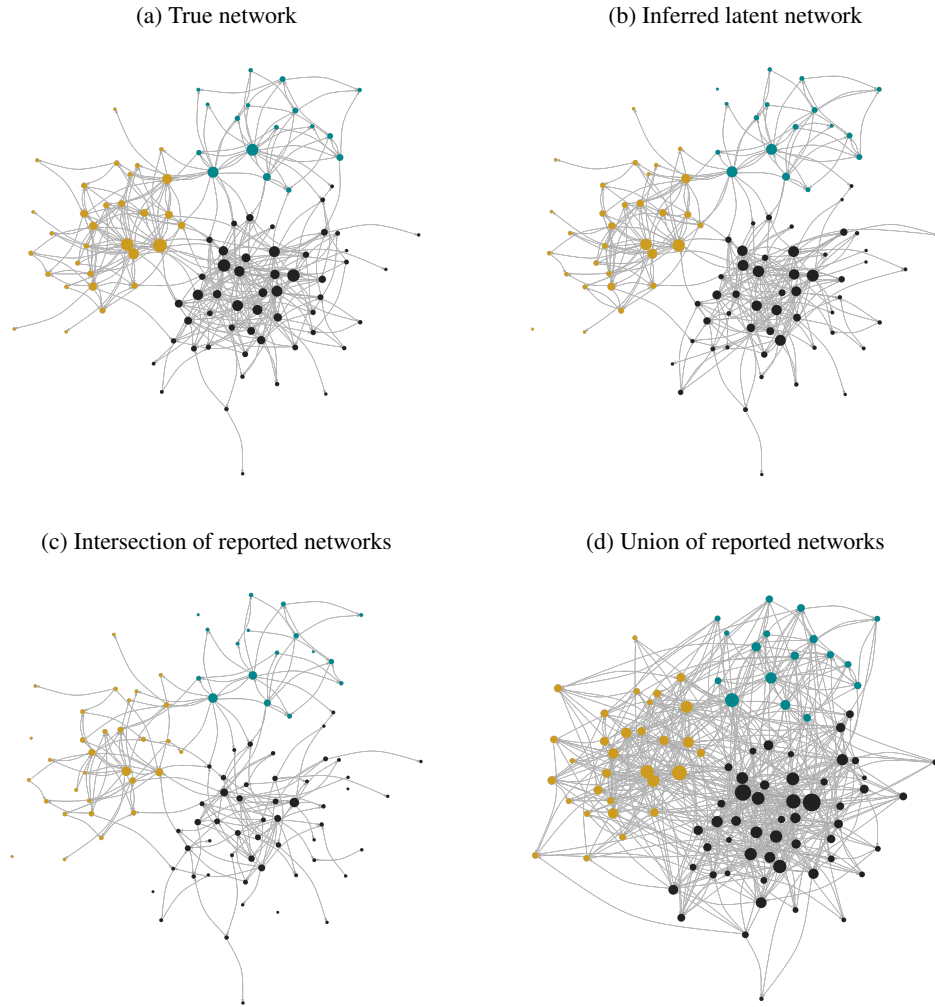


Fig. 1: In frame 1a, we plot a true network of transfers in a simulated group-structured population. Then, we simulate agents reporting on their transfer relations using a double-sampled design (i.e., reporting both who they gave to, and who they received from). In this case, we use a 1% false positive rate per node, and a 70% recall rate of true ties. In frame 1b, we plot the posterior median of our latent network. We also apply standard methods of reducing the double-sampled network to a single-layer network: in frame 1c we take the intersection (i.e., only classifying a dyadic tie as existing if both nodes reported the tie), and in frame 1d we take the union (i.e., classifying a dyadic tie as existing if either node reported the tie). We note, visually, that the latent network approach yields a better representation of the network—at least in terms of density—than the other methods. More formal tests are presented in the supplementary materials.

3.4. Simulation experiment of recency and frequency

Here, we replicate the data simulation protocol from section 3.3 exactly, but use the more complex model variant introduced in section 2.3 to analyze the data. That is, rather than analyzing pure self-report data, we now analyze 12 time points of sparse transfer data (in which transfers were accurately documented), in addition to the standard self report data. These sparse networks of “ground-truth” data might come from field observations (e.g., “scan sampling” or “spot-checks”, [Borgerhoff Mulder et al. 1985](#)), diary methods ([Paolisso and Hames 2010](#)), experimental games ([Ross and Redhead 2021](#)), video recordings ([DeTroy et al. 2021](#)), GPS tracking ([Davis et al. 2018](#)), proximity detection (e.g., through

cell phone data, [Urban 2021](#)), or a variety of other methods. Even with sparse, incomplete “ground-truth” data, network reconstruction can be improved, and covariate effects on network structure and reporting parameters disambiguated.

In Figure 5, we show that by integrating even sparse ground-truth data, we can recover directionally accurate estimates of the effect of status on both out-degree and recall of true ties. In the supplementary materials, we show that we recover the average effects of recency on recall rate and the flow rates of transfers along network ties. These rates were set to low levels (~ 0.12) per time step.

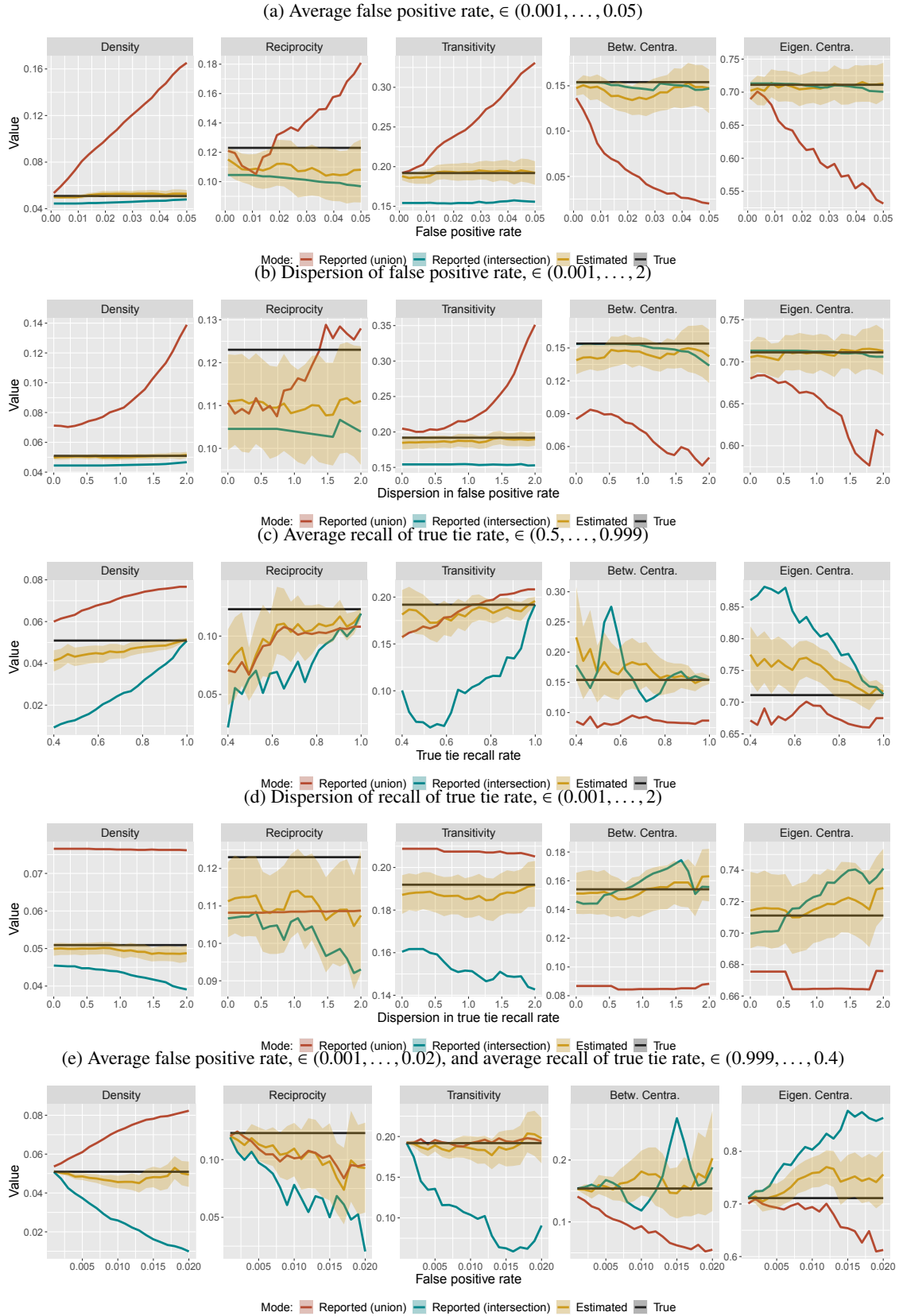


Fig. 2: Base model. Each frame plots an analysis of network-level properties. The true network and all associated parameters are held fixed, except for the parameters displayed in the column labels, which range over the indicated support. The levels of each outcome in the true network appear as horizontal black lines. The orange regions illustrate the posterior distributions of each outcome from the latent network model. The red and blue lines represent the outcomes resulting from application of either the union or intersection operator, respectively.

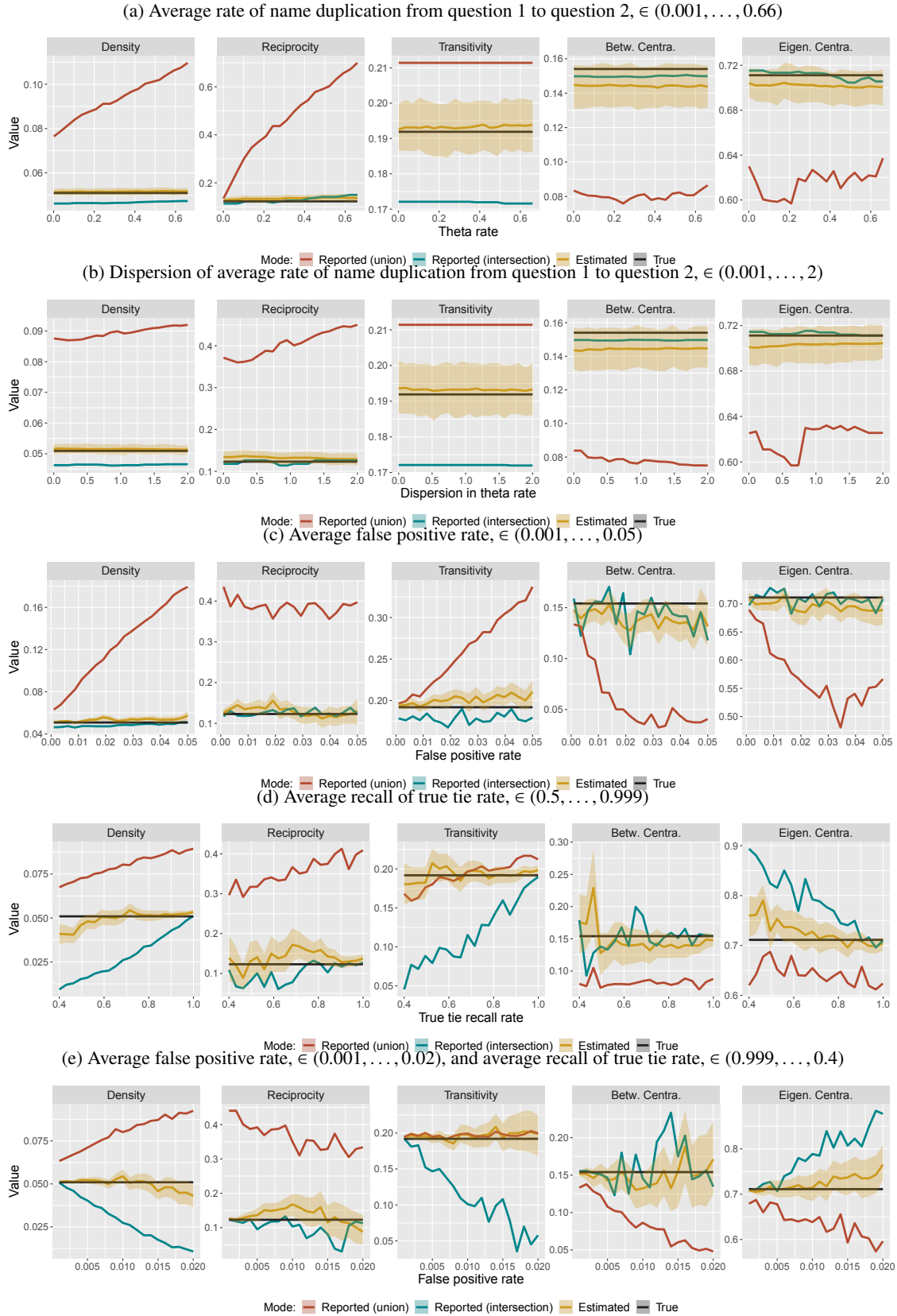


Fig. 3: Question order model. Rate of name duplication from question 1 to question 2: $\mu_\theta = 0.2$, unless otherwise noted. Each frame plots an analysis of network-level properties. The true network and all associated parameters are held fixed, except for the parameters displayed in the column labels, which range over the indicated support. The levels of each outcome in the true network appear as horizontal black lines. The orange regions illustrate the posterior distributions of each outcome from the latent network model. The red and blue lines represent the outcomes resulting from application of either the union or intersection operator, respectively.

4. The STRAND package

To make implementation of our models simple for end-users, we have released a fully documented R package, called STRAND, which can be used to both simulate realistic network data, and to analyze empirical network data using Bayesian methods. In this section, we introduce the STRAND package, and provide a step-by-step tutorial on its functionality.

4.1. Installation

Much of the functionality of STRAND is made possible by Stan (Stan Development Team 2021b) and CmdStanR (Stan Development Team 2021a). Users must install these programs prior to installing STRAND. Installation and loading of STRAND is then simple: just run three lines of code from R:

```
library(devtools)
install_github("ctross/STRAND")
library(STRAND)
```

4.2. Data simulation

After loading the library, STRAND can be used to run forward data simulations. We provide four key functions: `simulate_sbm_network`, `simulate_srm_network`, `simulate_sbm_plus_srm_network`, and finally the larger function `simulate_selfreport_network`. The first three functions are well described in the literature, referring to a stochastic block model, a social relations model, and the union of both models, respectively. The final function implements the full generative model introduced in section 2. The package documentation provides detailed descriptions of each function's arguments and outputs. We provide an example function call below:

```
s = simulate_selfreport_network(
  N_id = 100,
  N_groups = 3,
  group_probs = c(0.2, 0.5,
    0.3),
  in_block = 0.02,
  out_block = 0.002,
  false_positive_rate =
    c(0.01, 0.01, 0.01),
  recall_of_true_ties =
    c(0.85, 0.85, 0.99)
)
```

Many additional arguments—including covariate data and parameters to control their effects—can be supplied, in order to influence the structure of both the ‘true’ underlying network, and the observed reporting networks.

4.3. Data analysis

After data are simulated, they can be analyzed with simple, `lm`-style function calls. First, we use the function

`make_strand_data` to organize the simulated data into a format that can be read by Stan and later functions:

```
d = make_strand_data(self_report =
  list(
    s$reporting_network[, , 1],
    s$reporting_network[, , 2]),
  group_ids = factor(s$group_id),
  individual_covariates = NULL,
  dyadic_covariates = NULL)
```

The `make_strand_data` function requires `self_report` data to be supplied as a list of matrices; in cases where data come from a single-sampled design, the list will be of length 1, and in cases where data come from a double-sampled design, the list will be of length 2. Users then have the option to include further arguments, such as `group_ids` (a factor vector of group or block identification codes), `individual_covariates` (a data-frame of any number of covariates that must be matched by participant ID to the `self_report` data), and `dyadic_covariates` (a labeled list of matrices that must be the same dimensions as, and matched to, the `self_report` data). The `make_strand_data` function provides initial error checks, and identifies which network models available in STRAND can be fit given the input types—e.g., latent network models cannot be fit if `self_report` is only supplied as a single layer network.

An intercept only latent network model can then be fit using:

```
f = fit_latent_network_model(
  data = d,
  fpr_regression = ~ 1,
  rtt_regression = ~ 1,
  theta_regression = ~ 1,
  focal_regression = ~ 1,
  target_regression = ~ 1,
  dyad_regression = ~ 1,
  mode = "mcmc")
```

As before, there are numerous arguments that can be specified, and all are described in detail in the package documentation. Importantly, STRAND supports model fitting using all of Stan's modes: Markov Chain Monte Carlo, “mcmc”, variational Bayesian inference, “vb”, and optimization, “optim”. While “mcmc” is gold-standard for accurate posterior estimates, the other methods are useful for exploratory model fits, or for large networks that are computationally unfeasible with MCMC methods.

Parameter summaries are then returned and printed using:

```
r = summarize_strand_results(f)
```

4.4. An empirical example: Food/money sharing in Colombia

To explore how more complex models can be specified, we draw upon an empirical data-set from rural Colombia (Pisor et al. 2020). These data consist of a double-

sampled, self-report food/money sharing network, paired to a variety of individual-level and dyadic covariate data. The network and dyadic data are structured as adjacency matrices, and the individual-level data as a dataframe. As before, we first prepare the data:

```
net = list(
  TransferOut=Outgoing_Transfers,
  TransferIn=Incoming_Transfers
)

ind = data.frame(
  Age=center(Age),
  GoodsValues=center(GoodsValues),
  Male=Male,
  CantWork=CantWork,
  GripStrength=center(Grip),
  NoFood=NoFood,
  Depressed=Sad
)

dyad = list(
  Relatedness=Relatedness,
  Friends=Friends
)

model_dat = make_strand_data(
  self_report=net,
  group_ids=group_ids,
  individual_covariates=ind,
  dyadic_covariates=dyad
)
```

And then the model equations can be specified using `lm`-style syntax:

```
f = fit_latent_network_model(
  data=model_dat,
  focal_regression = ~ Age +
    GoodsValues + Male + NoFood +
    CantWork + GripStrength +
    Depressed,
  target_regression = ~ Age +
    GoodsValues + Male + NoFood +
    CantWork + GripStrength +
    Depressed,
  dyad_regression = ~ Relatedness
    + Friends,
  fpr_regression = ~ Age +
    GoodsValues + Depressed,
  rtt_regression = ~ Age +
    GoodsValues + Depressed,
  theta_regression = ~ 1,
  mode="mcmc",
  return_latent_network = TRUE
)
```

In Figure 6, we plot the networks implied by the actual nominations of outgoing transfers (frame 6a) and incoming transfers (frame 6b), followed by the networks implied by taking the intersection (frame 6c) or union (frame 6d) of these reported networks. Finally, we plot the posterior median realization from our latent network

Table 1: Network-level measurements as inferred by three network reconstruction methods. Values in parentheses are 89% highest posterior density intervals.

Metric	Intersect.	Union	Latent model
Edge density	0.001	0.01	0.007 (0.006, 0.008)
Reciprocity	0	0.496	0.08 (0.058, 0.124)
Transitivity	0.214	0.099	0.109 (0.076, 0.146)
Betw. Cent.	0.002	0.136	0.115 (0.092, 0.136)
Eigen. Cent.	0.983	0.962	0.954 (0.945, 0.963)

model in frame 6e. We find that use of the intersection leads to a very sparse reconstructed network, not representative of real social ties. The union produces a more plausible network structure. The latent network model settles on a similar structure to the union in this data set, but allows us to better quantify our uncertainty in network properties conditional on reporting biases (see Table 1). By disentangling true reciprocity from name duplication bias (i.e., across layer duplication in nominations within respondents), the implausibly high reciprocity rate of 0.49 (resulting from taking the union of reporting networks) is reduced to about 0.08.

In Figure 7, we plot the effects of various predictor variables on network properties (frame 7a) and reporting biases (frame 7b). As in Pisor et al. (2020), which used a simple social relations model, we find that transfers are more likely to flow between kin and friends. Likewise, those with food insecurity or depression are less likely to report sending outgoing food/money transfers to other members of their community, whereas those with high grip-strength (a proxy for health and physical well-being) are more likely to report making such transfers. There is little to no effect of covariates on reporting biases. Finally, in frame 7c, we see some baseline measures of reporting accuracy. We see that false positive rates appear quite low, with little variance across respondents. The recall rate of true ties appears to be substantially less than unity and variable across respondents; such recall appears higher in the network layer in which respondents commented on who they gave to, relative to who they received from. Lastly, there is evidence of a name duplication bias across network layers, with an approximately 25% rate of name duplication inflation across network layers.

5. Discussion

Within the broad field of social network research, self-report measurement instruments have remained a primary tool used for the collection of network data. Self-report research designs are logistically-feasible, and provide important information on (individuals' perceptions of) social relationships and interactions (Freeman 1992). However, the interpretation of such data remains an enduring problem as empirical evidence has suggested that: *i*) self-reports do not reliably reflect ground-truth data on interactions (Killworth and Bernard 1976), *ii*) relationships that are typically considered undirected (such as friendships) are characterized by low levels of internal consistency (Ball and Newman 2013), *iii*) certain at-

tributes may bias nominations (Simpson et al. 2011), and *iv*) question order may impact patterns of response (Pustejovsky and Spillane 2009). In an aim to detect and correct some of these issues, researchers have begun to collect multiple reports on each directed tie within a network (i.e., they collect double-sampled networks; Adams and Moody 2007; Nolin 2008). While such approaches produce data that can conceivably detect disagreements between multiple reports of a single social relationship (Ready and Power 2021), it has remained non-trivial to statistically disambiguate such conflicting reports.

Recently proposed latent network frameworks have provided promising avenues for examining, and accounting for, simple forms of measurement error (e.g., individual-level variation in the frequency of reporting false ties and/or forgetting real ties) in self-report social network data (Butts 2003; Peixoto 2018; Young et al. 2020). We build upon this architecture to incorporate a more realistic data generating process that can flexibly include recency and frequency biases, attribute-related biases, and the effects of question order. To ensure ease-of-use, we have developed a fully documented R package, STRAND, and have provided a step-by-step tutorial that outlines the STRAND workflow.

The STRAND R package offers social and behavioral scientists easy-to-use software to apply our latent network modeling framework (as well as other, more common tools for social network analysis; e.g., stochastic block models and social relations models). STRAND provides full Bayesian inference, with all models being implemented in the Stan programming language (Carpenter et al. 2017), without requiring end-users to have a detailed technical understanding of the underlying code. The latent network modeling framework that we propose here easily incorporates double sampled self-report network data and observational network data aimed at capturing the same type of underlying relationships, alongside any type of individual-level or dyadic-level covariate data. Users simply need to call the relevant functions within STRAND and supply formulas similar to those used in the popular `lm` syntax. Users who do not wish to implement our latent network model—or who do not have the appropriate data—may fit a social relations model (Kenny and La Voie 1984), stochastic block model (e.g., Peixoto 2019), or a combination of the two, instead.

The extensive simulation experiments that we present highlight the reliability of our latent network models. In realistic conditions—i.e., where there are moderate levels of false positive responses, question order effects, and attribute-related biases—our latent network models accurately recover individual-level and network-level properties. Our latent network models even remain reasonably accurate in highly unrealistic conditions (e.g., when false positives occur in a 3 to 1 ratio with true positives). In comparison, our experiments highlight how standard approaches for network reconstruction produce highly inaccurate inferences, even when data contain minimal amounts of the biases discussed above. For example, the standard approach of taking the union of double-sampled network data leads to enormous inflation of density, reciprocity and transitivity, and deflation of common centralization metrics in the presence of even minimal amounts

of false positives. Antithetically, use of the intersection of double-sampled data leads to substantial underestimation of density, reciprocity and transitivity, and overestimation of centralization metrics in the presence of false negatives. Contrasting with both approaches, our latent network model permits accurate recovery of an underlying network, without researchers needing to know—or guess—any reporting bias rates *a priori*.

The results of our simulation experiments highlight the need for greater consideration when running and interpreting social network models that do not adjust for measurement error. Our results build upon previous research, which has spotlighted the sensitivity of individual-level and network-level properties to informant false reporting and to survey design (Butts 2003; Kossinets 2006). Furthermore, when an individual-level attribute increases the probability of both true ties, and tendencies to make false reports or fail to recall true ties, identification of the effects of that attribute on either outcome is challenging. This problem is especially salient in social scientific applications, as the core aim of most applied network studies is to assess the effects of individual attributes on network structure. Estimates of the effects of individual attributes on network structure from models that do not simultaneously account for reporting biases may therefore be unreliable. This observation likely holds regardless of whether the network data are single- or double-sampled, and caution is thus needed when interpreting such estimates. Our latent network models can examine and adjust for these attribute-related biases, and—as highlighted by our simulation experiments—accurately recover the true effects of such attributes on both outcome types.

The latent network modeling framework that we advance here provides a platform for many fruitful avenues of future research. Currently, our package currently only allows for observed, single-membership stochastic block models. In future work, we will extend this package to incorporate models that allow individuals to be members of multiple, overlapping blocks (or communities; Newman and Leicht 2007). Similarly, we will develop models that allow block membership to itself be a latent unobserved variable (and thus be estimated; Wasserman and Anderson 1987). Alongside this, our framework provides an architecture which can be extended in future work to encompass multi-level (Lazega and Snijders 2015), and multi-layer cases (Kivelä et al. 2014). In sum, we hope that the latent network framework—and associated STRAND R package—will inspire and facilitate empirical applications designed to investigate and/or adjust for the biases associated with self-report social network data.

Acknowledgements. We thank Eleanor Power for helpful comments on an earlier version of the manuscript.

Funding. D.R. was supported by ESRC Grant number ES/V006495/1. D.R., R.M. and C.T.R. were supported by the Department of Human Behaviour, Ecology and Culture at the Max Planck Institute for Evolutionary Anthropology.

References

- Adams, J. and Moody, J. (2007). To tell the truth: Measuring concordance in multiply reported network data. *Social Networks*, 29(1):44–58.
- Ball, B. and Newman, M. E. (2013). Friendship networks and social status. *Network Science*, 1(1):16–30.
- Bandura, A. (2002). Social cognitive theory in cultural context. *Applied Psychology*, 51(2):269–290.
- Banerjee, A., Chandrasekhar, A. G., Duflo, E., and Jackson, M. O. (2013). The diffusion of microfinance. *Science*, 341(6144).
- Bansal, S., Read, J., Pourbohloul, B., and Meyers, L. A. (2010). The dynamic nature of contact networks in infectious disease epidemiology. *Journal of biological dynamics*, 4(5):478–489.
- Bell, D. C., Belli-McQueen, B., and Haider, A. (2007). Partner naming and forgetting: recall of network members. *Social networks*, 29(2):279–299.
- Bernard, H. R., Killworth, P. D., and Sailer, L. (1979). Informant accuracy in social network data iv: A comparison of clique-level structure in behavioral and cognitive network data. *Social Networks*, 2(3):191–218.
- Borgatti, S. P., Mehra, A., Brass, D. J., and Labianca, G. (2009). Network analysis in the social sciences. *science*, 323(5916):892–895.
- Borgerhoff Mulder, M., Caro, T. M., Chrisholm, J. S., Dumont, J.-P., Hall, R. L., Hinde, R. A., and Ohtsuka, R. (1985). The use of quantitative observational techniques in anthropology. *Current Anthropology*, 26(3):323–335.
- Brands, R. A. (2013). Cognitive social structures in social network research: A review. *Journal of Organizational Behavior*, 34(S1):S82–S103.
- Brewer, D. D. (2000). Forgetting in the recall-based elicitation of personal and social networks. *Social networks*, 22(1):29–43.
- Butts, C. T. (2003). Network inference, error, and informant (in) accuracy: a bayesian approach. *social networks*, 25(2):103–140.
- Buyalskaya, A., Gallo, M., and Camerer, C. F. (2021). The golden age of social science. *Proceedings of the National Academy of Sciences*, 118(5).
- Campbell, K. E. and Lee, B. A. (1991). Name generators in surveys of personal networks. *Social networks*, 13(3):203–221.
- Carley, K. (1986). An approach for relating social structure to cognitive structure. *Journal of Mathematical Sociology*, 12(2):137–189.
- Carpenter, B., Gelman, A., Hoffman, M. D., Lee, D., Goodrich, B., Betancourt, M., Brubaker, M., Guo, J., Li, P., and Riddell, A. (2017). Stan: A probabilistic programming language. *Journal of statistical software*, 76(1):1–32.
- Carrington, P. J., Scott, J., and Wasserman, S. (2005). *Models and methods in social network analysis*, volume 28. Cambridge university press.
- Casciaro, T. (1998). Seeing things clearly: Social structure, personality, and accuracy in social network perception. *Social Networks*, 20(4):331–351.
- Clauset, A., Moore, C., and Newman, M. E. (2008). Hierarchical structure and the prediction of missing links in networks. *Nature*, 453(7191):98–101.
- Crockett, W. H. (1982). Balance, agreement, and positivity in the cognition of small social structures. In *Advances in experimental social psychology*, volume 15, pages 1–57. Elsevier.
- Crofoot, M. C., Rubenstein, D. I., Maiya, A. S., and Berger-Wolf, T. Y. (2011). Aggression, grooming and group-level cooperation in white-faced capuchins (*cebus capucinus*): Insights from social networks. *American Journal of Primatology*, 73(8):821–833.
- Davis, G. H., Crofoot, M. C., and Farine, D. R. (2018). Estimating the robustness and uncertainty of animal social networks using different observational methods. *Animal Behaviour*, 141:29–44.
- De Ayala, R. J. (2013). *The theory and practice of item response theory*. Guilford Publications.
- De Soto, C. B. (1960). Learning a social structure. *The Journal of Abnormal and Social Psychology*, 60(3):417.
- DeTroy, S. E., Ross, C. T., Cronin, K. A., Van Leeuwen, E. J., and Haun, D. B. (2021). Cofeeding tolerance in chimpanzees depends on group composition: A longitudinal study across four communities. *Isience*, 24(3):102175.
- Dijkstra, J. K., Kretschmer, T., Pattiselanno, K., Franken, A., Harakeh, Z., Vollebergh, W., and Veenstra, R. (2015). Explaining adolescents’ delinquency and substance use: A test of the maturity gap: The snare study. *Journal of Research in Crime and Delinquency*, 52(5):747–767.
- Doreian, P. and Conti, N. (2012). Social context, spatial structure and social network structure. *Social networks*, 34(1):32–46.
- Eagle, D. E. and Proeschold-Bell, R. J. (2015). Methodological considerations in the use of name generators and interpreters. *Social Networks*, 40:75–83.
- Eagle, N., Pentland, A. S., and Lazer, D. (2009). Inferring friendship network structure by using mobile phone data. *Proceedings of the national academy of sciences*, 106(36):15274–15278.
- Farine, D. R., Strandburg-Peshkin, A., Berger-Wolf, T., Ziebart, B., Brugere, I., Li, J., and Crofoot, M. C. (2016). Both nearest neighbours and long-term affiliates predict individual locations during collective movement in wild baboons. *Scientific reports*, 6(1):1–10.
- Fischer, C. S. (2009). The 2004 gss finding of shrunken social networks: An artifact? *American Sociological Review*, 74(4):657–669.
- Fiske, A. P. (1995). Social schemata for remembering people: Relationships and person attributes in free recall of acquaintances. *Journal of Quantitative Anthropology*, 5(4):305–324.
- Fiske, S. T. and Taylor, S. E. (1991). *Social cognition*. McGraw-Hill Book Company.
- Flynn, F. J., Reagans, R. E., Amanatullah, E. T., and Ames, D. R. (2006). Helping one’s way to the top: self-monitors achieve status by helping others and knowing who helps whom. *Journal of personality and social psychology*, 91(6):1123.
- Flynn, F. J., Reagans, R. E., and Guillory, L. (2010). Do you two know each other? transitivity, homophily, and the need for (network) closure. *Journal of personality and social psychology*, 99(5):855.
- Freeman, L. C. (1992). Filling in the blanks: A theory of cognitive categories and the structure of social affiliation. *Social Psychology Quarterly*, pages 118–127.
- Freeman, L. C., Romney, A. K., and Freeman, S. C. (1987). Cognitive structure and informant accuracy. *American anthropologist*, 89(2):310–325.
- Furman, W. (1996). The measurement of friendship perceptions: Conceptual and methodological issues. *The company they keep: Friendship in childhood and adolescence*, pages 41–65.
- Gervais, M. M. (2017). Rich economic games for networked relationships and communities: development and preliminary validation in yasawa, fiji. *Field methods*, 29(2):113–129.
- Gin, B., Sim, N., Skrandal, A., and Rabe-Hesketh, S. (2020). A dyadic irt model. *Psychometrika*, 85(3):815–836.
- Grippa, F. and Gloor, P. A. (2009). You are who remembers you. detecting leadership through accuracy of recall. *Social networks*, 31(4):255–261.
- Guimerà, R. and Sales-Pardo, M. (2009). Missing and spurious interactions and the reconstruction of complex networks. *Proceedings of the National Academy of Sciences*, 106(52):22073–22078.
- Hammer, M. (1980). Some comments on the validity of network data. *Connections*, 3(1):13–15.
- Hammer, M. (1984). Explorations into the meaning of social network interview data. *Social networks*, 6(4):341–371.
- Harris, K. M. (2013). *The add health study: Design and accomplishments*. Chapel Hill: Carolina Population Center, University of North Carolina at Chapel Hill, pages 1–22.
- Heider, F. (1982). *The psychology of interpersonal relations*. Psychology Press.
- Holt-Lunstad, J., Smith, T. B., and Layton, J. B. (2010). Social relationships and mortality risk: a meta-analytic review. *PLoS medicine*, 7(7):e1000316.
- Jackson, M. O., Rogers, B. W., and Zenou, Y. (2017). The economic consequences of social-network structure. *Journal of Economic Literature*, 55(1):49–95.
- Keltner, D., Gruenfeld, D. H., and Anderson, C. (2003). Power, approach, and inhibition. *Psychological review*, 110(2):265.
- Kenny, D. A. and La Voie, L. (1984). The social relations model. *Advances in experimental social psychology*, 18:141–182.
- Killworth, P. and Bernard, H. (1976). Informant accuracy in social network data. *Human Organization*, 35(3):269–286.
- Killworth, P. D. and Bernard, H. R. (1979). Informant accuracy in social network data iii: A comparison of triadic structure in behavioral and cognitive data. *Social Networks*, 2(1):19–46.
- Kivelä, M., Arenas, A., Barthélemy, M., Gleeson, J. P., Moreno, Y., and Porter, M. A. (2014). Multilayer networks. *Journal of complex networks*, 2(3):203–271.
- Knecht, A. B., Burk, W. J., Weesie, J., and Steglich, C. (2011). Friendship and alcohol use in early adolescence: A multilevel social network approach. *Journal of research on adolescence*, 21(2):475–487.

- Kogovšek, T. and Ferligoj, A. (2005). Effects on reliability and validity of egocentered network measurements. *Social networks*, 27(3):205–229.
- Kossinets, G. (2006). Effects of missing data in social networks. *Social networks*, 28(3):247–268.
- Koster, J. M., Grote, M. N., and Winterhalder, B. (2013). Effects on household labor of temporary out-migration by male household heads in nicaragua and peru: an analysis of spot-check time allocation data using mixed-effects models. *Human Ecology*, 41(2):221–237.
- Koster, J. M. and Leckie, G. (2014). Food sharing networks in lowland nicaragua: an application of the social relations model to count data. *Social Networks*, 38:100–110.
- Krackhardt, D. (1987). Cognitive social structures. *Social networks*, 9(2):109–134.
- Krackhardt, D. and Kilduff, M. (1999). Whether close or far: Social distance effects on perceived balance in friendship networks. *Journal of personality and social psychology*, 76(5):770.
- Krause, J., Lusseau, D., and James, R. (2009). Animal social networks: an introduction. *Behavioral Ecology and Sociobiology*, 63(7):967–973.
- Kurvers, R. H., Krause, J., Croft, D. P., Wilson, A. D., and Wolf, M. (2014). The evolutionary and ecological consequences of animal social networks: emerging issues. *Trends in ecology & evolution*, 29(6):326–335.
- Lakey, B. and Cohen, S. (2000). Social support theory and measurement. In Cohen, S., Underwood, L. G., and Gottlieb, B. H., editors, *Social support measurement and intervention: A guide for health and social scientists*, pages 29–54. Oxford University Press.
- Lazega, E. and Snijders, T. A. (2015). *Multilevel network analysis for the social sciences: Theory, methods and applications*, volume 12. Springer.
- Lewandowski, D., Kurowicka, D., and Joe, H. (2009). Generating random correlation matrices based on vines and extended onion method. *Journal of multivariate analysis*, 100(9):1989–2001.
- Lewin, K. (1951). *Field theory in social science*. Harpers & Brothers Publishing: New York.
- Lin, N. (2002). *Social capital: A theory of social structure and action*. Number 19. Cambridge university press.
- Lusher, D., Koskinen, J., and Robins, G. (2013). *Exponential random graph models for social networks: Theory, methods, and applications*, volume 35. Cambridge University Press.
- Marin, A. (2004). Are respondents more likely to list alters with certain characteristics?: Implications for name generator data. *Social networks*, 26(4):289–307.
- Marin, A. and Hampton, K. N. (2007). Simplifying the personal network name generator: Alternatives to traditional multiple and single name generators. *Field methods*, 19(2):163–193.
- Marineau, J. E., Labianca, G. J., Brass, D. J., Borgatti, S. P., and Vecchi, P. (2018). Individuals’ power and their social network accuracy: A situated cognition perspective. *Social Networks*, 54:145–161.
- Marsden, P. V. (1990). Network data and measurement. *Annual review of sociology*, pages 435–463.
- Marsden, P. V. (2005). Recent developments in network measurement. *Models and methods in social network analysis*, 8:30.
- McPherson, M., Smith-Lovin, L., and Brashears, M. E. (2006). Social isolation in america: Changes in core discussion networks over two decades. *American sociological review*, 71(3):353–375.
- Newman, M. E. (2018). Estimating network structure from unreliable measurements. *Physical Review E*, 98(6):062321.
- Newman, M. E. and Leicht, E. A. (2007). Mixture models and exploratory analysis in networks. *Proceedings of the National Academy of Sciences*, 104(23):9564–9569.
- Nolin, D. A. (2008). *Food-sharing networks in Lamalera, Indonesia: Tests of adaptive hypotheses*. University of Washington.
- Nolin, D. A. (2012). Food-sharing networks in lamalera, indonesia: status, sharing, and signaling. *Evolution and Human Behavior*, 33(4):334–345.
- Paik, A. and Sanchagrin, K. (2013). Social isolation in america: An artifact. *American Sociological Review*, 78(3):339–360.
- Paolisso, M. and Hames, R. (2010). Time diary versus instantaneous sampling: A comparison of two behavioral research methods. *Field Methods*, 22(4):357–377.
- Peixoto, T. P. (2018). Reconstructing networks with unknown and heterogeneous errors. *Physical Review X*, 8(4):041011.
- Peixoto, T. P. (2019). Bayesian stochastic blockmodeling. *Advances in network clustering and blockmodeling*, pages 289–332.
- Pinter-Wollman, N., Hobson, E. A., Smith, J. E., Edelman, A. J., Shizuka, D., De Silva, S., Waters, J. S., Prager, S. D., Sasaki, T., Wittemyer, G., et al. (2014). The dynamics of animal social networks: analytical, conceptual, and theoretical advances. *Behavioral Ecology*, 25(2):242–255.
- Pisor, A. and Ross, C. T. (2021). How generalizable are patterns of parochial altruism in humans?
- Pisor, A. C., Gervais, M. M., Purzycki, B. G., and Ross, C. T. (2020). Preferences and constraints: the value of economic games for studying human behaviour. *Royal Society open science*, 7(6):192090.
- Power, E. A. (2017). Social support networks and religiosity in rural south india. *Nature Human Behaviour*, 1(3):1–6.
- Power, E. A. and Ready, E. (2018). Building bigness: reputation, prominence, and social capital in rural south india. *American Anthropologist*, 120(3):444–459.
- Press, A. N., Crockett, W. H., and Rosenkrantz, P. S. (1969). Cognitive complexity and the learning of balanced and unbalanced social structures 1. *Journal of Personality*, 37(4):541–553.
- Pustejovsky, J. E. and Spillane, J. P. (2009). Question-order effects in social network name generators. *Social networks*, 31(4):221–229.
- Ready, E., Habecker, P., Abadie, R., Khan, B., and Dombrowski, K. (2020). Competing forces of withdrawal and disease avoidance in the risk networks of people who inject drugs. *PloS one*, 15(6):e0235124.
- Ready, E. and Power, E. (2021). Measuring reciprocity: Double sampling, concordance, and network construction.
- Redhead, D., Cheng, J., and O’Gorman, R. (2018). Higher status in group. In *Encyclopedia of evolutionary Psychological Science*. Springer.
- Redhead, D. and Power, E. A. (2021). Social hierarchies and social networks in humans. *Philosophical Transaction of the Royal Society B*.
- Redhead, D. and von Rueden, C. R. (2021). Coalitions and conflict: A longitudinal analysis of men’s politics. *Evolutionary Human Sciences*, 3.
- Ross, C. T. and Redhead, D. (2021). DieTryin: An R package for data collection, automated data entry, and post-processing of network-structured economic games, social networks, and other roster-based dyadic data. *Behavior Research Methods*, pages 1–21.
- Schwarz, N. (1999). Self-reports: how the questions shape the answers. *American psychologist*, 54(2):93.
- Selfhout, M., Burk, W., Branje, S., Denissen, J., Van Aken, M., and Meeus, W. (2010). Emerging late adolescent friendship networks and big five personality traits: A social network approach. *Journal of personality*, 78(2):509–538.
- Shakya, H. B., Christakis, N. A., and Fowler, J. H. (2017). An exploratory comparison of name generator content: Data from rural india. *Social networks*, 48:157–168.
- Simpson, B., Markovsky, B., and Steketee, M. (2011). Power and the perception of social networks. *Social Networks*, 33(2):166–171.
- Smith, E. B., Brands, R. A., Brashears, M. E., and Kleinbaum, A. M. (2020). Social networks and cognition. *Annual Review of Sociology*, 46(1):159–174.
- Smith, K. P. and Christakis, N. A. (2008). Social networks and health. *Annu. Rev. Sociol.*, 34:405–429.
- Snijders, T. A. (2017). Stochastic actor-oriented models for network dynamics. *Annual Review of Statistics and Its Application*, 4:343–363.
- Snijders, T. A., Pattison, P. E., Robins, G. L., and Handcock, M. S. (2006). New specifications for exponential random graph models. *Sociological methodology*, 36(1):99–153.
- Stan Development Team (2021a). Cmdstanr: A lightweight interface to stan for r users.
- Stan Development Team (2021b). Stan modeling language users guide and reference manual, version 2.21.
- Tourangeau, R. and Rasinski, K. A. (1988). Cognitive processes underlying context effects in attitude measurement. *Psychological bulletin*, 103(3):299.
- Urban, M. (2021). „die hoffnung, informiert zu sein“. effekte der corona-warn-app. *Prävention und Gesundheitsförderung*, pages 1–6.
- Van Duijn, M. A. and Vermunt, J. K. (2006). What is special about social network analysis? *Methodology: European Journal of Research Methods for the Behavioral and Social Sciences*, 2(1):2.
- von Rueden, C. R., Redhead, D., O’Gorman, R., Kaplan, H., and Gurven, M. (2019). The dynamics of men’s cooperation and social status in a small-scale society. *Proceedings of the Royal Society B*, 286(1908):20191367.

- Walker, C. J. (1976). The employment of vertical and horizontal social schemata in the learning of a social structure. *Journal of Personality and Social Psychology*, 33(2):132.
- Wasserman, S. and Anderson, C. (1987). Stochastic a posteriori block-models: Construction and assessment. *Social networks*, 9(1):1–36.
- Young, J.-G., Cantwell, G. T., and Newman, M. (2020). Robust bayesian inference of network structure from unreliable data. *arXiv preprint arXiv:2008.03334*.
- Yousefi-Nooraie, R., Marin, A., Hanneman, R., Pullenayegum, E., Lohfeld, L., and Dobbins, M. (2019). The relationship between the position of name generator questions and responsiveness in multiple name generator surveys. *Sociological Methods & Research*, 48(2):243–262.
- Zhu, X., Woo, S. E., Porter, C., and Brzezinski, M. (2013). Pathways to happiness: From personality to social networks and perceived support. *Social networks*, 35(3):382–393.

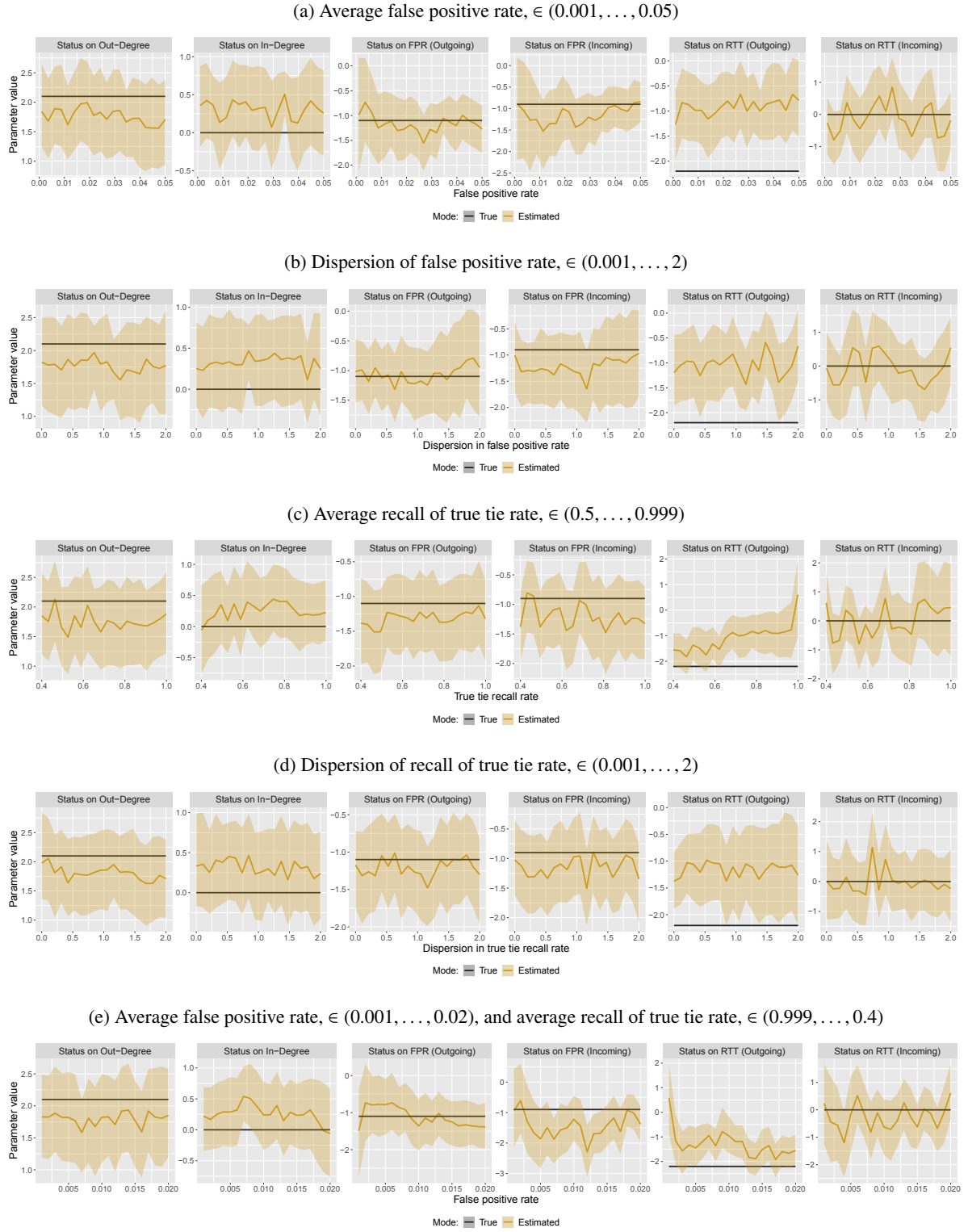


Fig. 4: Model with self-report data only. Each frame plots our recovery of the effect of status on different individual-level properties. The true network and all associated parameters are held fixed, except for the parameters displayed in the column labels, which range over the indicated support. The levels of each outcome in the true network appear as horizontal black lines. The orange regions illustrate the posterior distribution of each outcome from the latent network model.

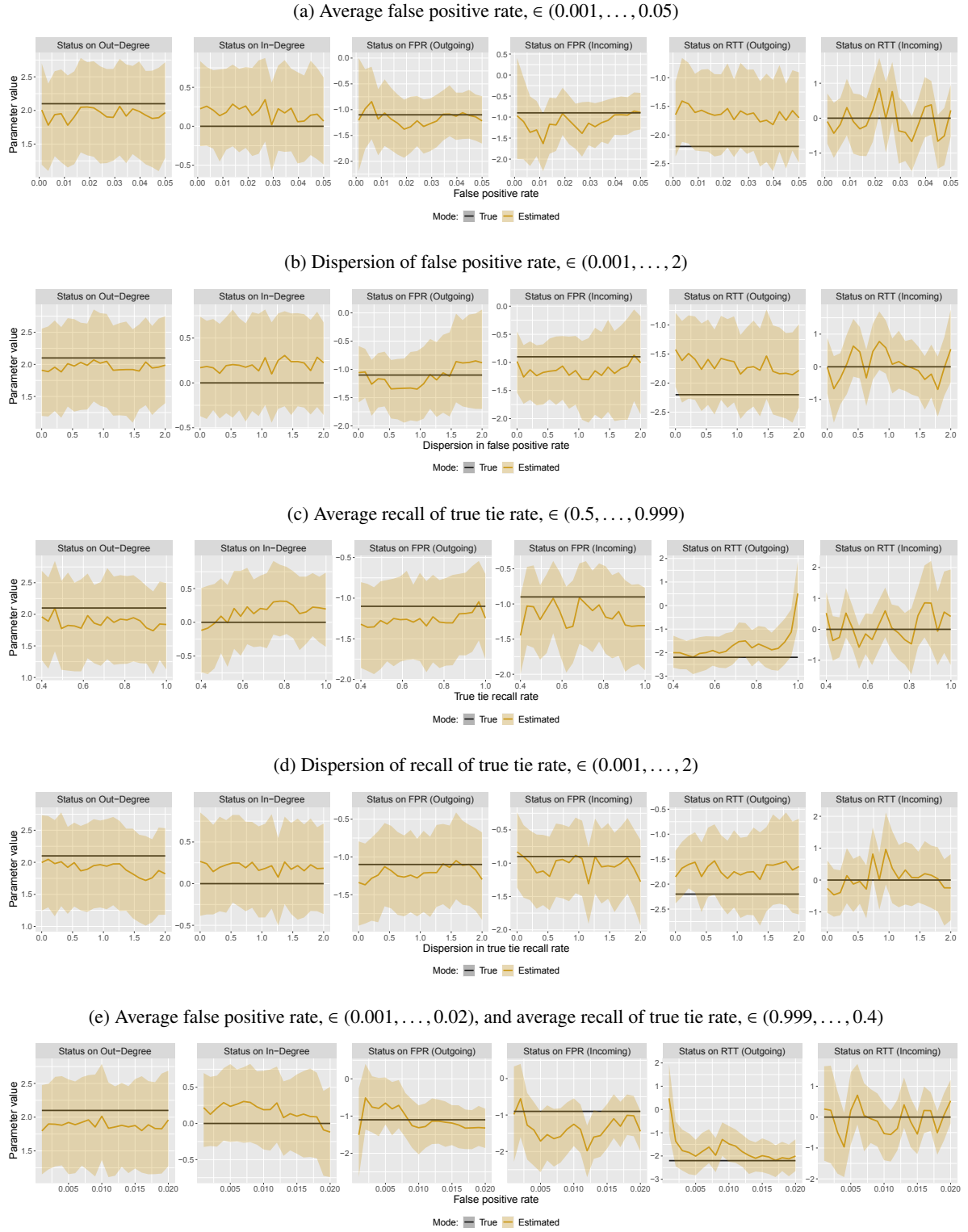


Fig. 5: Model with self-report and ground-truth data. Each frame plots our recovery of the effect of status on different individual-level properties. The true network and all associated parameters are held fixed, except for the parameters displayed in the column labels, which range over the indicated support. The levels of each outcome in the true network appear as horizontal black lines. The orange regions illustrate the posterior distribution of each outcome from the latent network model.

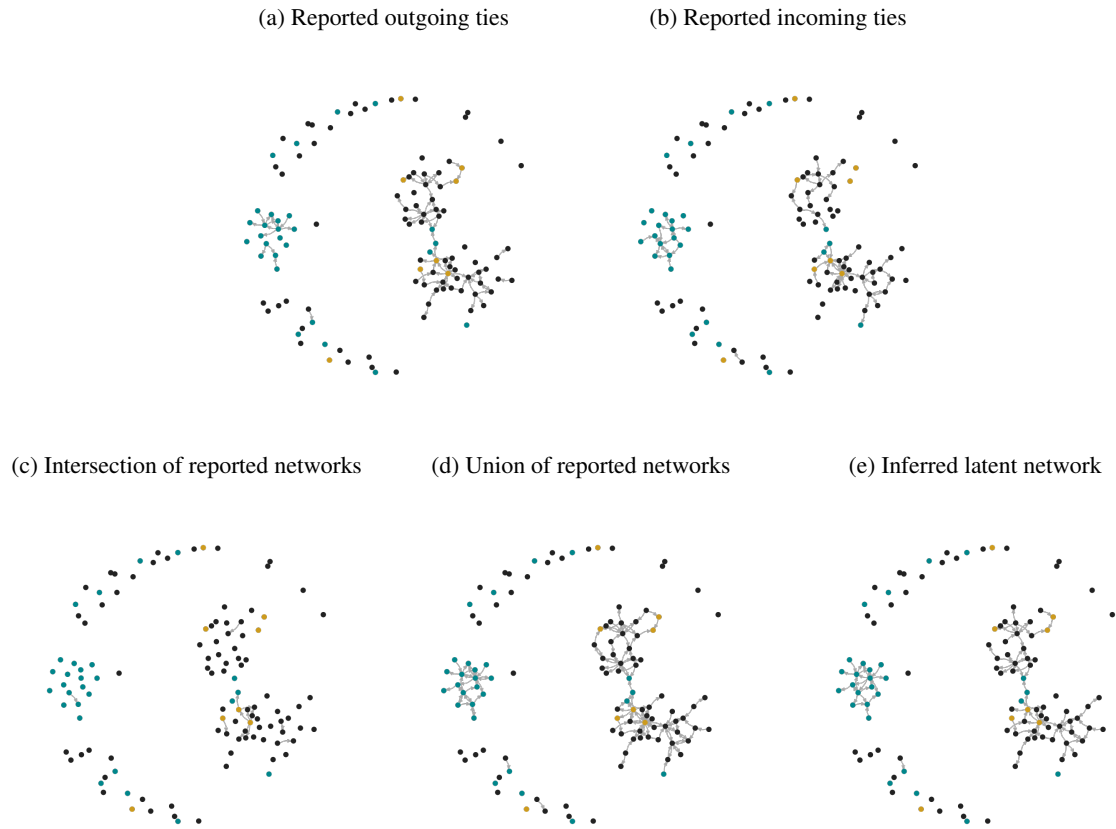


Fig. 6: In frames 6a and 6b, we plot the reported networks of food/money transfers in a rural Colombian population. Then, in frames 6c and 6d, we plot the intersection and union, respectively, of these networks. Finally, in frame 6e, we plot the posterior median latent network as inferred by our model. We note, visually, that the intersection appears too sparse. The latent network model settles on a structure similar to the union in this case, but corrects for reciprocity inflation arising from question order effects; see Table 1.

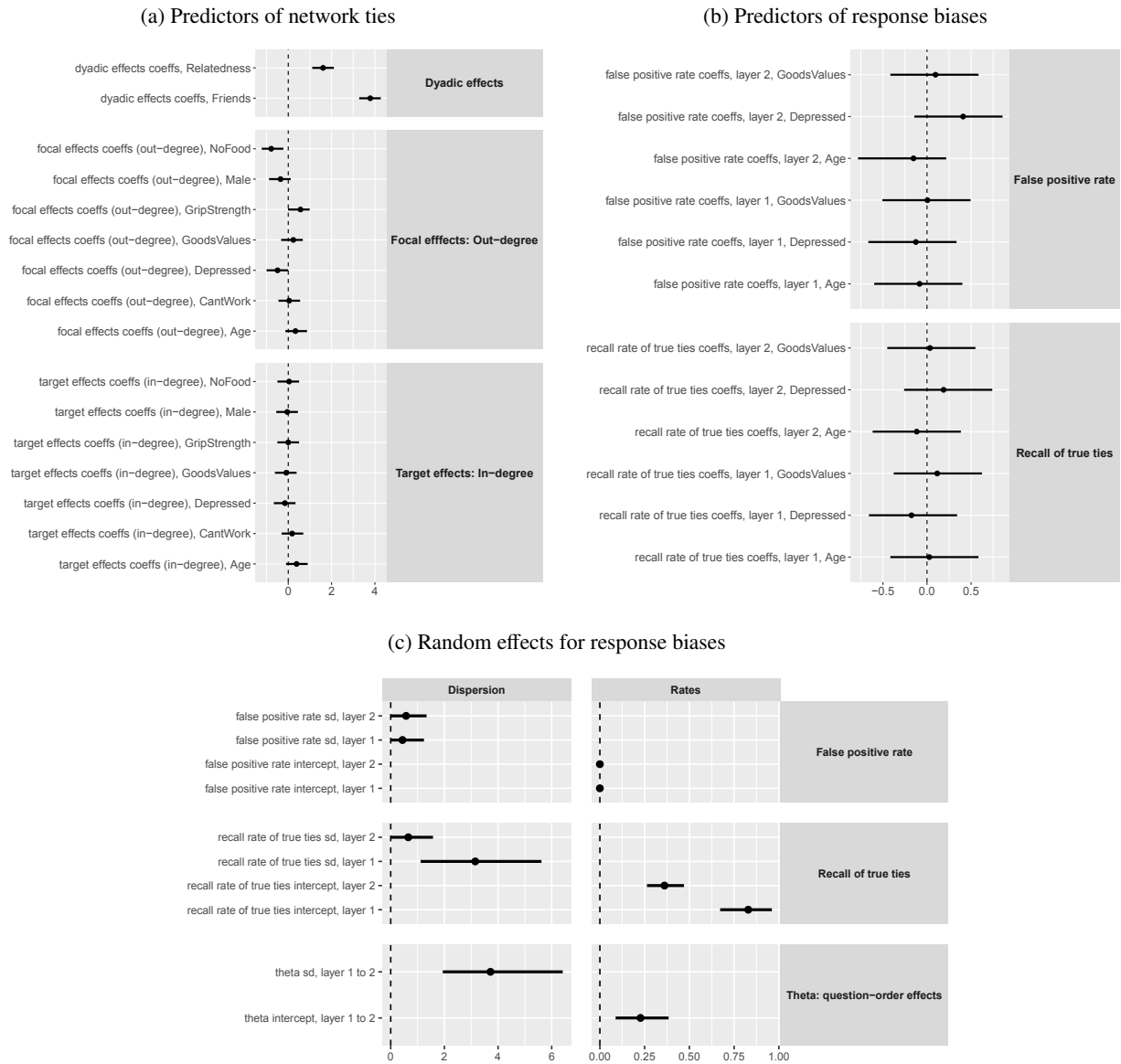


Fig. 7: Posterior median estimates (and 89% highest posterior density intervals) of various effects from our regression models. In frame 7a, we plot the standardized effects of various covariates on network structure. In frame 7b, we plot the standardized effects of various covariates on false positive and true tie recall rates. Finally, in frame 7c, we plot the unstandardized parameters controlling the mean and cross-individual variation in the random effects measuring reporting biases.