

Online Appendices, Not for Publication

APPENDIX A. NETWORK DEFINITIONS

In this section, we provide basic definitions and interpretations for the different network characteristics that we consider. See [Jackson \(2008\)](#) for details. At the household level, we study:

- Degree: the number of links that a household has. This is a measure of how well connected a node is in the graph.
- Clustering coefficient: the fraction of a household’s neighbors that are themselves neighbors. This is a measure of how interwoven a household’s neighborhood is.
- Eigenvector centrality: recursively defined notion of importance. A household’s importance is defined to be proportional to the sum of its neighbors’ importances. It corresponds to the i^{th} entry of the eigenvector corresponding to the maximal eigenvalue of the adjacency matrix. This is a measure of how important a node is, in the sense of information flow. We take the eigenvector normalized with $\|\cdot\|_2 = 1$.
- Reachability and distance: we say two households i and j are reachable if there exists a path through the network that connects them. The distance is the length of the shortest such path.

At the hamlet level, we consider:

- Average degree: the mean number of links that a household has in the hamlet. A network with higher average degree has more edges on which to transmit information.
- Average clustering: the mean clustering coefficient of households in the hamlet. This measures how interwoven the network is.
- Average path length: the mean length of the shortest path between any two households in the hamlet which are connected. Shorter average path length means information has to travel less (on average) to get from household i to household j .
- First eigenvalue: the largest eigenvalue of the adjacency matrix. This is a measure of how diffusive the network is. A higher first eigenvalue tends to mean that information is generally transmitted more.
- Fraction of nodes in the giant component: the share of nodes in the graph that are in the largest connected component. Typically, realistic graphs have a “giant” component with almost all nodes in it. Thus, the measure should be approaching one. For a network that is sampled, this number can be significantly lower. In particular, networks which were tenuously or sparsely connected to begin with, may “shatter” under sampling and therefore the giant component may no longer be giant after sampling. In turn, it becomes a useful measure of how interwoven the underlying network is.
- Link density: the average share of connections (out of potential connections) that a household has. This measure looks at the rate of edge formation in a graph.

APPENDIX B. KALMAN FILTER, ESTIMATION AND SIMULATION

This section develops the formal algorithm for the model and discusses estimation.

B.1. Model.

Setup. Without loss of generality, fix node j about whose wealth the remainder of the nodes are learning. Wealth follows an AR(1) process given by

$$w_{j,t} = c + \rho w_{j,t-1} + \epsilon_{j,t}.$$

Individuals $i \in V \setminus \{j\}$ want to guess $w_{j,t}$ when surveyed at period t , given an information set $\mathcal{F}_{i,t-1}^j$ that is informed from social learning.

The model will have a transmission error in every step when an individual speaks to another individual. For instance, if l communicates with i in period r , this communication can be disturbed by some $u_r^{l \rightarrow i}$. We will assume that every $u_r^{l \rightarrow i}$ is independently and identically distributed, $\mathcal{N}(0, \sigma_u^2)$ with common mean and variance.

Individuals communicate with each others as follows. At period t we look at what node i receives from others and consider her updating problem:

- *Signals from the source j :* Every period, the source j generates a signal about her $t - 1$ wealth that she transmits to each of her neighbors, $i \in N_j$.

$$S_{t-1}^{j \rightarrow i} = w_{j,t-1} + u_{t-1}^{j \rightarrow i}.$$

- *Signals from arbitrary node l to i :* Every period, a node l noisily transmits the most recent piece of gossip she has heard to each of her neighbors, independently.
 - Let $k^* := k^*(l, j)$ be the neighbor of l that is closest to j .³⁶ The signal that l received from k^* the previous period is what will then be passed on.
 - Passing only occurs if l is sufficiently sure enough about the quality of this information. This is mapped to some threshold τ such that if $d(k^*(l, j), j) \leq \tau$, then l passes information to each of her neighbors. If $d(k^*(l, j), j) > \tau$, then no information is passed.
 - When l passes information, it is

$$S_{t-d(l,j)}^{l \rightarrow i} = S_{t-1-d(k^*,j)}^{k^* \rightarrow l} + u_{t-d(l,j)}^{l \rightarrow i}.$$

The above protocol defines a signal generation process. Thus, in every period t , a generic node i in the graph has received a vector of signals

$$\mathbf{s}^{i,t} := \left(s_1^{i,t}, \dots, s_{t-d(i,j)}^{i,t} \right).$$

This vector of signals encodes the information that i has about each $w_{j,1}, \dots, w_{j,t-d(i,j)}$. Note that the signal vector $\mathbf{s}^{i,t}$ is double-indexed since it can have time-varying elements.

³⁶For presentation purposes we assume this is unique. If it is not unique, and there are two such closest signals, then we assume that l passes the average.

- The signal vector can be treated as a collection of independent draws (conditional on the wealth sequence) with

$$s_r^{i,t} \sim \mathcal{N}(w_{j,r}, \sigma_{r,t,i}^2)$$

where i 's t th period signal about $w_{j,r}$ can only come from neighbors that are close enough to j ,

$$s_r^{i,t} = \sum_{l \in N_i} \omega_{l,r,t,i} \cdot s_r^{l \rightarrow i},$$

and $\omega_{l,r,t,i} = \frac{\mathbf{1}\{t-r \geq d(l,j)+1\} / [\sigma_u^2 d(l,j)]}{\sum_{k \in N_i} \mathbf{1}\{t-r \geq d(k,j)+1\} / [\sigma_u^2 d(k,j)]}$ is the weight that i puts on l in period t about l 's estimate of $w_{j,r}$.

This is because only neighbors of i that are within $t - r - 1$ steps of j can reveal an estimate of $w_{j,r}$ to i by period t . Every time the signal is transferred across individuals, it is disturbed by a shock with variance σ_u^2 , leading to a variance of $\sigma_u^2 \cdot d(l, j)$ for a signal that has traveled $d(l, j)$ steps.

In this case, we can compute i 's period t variance about $w_{j,r}$ as

$$\sigma_{r,t,i}^2 = \sum_{l \in N_i} \omega_{l,r,t,i}^2 \cdot \sigma_u^2 d(l, j).$$

- Given $\mathbf{s}^{i,t}$, node i applies the Kalman filter to obtain the posterior mean and variance over $w_{j,t}$.

Kalman Filter. The Kalman filter is as follows. In what follows, we reserve r to index time (and describe the process only for nodes that are speaking). At period t a node i makes the following computation. She treats the system as the t th period of a linear Gauss-Markov dynamical system with

- state equation is given by

$$w_{j,r} = c + \rho w_{j,r-1} + \epsilon_{j,r}, \quad r = 1, \dots, t+1.$$

- measurement equation given by

$$s_r^{i,t} = w_{j,r} + v_r^{i,t},$$

where $v_r^{i,t} \sim \mathcal{N}(0, \sigma_{r,t,i}^2)$.

Then the computation of the Kalman filter is entirely standard given the vector $\mathbf{s}^{i,t}$ of measurements and knowledge of parameters $c, \rho, \sigma_\epsilon^2, \sigma_u^2$ and $d(k, j) \quad \forall k \in N_i$. The crucial equations are how to do a time update given prior information and how to incorporate the new measurements to correct the system:

- Time update equations:

$$\begin{aligned} \hat{w}_r^- &= \rho \hat{w}_{r-1} + c \\ P_r^- &= \rho^2 P_{r-1} + \sigma_\epsilon^2. \end{aligned}$$

- Measurement update equations:

$$\begin{aligned} K_r &= \frac{P_r^-}{P_r^- + \sigma_{r,t,i}^2} \\ \hat{w}_r &= \hat{w}_r^- + K_r (s_r^{i,t} - \hat{w}_r^-) \\ P_r &= (1 - K_r) P_r^- \end{aligned}$$

The initialization is at the mean of the invariant distribution $w_0 = \frac{c}{1-\rho}$ and the variance $P_0 = \frac{\sigma_\epsilon^2}{1-\rho^2}$.

B.2. Estimation. Before conducting our simulated method of moments, we first estimate some preliminary parameters.

- (1) Auto-correlation parameter (ρ): We use data from Indonesia Family Life Survey. We construct a panel data for 1993, 1997, 2000, and 2007. The sample used contains only those households that were surveyed in all the years. We use real total expenditures as our variable of interest.³⁷ As there was a financial crisis in 1997-1998, we omit the 1997-2000 period from the estimation. Given that the gap between the years is long and variable, we use the mean gap to compute an approximate yearly ρ . The mean gap is 5.5 years so we obtain ρ using $(\rho)^k = \hat{\rho}_{Panel}$, for $k = 5.5$, and its distribution is derived using the delta method. We estimate $\hat{\rho} = 0.86$ and because of the size of the panel, the parameter is tightly estimated (a t -statistic of 12.33); thus, we view it as super-consistent relative to the structural parameters in our model.
- (2) Variance of the shock term (σ_ϵ^2): We obtain this variable using the stationary variance of the consumption process $\sigma_w^2 = \frac{\sigma_\epsilon^2}{(1-\rho^2)}$. Again, given the size of the data set this can be viewed as super-consistent.

B.3. Simulated Method of Moments. Equipped with a collection of over 600 graphs, ρ , and σ_ϵ^2 , we estimate our model via simulated method of moments. The two parameters we are interested in are (α, τ) where $\alpha := \frac{\sigma_y^2}{\sigma_\epsilon^2}$ and τ is the maximal distance away from the source for an individual to be confident enough to pass information to her neighbors.

Our approach is a grid-based simulated method of moments which allows us to conduct inference on a large simulation quite easily (Banerjee et al., 2013). We let Θ be the parameter space and Ξ be a grid on Θ , which we describe below. We put $\psi(\cdot)$ as the moment function and let $z_r = (y_r, x_r)$ denote the empirical data for network r with a vector of wealth ranking decisions for each surveyed individual, y_r , as well as data, x_r , which includes expenditure data and the graph G_r .

Define $m_{emp,r} := \psi(z_r)$ as the empirical moment for hamlet r and $m_{sim,r}(s, \theta) := \psi(z_r^s(\theta)) = \psi(y_r^s(\theta), x_r)$ as the s th simulated moment for hamlet r at parameter value θ . Finally, put B as the number of bootstraps and S as the number of simulations used to construct the simulated moment. This nests the case with $B = 1$ in which we just find the minimizer of the objective function.

- (1) Pick lattice $\Xi \subset \Theta$. For $\xi \in \Xi$ on the grid:

³⁷Expenditures were converted in real terms using the CPI published by the Central Bank.

(a) For each network $r \in [R]$, compute

$$d(r, \xi) := \frac{1}{S} \sum_{s \in [S]} m_{sim,r}(s, \theta) - m_{emp,r}.$$

(b) For each $b \in [B]$, compute

$$D(b, \xi) := \frac{1}{R} \sum_{r \in [R]} \omega_r^b \cdot d(r, \xi)$$

where $\omega_r^b = e_{br}/\bar{e}_r$, with e_{br} iid $\exp(1)$ random variables and $\bar{e}_r = \frac{1}{R} \sum e_{br}$ if we are conducting bootstrap, and $\omega_r^b = 1$ if we are just finding the minimizer.

(c) Find $\xi^{*b} = \operatorname{argmin} Q^{*b}(\xi)$, with $Q^{*b}(\xi) = D(b, \xi)'D(b, \xi)$.³⁸

(2) Obtain $\{\xi^{*b}\}_{b \in B}$.

(3) For conservative inference on $\hat{\theta}_j$, the j^{th} component, consider the $1 - \alpha/2$ and $\alpha/2$ quantiles of the ξ_j^{*b} marginal empirical distribution.

In all simulations we use $B = 10000$, $S = 50$. We set $\Xi = [0.1 : 0.033 : 0.85] \times \{1, \dots, 7\}$.

B.4. Simulations for Regressions. To generate our synthetic data we fix our parameters $(\hat{\alpha}, \hat{\tau})$ and generate 50 draws. We then compute

$$\overline{Error}_{ijk r}^{SIM} = \sum_s Error_{ijk r}^s / 50.$$

This allows us to aggregate the errors to any level we need. For instance by integrating over all the triples in our sample, we can compute $\overline{Error}_r^{SIM}$, the simulated error rate for hamlet r . We then conduct our regression analysis with these simulated outcomes.

³⁸Because we are just identified we do not need to weight the moments.

APPENDIX C. DETAILS ON POVERTY TARGETING PROCEDURES

This appendix briefly describes the poverty targeting procedures used to allocate the transfer program to households. Additional details can be found in [Alatas et al. \(2012\)](#).

- **PMT Treatment:** the government created formulas that mapped 49 easily observable household characteristics into a single index using regression techniques.³⁹ Government enumerators collected these indicators from all households in the PMT hamlets by conducting a door-to-door survey. These data were then used to calculate a computer-generated predicted consumption score for each household using a district-specific PMT formula. A list of beneficiaries was generated by selecting the pre-determined number of households with the lowest scores in each hamlet, based on quotas determined by a geographic targeting procedure.
- **Community Treatment:** To start, a local facilitator visited each hamlet to publicize the program and invite individuals to a community meeting.⁴⁰ At the meeting, the facilitator first explained the program. Next, he or she displayed the list of all households in the hamlet (which came from the baseline survey). The facilitator then spent about 15 minutes having the community brainstorm a list of characteristics that differentiate the poor from the wealthy households in their community. The facilitator then proceeded with the ranking exercise using a set of randomly-ordered index cards that displayed the names of each household in the neighborhood. He or she hung a string from wall to wall, with one end labeled as “most well-off” (paling mampu) and the other side labeled as “poorest” (paling miskin). Then, he or she started by holding up the first two name cards from the randomly-ordered stack and asking the community, “Which of these two households is better off?” Based on the community’s response, he or she attached the cards along the string, with the poorer household placed closer to the “poorest” end. Next, the facilitator displayed the third card and asked how this household ranked relative to the first two households. The activity continued with each card being positioned relative to the already-ranked households one-by-one until complete. Before the final ranking was recorded, the facilitator read the ranking aloud so adjustments could be made if necessary. After all meetings were complete, the facilitators were provided with “beneficiary quotas” for each hamlet based on the geographic targeting procedure. Households ranked below the quota were deemed eligible.
- **Hybrid Treatment:** This method combines the community ranking procedure with a subsequent PMT verification. The ranking exercise, described above, was implemented first.

³⁹The chosen indicators encompassed the household’s home attributes (wall type, roof type, etc), assets (TV, motor-bike, etc), household composition, and household head’s education and occupation. The formulas were derived using pre-existing survey data: specifically, the government estimated the relationship between the variables of interest and household per capita consumption. While the same indicators were considered across regions, the government estimated district-specific formulas due to the perceived high variance in the best predictors of poverty across regions

⁴⁰On average, 45 percent of households attended the meeting. Note, however, that we only invited the full community in half of the community treatment hamlets. In the other half (randomly selected), only local elites were invited, so that we can see whether elites are more likely to capture the community process when they have control over the process.

However, there was one key difference: at the start of these meetings, the facilitator announced that the lowest-ranked households would be independently checked by the government enumerators before the beneficiary list was finalized. After the community meetings were complete, the government enumerators indeed visited the lowest-ranked households to collect the data needed to calculate their PMT score. The number of households to be visited was computed by multiplying the “beneficiary quotas” by 150 percent. Households were ranked by their PMT score, and those below the village quota became beneficiaries of the program. Thus, it was possible that some households could become beneficiaries even if they were ranked as slightly wealthier than the beneficiary quota cutoff line on the community list. Conversely, some relatively poor-ranked households on the community list might become ineligible.

APPENDIX D. TABLES WITHOUT DON'T KNOWS

TABLE D.1. The Correlation between Household Network Characteristics and the Error Rate in Ranking Income Status of Household, Conditional on Offering Assessments

	(1)	(2)	(3)	(4)
<i>Outcome variable: Error rate conditional on reporting</i>				
<i>Panel A: Consumption Metric</i>				
Degree	-0.00280*** (0.000713)			-0.00175 (0.00116)
Clustering		-0.0117 (0.00893)		-0.00830 (0.00992)
Eigenvector Centrality			-0.0857*** (0.0233)	-0.0434 (0.0383)
R-squared	0.664	0.663	0.665	0.665
<i>Panel B: Self-Assessment Metric</i>				
Degree	-0.00384*** (0.000713)			-0.00266** (0.00122)
Clustering		-0.00248 (0.0100)		0.000797 (0.0109)
Eigenvector Centrality			-0.104*** (0.0248)	-0.0471 (0.0410)
R-squared	0.671	0.669	0.671	0.671
Hamlet Fixed Effect	Yes	Yes	Yes	Yes

Notes: This table provides estimates of the correlation between a household's network characteristics and its ability to accurately rank the poverty status of other members of the hamlet. The sample comprises 5,633 households. The mean of the dependent variable in Panel A (a household's error rate in ranking others in the hamlet based on consumption) is 0.52, while the mean of the dependent variable in Panel B (a household's error rate in ranking others in the hamlet based on a household's own self-assessment of poverty status) is 0.46. "Don't know" answers are dropped. Standard errors are clustered by hamlet and are listed in parentheses.

*** p<0.01, ** p<0.05, * p<0.1.

TABLE D.2. The Correlation between Household Network Characteristics and the Error Rate in Ranking Income Status of Household, Controlling for Household Characteristics, Conditional on Offering Assessments

	(1)	(2)	(3)	(4)
<i>Outcome variable: Error rate conditional on reporting</i>				
<i>Panel A: Consumption Metric</i>				
Degree	-0.00215*** (0.000698)			-0.00122 (0.00112)
Clustering		-0.0115 (0.00879)		-0.00828 (0.00979)
Eigenvector Centrality			-0.0694*** (0.0231)	-0.0387 (0.0375)
R-squared	0.668	0.667	0.668	0.668
<i>Panel B: Self-Assessment Metric</i>				
Degree	-0.00301*** (0.000697)			-0.00200* (0.00119)
Clustering		-0.00239 (0.00973)		0.000599 (0.0107)
Eigenvector Centrality			-0.0828*** (0.0243)	-0.0406 (0.0401)
R-squared	0.676	0.674	0.676	0.676
Hamlet Fixed Effect	Yes	Yes	Yes	Yes

Notes: This table provides estimates of the correlation between a household's network characteristics and its ability to accurately rank the poverty status of other members of the hamlet, controlling for the household's characteristics as in Table 3. The sample comprises 5,630 households for panel. The mean of the dependent variable in Panel A (a household's error rate in ranking others in the hamlet based on consumption) is 0.52, while the mean of the dependent variable in Panel B (a household's error rate in ranking others in the hamlet based on a household's own self-assessment of poverty status) is 0.46. "Don't know" answers are dropped. Standard errors are clustered by hamlet and are listed in parentheses. *** p<0.01, ** p<0.05, * p<0.1.

TABLE D.3. The Correlation Between Inaccuracy in Ranking a Pair of Households in a Hamlet and the Average Inverse Distance to Rankees, Conditional on Offering Assessments

	(1)	(2)	(3)	(4)
<i>Outcome variable: Error rate conditional on reporting</i>				
<i>Panel A: Consumption Metric</i>				
Average Inverse Distance	-0.00645 (0.00662)	-0.0178*** (0.00660)	-0.00886* (0.00491)	-0.0215 (0.0132)
Average Degree		0.00139 (0.00140)	0.00571* (0.00307)	0.00628* (0.00320)
Average Clustering Coefficient		0.0252 (0.0228)	0.0618** (0.0266)	0.0693** (0.0280)
Average Eigenvector Centrality		-0.0480 (0.0547)	-0.177* (0.1000)	-0.153 (0.106)
R-squared	0.000	0.008	0.082	0.139
<i>Panel B: Self-Assessment Metric</i>				
Average Inverse Distance	-0.00802 (0.00689)	-0.0143** (0.00694)	-0.00458 (0.00507)	0.00230 (0.0131)
Average Degree		0.000997 (0.00146)	0.00175 (0.00297)	0.00118 (0.00312)
Average Clustering Coefficient		-0.0216 (0.0237)	0.0228 (0.0289)	0.0242 (0.0300)
Average Eigenvector Centrality		0.0298 (0.0549)	0.0304 (0.0993)	0.0346 (0.104)
R-squared	0.000	0.002	0.092	0.168
Demographic Controls	No	Yes	Yes	Yes
Hamlet Fixed Effect	No	No	Yes	Yes
Ranker Fixed Effect	No	No	No	Yes

Notes: This table provides an estimate of the correlation between the accuracy in ranking a pair of households in a hamlet and the characteristics of the households that are being ranked. In Panel A, the dependent variable is a dummy variable for whether household i ranks household j versus household k incorrectly based on using consumption as the metric of truth (the sample mean is 0.52). In Panel B, the self-assessment variable is the metric of truth (the sample mean is 0.46). "Don't know" answers are dropped. The sample is comprised of 80,380 ranked pairs in Panel A and 116,338 in Panel B. Standard errors are clustered by hamlet and are listed in parentheses. *** $p < 0.01$, ** $p < 0.05$, * $p < 0.1$.

APPENDIX E. EXTENDED MICRO TABLES USING SIMULATIONS

TABLE E.1. The Correlation between Household Network Characteristics and the Error Rate in Ranking Income Status of Households, Including Simulations

	(1)	(2)	(3)	(4)	(5)	(6)	(7)	(8)
	Outcome variable: Error rate				Outcome variable: Share of don't knows			
Panel A: Consumption Metric								
Degree	-0.00281*** (0.000712)			-0.00182 (0.00116)	-0.00306*** (0.000666)			-0.00152 (0.00115)
Clustering		-0.0118 (0.00889)		-0.00859 (0.00986)		0.00297 (0.0110)		0.00572 (0.0120)
Eigenvector Centrality			-0.0847*** (0.0232)	-0.0409 (0.0380)			-0.0937*** (0.0255)	-0.0619 (0.0414)
R-squared	0.667	0.666	0.667	0.668	0.721	0.720	0.721	0.722
Panel B: Self-Assessment Metric								
Degree	-0.00386*** (0.000712)			-0.00276** (0.00122)	-0.00306*** (0.000666)			-0.00152 (0.00115)
Clustering		-0.00283 (0.00999)		0.000172 (0.0108)		0.00297 (0.0110)		0.00572 (0.0120)
Eigenvector Centrality			-0.103*** (0.0247)	-0.0439 (0.0408)			-0.0937*** (0.0255)	-0.0619 (0.0414)
R-squared	0.674	0.672	0.673	0.674	0.721	0.720	0.721	0.722
Panel C: Simulations								
Degree	-0.00836*** (0.000576)			-0.00465*** (0.000781)	-0.0166*** (0.00113)			-0.00945*** (0.00153)
Clustering		-0.0780*** (0.00741)		-0.0706*** (0.00714)		-0.157*** (0.0146)		-0.143*** (0.0141)
Eigenvector Centrality			-0.310*** (0.0187)	-0.176*** (0.0264)			-0.613*** (0.0368)	-0.341*** (0.0518)
R-squared	0.896	0.887	0.905	0.913	0.902	0.893	0.912	0.921
Hamlet Fixed Effect	Yes	Yes	Yes	Yes	Yes	Yes	Yes	Yes

Notes: This table provides estimates of the correlation between a household's network characteristics and its ability to accurately rank the poverty status of other members of the hamlet. The sample comprises 5,633 households. The mean of the dependent variable in Panel A (a household's error rate in ranking others in the hamlet based on consumption) is 0.52, while the mean of the dependent variable in Panel B (a household's error rate in ranking others in the hamlet based on a household's own self-assessment of poverty status) is 0.46. Details of the simulation procedure for Panel C are contained in Appendix B. Standard errors are clustered by hamlet and are listed in parentheses.

*** p<0.01, ** p<0.05, * p<0.1.

TABLE E.2. The Correlation between Household Network Characteristics and the Error Rate in Ranking Income Status of Households, Controlling for Household Characteristics, Including Simulations

	(1)	(2)	(3)	(4)	(5)	(6)	(7)	(8)
	Outcome variable: Error rate				Outcome variable: Share of don't knows			
	Panel A: Consumption Metric							
Degree	-0.00215*** (0.000697)			-0.00127 (0.00112)	-0.00238*** (0.000657)			-0.000913 (0.00113)
Clustering		-0.0116 (0.00877)		-0.00860 (0.00974)		0.00185 (0.0108)		0.00522 (0.0118)
Eigenvector Centrality			-0.0684*** (0.0230)	-0.0363 (0.0372)			-0.0774*** (0.0252)	-0.0594 (0.0406)
R-squared	0.671	0.670	0.671	0.671	0.725	0.724	0.725	0.725
	Panel B: Self-Assessment Metric							
Degree	-0.00302*** (0.000697)			-0.00209* (0.00118)	-0.00238*** (0.000657)			-0.000913 (0.00113)
Clustering		-0.00279 (0.00972)		-5.26e-05 (0.0106)		0.00185 (0.0108)		0.00522 (0.0118)
Eigenvector Centrality			-0.0819*** (0.0243)	-0.0376 (0.0398)			-0.0774*** (0.0252)	-0.0594 (0.0406)
R-squared	0.679	0.677	0.678	0.679	0.725	0.724	0.725	0.725
	Panel C: Simulations							
Degree	-0.00834*** (0.000576)			-0.00461*** (0.000779)	-0.0166*** (0.00113)			-0.00940*** (0.00152)
Clustering		-0.0782*** (0.00735)		-0.0703*** (0.00710)		-0.158*** (0.0145)		-0.143*** (0.0141)
Eigenvector Centrality			-0.309*** (0.0187)	-0.176*** (0.0264)			-0.612*** (0.0368)	-0.341*** (0.0517)
R-squared	0.896	0.887	0.905	0.913	0.903	0.894	0.912	0.921
Hamlet Fixed Effect	Yes	Yes	Yes	Yes	Yes	Yes	Yes	Yes

Notes: This table provides estimates of the correlation between a household's network characteristics and its ability to accurately rank the poverty status of other members of the hamlet, controlling for the household's characteristics as in Table 3. The sample comprises 5,630 households for panel. The mean of the dependent variable in Panel A (a household's error rate in ranking others in the hamlet based on consumption) is 0.52, while the mean of the dependent variable in Panel B (a household's error rate in ranking others in the hamlet based on a household's own self-assessment of poverty status) is 0.46. Details of the simulation procedure for Panel C are contained in Appendix B. Standard errors are clustered by hamlet and are listed in parentheses. *** p<0.01, ** p<0.05, * p<0.1.

TABLE E.3. The Correlation between Inaccuracy in Ranking a Pair of Households in a Hamlet and the Average Distance to Rankees, Including Simulations

	(1)	(2)	(3)	(4)	(5)	(6)	(7)	(8)
	Outcome variable: Error rate				Outcome variable: Share of don't knows			
	Panel A: Consumption Metric							
Average Inverse Distance	-0.0576*** (0.00847)	-0.0383*** (0.00835)	-0.0220*** (0.00565)	-0.0159 (0.0127)	-0.0737*** (0.00950)	-0.0414*** (0.00992)	-0.0280*** (0.00707)	-0.00756 (0.0132)
Average Degree		-0.00500*** (0.00176)	0.00243 (0.00318)	0.00258 (0.00323)		-0.00961*** (0.00212)	-0.00257 (0.00309)	-0.00270 (0.00309)
Average Clustering Coefficient		0.00200 (0.0256)	0.0322 (0.0275)	0.0339 (0.0279)		-0.0298 (0.0307)	-0.0132 (0.0288)	-0.0144 (0.0286)
Average Eigenvector Centrality		0.0470 (0.0675)	-0.0855 (0.0922)	-0.109 (0.0956)		0.129 (0.0825)	0.0881 (0.0985)	0.0232 (0.105)
R-squared	0.007	0.011	0.137	0.202	0.019	0.061	0.330	0.443
	Panel B: Self-Assessment Metric							
Average Inverse Distance	-0.0661*** (0.00951)	-0.0387*** (0.00918)	-0.0221*** (0.00607)	-0.00615 (0.0137)	-0.0737*** (0.00950)	-0.0414*** (0.00992)	-0.0280*** (0.00707)	-0.00756 (0.0132)
Average Degree		-0.00614*** (0.00194)	0.000118 (0.00340)	-0.000378 (0.00349)		-0.00961*** (0.00212)	-0.00257 (0.00309)	-0.00270 (0.00309)
Average Clustering Coefficient		-0.0357 (0.0275)	0.00741 (0.0304)	0.00846 (0.0304)		-0.0298 (0.0307)	-0.0132 (0.0288)	-0.0144 (0.0286)
Average Eigenvector Centrality		0.110 (0.0757)	0.0407 (0.105)	0.00455 (0.108)		0.129 (0.0825)	0.0881 (0.0985)	0.0232 (0.105)
R-squared	0.009	0.019	0.166	0.247	0.019	0.061	0.330	0.443
	Panel C: Simulations							
Average Inverse Distance	-0.246*** (0.00479)	-0.210*** (0.00595)	-0.194*** (0.00736)	-0.222*** (0.0126)	-0.516*** (0.0147)	-0.443*** (0.00915)	-0.386*** (0.0114)	-0.449*** (0.0239)
Average Degree		-0.00640*** (0.00160)	-0.00818*** (0.00277)	-0.00719** (0.00299)		-0.0112*** (0.00118)	-0.00897*** (0.00304)	-0.00694** (0.00295)
Average Clustering Coefficient		-0.121*** (0.0209)	-0.159*** (0.0220)	-0.162*** (0.0239)		-0.192*** (0.0320)	-0.264*** (0.0329)	-0.270*** (0.0328)
Average Eigenvector Centrality		-0.0148 (0.0531)	-0.110 (0.0716)	-0.0515 (0.0739)		-0.172 (0.111)	-0.399*** (0.103)	-0.264** (0.115)
R-squared	0.127	0.133	0.213	0.233	0.578	0.613	0.804	0.833
Demographic Controls	No	Yes	Yes	Yes	No	Yes	Yes	Yes
Hamlet Fixed Effects	No	No	Yes	Yes	No	No	Yes	Yes
Ranker Fixed Effects	No	No	No	Yes	No	No	No	Yes

Notes: This table provides an estimate of the correlation between the accuracy in ranking a pair of households in a hamlet and the characteristics of the households that are being ranked. In Panel A, the dependent variable is a dummy variable for whether household i ranks household j versus household k incorrectly based on using consumption as the metric of truth (the sample mean is 0.52). In Panel B, the self-assessment variable is the metric of truth (the sample mean is 0.46). The sample is comprised of 104,445 ranked pairs in Panel A and 103,425 in Panel B. Details of the simulation procedure for Panel C are contained in Appendix B. Standard errors are clustered by hamlet and are listed in parentheses. *** $p < 0.01$, ** $p < 0.05$, * $p < 0.1$.

APPENDIX F. TABLES WITHOUT DEMOGRAPHIC COVARIATES

TABLE F.1. Without Controls: Numerical Predictions on Correlation between Hamlet Network Characteristics and Hamlet Level Error Rates

	(1)	(2)	(3)	(4)	(5)	(6)	(7)
Average Degree	-0.0283*** (0.00233)						0.0802*** (0.00792)
Average Clustering		-0.349*** (0.0362)					0.430*** (0.0754)
Number of Households			0.00108*** (0.000263)				0.000610** (0.000283)
First eigenvalue $\lambda_1(A)$				-0.0279*** (0.00215)			-0.0546*** (0.00469)
Fraction of Nodes in Giant Component					-0.382*** (0.0240)		-0.749*** (0.0514)
Link Density						-0.543*** (0.0677)	-0.545*** (0.0922)
R-squared	0.182	0.126	0.028	0.245	0.281	0.107	0.483

Notes: Same as Table 7, without demographic controls. It reports the relationship between hamlet network characteristics and the error rate in ranking others in the hamlet. Columns 1-6 show univariate regressions, while column 7 reports the results from a multivariate regression. The sample comprises 631 hamlets. Results for error rates using simulated data, as described in Appendix B. Robust standard errors in parentheses, *** p<0.01, ** p<0.05, * p<0.1.

TABLE F.2. Without Controls: Numerical Predictions on Correlation between Hamlet Network Characteristics and Hamlet Level Error Rates

	(1)	(2)	(3)	(4)	(5)	(6)	(7)
<i>Panel A: Consumption Metric</i>							
Average Degree	-0.0200*** (0.00274)						0.0356*** (0.0112)
Average Clustering		-0.361*** (0.0406)					-0.359*** (0.0953)
Number of Households			0.000892*** (0.000301)				0.000305 (0.000391)
First eigenvalue $\lambda_1(A)$				-0.0168*** (0.00217)			-0.0211*** (0.00578)
Fraction of Nodes in Giant Component					-0.264*** (0.0300)		-0.205*** (0.0699)
Link Density						-0.349*** (0.0780)	0.108 (0.138)
R-squared	0.076	0.114	0.016	0.075	0.113	0.037	0.153
<i>Panel B: Self-Assessment Metric</i>							
Average Degree	-0.0276*** (0.00294)						0.0294** (0.0124)
Average Clustering		-0.495*** (0.0431)					-0.476*** (0.106)
Number of Households			0.00135*** (0.000337)				0.000266 (0.000418)
First eigenvalue $\lambda_1(A)$				-0.0206*** (0.00251)			-0.0165** (0.00660)
Fraction of Nodes in Giant Component					-0.355*** (0.0319)		-0.219*** (0.0779)
Link Density						-0.524*** (0.0816)	0.163 (0.148)
R-squared	0.115	0.170	0.029	0.090	0.161	0.066	0.198

Notes: Same as Table 9, without demographic controls. It reports the relationship between hamlet network characteristics and the error rate in ranking others in the hamlet. Columns 1-6 show univariate regressions, while column 7 reports the results from a multivariate regression. The sample comprises 631 hamlets. Panel A presents results for error rates using the consumption metric. Panel B presents results for error rates using the self-assessment metric. Robust standard errors in parentheses, *** p<0.01, ** p<0.05, * p<0.1.

APPENDIX G. ALTERNATIVE PARAMETERS

TABLE G.1. Numerical Predictions on Stochastic Dominance with Alternative Parameters

	(1)	(2)	(3)	(4)	(5)	(6)	(7)	(8)
	<i>Panel A: ($\alpha = 0.04, \tau = 3$)</i>				<i>Panel C: ($a = 0.36, \tau = 3$)</i>			
I fofd J	-0.116*** (0.0157)	-0.235*** (0.0243)	-0.119*** (0.0157)	-0.225*** (0.0247)	-0.124*** (0.0149)	-0.253*** (0.0230)	-0.125*** (0.0151)	-0.239*** (0.0230)
J fofd I	0.121*** (0.0172)		0.123*** (0.0174)		0.129*** (0.0165)		0.129*** (0.0162)	
Observations	199,396	147,460	199,396	147,460	199,396	147,460	199,396	147,460
	<i>Panel B: ($\alpha = 0.04, \tau = 5$)</i>				<i>Panel D: ($\alpha = 0.36, \tau = 5$)</i>			
I fofd J	-0.145*** (0.0163)	-0.281*** (0.0251)	-0.148*** (0.0167)	-0.277*** (0.0260)	-0.123*** (0.0162)	-0.227*** (0.0248)	-0.131*** (0.0163)	-0.226*** (0.0253)
J fofd I	0.134*** (0.0178)		0.141*** (0.0182)		0.0985*** (0.0180)		0.106*** (0.0179)	
Observations	199,396	147,460	199,396	147,460	199,396	147,460	199,396	147,460
Non-Comparable	Yes	No	Yes	No	Yes	No	Yes	No
Demographic Controls	No	No	Yes	Yes	No	No	Yes	Yes
Stratification Group FE	Yes	Yes	Yes	Yes	Yes	Yes	Yes	Yes

Notes: Same as Table 6, with alternative parameters generating the simulations.

TABLE G.2. Numerical Predictions on Stochastic Dominance with Alternative Parameters

	(1)	(2)	(3)	(4)	(5)	(6)	(7)	(8)	(9)	(10)	(11)	(12)	(13)	(14)
	Panel A: ($\alpha = 0.04$, $t = 3$)			Panel C: ($\alpha = 0.36$, $t = 3$)										
Average Degree	-0.0304*** (0.00267)						0.0836*** (0.00810)	-0.0312*** (0.00272)						0.0782*** (0.00814)
Average Clustering		-0.362*** (0.0433)					0.326*** (0.0753)		-0.374*** (0.0433)					0.319*** (0.0761)
Number of Households			0.000842*** (0.000251)				0.000520* (0.000272)			0.000963*** (0.000251)				0.000510* (0.000278)
First eigenvalue $\lambda_1(A)$				-0.0294*** (0.00236)			-0.0580*** (0.00503)			-0.0297*** (0.00233)				-0.0558*** (0.00499)
Fraction of Nodes in Giant Component					-0.422*** (0.0273)		-0.689*** (0.0528)				-0.429*** (0.0271)			-0.680*** (0.0531)
Link Density						-0.534*** (0.0730)	-0.633*** (0.0900)					-0.540*** (0.0760)		-0.573*** (0.0923)
R-squared	0.242	0.181	0.106	0.309	0.322	0.173	0.510	0.248	0.184	0.103	0.311	0.327	0.171	0.504
	Panel B: ($\alpha = 0.04$, $t = 5$)			Panel D: ($\alpha = 0.36$, $t = 5$)										
Average Degree	-0.0295*** (0.00278)						0.0752*** (0.00782)	-0.0290*** (0.00278)						0.0780*** (0.00786)
Average Clustering		-0.359*** (0.0431)					0.374*** (0.0722)		-0.355*** (0.0434)					0.372*** (0.0730)
Number of Households			0.000879*** (0.000252)				0.000253 (0.000300)			0.000897*** (0.000249)				0.000298 (0.000282)
First eigenvalue $\lambda_1(A)$				-0.0281*** (0.00241)			-0.0508*** (0.0462)			-0.0278*** (0.00238)				-0.0519*** (0.00455)
Fraction of Nodes in Giant Component					-0.446*** (0.0269)		-0.811*** (0.0490)				-0.442*** (0.0269)			-0.810*** (0.0493)
Link Density						-0.476*** (0.0734)	-0.424*** (0.0862)					-0.475*** (0.0757)		-0.463*** (0.0871)
R-squared	0.232	0.176	0.094	0.292	0.356	0.150	0.546	0.224	0.171	0.092	0.285	0.350	0.146	0.543

Notes: Same as Table 7, with alternative parameters generating the simulations.

APPENDIX H. GRAPHICAL RECONSTRUCTION

We now conduct a graphical reconstruction exercise where we integrate over the missing data in each network as described in [Chandrasekhar and Lewis \(2012\)](#). Before we get started, it is useful to establish some notation. Let $\mathbf{A}$ denote the adjacency matrix which is composed of a kin adjacency matrix, $\mathbf{K}$, and a social group matrix, $\mathbf{S}$. That is, entrywise we have

$$\mathbf{A} = 1 \{ \mathbf{K} + \mathbf{S} > 0 \}.$$

Next, let Γ denote the set of surveyed nodes and let Δ denote those on whom we have full kinship data. This means $\Gamma \subset \Delta$ but also that Δ includes all informal and formal leaders as well as those that were nominated by any of the surveyed households to be in the top or bottom 5 of the wealth distribution. This implies that we only fail to know K_{ij} if both $i, j \notin \Delta$. Meanwhile, we know S_{ij} so long as $i, j \in \Gamma$.

To describe what our current data looks like, let $\bar{\mathbf{A}}$, $\bar{\mathbf{K}}$ and $\bar{\mathbf{S}}$ denote the adjacency matrices we use in our analysis above. Let us take the example of the kinship network. Notice that

$$\bar{K}_{ij} = \begin{cases} 1 & \text{if } i \in \Delta \text{ or } j \in \Delta \text{ and } K_{ij} = 1 \\ 0 & \text{o.w.} \end{cases}$$

Crucially, this means that when $i, j \notin \Delta$ one cannot determine whether $K_{ij} = 1$ or $K_{ij} = 0$. (An analogous statement is true for $\bar{S}$ and Γ .)

Our goal is to now construct estimates of the regressors conditional on the observed data. Let $\mathbf{A}_r^{obs}$ denote the observed part of the adjacency matrix (where we definitively know whether there is a link present or not). Our goal is to construct

$$\mathbb{E} \left[W(\mathbf{A}_r) | \mathbf{A}_r^{obs} \right]$$

the expectation of the regressor W of interest (which can be an attribute of a node or the entire network) given the observed part of the data. Across independent networks, the estimator of the regression coefficient will be consistent under a correctly specified model.

To implement this, we do the following. For each hamlet r , we assume a 2-parameter model for the missing links: $(\hat{p}_r^{kin}, \hat{p}_r^{social})$. This is as if we are allowing $\mathbf{K}$ and $\mathbf{S}$ to be different Erdos-Renyi graphs ($\mathbf{A}$ is its union), with parameters that can vary hamlet-by-hamlet to allow for “hamlet fixed effects”. After estimating these parameters, we will use the parameters to reconstruct potential values of the missing data and average over them.

(1) Construct estimates $(\hat{p}_r^{kin}, \hat{p}_r^{social})$:

- Kinship network: for each hamlet we can use the rows of $\bar{K}_{i\cdot}$ for $i \in \Gamma$. This is a randomly selected set of nodes and therefore we can use this to estimate $\hat{p}_r^{kin}$.
- Social network: for each hamlet we can use $\bar{S}_{i,j}$ for $i, j \in \Gamma$. Again, this is a randomly selected set of pairs of nodes and therefore we can use the share of these that are linked to establish $\hat{p}_r^{social}$.

Once equipped with $(\hat{p}_r^{kin}, \hat{p}_r^{social})$, we proceed to integrating over the missing data.

(2) Averaging over missing data:

- (a) Construct a sequence of 500 matrices $\{\mathbf{A}^{*b}, \mathbf{K}^{*b}, \mathbf{S}^{*b}\}_{b=1}^B$ where $\mathbf{A}^{*b} = 1 \{\mathbf{K}^{*b} + \mathbf{S}^{*b} > 0\}$ and $B = 500$.
- (b) Construct regressors

$$\hat{\mathbb{E}} \left[W(\mathbf{A}_r) | \mathbf{A}_r^{obs}, \hat{p}_r^{kin}, \hat{p}_r^{social} \right] = \frac{1}{B} \sum_b W(\mathbf{A}_r^{*b}).$$

Note that these can be regressors for the within-village analysis (such as centralities of nodes or distances between nodes) or regressors for the network-level analysis (where they are features such as degree, the degree distribution, clustering, etc).

- (3) Run within-village analysis using $\hat{\mathbb{E}} \left[W(\mathbf{A}) | \mathbf{A}^{obs} \right]$ as our regressor.
 - (4) Estimate our GMM model.
 - For each $b = 1, \dots, B$ run S draws of the learning process
 - Begin with a graph $\mathbf{A}_r^{*b}$ as the underlying graph for network r . Compute empirical moments $m_{emp,b,r}$ for each b for each network r .
 - Generate $S = 50$ simulations as described in Appendix B.
 - Compute the expected deviation of the simulated method from the empirical moment
- $$d(r, \xi) := \frac{1}{B} \sum_{b \in [B]} \left\{ \frac{1}{S} \sum_{s \in [S]} m_{sim,b,r}(s, \theta) - m_{emp,b,r} \right\}.$$
- Minimize the objective function, which is the quadratic form of this deviation, as described in Appendix B. Standard errors are as described there, via the Bayesian bootstrap.
 - Generate synthetic outcome data.
 - For each $b = 1, \dots, B$ run $S = 50$ draws of the learning process at estimated parameters $(\hat{\alpha}, \hat{\tau})$ over $\mathbf{A}_r^{*b}$.
 - Compute error rates by averaging over the $B \times S$ draws per hamlet.
- (5) Run our village-level analysis using $\hat{\mathbb{E}} \left[W(\mathbf{A}) | \mathbf{A}^{obs} \right]$ as our regressor where our outcome data are either the empirical data or the synthetic data described above.

After conducting this exercise, we find that the results are broadly consistent with our original results. Tables H.1-H.7 report the results. Let us summarize the main findings.

- Tables H.1/H.2:

Higher degree and more central nodes tend to have lower error rates and are less likely to report not having an opinion. (The latter is true for clustering as well.) These results are mostly robust to the inclusion of hamlet fixed effects as well as a large set of demographic covariates. Overall, this is consistent with the simulated outcomes.
- Table H.3:

Nodes that are further from those that they are ranking are more likely to make mistakes and more likely to not know (not have an opinion). This is broadly true irrespective of demographic controls and hamlet fixed effects. The results are underpowered with ranker fixed effects. Overall, this is consistent with the simulated outcomes.

- Tables [H.4](#) and [H.6](#): Networks that tend to dominate other networks (in expectation now since we are integrating over missing data) tend to have lower error rates, both when we look at the synthetic outcome data as well as in the empirical data.
- Tables [H.5](#) and [H.7](#): Both in the simulations and in the empirical data we see that a higher degree reduces error rates, more clustering reduces error rates, a higher first eigenvalue reduces error rates, and a higher density reduces error rates. When looking at the moments conditional on each other, the average degree and the share of nodes in the giant component robustly matter in the simulations and in the data. Note a difference once we correct for sampling: the average degree no longer has the “wrong” sign.
- Tables [H.8](#) and [H.9](#): Both show that community targeting is differentially more effective relative to the PMT in more diffusive hamlets. Therefore, the main result of Section 6 is largely unchanged when we integrate over the missing data.

In sum, correcting for missing data in this way leaves the main results of the paper intact.

TABLE H.1. Graphical Reconstruction: The Correlation between Household Network Characteristics and the Error Rate in Ranking Income Status of Household

	(1)	(2)	(3)	(4)	(5)	(6)	(7)	(8)
<i>Panel A: Consumption Metric, Error Rate</i>								
Degree	-0.00790*** (0.00112)			-0.00823*** (0.00133)	-0.00175*** (0.000627)			-0.00161 (0.000982)
Clustering		-0.0668*** (0.0169)		-0.0668*** (0.0161)		-0.00963 (0.0105)		-0.0126 (0.0115)
Eigenvector Centrality			-0.0949* (0.0501)	0.0611 (0.0586)			-0.0632** (0.0272)	-0.0146 (0.0409)
R-squared	0.027	0.005	0.002	0.033	0.668	0.667	0.668	0.668
<i>Panel B: Self-Assessment Metric, Error Rate</i>								
Degree	-0.00986*** (0.00129)			-0.0101*** (0.00154)	-0.00270*** (0.000745)			-0.00205 (0.00134)
Clustering		-0.0829*** (0.0184)		-0.0807*** (0.0174)		0.00411 (0.0125)		0.000974 (0.0136)
Eigenvector Centrality			-0.135** (0.0568)	0.0560 (0.0663)			-0.0907*** (0.0322)	-0.0336 (0.0548)
R-squared	0.034	0.007	0.004	0.040	0.671	0.670	0.671	0.671
<i>Panel C: Share of Don't Knows</i>								
Degree	-0.0118*** (0.00133)			-0.0122*** (0.00157)	-0.00275*** (0.000652)			-0.00202* (0.00106)
Clustering		-0.0985*** (0.0218)		-0.0963*** (0.0205)		-0.00115 (0.0130)		-0.00381 (0.0139)
Eigenvector Centrality			-0.142** (0.0575)	0.0903 (0.0667)			-0.0978*** (0.0304)	-0.0409 (0.0466)
R-squared	0.051	0.010	0.004	0.060	0.712	0.711	0.712	0.712
Hamlet Fixed Effect	No	No	No	No	Yes	Yes	Yes	Yes

Notes: Same as Table 2.

TABLE H.2. Graphical Reconstruction: The Correlation between Household Network Characteristics and the Error Rate in Ranking Income Status of Households

	(1)	(2)	(3)	(4)	(5)	(6)	(7)	(8)
<i>Panel A: Consumption Metric, Error Rate</i>								
Degree	-0.00687*** (0.00108)			-0.00713*** (0.00130)	-0.00130** (0.000628)			-0.00118 (0.000976)
Clustering		-0.0523*** (0.0164)		-0.0534*** (0.0159)		-0.00888 (0.0104)		-0.0109 (0.0114)
Eigenvector Centrality			-0.0924* (0.0492)	0.0410 (0.0581)			-0.0486* (0.0272)	-0.0125 (0.0407)
R-squared	0.044	0.028	0.026	0.048	0.671	0.671	0.671	0.671
<i>Panel B: Self-Assessment Metric, Error Rate</i>								
Degree	-0.00821*** (0.00125)			-0.00831*** (0.00149)	-0.00209*** (0.000737)			-0.00149 (0.00132)
Clustering		-0.0603*** (0.0178)		-0.0586*** (0.0171)		0.00487 (0.0122)		0.00287 (0.0134)
Eigenvector Centrality			-0.130** (0.0551)	0.0231 (0.0652)			-0.0708** (0.0316)	-0.0301 (0.0541)
R-squared	0.067	0.047	0.047	0.070	0.676	0.675	0.676	0.676
<i>Panel C: Share of Don't Knows</i>								
Degree	-0.00953*** (0.00119)			-0.00992*** (0.00146)	-0.00198*** (0.000600)			-0.00111 (0.00103)
Clustering		-0.0581*** (0.0198)		-0.0597*** (0.0191)		0.00524 (0.0126)		0.00457 (0.0135)
Eigenvector Centrality			-0.119** (0.0535)	0.0584 (0.0645)			-0.0727** (0.0297)	-0.0431 (0.0474)
R-squared	0.096	0.067	0.066	0.100	0.722	0.722	0.722	0.722
Hamlet Fixed Effect	No	No	No	No	Yes	Yes	Yes	Yes

Notes: Same as Table 3.

TABLE H.3. Graphical Reconstruction: The Correlation Between Inaccuracy in Ranking a Pair of Households in a Hamlet and the Average Distance to Rankees

	(1)	(2)	(3)	(4)
<i>Panel A: Consumption Metric, Error Rate</i>				
Inverse of the Distance	-0.0553*** (0.00919)	-0.0331*** (0.00822)	-0.0174*** (0.00554)	-0.01000 (0.00894)
Average Degree		-0.00573*** (0.00165)	0.00551* (0.00330)	0.00543 (0.00337)
Average Clustering Coefficient		-0.0384 (0.0284)	0.0314 (0.0305)	0.0292 (0.0308)
Average Eigenvector Centrality		0.0648 (0.0797)	-0.247** (0.119)	-0.270** (0.122)
R-squared	0.005	0.011	0.137	0.203
<i>Panel B: Self-Assessment Metric, Error Rate</i>				
Inverse of the Distance	-0.0625*** (0.0102)	-0.0335*** (0.00903)	-0.0136** (0.00589)	0.000397 (0.00990)
Average Degree		-0.00603*** (0.00184)	0.000466 (0.00358)	-0.000308 (0.00366)
Average Clustering Coefficient		-0.0763*** (0.0287)	-0.000133 (0.0327)	0.000232 (0.0329)
Average Eigenvector Centrality		0.0772 (0.0871)	-0.000318 (0.140)	-0.0168 (0.142)
R-squared	0.006	0.018	0.166	0.247
<i>Panel C: Share of Don't Knows</i>				
Inverse of the Distance	-0.0665*** (0.0100)	-0.0345*** (0.0101)	-0.0226*** (0.00658)	-0.00945 (0.00975)
Average Degree		-0.00946*** (0.00210)	-0.00186 (0.00305)	-0.00225 (0.00308)
Average Clustering Coefficient		-0.0954*** (0.0344)	-0.0194 (0.0329)	-0.0200 (0.0328)
Average Eigenvector Centrality		0.112 (0.0981)	0.0530 (0.129)	0.0132 (0.132)
R-squared	0.012	0.060	0.331	0.445
Demographic Controls	No	Yes	Yes	Yes
Hamlet Fixed Effects	No	No	Yes	Yes
Ranker Fixed Effects	No	No	No	Yes

Notes: Same as Table 4.

TABLE H.4. Graphical Reconstruction: Numerical Predictions on Stochastic Dominance

	(1)	(2)	(3)	(4)
I fofd J	-0.126*** (0.00379)	-0.224*** (0.00281)	-0.105*** (0.00378)	-0.190*** (0.00283)
J fofd I	0.0938*** (0.00411)		0.0971*** (0.00417)	
Observations	199,396	147,460	199,396	147,460
Non-Comparable	Yes	No	Yes	No
Demographic Controls	No	No	Yes	Yes
Stratification Group FE	Yes	Yes	Yes	Yes

Notes: Same as Table 6.

TABLE H.5. Graphical Reconstruction: Numerical Predictions on Correlation between Hamlet Network Characteristics and Hamlet Level Error Rate

	(1)	(2)	(3)	(4)	(5)	(6)	(7)
Average Degree	-0.0304*** (0.00341)						-0.0327*** (0.00711)
Average Clustering		-0.214*** (0.0666)					0.0904 (0.0699)
Number of Households			0.000826** (0.000345)				0.00250*** (0.000399)
First eigenvalue $\lambda_1(A)$				-0.0209*** (0.00356)			0.00493 (0.00709)
Fraction of Nodes in Giant Component					-1.292*** (0.0802)		-1.051*** (0.0745)
Link Density						-0.0286*** (0.00435)	0.00643 (0.00505)
R-squared	0.236	0.023	0.010	0.154	0.475	0.122	0.598

Notes: Same as Table 7.

TABLE H.6. Graphical Reconstruction: Empirical Results on Stochastic Dominance

	(1)	(2)	(3)	(4)
<i>Panel A: Consumption Metric</i>				
I fofd J	-0.0986*** (0.0261)	-0.118*** (0.0266)	-0.102*** (0.0260)	-0.125*** (0.0264)
J fofd I	0.0141 (0.0243)		0.0349 (0.0231)	
Observations	196,878	145,697	196,878	145,697
<i>Panel B: Self-Assessment Metric</i>				
I fofd J	-0.0814*** (0.0225)	-0.106*** (0.0237)	-0.0769*** (0.0228)	-0.0997*** (0.0239)
J fofd I	0.00517 (0.0225)		0.0236 (0.0214)	
Observations	196,878	145,697	196,878	145,697
Non-Comparable	Yes	No	Yes	No
Demographic Controls	No	No	Yes	Yes
Stratification Group FE	Yes	Yes	Yes	Yes

Notes: Same as Table 8.

TABLE H.7. Graphical Reconstruction: Empirical Results on Correlation between Hamlet Network Characteristics and Hamlet Level Error Rate

	(1)	(2)	(3)	(4)	(5)	(6)	(7)
<i>Panel A: Consumption Metric</i>							
Average Degree	-0.0107*** (0.00204)						-0.0215*** (0.00818)
Average Clustering		-0.136*** (0.0504)					0.0916 (0.0732)
Number of Households			0.000854*** (0.000294)				0.00182*** (0.000421)
First eigenvalue $\lambda_1(A)$				-0.00727*** (0.00165)			0.00515 (0.00782)
Fraction of Nodes in Giant Component					-0.270*** (0.0496)		-0.119** (0.0549)
Link Density						-0.0110*** (0.00251)	0.00202 (0.00534)
R-squared	0.128	0.102	0.108	0.115	0.119	0.114	0.179
<i>Panel B: Self-Assessment Metric</i>							
Average Degree	-0.0103*** (0.00232)						-0.0143 (0.00899)
Average Clustering		-0.205*** (0.0543)					0.0456 (0.0791)
Number of Households			0.00119*** (0.000325)				0.00193*** (0.000436)
First eigenvalue $\lambda_1(A)$				-0.00726*** (0.00187)			0.00225 (0.00854)
Fraction of Nodes in Giant Component					-0.319*** (0.0514)		-0.206*** (0.0563)
Link Density						-0.0131*** (0.00273)	-0.00178 (0.00572)
R-squared	0.169	0.162	0.167	0.161	0.173	0.168	0.228

Notes: Same as Table 9.

TABLE H.8. Graphical Reconstruction: Rank Correlation on Targeting Type Interacted with Diffusiveness (Principal Component)

	(1)	(2)	(3)	(4)	(5)	(6)
<i>Panel A: Rank Correlation (Consumption)</i>						
Community x Diffusiveness		-0.0680 (0.119)	-0.0686 (0.119)	-0.0940 (0.124)	-0.0936 (0.126)	
Hybrid x Diffusiveness		-0.0425 (0.116)	-0.0452 (0.117)	-0.0632 (0.126)		
Community	-0.0588* (0.0319)	-0.0217 (0.0652)	-0.0176 (0.0651)	-0.00804 (0.0677)	-0.00370 (0.0678)	
Hybrid	-0.0614* (0.0327)	-0.0362 (0.0683)	-0.0314 (0.0688)	-0.0180 (0.0760)		
Diffusiveness		-0.0143 (0.0755)	0.0119 (0.0781)	0.0730 (0.0929)	0.0973 (0.105)	0.0727 (0.0927)
(Community or Hybrid) x Diffusiveness						-0.0766 (0.104)
(Community or Hybrid)						-0.0134 (0.0597)
R-squared	0.014	0.009	0.013	0.095	0.152	0.095
<i>Panel B: Rank Correlation (Self-Assessment)</i>						
Community x Diffusiveness		0.198* (0.113)	0.197* (0.113)	0.156 (0.119)	0.143 (0.121)	
Hybrid x Diffusiveness		0.184 (0.114)	0.181 (0.115)	0.152 (0.123)		
Community	0.108*** (0.0321)	0.0181 (0.0682)	0.0250 (0.0676)	0.0464 (0.0710)	0.0503 (0.0724)	
Hybrid	0.0839** (0.0331)	-0.00459 (0.0702)	0.00276 (0.0703)	0.0185 (0.0775)		
Diffusiveness		-0.195** (0.0808)	-0.149* (0.0832)	-0.164 (0.101)	-0.158 (0.110)	-0.162 (0.100)
(Community or Hybrid) x Diffusiveness						0.148 (0.105)
(Community or Hybrid)						0.0355 (0.0643)
R-squared	0.033	0.030	0.044	0.137	0.176	0.135
Stratification Group FE	No	No	No	Yes	Yes	Yes
Demographic Covariates	No	No	No	Yes	Yes	Yes

Notes: Same as Table 10.

TABLE H.9. Graphical Reconstruction: Rank Correlation on Targeting Type Interacted with Diffusiveness (1 - Simulated Error Rate)

	(1)	(2)	(3)	(4)	(5)	(6)
<i>Panel A: Rank Correlation (Consumption)</i>						
Community x Diffusiveness		0.0472 (0.110)	0.0228 (0.114)	0.0224 (0.116)	0.0506 (0.104)	
Hybrid x Diffusiveness		0.0249 (0.109)	-0.0352 (0.112)	-0.0476 (0.114)		
Community	-0.0588* (0.0319)	-0.0744 (0.0636)	-0.0672 (0.0665)	-0.0722 (0.0686)	-0.0587 (0.0618)	
Hybrid	-0.0614* (0.0327)	-0.0708 (0.0592)	-0.0422 (0.0619)	-0.0296 (0.0638)		
Diffusiveness		-0.0736 (0.0733)	-0.0132 (0.0841)	-0.00715 (0.0854)	-0.0380 (0.0700)	-0.00746 (0.0851)
(Community or Hybrid) x Diffusiveness						-0.0135 (0.0977)
(Community or Hybrid)						-0.0503 (0.0558)
R-squared	0.014	0.015	0.083	0.091	0.086	0.090
<i>Panel B: Rank Correlation (Self-Assessment)</i>						
Community x Diffusiveness		0.256** (0.118)	0.260** (0.121)	0.280** (0.124)	0.174* (0.105)	
Hybrid x Diffusiveness		0.246** (0.118)	0.220* (0.117)	0.195 (0.119)		
Community	0.108*** (0.0321)	-0.0170 (0.0662)	-0.0209 (0.0682)	-0.0305 (0.0706)	-0.0264 (0.0598)	
Hybrid	0.0839** (0.0331)	-0.0311 (0.0653)	-0.0199 (0.0665)	-0.00324 (0.0687)		
Diffusiveness		-0.269*** (0.0837)	-0.257*** (0.0984)	-0.251** (0.102)	-0.140* (0.0775)	-0.250** (0.101)
(Community or Hybrid) x Diffusiveness						0.236** (0.106)
(Community or Hybrid)						-0.0162 (0.0608)
R-squared	0.033	0.050	0.115	0.135	0.117	0.133
Stratification Group FE	No	No	No	Yes	Yes	Yes
Demographic Covariates	No	No	No	Yes	Yes	Yes

Notes: Same as Table 11.

APPENDIX I. AUGMENTING NETWORK DATA

Another approach to investigating the degree to which the missing data is an issue is to collect new data. Since one might be concerned about the missing 10 percent of kinship links, particularly since they are a non-random set, we decided to collect new data.

To investigate the degree to which this is an issue, we randomly selected 10 villages in Java, and went back and revisited these villages in the field. In these 10 villages, we obtained data on links in the study hamlet through a key informant survey: we sat with the neighborhood head, his spouse, and/or another local leaders who the neighborhood head viewed as knowledgeable about the community. We then walked them through the full list of households one by one, and asked him to enumerate the full set of kin links for the entire network in the hamlet. On average, this procedure took 1.6 hours per village.

Why is this useful? From our previous discussion, we know that we are missing about 10% of our kinship links. Returning to the field allows us to attempt to “fill in” these links, though we do again stress that we likely have information on 90% of the potential kinship links.

Of course, since we enlist the neighborhood head to go through and enumerate everyone’s kin, surely he wouldn’t do a complete enumeration. He does a good job though: 74% of the kin he named were indeed kin in our baseline and at the same time, for each individual who we directly surveyed, he named about 1/3 of their kin. This suggests that he isn’t adding much noise but at the same time will help us get about a 1/3 of the missing 10%. In short, this exercise puts us at having 93% of our kinship data.

For these 10 hamlets, we now augment our old kinship data with the new data. We then conduct two exercises. First, we use this augmented network directly. That is, we add any links that were reported in 2015 but not in 2007. Second, we compute the network using this augmented data but only using the augmented information on those individuals who we would have observed under our old sampling scheme. This means that we are not updating the links for someone who was not directly survey nor named in any of the listing exercises.

Taken together, both of these exercises holds the data fixed (augments the 2007 sample with the 2015 update) but varies the sampling scheme. We then compare the results to the original datasets. We expect that little should change given that we went from having about 90% of links to 93% of links.

Specifically, we replicate the within-village regressions (Tables 2, 3, and 4) using these two sets of networks. We do not replicate the cross-village regressions since we only have the new complete data for 10 hamlets. These are shown below, in Tables I.1-I.3.

What is apparent from these tables is that, holding the data fixed, sampling makes virtually no difference.

Let us discuss Table I.1 in detail, but the reader can easily verify that the same is true turning to Table I.2 (which just includes demographic covariates) and Table I.3 (which looks at distance of ranker to rankees). Columns 1-8 use the 2015-augmented kinship data, Columns 7-16 use the 2015-augmented data, restricted to our previous sample, and Columns 17-24 use our original 2007 data. For example, when we compare columns 1, 9 and 17 we see that the results are nearly identical when we look at the degree of the ranking household. Similarly, looking at columns 4, 12,

and 20 shows that even when we take degree, clustering and centrality conditional on each other, the results are qualitatively (and quantitatively) very similar.

In general, by comparing columns x , $x+8$ and $x+16$ in Tables I.1 and I.2 or columns x , $x+4$ and $x+8$ in Table I.3, we can see that the results are essentially identical when we use the full dataset and when we use the same data but only information on the nodes we sample and very similar to our original data. As discussed above, we believe that at an intuitive level the reason the sampling does not change the results is that (i) we already had complete kinship data on almost 70 percent of the network, (ii) for the nodes we are particularly interested in, we had an even higher proportion of the links, and (iii) we nearly had 90 percent of all kinship links overall.

TABLE I.1. The Correlation between Household Network Characteristics and the Error Rate in Ranking Income Status of Households

	(1)	(2)	(3)	(4)	(5)	(6)	(7)	(8)	(9)	(10)	(11)	(12)	(13)	(14)	(15)	(16)
	Outcome Variable: Consumption Metric, Error Rate								Outcome Variable: Self-Assessment Metric, Error Rate							
	Panel A. New Data															
Degree	-0.0185* (0.00922)			-0.0448*** (0.0141)	-0.000264 (0.00559)			-0.00189 (0.0173)	-0.0173* (0.00844)			-0.0357* (0.0171)	0.000264 (0.00461)			-0.00113 (0.0158)
Clustering		-0.00917 (0.103)		-0.432*** (0.106)		-0.0299 (0.0604)		-0.0583 (0.129)		0.0824 (0.116)		-0.243 (0.139)		-0.0127 (0.0853)		-0.0237 (0.186)
Eigenvector Centrality			-0.257 (0.405)	0.544 (0.596)			-0.0848 (0.207)	-0.0840 (0.450)			-0.0957 (0.550)	0.608 (0.808)			-0.00939 (0.160)	0.00293 (0.429)
R-squared	0.048	0.000	0.005	0.109	0.811	0.811	0.811	0.812	0.043	0.004	0.001	0.071	0.819	0.819	0.819	0.819
	Panel B. New Data, Old Sampling															
Degree	-0.0185* (0.00922)			-0.0451*** (0.0140)	-0.000264 (0.00559)			-0.00209 (0.0172)	-0.0173* (0.00844)			-0.0359* (0.0170)	0.000264 (0.00461)			-0.000769 (0.0158)
Clustering		-0.00959 (0.103)		-0.435*** (0.105)		-0.0303 (0.0605)		-0.0601 (0.129)		0.0825 (0.116)		-0.244 (0.138)		-0.0118 (0.0853)		-0.0200 (0.186)
Eigenvector Centrality			-0.250 (0.403)	0.557 (0.591)			-0.0848 (0.207)	-0.0790 (0.448)			-0.0883 (0.546)	0.619 (0.802)			-0.00921 (0.160)	-0.00490 (0.428)
R-squared	0.048	0.000	0.004	0.110	0.811	0.811	0.811	0.812	0.043	0.004	0.001	0.072	0.819	0.819	0.819	0.819
	Panel C. Old Data															
Degree	-0.0239** (0.00864)			-0.0520*** (0.0128)	-0.00230 (0.00455)			-0.000942 (0.0150)	-0.0199** (0.00801)			-0.0452*** (0.0124)	0.0000511 (0.00525)			-0.00382 (0.0164)
Clustering		0.0169 (0.120)		-0.333*** (0.0872)		-0.00630 (0.0454)		-0.0189 (0.0980)		0.0727 (0.109)		-0.226** (0.101)		-0.0031 (0.0642)		-0.0246 (0.137)
Eigenvector Centrality			-0.266 (0.339)	0.916 (0.597)			-0.128 (0.180)	-0.110 (0.488)			0.0023 (0.517)	1.046 (0.766)			0.0123 (0.176)	0.115 (0.54)
R-squared	0.066	0.000	0.005	0.123	0.811	0.811	0.812	0.812	0.047	0.004	0.000	0.094	0.819	0.819	0.819	0.819
Hamlet Fixed Effect	No	No	No	No	Yes	Yes	Yes	Yes	No	No	No	No	Yes	Yes	Yes	Yes

Notes: Same as Table 2. The sample is restricted to the 10 villages where we collected new data in 2015 to augment our old data.

TABLE I.2. The Correlation between Household Network Characteristics and the Error Rate in Ranking Income Status of Households

	(1)	(2)	(3)	(4)	(5)	(6)	(7)	(8)	(9)	(10)	(11)	(12)	(13)	(14)	(15)	(16)
	Outcome Variable: Consumption Metric, Error Rate								Outcome Variable: Self-Assessment Metric, Error Rate							
	Panel A. New Data															
Degree	-0.0111 (0.00701)			-0.0275 (0.0270)	-0.00113 (0.00558)			-0.00397 (0.0133)	-0.00828 (0.00690)			0.00215 (0.0286)	-9.58e-05 (0.00430)			-0.00246 (0.0154)
Clustering		-0.0556 (0.0786)		-0.318 (0.210)		-0.0342 (0.0578)		-0.0834 (0.0954)		0.0411 (0.0933)		0.0206 (0.234)		-0.0131 (0.0770)		-0.0354 (0.168)
Eigenvector Centrality			-0.587* (0.285)	-0.00436 (0.817)			-0.144 (0.217)	-0.0973 (0.393)			-0.514 (0.494)	-0.562 (1.092)			-0.0265 (0.148)	0.0186 (0.422)
R-squared	0.146	0.134	0.154	0.177	0.819	0.819	0.820	0.821	0.170	0.163	0.179	0.179	0.823	0.823	0.823	0.823
	Panel B. New Data, Old Sampling															
Degree	-0.0111 (0.00701)			-0.0281 (0.0269)	-0.00113 (0.00558)			-0.00419 (0.0131)	-0.00828 (0.00690)			0.00181 (0.0286)	-9.58e-05 (0.00430)			-0.00209 (0.0153)
Clustering		-0.0558 (0.0786)		-0.323 (0.210)		-0.0346 (0.0580)		-0.0854 (0.0948)		0.0416 (0.0932)		0.0194 (0.235)		-0.0121 (0.0771)		-0.0316 (0.168)
Eigenvector Centrality			-0.581* (0.282)	0.0194 (0.809)			-0.143 (0.217)	-0.0914 (0.392)			-0.506 (0.489)	-0.546 (1.086)			-0.0262 (0.148)	0.0106 (0.421)
R-squared	0.146	0.135	0.154	0.177	0.819	0.819	0.820	0.821	0.170	0.163	0.179	0.179	0.823	0.823	0.823	0.823
	Panel C. Old Data															
Degree	-0.0166* (0.00890)			-0.0387 (0.0402)	-0.00302 (0.00455)			-0.00315 (0.0139)	-0.0100 (0.00790)			-0.0109 (0.0386)	-0.000167 (0.00444)			-0.00580 (0.0162)
Clustering		-0.0365 (0.104)		-0.274 (0.212)		-0.0164 (0.0474)		-0.0438 (0.0945)		0.0174 (0.0873)		-0.0555 (0.227)		-0.00659 (0.0525)		-0.0405 (0.117)
Eigenvector Centrality			-0.589** (0.220)	0.449 (1.300)			-0.181 (0.190)	-0.108 (0.511)			-0.390 (0.467)	-0.0910 (1.416)			0.00263 (0.145)	0.158 (0.512)
R-squared	0.158	0.134	0.155	0.182	0.819	0.819	0.820	0.821	0.172	0.163	0.172	0.174	0.823	0.823	0.823	0.823
Hamlet Fixed Effect	No	No	No	No	Yes	Yes	Yes	Yes	No	No	No	No	Yes	Yes	Yes	Yes

Notes: Same as Table 3. The sample is restricted to the 10 villages where we collected new data in 2015 to augment our old data.

Notes: Same as Table 3. The sample is restricted to the 10 villages where we collected new data in 2015 to augment our old data.

TABLE I.3. The Correlation Between Inaccuracy in Ranking a Pair of Households in a Hamlet and the Average Inverse Distance to Ranker

(1)	(2)	(3)	(4)	(5)	(6)	(7)	(8)	(9)	(10)	(11)	(12)
<i>Panel A: Consumption Metric, Error Rate</i>											
Average Inverse Distance	-0.165* (0.0784)	<i>New Data</i>		-0.165* (0.0782)	<i>New Data, Old Sampling</i>		-0.123 (0.115)	-0.148** (0.0564)	<i>Old Data</i>		-0.150 (0.0877)
		-0.142* (0.0748)	-0.0370 (0.0527)		-0.144* (0.0739)	-0.0376 (0.0523)			-0.141** (0.0493)	-0.0405 (0.0311)	
Average Degree	-0.000230 (0.0370)	0.349 (0.0658)	0.0312 (0.0654)	-0.000313 (0.0372)	<i>New Data, Old Sampling</i>		0.0323 (0.0645)	-0.00288 (0.0501)	<i>Old Data</i>		0.0625 (0.0551)
		0.237 (0.395)	0.603 (0.488)		0.0739 (0.399)	0.0523 (0.483)			0.0493 (0.312)	0.0625 (0.300)	
Average Clustering Coefficient	-0.396 (0.849)	-0.225 (1.488)	0.184 (1.527)	-0.374 (0.847)	<i>New Data, Old Sampling</i>		0.154 (1.504)	-0.0846 (1.492)	<i>Old Data</i>		-0.722 (1.964)
		0.012 (0.136)	0.202 (0.202)		0.0739 (0.399)	0.0523 (0.483)			0.0493 (0.312)	0.0625 (0.300)	
R-squared	0.007			0.053	0.058	0.193	0.231	0.054	0.055	0.190	0.228
<i>Panel B: Self-Assessment Metric, Error Rate</i>											
Average Inverse Distance	-0.104 (0.0869)	<i>New Data</i>		-0.103 (0.0867)	<i>New Data, Old Sampling</i>		-0.0839 (0.285)	-0.0664 (0.0746)	<i>Old Data</i>		-0.134 (0.197)
		-0.0929 (0.0668)	-0.00445 (0.0705)		-0.0912 (0.0682)	-0.00370 (0.0707)			-0.0677 (0.0540)	0.0126 (0.0359)	
Average Degree	0.0294 (0.0253)	0.0209 (0.0512)	0.0156 (0.0526)	0.0281 (0.0254)	<i>New Data, Old Sampling</i>		0.0128 (0.0540)	-0.0284 (0.0201)	<i>Old Data</i>		0.0625 (0.0364)
		0.608 (0.369)	0.584 (0.671)		0.0281 (0.374)	0.0180 (0.682)			0.0284 (0.242)	0.0631 (0.341)	
Average Clustering Coefficient	-1.310 (0.945)	-0.0120 (1.351)	0.430 (1.844)	-1.278 (0.949)	<i>New Data, Old Sampling</i>		0.487 (1.851)	-0.0202 (0.983)	<i>Old Data</i>		-0.546 (1.792)
		0.079 (0.196)	0.234 (0.234)		0.079 (0.196)	0.0507 (1.367)			0.0202 (0.983)	0.190 (1.687)	
R-squared	0.063			0.063	0.078	0.196	0.234	0.002	0.020	0.190	0.228
Demographic Controls	No	Yes	Yes	No	Yes	Yes	Yes	No	Yes	Yes	Yes
Hamlet Fixed Effects	No	No	Yes	No	No	Yes	Yes	No	No	Yes	Yes
Ranker Fixed Effects	No	No	No	No	No	No	Yes	No	No	No	Yes

Notes: Same as Table 4. The sample is restricted to the 10 villages where we collected new data in 2015 to augment our old data.

APPENDIX J. DROPPING FURTHER DATA

Here we show that, qualitatively, our results do not appear to get much worse if we were to use an even sparser network sampled in the same way. Specifically, we show that neither dropping 25% of links randomly nor sampling 6 people instead of 8 substantially alters our results. This, again, suggests that our results are not too sensitive to having a partial network structure.

In order to operationalize this, for each network we do the following.

- For $b = 1, \dots, B$
 - Exercise 1: drop 25% of links uniformly at random. This generates a new adjacency matrix $\mathbf{A}_r^b$ for each network r .
 - Exercise 2: select two households (out of the surveyed households) uniformly at random and drop their links. Also drop their survey responses where they nominate the 5 poorest, 5 richest, and elites as well as all of the kin of these folk. This generates a new adjacency matrix $\mathbf{A}_r^b$ for each network r .
- For both exercises, construct regressors $W(\mathbf{A}_r^b)$ for each draw for each network and then construct an average, integrating over the missing data

$$\hat{\mathbb{E}}[W(\mathbf{A}_r)] := \frac{1}{B} \sum_{b=1}^B W(\mathbf{A}_r^b)$$

which we then use in our regressions.

In this way, we then rerun our analysis from the main part of the paper, constructing regressors from this data, simulating the network learning process on this subgraph, etc. The goal is to document that the qualitative results, in this case, are robust to this procedure. In practice we set $B = 100$.

J.1. Dropping 25% of links uniformly at random.

- Tables J.1/J.2:
We see that higher degree and more (eigenvector) central nodes have lower error rates. Further, the results typically hold even when adding demographic controls and are mostly robust to the inclusion of hamlet fixed effects. This is consistent with our main results.
- Table J.3:
When nodes are further on average from those whom they are ranking, they are more likely to make a mistake. This is robust to including demographic controls and hamlet fixed effects. Again the results are underpowered with ranker fixed effects. This is consistent with our main results.
- Table J.4:
Networks that have degree distributions that (first order stochastically) dominate other networks tend to have lower error rates. Again, this is consistent with our main results.
- Table J.5:
We find that a higher degree, more clustering, a higher first eigenvalue, and a higher density, all are associated with a lower error rate. This is consistent with our main results.
- Tables J.6 and J.7:

Here we look at whether community targeting does better relative to PMT in more “diffusive” hamlets where this is computed either via principal components (Table J.6) or by using the simulated error rate in these hamlets (Table J.7). We find that being in a more diffusive hamlet (or less error-prone hamlet under our model) corresponds to community targeting being more effective than the PMT. This is consistent with our main results.

TABLE J.1. After Dropping 25% of links: The Correlation between Household Network Characteristics and the Error Rate in Ranking Income Status of Household

	(1)	(2)	(3)	(4)	(5)	(6)	(7)	(8)
<i>Panel A: Consumption Metric, Error Rate</i>								
Degree	-0.0163*** (0.00257)			-0.0226*** (0.00373)	-0.00319** (0.00145)			-0.00220 (0.00213)
Clustering		-0.0314 (0.0193)		-0.0104 (0.0192)		-0.00973 (0.00945)		-0.00453 (0.00998)
Eigenvector Centrality			-0.0419** (0.0205)	0.100*** (0.0329)			-0.0255* (0.0132)	-0.0101 (0.0202)
R-squared	0.016	0.001	0.001	0.019	0.668	0.668	0.668	0.668
<i>Panel B: Self-Assessment Metric, Error Rate</i>								
Degree	-0.0158*** (0.00289)			-0.0212*** (0.00432)	-0.00355** (0.00154)			-0.00227 (0.00237)
Clustering		-0.0366* (0.0196)		-0.0160 (0.0195)		-0.0207* (0.0105)		-0.0154 (0.0112)
Eigenvector Centrality			-0.0479** (0.0224)	0.0878** (0.0370)			-0.0302** (0.0142)	-0.0104 (0.0225)
R-squared	0.056	0.046	0.046	0.058	0.676	0.676	0.676	0.676
<i>Panel C: Share of Don't Knows</i>								
Degree	-0.0244*** (0.00298)			-0.0354*** (0.00440)	-0.00508*** (0.00137)			-0.00389* (0.00207)
Clustering		-0.0254 (0.0212)		0.00652 (0.0205)		-0.00896 (0.00968)		-0.000934 (0.0101)
Eigenvector Centrality			-0.0529** (0.0224)	0.161*** (0.0371)			-0.0379*** (0.0134)	-0.0133 (0.0207)
R-squared	0.029	0.001	0.001	0.037	0.715	0.714	0.715	0.715
Hamlet Fixed Effect	No	No	No	No	Yes	Yes	Yes	Yes

Notes: Same as Table 2.

TABLE J.2. After Dropping 25% of links: The Correlation between Household Network Characteristics and the Error Rate in Ranking Income Status of Households

	(1)	(2)	(3)	(4)	(5)	(6)	(7)	(8)
<i>Panel A: Consumption Metric, Error Rate</i>								
Degree	-0.00666*** (0.00164)			-0.00750*** (0.00225)	-0.00245** (0.000981)			0.00104 (0.00215)
Clustering		0.00889 (0.0240)		-0.0277 (0.0264)		-0.00399 (0.0154)		0.00308 (0.0189)
Eigenvector Centrality			-0.0432 (0.0641)	0.0674 (0.0843)			-0.128*** (0.0385)	-0.161** (0.0790)
R-squared	0.033	0.016	0.016	0.034	0.688	0.687	0.689	0.689
<i>Panel B: Self-Assessment Metric, Error Rate</i>								
Degree	-0.00752*** (0.00194)			-0.00851*** (0.00262)	-0.00345*** (0.00105)			-0.000727 (0.00258)
Clustering		0.0202 (0.0249)		-0.0228 (0.0266)		0.00476 (0.0175)		0.00322 (0.0208)
Eigenvector Centrality			-0.0258 (0.0695)	0.0962 (0.0943)			-0.148*** (0.0494)	-0.125 (0.107)
R-squared	0.098	0.066	0.069	0.100	0.679	0.677	0.678	0.679
<i>Panel B: Self-Assessment Metric, Error Rate</i>								
Degree	-0.0185*** (0.00274)			-0.0267*** (0.00414)	-0.00379*** (0.00131)			-0.00252 (0.00200)
Clustering		-0.0154 (0.0194)		0.00908 (0.0191)		-0.00716 (0.00976)		-0.000553 (0.0100)
Eigenvector Centrality			-0.0431** (0.0210)	0.117*** (0.0357)			-0.0304** (0.0132)	-0.0144 (0.0203)
R-squared	0.081	0.065	0.065	0.086	0.725	0.724	0.725	0.725
Hamlet Fixed Effect	No	No	No	No	Yes	Yes	Yes	Yes

Notes: Same as Table 3.

TABLE J.3. After Dropping 25% of links: The Correlation Between Inaccuracy in Ranking a Pair of Households in a Hamlet and the Average Inverse Distance to Rankees

	(1)	(2)	(3)	(4)
<i>Panel A: Consumption Metric, Error Rate</i>				
Average Inverse Distance	-0.0418*** (0.00898)	-0.0177** (0.00839)	-0.00734 (0.00533)	0.000479 (0.00888)
Average Degree		-0.0175*** (0.00601)	0.00545 (0.00790)	0.00446 (0.00815)
Average Clustering Coefficient		0.0178 (0.0398)	0.00646 (0.0401)	0.00364 (0.0412)
Average Eigenvector Centrality		0.0390 (0.0587)	-0.0962 (0.0689)	-0.104 (0.0712)
R-squared	0.003	0.009	0.138	0.208
<i>Panel B: Self-Assessment Metric, Error Rate</i>				
Average Inverse Distance	-0.0478*** (0.00965)	-0.0179** (0.00841)	-0.00761 (0.00539)	0.00522 (0.00882)
Average Degree		-0.0184*** (0.00660)	0.00355 (0.00791)	0.00119 (0.00813)
Average Clustering Coefficient		0.0312 (0.0403)	0.0105 (0.0412)	0.00423 (0.0425)
Average Eigenvector Centrality		0.0460 (0.0667)	-0.0734 (0.0733)	-0.0836 (0.0757)
R-squared	0.004	0.015	0.167	0.253
<i>Panel C: Share of Don't Knows</i>				
Average Inverse Distance	-0.0572*** (0.00916)	-0.0261*** (0.00971)	-0.0121** (0.00576)	-0.00660 (0.00915)
Average Degree		-0.0289*** (0.00738)	0.000929 (0.00757)	1.37e-05 (0.00779)
Average Clustering Coefficient		0.0675 (0.0461)	0.0310 (0.0410)	0.0310 (0.0417)
Average Eigenvector Centrality		0.116 (0.0737)	-0.0553 (0.0709)	-0.0674 (0.0732)
R-squared	0.010	0.054	0.339	0.457
Demographic Controls	No	Yes	Yes	Yes
Hamlet Fixed Effect	No	No	Yes	Yes
Ranker Fixed Effect	No	No	No	Yes

Notes: Same as Table 4.

TABLE J.4. After Dropping 25% of links: Empirical Results on Stochastic Dominance

	(1)	(2)	(3)	(4)
<i>Panel A: Consumption Metric</i>				
I fofd J	-0.0968*** (0.0192)	-0.141*** (0.0297)	-0.0906*** (0.0191)	-0.124*** (0.0280)
J fofd I	0.0471** (0.0184)		0.0484*** (0.0179)	
<i>Panel B: Self-Assessment Metric</i>				
I fofd J	-0.102*** (0.0178)	-0.172*** (0.0265)	-0.0772*** (0.0181)	-0.125*** (0.0262)
J fofd I	0.0735*** (0.0168)		0.0593*** (0.0168)	
Non-Comparable	Yes	No	Yes	No
Demographic Controls	No	No	Yes	Yes
Stratification Group FE	Yes	Yes	Yes	Yes

Notes: Same as Table 8.

TABLE J.5. After Dropping 25% of links: Empirical Results on Correlation between Hamlet Network Characteristics and Hamlet Level Error Rate

	(1)	(2)	(3)	(4)	(5)	(6)	(7)
<i>Panel A: Consumption Metric</i>							
Average Degree	-0.0313** (0.0134)						0.0231 (0.0393)
Average Clustering		-0.223 (0.138)					0.0772 (0.258)
Number of Households			0.000762* (0.000398)				0.000978** (0.000422)
First eigenvalue $\lambda_1(A)$				-0.0222** (0.00966)			-0.0407* (0.0220)
Fraction of Nodes in Giant Component					-0.106** (0.0443)		-0.0264 (0.0972)
Link Density						-0.0355* (0.0188)	0.00856 (0.0329)
R-squared	0.248	0.244	0.249	0.249	0.249	0.247	0.257
<i>Panel B: Self-Assessment Metric</i>							
Average Degree	-0.0376*** (0.0131)						0.0228 (0.0453)
Average Clustering		-0.0424 (0.132)					0.515*** (0.188)
Number of Households			0.00118*** (0.000419)				0.00137*** (0.000449)
First eigenvalue $\lambda_1(A)$				-0.0215** (0.0101)			-0.0385 (0.0238)
Fraction of Nodes in Giant Component					-0.147*** (0.0437)		-0.0908 (0.126)
Link Density						-0.0327* (0.0193)	-0.0124 (0.0367)
R-squared	0.311	0.303	0.319	0.309	0.315	0.307	0.329

Notes: Same as Table 9.

TABLE J.6. After Dropping 25% of links: Rank Correlation on Targeting Type Interacted with Diffusiveness (Principal Component)

	(1)	(2)	(3)	(4)	(5)	(6)
<i>Panel A: Rank Correlation (Consumption)</i>						
Community x Diffusiveness		-0.0443 (0.119)	-0.0511 (0.118)	-0.0946 (0.124)	-0.0866 (0.128)	
Hybrid x Diffusiveness		-0.0584 (0.114)	-0.0601 (0.114)	-0.0986 (0.123)		
Community	-0.0588* (0.0319)	-0.0393 (0.0659)	-0.0322 (0.0655)	-0.0160 (0.0686)	-0.0162 (0.0698)	
Hybrid	-0.0614* (0.0327)	-0.0349 (0.0671)	-0.0307 (0.0674)	-0.00670 (0.0739)		
Diffusiveness		-0.0354 (0.0788)	-0.0115 (0.0803)	0.0530 (0.0933)	0.0765 (0.102)	0.0525 (0.0931)
(Community or Hybrid) x Diffusiveness						-0.0954 (0.105)
(Community or Hybrid)						-0.0120 (0.0600)
R-squared	0.014	0.012	0.016	0.095	0.151	0.094
<i>Panel B: Rank Correlation (Self-Assessment)</i>						
Community x Diffusiveness		0.278** (0.110)	0.266** (0.108)	0.237** (0.116)	0.241** (0.120)	
Hybrid x Diffusiveness		0.326*** (0.111)	0.325*** (0.111)	0.316*** (0.118)		
Community	0.108*** (0.0321)	-0.0308 (0.0659)	-0.0185 (0.0649)	-0.00424 (0.0696)	-0.00846 (0.0718)	
Hybrid	0.0839** (0.0331)	-0.0842 (0.0676)	-0.0777 (0.0674)	-0.0731 (0.0735)		
Diffusiveness		-0.267*** (0.0777)	-0.225*** (0.0785)	-0.222** (0.0898)	-0.246** (0.0964)	-0.220** (0.0897)
(Community or Hybrid) x Diffusiveness						0.273*** (0.102)
(Community or Hybrid)						-0.0358 (0.0621)
R-squared	0.033	0.029	0.043	0.127	0.161	0.125
Stratification Group FE	No	No	No	Yes	Yes	Yes
Demographic Covariates	No	No	No	Yes	Yes	Yes

Notes: Same as Table 10.

TABLE J.7. After Dropping 25% of links: Rank Correlation on Targeting Type Interacted with Diffusiveness (1-Simulated Error Rate)

	(1)	(2)	(3)	(4)	(5)	(6)
<i>Panel A: Rank Correlation (Consumption)</i>						
Community x Diffusiveness		0.110 (0.112)	0.0831 (0.117)	0.0747 (0.119)	0.100 (0.102)	
Hybrid x Diffusiveness		0.0126 (0.114)	-0.0252 (0.116)	-0.0355 (0.118)		
Community	-0.0588* (0.0319)	-0.109* (0.0617)	-0.0967 (0.0647)	-0.0965 (0.0661)	-0.0817 (0.0596)	
Hybrid	-0.0614* (0.0327)	-0.0649 (0.0592)	-0.0464 (0.0609)	-0.0351 (0.0627)		
Diffusiveness		-0.0993 (0.0793)	-0.0787 (0.0953)	-0.0763 (0.0958)	-0.105 (0.0769)	-0.0758 (0.0955)
(Community or Hybrid) x Diffusiveness						0.0163 (0.103)
(Community or Hybrid)						-0.0642 (0.0543)
R-squared	0.014	0.018	0.088	0.096	0.092	0.094
<i>Panel B: Rank Correlation (Self-Assessment)</i>						
Community x Diffusiveness		0.187 (0.116)	0.224* (0.123)	0.230* (0.124)	0.114 (0.106)	
Hybrid x Diffusiveness		0.213* (0.118)	0.232* (0.122)	0.199 (0.122)		
Community	0.108*** (0.0321)	0.0261 (0.0650)	0.00938 (0.0678)	0.00671 (0.0691)	0.0135 (0.0605)	
Hybrid	0.0839** (0.0331)	-0.0128 (0.0640)	-0.0230 (0.0666)	-0.00365 (0.0679)		
Diffusiveness		-0.248*** (0.0833)	-0.284*** (0.102)	-0.281*** (0.105)	-0.162** (0.0809)	-0.279*** (0.104)
(Community or Hybrid) x Diffusiveness						0.213** (0.107)
(Community or Hybrid)						0.00193 (0.0588)
R-squared	0.033	0.049	0.115	0.132	0.115	0.131
Stratification Group FE	No	No	No	Yes	Yes	Yes
Demographic Covariates	No	No	No	Yes	Yes	Yes

Notes: Same as Table 11.

J.2. Dropping 2 sampled households and all of their information.

- Tables J.8/J.9:

We see that higher degree and more (eigenvector) central nodes have lower error rates. Further, the results typically hold even when adding demographic controls and are mostly robust to the inclusion of hamlet fixed effects. This is consistent with our main results.

- Table J.10:

When nodes are further on average from those whom they are ranking, they are more likely to make a mistake. This is robust to including demographic controls and hamlet fixed effects. Again the results are underpowered with ranker fixed effects. This is consistent with our main results.

- Table J.11:

Networks that have degree distributions that (first order stochastically) dominate other networks tend to have lower error rates. Again, this is consistent with our main results.

- Table J.12:

We find that a higher degree reduces error rates, more clustering reduces error rates, a higher first eigenvalue reduces error rates, and a higher density reduces error rates. This is consistent with our main results.

- Tables J.13 and J.14:

Here we look at whether community targeting does better relative to PMT in more “diffusive” villages where this is computed either via principal components (Table J.13) or by using the simulated error rate in these hamlets (Table J.14). We find that being in a more diffusive hamlet (or less error-prone hamlet under our model) corresponds to community targeting being more effective than the PMT. This is consistent with our main results.

TABLE J.8. Dropping Two Households: The Correlation between Household Network Characteristics and the Error Rate in Ranking Income Status of Household

	(1)	(2)	(3)	(4)	(5)	(6)	(7)	(8)
<i>Panel A: Consumption Metric, Error Rate</i>								
Degree	-0.0116*** (0.00134)			-0.0138*** (0.00199)	-0.00322*** (0.000899)			-0.00186 (0.00172)
Clustering		-0.0357** (0.0161)		-0.0363** (0.0165)		-0.00945 (0.0103)		-0.00727 (0.0137)
Eigenvector Centrality			-0.173*** (0.0378)	0.126** (0.0584)			-0.0871*** (0.0273)	-0.0490 (0.0519)
R-squared	0.045	0.002	0.010	0.048	0.689	0.688	0.690	0.690
<i>Panel B: Self-Assessment Metric, Error Rate</i>								
Degree	-0.0134*** (0.00153)			-0.0166*** (0.00227)	-0.00320*** (0.00103)			-0.00163 (0.00197)
Clustering		-0.0243 (0.0173)		-0.0284* (0.0172)		-0.00643 (0.0113)		-0.00338 (0.0143)
Eigenvector Centrality			-0.175*** (0.0420)	0.171*** (0.0645)			-0.0867*** (0.0299)	-0.0549 (0.0562)
R-squared	0.048	0.001	0.008	0.053	0.686	0.685	0.686	0.686
<i>Panel C: Share of Don't Knows</i>								
Degree	-0.0161*** (0.00160)			-0.0188*** (0.00228)	-0.00451*** (0.000931)			0.000510 (0.00159)
Clustering		-0.0251 (0.0200)		-0.0188 (0.0194)		-0.000908 (0.0121)		0.0156 (0.0144)
Eigenvector Centrality			-0.240*** (0.0457)	0.141** (0.0657)			-0.147*** (0.0284)	-0.165*** (0.0481)
R-squared	0.072	0.001	0.016	0.075	0.732	0.729	0.733	0.733
Hamlet Fixed Effect	No	No	No	No	Yes	Yes	Yes	Yes

Notes: Same as Table 2.

TABLE J.9. Dropping Two Households: The Correlation between Household Network Characteristics and the Error Rate in Ranking Income Status of Households

	(1)	(2)	(3)	(4)	(5)	(6)	(7)	(8)
<i>Panel A: Consumption Metric, Error Rate</i>								
Degree	-0.0103*** (0.00132)			-0.0121*** (0.00196)	-0.00241*** (0.000902)			-0.00126 (0.00169)
Clustering		-0.0376** (0.0156)		-0.0369** (0.0164)		-0.0103 (0.0102)		-0.00785 (0.0135)
Eigenvector Centrality			-0.157*** (0.0373)	0.107* (0.0578)			-0.0697** (0.0275)	-0.0425 (0.0513)
R-squared	0.057	0.026	0.032	0.059	0.692	0.691	0.692	0.692
<i>Panel B: Self-Assessment Metric, Error Rate</i>								
Degree	-0.0112*** (0.00151)			-0.0138*** (0.00221)	-0.00220** (0.00102)			-0.000888 (0.00195)
Clustering		-0.0273* (0.0165)		-0.0292* (0.0169)		-0.00741 (0.0112)		-0.00408 (0.0141)
Eigenvector Centrality			-0.151*** (0.0414)	0.138** (0.0632)			-0.0647** (0.0299)	-0.0463 (0.0558)
R-squared	0.098	0.066	0.069	0.100	0.679	0.677	0.678	0.679
<i>Panel B: Self-Assessment Metric, Error Rate</i>								
Degree	-0.0131*** (0.00146)			-0.0151*** (0.00216)	-0.00322*** (0.000905)			0.00188 (0.00157)
Clustering		-0.0184 (0.0178)		-0.0136 (0.0184)		0.000576 (0.0116)		0.0192 (0.0140)
Eigenvector Centrality			-0.192*** (0.0409)	0.109* (0.0637)			-0.122*** (0.0281)	-0.167*** (0.0479)
R-squared	0.074	0.043	0.048	0.077	0.690	0.689	0.690	0.690
Hamlet Fixed Effect	No	No	No	No	Yes	Yes	Yes	Yes

Notes: Same as Table 2.

TABLE J.10. Dropping Two Households: The Correlation Between Inaccuracy in Ranking a Pair of Households in a Hamlet and the Average Inverse Distance to Rankees

	(1)	(2)	(3)	(4)
<i>Panel A: Consumption Metric, Error Rate</i>				
Average Inverse Distance	-0.0559*** (0.00924)	-0.0348*** (0.0103)	-0.0160** (0.00749)	-0.00514 (0.0155)
Average Degree		-0.0106*** (0.00322)	-0.00153 (0.00516)	-0.00140 (0.00568)
Average Clustering Coefficient		0.0110 (0.0340)	0.0277 (0.0406)	0.0296 (0.0439)
Average Eigenvector Centrality		0.180* (0.0997)	0.166 (0.120)	0.118 (0.129)
R-squared	0.007	0.013	0.175	0.269
<i>Panel B: Self-Assessment Metric, Error Rate</i>				
Average Inverse Distance	-0.0708*** (0.00990)	-0.0336*** (0.0105)	-0.0211*** (0.00796)	0.00503 (0.0146)
Average Degree		-0.0134*** (0.00333)	-0.00736 (0.00537)	-0.00734 (0.00577)
Average Clustering Coefficient		0.0181 (0.0363)	0.0304 (0.0405)	0.0325 (0.0432)
Average Eigenvector Centrality		0.123 (0.107)	0.176 (0.128)	0.0587 (0.138)
R-squared	0.011	0.023	0.211	0.322
<i>Panel C: Share of Don't Knows</i>				
Average Inverse Distance	-0.0707*** (0.00964)	-0.0406*** (0.0118)	-0.0186** (0.00845)	0.00521 (0.0138)
Average Degree		-0.0154*** (0.00375)	-0.00492 (0.00469)	-0.00427 (0.00491)
Average Clustering Coefficient		0.0322 (0.0391)	0.0102 (0.0379)	0.0143 (0.0391)
Average Eigenvector Centrality		0.235** (0.120)	0.207* (0.123)	0.0672 (0.129)
R-squared	0.020	0.063	0.387	0.525
Demographic Controls	No	Yes	Yes	Yes
Hamlet Fixed Effect	No	No	Yes	Yes
Ranker Fixed Effect	No	No	No	Yes

Notes: Same as Table 4.

TABLE J.11. Dropping Two Households: Empirical Results on Stochastic Dominance

	(1)	(2)	(3)	(4)
<i>Panel A: Consumption Metric</i>				
I fofd J	-0.0987*** (0.0202)	-0.139*** (0.0278)	-0.0942*** (0.0201)	-0.120*** (0.0258)
J fofd I	0.0530*** (0.0188)		0.0577*** (0.0183)	
<i>Panel B: Self-Assessment Metric</i>				
I fofd J	-0.108*** (0.0189)	-0.162*** (0.0249)	-0.0829*** (0.0188)	-0.111*** (0.0235)
J fofd I	0.0730*** (0.0173)		0.0601*** (0.0175)	
Non-Comparable	Yes	No	Yes	No
Demographic Controls	No	No	Yes	Yes
Stratification Group FE	Yes	Yes	Yes	Yes

Notes: Same as Table 8.

TABLE J.12. Dropping Two Households: Empirical Results on Correlation between Hamlet Network Characteristics and Hamlet Level Error Rate

	(1)	(2)	(3)	(4)	(5)	(6)	(7)
<i>Panel A: Consumption Metric</i>							
Average Degree	-0.0171** (0.00666)						0.0440*** (0.0152)
Average Clustering		-0.213*** (0.0701)					-0.162 (0.111)
Number of Households			0.000762* (0.000398)				0.000738* (0.000434)
First eigenvalue $\lambda_1(A)$				-0.0145*** (0.00506)			-0.0328*** (0.00983)
Fraction of Nodes in Giant Component					-0.182*** (0.0502)		-0.186** (0.0746)
Link Density						-0.00888 (0.00570)	0.0118 (0.00807)
R-squared	0.248	0.253	0.249	0.249	0.261	0.238	0.275
<i>Panel B: Self-Assessment Metric</i>							
Average Degree	-0.0222*** (0.00602)						0.0289* (0.0154)
Average Clustering		-0.256*** (0.0655)					-0.108 (0.115)
Number of Households			0.00118*** (0.000419)				0.00102* (0.000510)
First eigenvalue $\lambda_1(A)$				-0.0134*** (0.00448)			-0.0261** (0.0111)
Fraction of Nodes in Giant Component					-0.214*** (0.0465)		-0.166** (0.0706)
Link Density						-0.00779 (0.00527)	0.0136 (0.00947)
R-squared	0.313	0.316	0.319	0.305	0.324	0.297	0.327

Notes: Same as Table 9.

TABLE J.13. Dropping Two Households: Rank Correlation on Targeting Type Interacted with Diffusiveness (Principal Component)

	(1)	(2)	(3)	(4)	(5)	(6)
<i>Panel A: Rank Correlation (Consumption)</i>						
Community x Diffusiveness		-0.0443 (0.119)	-0.0511 (0.118)	-0.0946 (0.124)	-0.0866 (0.128)	
Hybrid x Diffusiveness		-0.0584 (0.114)	-0.0601 (0.114)	-0.0986 (0.123)		
Community	-0.0588* (0.0319)	-0.0393 (0.0659)	-0.0322 (0.0655)	-0.0160 (0.0686)	-0.0162 (0.0698)	
Hybrid	-0.0614* (0.0327)	-0.0349 (0.0671)	-0.0307 (0.0674)	-0.00670 (0.0739)		
Diffusiveness		-0.0354 (0.0788)	-0.0115 (0.0803)	0.0530 (0.0933)	0.0765 (0.102)	0.0525 (0.0931)
(Community or Hybrid) x Diffusiveness						-0.0954 (0.105)
(Community or Hybrid)						-0.0120 (0.0600)
R-squared	0.014	0.012	0.016	0.095	0.151	0.094
<i>Panel B: Rank Correlation (Self-Assessment)</i>						
Community x Diffusiveness		0.262** (0.112)	0.250** (0.112)	0.243** (0.118)	0.239** (0.120)	
Hybrid x Diffusiveness		0.326*** (0.110)	0.319*** (0.110)	0.321*** (0.116)		
Community	0.108*** (0.0321)	-0.0201 (0.0666)	-0.00790 (0.0661)	-0.00429 (0.0695)	-0.00443 (0.0712)	
Hybrid	0.0839** (0.0331)	-0.0841 (0.0669)	-0.0746 (0.0671)	-0.0763 (0.0726)		
Diffusiveness		-0.245*** (0.0772)	-0.202** (0.0796)	-0.226** (0.103)	-0.252** (0.117)	-0.221** (0.103)
(Community or Hybrid) x Diffusiveness						0.278*** (0.102)
(Community or Hybrid)						-0.0372 (0.0613)
R-squared	0.033	0.035	0.047	0.132	0.167	0.130
Stratification Group FE	No	No	No	Yes	Yes	Yes
Demographic Covariates	No	No	No	Yes	Yes	Yes

Notes: Same as Table 10.

TABLE J.14. Dropping Two Households: Rank Correlation on Targeting Type Interacted with Diffusiveness (1-Simulated Error Rate)

	(1)	(2)	(3)	(4)	(5)	(6)
<i>Panel A: Rank Correlation (Consumption)</i>						
Community x Diffusiveness		0.117 (0.115)	0.106 (0.122)	0.115 (0.123)	0.144 (0.113)	
Hybrid x Diffusiveness		-0.0145 (0.116)	-0.0363 (0.119)	-0.0413 (0.122)		
Community	-0.0588* (0.0319)	-0.110 (0.0671)	-0.109 (0.0703)	-0.118 (0.0724)	-0.104 (0.0674)	
Hybrid	-0.0614* (0.0327)	-0.0555 (0.0613)	-0.0452 (0.0631)	-0.0360 (0.0651)		
Diffusiveness		-0.0882 (0.0755)	-0.0543 (0.0892)	-0.0533 (0.0913)	-0.0846 (0.0787)	-0.0511 (0.0909)
(Community or Hybrid) x Diffusiveness						0.0297 (0.103)
(Community or Hybrid)						-0.0731 (0.0568)
R-squared	0.014	0.017	0.083	0.090	0.086	0.087
<i>Panel B: Rank Correlation (Self-Assessment)</i>						
Community x Diffusiveness		0.293** (0.116)	0.330*** (0.124)	0.326*** (0.124)	0.184* (0.105)	
Hybrid x Diffusiveness		0.248** (0.116)	0.262** (0.120)	0.244** (0.121)		
Community	0.108*** (0.0321)	-0.0360 (0.0656)	-0.0587 (0.0685)	-0.0551 (0.0689)	-0.0342 (0.0608)	
Hybrid	0.0839** (0.0331)	-0.0309 (0.0633)	-0.0415 (0.0663)	-0.0282 (0.0676)		
Diffusiveness		-0.209** (0.0854)	-0.177* (0.0990)	-0.187* (0.102)	-0.0440 (0.0760)	-0.185* (0.102)
(Community or Hybrid) x Diffusiveness						0.282*** (0.108)
(Community or Hybrid)						-0.0400 (0.0586)
R-squared	0.033	0.043	0.114	0.132	0.113	0.131
Stratification Group FE	No	No	No	Yes	Yes	Yes
Demographic Covariates	No	No	No	Yes	Yes	Yes

Notes: Same as Table 11.

APPENDIX K. SMALL HAMLETS

Because we fully survey 8 households at random as well as the leader, the amount of information we have varies with the number of households in the hamlet. Specifically, we have better network data for smaller networks. In this section we show that both our within-village and across-village results are robust to looking at hamlets with households below the median number of households. In all specifications we look at only hamlets with below the median number of households and, further, even when not controlling for hamlet fixed effects in our household level analysis we always control for hamlet size.

TABLE K.1. Small Hamlets: The Correlation between Household Network Characteristics and the Error Rate in Ranking Income Status of Household

	(1)	(2)	(3)	(4)	(5)	(6)	(7)	(8)
<i>Panel A: Consumption Metric, Error Rate</i>								
Degree	-0.00776*** (0.00164)			-0.00854*** (0.00215)	-0.00162* (0.000878)			-0.000617 (0.00170)
Clustering		-0.0435* (0.0222)		-0.0379* (0.0214)		-0.00209 (0.0124)		0.00159 (0.0150)
Eigenvector Centrality			-0.0961* (0.0550)	0.0796 (0.0749)			-0.0533* (0.0291)	-0.0402 (0.0570)
R-squared	0.030	0.007	0.003	0.032	0.707	0.707	0.707	0.707
<i>Panel B: Self-Assessment Metric, Error Rate</i>								
Degree	-0.0111*** (0.00179)			-0.0136*** (0.00235)	-0.00203** (0.000913)			-0.00243 (0.00189)
Clustering		-0.0417* (0.0249)		-0.0492** (0.0224)		-0.00156 (0.0136)		-0.00508 (0.0162)
Eigenvector Centrality			-0.0741 (0.0583)	0.213*** (0.0814)			-0.0431 (0.0334)	0.0153 (0.0659)
R-squared	0.053	0.002	0.009	0.061	0.706	0.705	0.705	0.706
<i>Panel C: Share of Don't Knows</i>								
Degree	-0.0103*** (0.00174)			-0.0127*** (0.00235)	-0.00154** (0.000741)			-0.000888 (0.00151)
Clustering		-0.00599 (0.0246)		-0.0173 (0.0235)		0.00507 (0.0129)		0.00573 (0.0144)
Eigenvector Centrality			-0.0333 (0.0581)	0.202** (0.0816)			-0.0445 (0.0312)	-0.0255 (0.0579)
R-squared	0.043	0.000	0.000	0.052	0.807	0.806	0.807	0.807
Hamlet Fixed Effect	No	No	No	No	Yes	Yes	Yes	Yes

Notes: Same as Table 2. The sample is restricted to villages with below the median number of households.

TABLE K.2. Small Hamlets: The Correlation between Household Network Characteristics and the Error Rate in Ranking Income Status of Households

	(1)	(2)	(3)	(4)	(5)	(6)	(7)	(8)
<i>Panel A: Consumption Metric, Error Rate</i>								
Degree	-0.00613*** (0.00150)			-0.00670*** (0.00203)	-0.00145* (0.000873)			-0.000462 (0.00168)
Clustering		-0.0357* (0.0200)		-0.0320 (0.0205)		-0.00214 (0.0125)		0.00158 (0.0151)
Eigenvector Centrality			-0.0865* (0.0512)	0.0572 (0.0736)			-0.0496* (0.0296)	-0.0400 (0.0573)
R-squared	0.045	0.032	0.033	0.047	0.708	0.708	0.708	0.708
<i>Panel B: Self-Assessment Metric, Error Rate</i>								
Degree	-0.00857*** (0.00160)			-0.0108*** (0.00219)	-0.00177* (0.000905)			-0.00222 (0.00187)
Clustering		-0.0343 (0.0210)		-0.0410* (0.0212)		-0.00215 (0.0135)		-0.00572 (0.0161)
Eigenvector Centrality			-0.0439 (0.0543)	0.180** (0.0786)			-0.0367 (0.0334)	0.0171 (0.0657)
R-squared	0.082	0.059	0.058	0.088	0.707	0.706	0.707	0.707
<i>Panel C: Share of Don't Knows</i>								
Degree	-0.00778*** (0.00156)			-0.00973*** (0.00218)	-0.00147** (0.000742)			-0.000754 (0.00151)
Clustering		-0.00226 (0.0217)		-0.0109 (0.0218)		0.00327 (0.0127)		0.00432 (0.0141)
Eigenvector Centrality			-0.0259 (0.0538)	0.151* (0.0782)			-0.0450 (0.0312)	-0.0287 (0.0580)
R-squared	0.083	0.060	0.060	0.088	0.807	0.807	0.807	0.807
Hamlet Fixed Effect	No	No	No	No	Yes	Yes	Yes	Yes

Notes: Same as Table 3. The sample is restricted to villages with below the median number of households.

TABLE K.3. Small Hamlets: The Correlation Between Inaccuracy in Ranking a Pair of Households in a Hamlet and the Average Inverse Distance to Rankees

	(1)	(2)	(3)	(4)
<i>Panel A: Consumption Metric, Error Rate</i>				
Average Inverse Distance	-0.0367*** (0.0122)	-0.0395*** (0.0117)	-0.0212*** (0.00788)	-0.0150 (0.0208)
Average Degree		-0.00293 (0.00281)	0.00131 (0.00528)	0.00189 (0.00533)
Average Clustering Coefficient		0.0111 (0.0394)	-0.0174 (0.0456)	-0.0155 (0.0459)
Average Eigenvector Centrality		0.142 (0.102)	0.00693 (0.177)	-0.0195 (0.190)
R-squared	0.007	0.008	0.136	0.189
<i>Panel B: Self-Assessment Metric, Error Rate</i>				
Average Inverse Distance	-0.0338*** (0.0127)	-0.0363*** (0.0127)	-0.0188** (0.00853)	-0.0132 (0.0220)
Average Degree		-0.00617** (0.00287)	0.00321 (0.00594)	0.00300 (0.00603)
Average Clustering Coefficient		0.0122 (0.0399)	0.0201 (0.0510)	0.0216 (0.0504)
Average Eigenvector Centrality		0.271** (0.111)	0.0262 (0.212)	0.0263 (0.220)
R-squared	0.011	0.014	0.154	0.218
<i>Panel C: Share of Don't Knows</i>				
Average Inverse Distance	-0.0509*** (0.0143)	-0.0310** (0.0148)	-0.0192** (0.00963)	-0.0137 (0.0216)
Average Degree		-0.0102*** (0.00330)	-0.00778 (0.00526)	-0.00696 (0.00521)
Average Clustering Coefficient		0.0297 (0.0450)	-0.0234 (0.0461)	-0.0166 (0.0453)
Average Eigenvector Centrality		0.260** (0.122)	0.300 (0.195)	0.263 (0.210)
R-squared	0.008	0.044	0.339	0.414
Demographic Controls	No	Yes	Yes	Yes
Hamlet Fixed Effect	No	No	Yes	Yes
Ranker Fixed Effect	No	No	No	Yes

Notes: Same as Table 4. The sample is restricted to villages with below the median number of households.

TABLE K.4. Small Hamlets: Empirical Results on Stochastic Dominance

	(1)	(2)	(3)	(4)
<i>Panel A: Consumption Metric</i>				
I fofd J	-0.0476*** (0.00538)	-0.0145** (0.00660)	-0.0630*** (0.00555)	-0.0395*** (0.00704)
J fofd I	-0.0254*** (0.00535)		-0.00889* (0.00510)	
Observations	52,650	35,753	52,650	35,753
<i>Panel B: Self-Assessment Metric</i>				
I fofd J	-0.0773*** (0.00494)	-0.0853*** (0.00619)	-0.0890*** (0.00493)	-0.102*** (0.00682)
J fofd I	0.0145*** (0.00548)		0.0253*** (0.00545)	
Observations	52,650	35,753	52,650	35,753
Non-Comparable	Yes	No	Yes	No
Demographic Controls	No	No	Yes	Yes
Stratification Group FE	Yes	Yes	Yes	Yes

Notes: Same as Table 8. The sample is restricted to villages with below the median number of households.

TABLE K.5. Small Hamlets: Empirical Results on Correlation between Hamlet Network Characteristics and Hamlet Level Error Rate

	(1)	(2)	(3)	(4)	(5)	(6)	(7)
<i>Panel A: Consumption Metric</i>							
Average Degree	-0.0136*** (0.00387)						0.0330* (0.0173)
Average Clustering		-0.314*** (0.0652)					-0.384*** (0.128)
Number of Households			-0.000390 (0.00120)				0.000752 (0.00191)
First eigenvalue $\lambda_1(A)$				-0.0128*** (0.00343)			-0.0212* (0.0108)
Fraction of Nodes in Giant Component					-0.239*** (0.0506)		-0.195** (0.0918)
Link Density						-0.195** (0.0905)	0.185 (0.232)
R-squared	0.042	0.082	0.006	0.045	0.079	0.020	0.120
<i>Panel B: Self-Assessment Metric</i>							
Average Degree	-0.0219*** (0.00395)						0.0204 (0.0183)
Average Clustering		-0.464*** (0.0674)					-0.517*** (0.140)
Number of Households			-0.000503 (0.00137)				-0.000467 (0.00195)
First eigenvalue $\lambda_1(A)$				-0.0188*** (0.00382)			-0.0130 (0.0116)
Fraction of Nodes in Giant Component					-0.333*** (0.0523)		-0.170* (0.0981)
Link Density						-0.344*** (0.0912)	0.211 (0.230)
R-squared	0.092	0.153	0.011	0.083	0.133	0.051	0.182

Notes: Same as Table 9. The sample is restricted to villages with below the median number of households.

TABLE K.6. Small Hamlets: Rank Correlation on Targeting Type Interacted with Diffusiveness (Principal Component)

	(1)	(2)	(3)	(4)	(5)	(6)
<i>Panel A: Rank Correlation (Consumption)</i>						
Community x Diffusiveness		-0.168 (0.163)	-0.141 (0.165)	-0.0942 (0.183)	-0.104 (0.205)	
Hybrid x Diffusiveness		0.0358 (0.176)	0.0534 (0.177)	0.0382 (0.194)		
Community	-0.0471 (0.0486)	0.0552 (0.0997)	0.0380 (0.101)	0.00983 (0.113)	0.0445 (0.126)	
Hybrid	-0.0311 (0.0452)	-0.0518 (0.122)	-0.0632 (0.122)	-0.0368 (0.136)		
Diffusiveness		-0.134 (0.112)	-0.100 (0.117)	-0.105 (0.132)	-0.103 (0.170)	-0.101 (0.132)
(Community or Hybrid) x Diffusiveness						-0.00939 (0.161)
(Community or Hybrid)						-0.0204 (0.107)
R-squared	0.020	0.022	0.030	0.200	0.291	0.197
<i>Panel B: Rank Correlation (Self-Assessment)</i>						
Community x Diffusiveness		-0.0352 (0.177)	-0.000627 (0.178)	-0.0290 (0.210)	-0.0320 (0.224)	
Hybrid x Diffusiveness		0.0391 (0.189)	0.0613 (0.190)	-0.0113 (0.215)		
Community	0.182*** (0.0448)	0.204 (0.124)	0.182 (0.124)	0.224 (0.146)	0.228 (0.158)	
Hybrid	0.155*** (0.0462)	0.128 (0.137)	0.114 (0.137)	0.185 (0.158)		
Diffusiveness		-0.0362 (0.134)	0.00625 (0.137)	0.000362 (0.164)	-0.00402 (0.183)	-0.000306 (0.163)
(Community or Hybrid) x Diffusiveness						-0.0256 (0.188)
(Community or Hybrid)						0.207 (0.135)
R-squared	0.068	0.055	0.069	0.249	0.391	0.248
Stratification Group FE	No	No	No	Yes	Yes	Yes
Demographic Covariates	No	No	No	Yes	Yes	Yes

Notes: Same as Table 10. The sample is restricted to villages with below the median number of households.

TABLE K.7. Small Hamlets: Rank Correlation on Targeting Type Interacted with Diffusiveness (1-Simulated Error Rate)

	(1)	(2)	(3)	(4)	(5)	(6)
<i>Panel A: Rank Correlation (Consumption)</i>						
Community x Diffusiveness		-0.0966 (0.166)	-0.0608 (0.186)	-0.0623 (0.187)	-0.0481 (0.158)	
Hybrid x Diffusiveness		-0.0315 (0.150)	-0.0167 (0.163)	-0.0175 (0.164)		
Community	-0.0471 (0.0486)	0.00818 (0.0984)	-0.00632 (0.109)	-0.0173 (0.110)	-0.0122 (0.101)	
Hybrid	-0.0311 (0.0452)	-0.0115 (0.0889)	-0.00276 (0.0912)	-0.0124 (0.0950)		
Diffusiveness		-0.0172 (0.107)	0.0598 (0.145)	0.0942 (0.148)	0.0755 (0.107)	0.0960 (0.148)
(Community or Hybrid) x Diffusiveness						-0.0341 (0.150)
(Community or Hybrid)						-0.0159 (0.0841)
R-squared	0.020	0.023	0.177	0.198	0.197	0.196
<i>Panel B: Rank Correlation (Self-Assessment)</i>						
Community x Diffusiveness		-0.0499 (0.157)	0.00950 (0.181)	0.00793 (0.188)	-0.0493 (0.160)	
Hybrid x Diffusiveness		0.0219 (0.169)	0.0543 (0.178)	0.0466 (0.181)		
Community	0.182*** (0.0448)	0.213** (0.0919)	0.210** (0.102)	0.208* (0.107)	0.133 (0.0969)	
Hybrid	0.155*** (0.0462)	0.147 (0.0967)	0.156 (0.100)	0.159 (0.101)		
Diffusiveness		-0.0675 (0.117)	-0.0983 (0.148)	-0.121 (0.155)	-0.0233 (0.122)	-0.123 (0.155)
(Community or Hybrid) x Diffusiveness						0.0291 (0.161)
(Community or Hybrid)						0.181** (0.0881)
R-squared	0.068	0.072	0.233	0.252	0.209	0.251
Stratification Group FE	No	No	No	Yes	Yes	Yes
Demographic Covariates	No	No	No	Yes	Yes	Yes

Notes: Same as Table 11. The sample is restricted to villages with below the median number of households.

APPENDIX L. RURAL SAMPLE

In this section we restrict our sample to only rural villages.

TABLE L.1. Rural Sample: The Correlation between Household Network Characteristics and the Error Rate in Ranking Income Status of Household

	(1)	(2)	(3)	(4)	(5)	(6)	(7)	(8)
<i>Panel A: Consumption Metric, Error Rate</i>								
Degree	-0.00546*** (0.00130)			-0.00554*** (0.00166)	-0.000987 (0.000819)			-0.00125 (0.00141)
Clustering		-0.0183 (0.0201)		-0.0174 (0.0195)		0.00451 (0.0126)		0.00158 (0.0133)
Eigenvector Centrality			-0.113** (0.0538)	0.00720 (0.0679)			-0.0200 (0.0348)	0.0125 (0.0580)
R-squared	0.017	0.001	0.004	0.017	0.651	0.651	0.651	0.651
<i>Panel B: Self-Assessment Metric, Error Rate</i>								
Degree	-0.00818*** (0.00148)			-0.00859*** (0.00190)	-0.00220** (0.000872)			-0.00207 (0.00152)
Clustering		-0.0198 (0.0215)		-0.0209 (0.0201)		0.0165 (0.0137)		0.0125 (0.0144)
Eigenvector Centrality			-0.149** (0.0604)	0.0337 (0.0756)			-0.0519 (0.0363)	-0.000565 (0.0591)
R-squared	0.030	0.000	0.006	0.031	0.648	0.647	0.647	0.648
<i>Panel C: Share of Don't Knows</i>								
Degree	-0.00810*** (0.00124)			-0.00841*** (0.00157)	-0.00214*** (0.000725)			-0.000783 (0.00123)
Clustering		-0.0302 (0.0211)		-0.0318 (0.0205)		0.0101 (0.0142)		0.0108 (0.0156)
Eigenvector Centrality			-0.158*** (0.0523)	0.0254 (0.0642)			-0.0755** (0.0323)	-0.0573 (0.0501)
R-squared	0.042	0.002	0.010	0.044	0.708	0.707	0.708	0.709
Hamlet Fixed Effect	No	No	No	No	Yes	Yes	Yes	Yes

Notes: Same as Table 2. The sample is restricted to rural villages.

TABLE L.2. Rural Sample: The Correlation between Household Network Characteristics and the Error Rate in Ranking Income Status of Households

	(1)	(2)	(3)	(4)	(5)	(6)	(7)	(8)
<i>Panel A: Consumption Metric, Error Rate</i>								
Degree	-0.00537*** (0.00131)			-0.00545*** (0.00168)	-0.000801 (0.000810)			-0.000949 (0.00140)
Clustering		-0.0193 (0.0201)		-0.0183 (0.0195)		0.00341 (0.0125)		0.00134 (0.0133)
Eigenvector Centrality			-0.110** (0.0540)	0.00790 (0.0683)			-0.0172 (0.0348)	0.00731 (0.0580)
R-squared	0.020	0.005	0.009	0.021	0.653	0.653	0.653	0.653
<i>Panel B: Self-Assessment Metric, Error Rate</i>								
Degree	-0.00792*** (0.00148)			-0.00832*** (0.00191)	-0.00196** (0.000874)			-0.00171 (0.00154)
Clustering		-0.0220 (0.0213)		-0.0229 (0.0200)		0.0156 (0.0136)		0.0128 (0.0142)
Eigenvector Centrality			-0.143** (0.0603)	0.0335 (0.0760)			-0.0477 (0.0365)	-0.00651 (0.0597)
R-squared	0.036	0.009	0.014	0.037	0.650	0.649	0.649	0.650
<i>Panel B: Self-Assessment Metric, Error Rate</i>								
Degree	-0.00776*** (0.00124)			-0.00805*** (0.00158)	-0.00180** (0.000722)			-0.000464 (0.00125)
Clustering		-0.0324 (0.0209)		-0.0338* (0.0203)		0.00786 (0.0142)		0.00937 (0.0156)
Eigenvector Centrality			-0.151*** (0.0518)	0.0242 (0.0643)			-0.0669** (0.0316)	-0.0570 (0.0495)
R-squared	0.055	0.019	0.026	0.057	0.712	0.711	0.712	0.712
Hamlet Fixed Effect	No	No	No	No	Yes	Yes	Yes	Yes

Notes: Same as Table 3. The sample is restricted to rural villages.

TABLE L.3. Rural Sample: The Correlation Between Inaccuracy in Ranking a Pair of Households in a Hamlet and the Average Inverse Distance to Rankees

	(1)	(2)	(3)	(4)
<i>Panel A: Consumption Metric, Error Rate</i>				
Average Inverse Distance	-0.0452*** (0.0113)	-0.0346*** (0.0113)	-0.0157** (0.00713)	-0.00819 (0.0155)
Average Degree		-0.00410* (0.00210)	0.00443 (0.00368)	0.00442 (0.00372)
Average Clustering Coefficient		-0.00974 (0.0362)	0.00316 (0.0358)	0.00263 (0.0358)
Average Eigenvector Centrality		0.0661 (0.0918)	-0.0687 (0.122)	-0.0852 (0.127)
R-squared	0.007	0.007	0.125	0.186
<i>Panel B: Self-Assessment Metric, Error Rate</i>				
Average Inverse Distance	-0.0542*** (0.0126)	-0.0370*** (0.0130)	-0.0160** (0.00806)	-0.0110 (0.0176)
Average Degree		-0.00638*** (0.00226)	0.00136 (0.00436)	0.00102 (0.00442)
Average Clustering Coefficient		-0.0446 (0.0400)	-0.0143 (0.0437)	-0.00880 (0.0440)
Average Eigenvector Centrality		0.123 (0.107)	0.0857 (0.161)	0.0732 (0.161)
R-squared	0.010	0.012	0.153	0.230
<i>Panel C: Share of Don't Knows</i>				
Average Inverse Distance	-0.0651*** (0.0142)	-0.0364*** (0.0132)	-0.0240*** (0.00878)	-0.00811 (0.0161)
Average Degree		-0.00955*** (0.00230)	-0.00112 (0.00368)	-0.00104 (0.00361)
Average Clustering Coefficient		-0.0636 (0.0425)	-0.0434 (0.0397)	-0.0421 (0.0394)
Average Eigenvector Centrality		0.145 (0.110)	0.103 (0.143)	0.0481 (0.146)
R-squared	0.029	0.037	0.310	0.418
Demographic Controls	No	Yes	Yes	Yes
Hamlet Fixed Effects	No	No	Yes	Yes
Ranker Fixed Effects	No	No	No	Yes

Notes: Same as Table 4. The sample is restricted to rural villages.

TABLE L.4. Rural Sample: Empirical Results on Stochastic Dominance

	(1)	(2)	(3)	(4)
<i>Panel A: Consumption Metric</i>				
I fofd J	-0.0352 (0.0278)	-0.0557 (0.0412)	-0.0562** (0.0270)	-0.0833** (0.0386)
J fofd I	0.0177 (0.0256)		0.0367 (0.0240)	
<i>Panel B: Self-Assessment Metric</i>				
I fofd J	-0.0625** (0.0252)	-0.122*** (0.0367)	-0.0732*** (0.0250)	-0.135*** (0.0365)
J fofd I	0.0626*** (0.0235)		0.0757*** (0.0231)	
Non-Comparable	Yes	No	Yes	No
Demographic Controls	No	No	Yes	Yes
Stratification Group FE	Yes	Yes	Yes	Yes

Notes: Same as Table 8. The sample is restricted to rural villages.

TABLE L.5. Rural Sample: Empirical Results on Correlation between Hamlet Network Characteristics and Hamlet Level Error Rate

	(1)	(2)	(3)	(4)	(5)	(6)	(7)
<i>Panel A: Consumption Metric</i>							
Average Degree	-0.00599* (0.00342)						0.00207 (0.0206)
Average Clustering		-0.105 (0.0686)					-0.0626 (0.155)
Number of Households			-7.32e-05 (0.000523)				-0.000155 (0.000493)
First eigenvalue $\lambda_1(A)$				-0.00534 (0.00331)			-0.00371 (0.0120)
Fraction of Nodes in Giant Component					-0.0988* (0.0554)		-0.122 (0.108)
Link Density						-0.0347 (0.0983)	0.168 (0.198)
R-squared	0.173	0.173	0.167	0.174	0.177	0.168	0.170
<i>Panel B: Self-Assessment Metric</i>							
Average Degree	-0.0118*** (0.00393)						0.000302 (0.0214)
Average Clustering		-0.239*** (0.0842)					-0.263 (0.192)
Number of Households			0.000264 (0.000568)				-6.49e-05 (0.000607)
First eigenvalue $\lambda_1(A)$				-0.00882*** (0.00323)			-0.00361 (0.0117)
Fraction of Nodes in Giant Component					-0.171*** (0.0614)		-0.0917 (0.113)
Link Density						-0.135 (0.0891)	0.293 (0.177)
R-squared	0.239	0.245	0.222	0.236	0.244	0.225	0.228

Notes: Same as Table 9. The sample is restricted to rural villages.

TABLE L.6. Rural Sample: Rank Correlation on Targeting Type Interacted with Diffusiveness (Principal Component)

	(1)	(2)	(3)	(4)	(5)	(6)
<i>Panel A: Rank Correlation (Consumption)</i>						
Community x Diffusiveness		-0.118 (0.194)	-0.118 (0.194)	-0.106 (0.201)	-0.164 (0.206)	
Hybrid x Diffusiveness		-0.265 (0.180)	-0.265 (0.180)	-0.306 (0.204)		
Community	-0.0801 (0.0502)	-0.0128 (0.126)	-0.0128 (0.126)	-0.0333 (0.136)	0.0188 (0.139)	
Hybrid	-0.0753 (0.0481)	0.0869 (0.122)	0.0869 (0.122)	0.0925 (0.141)		
Diffusiveness		0.0498 (0.133)	0.0498 (0.133)	0.0925 (0.152)	0.136 (0.172)	0.0922 (0.152)
(Community or Hybrid) x Diffusiveness						-0.209 (0.178)
(Community or Hybrid)						0.0294 (0.123)
R-squared	0.010	0.021	0.021	0.168	0.211	0.164
<i>Panel B: Rank Correlation (Self-Assessment)</i>						
Community x Diffusiveness		0.205 (0.182)	0.205 (0.182)	0.120 (0.186)	0.0377 (0.197)	
Hybrid x Diffusiveness		0.113 (0.186)	0.113 (0.186)	0.111 (0.185)		
Community	0.109** (0.0472)	-0.00478 (0.125)	-0.00478 (0.125)	0.0395 (0.127)	0.0890 (0.137)	
Hybrid	0.0817* (0.0482)	0.0146 (0.129)	0.0146 (0.129)	0.0195 (0.129)		
Diffusiveness		-0.0585 (0.137)	-0.0585 (0.137)	0.0157 (0.146)	0.0765 (0.161)	0.0159 (0.145)
(Community or Hybrid) x Diffusiveness						0.113 (0.161)
(Community or Hybrid)						0.0298 (0.112)
R-squared	0.017	0.023	0.023	0.172	0.213	0.171
Stratification Group FE	No	No	No	Yes	Yes	Yes
Demographic Covariates	No	No	No	Yes	Yes	Yes

Notes: Same as Table 10. The sample is restricted to rural villages.

TABLE L.7. Rural Sample: Rank Correlation on Targeting Type Interacted with Diffusiveness (1-Simulated Error Rate)

	(1)	(2)	(3)	(4)	(5)	(6)
<i>Panel A: Rank Correlation (Consumption)</i>						
Community x Diffusiveness		0.256 (0.168)	0.238 (0.186)	0.270 (0.193)	0.233 (0.163)	
Hybrid x Diffusiveness		0.147 (0.158)	0.0769 (0.182)	0.0809 (0.187)		
Community	-0.0801 (0.0502)	-0.205** (0.0982)	-0.213** (0.107)	-0.240** (0.112)	-0.168* (0.0994)	
Hybrid	-0.0753 (0.0481)	-0.143 (0.0970)	-0.127 (0.105)	-0.140 (0.107)		
Diffusiveness		-0.165 (0.111)	-0.0712 (0.142)	-0.0611 (0.147)	-0.0494 (0.107)	-0.0630 (0.146)
(Community or Hybrid) x Diffusiveness						0.168 (0.166)
(Community or Hybrid)						-0.186** (0.0935)
R-squared	0.010	0.016	0.139	0.168	0.156	0.164
<i>Panel B: Rank Correlation (Self-Assessment)</i>						
Community x Diffusiveness		0.175 (0.179)	0.132 (0.185)	0.108 (0.193)	0.188 (0.160)	
Hybrid x Diffusiveness		-0.0935 (0.187)	-0.0847 (0.195)	-0.160 (0.191)		
Community	0.109** (0.0472)	0.0220 (0.104)	0.0544 (0.103)	0.0589 (0.110)	-0.0346 (0.0926)	
Hybrid	0.0817* (0.0482)	0.147 (0.105)	0.146 (0.109)	0.182* (0.106)		
Diffusiveness		-0.0818 (0.136)	-0.0434 (0.160)	0.00794 (0.161)	-0.0457 (0.128)	0.00448 (0.160)
(Community or Hybrid) x Diffusiveness						-0.0362 (0.170)
(Community or Hybrid)						0.125 (0.0949)
R-squared	0.017	0.028	0.146	0.187	0.172	0.179
Stratification Group FE	No	No	No	Yes	Yes	Yes
Demographic Covariates	No	No	No	Yes	Yes	Yes

Notes: Same as Table 11. The sample is restricted to rural villages.

APPENDIX M. ALTERNATIVE MEASURE OF INEQUALITY

In this section we use an alternative measure of inequality and control for it in our cross-village regressions. For 542 villages in our sample we observe from the 2003 Indonesia agricultural census the complete distribution of land holdings in each village in Indonesia. We use this data to construct an inequality measure for each village, and because this measure is based on a 100% sample of the census, it is measured with relatively little measurement error. We repeat all hamlet level regressions controlling for this new measure of inequality, reported below.

TABLE M.1. Alternate Measure of Inequality: Numerical Predictions on Stochastic Dominance

	(1)	(2)	(3)	(4)
I fofd J	-0.118*** (0.0161)	-0.243*** (0.0250)	-0.101*** (0.0176)	-0.222*** (0.0293)
J fofd I	0.132*** (0.0182)		0.134*** (0.0206)	
Non-Comparable	Yes	No	Yes	No
Demographic Controls	No	No	Yes	Yes
Stratification Group FE	Yes	Yes	Yes	Yes

Notes: Same as Table 6. The sample restricted to the 542 villages where we have the alternative inequality measure.

TABLE M.2. Alternate Measure of Inequality: Numerical Predictions on Correlation between Hamlet Network Characteristics and Hamlet Level Error Rate

	(1)	(2)	(3)	(4)	(5)	(6)	(7)
Average Degree	-0.0305*** (0.00599)						0.0521*** (0.0141)
Average Clustering		-0.349*** (0.0742)					0.436*** (0.148)
Number of Households			0.000396 (0.000357)				-4.58e-05 (0.000518)
First eigenvalue $\lambda_1(A)$				-0.0265*** (0.00486)			-0.0414*** (0.00741)
Fraction of Nodes in Giant Component					-0.481*** (0.0703)		-0.848*** (0.113)
Link Density						-0.426*** (0.121)	-0.271* (0.147)
R-squared	0.589	0.560	0.530	0.599	0.629	0.549	0.683

Notes: Same as Table 7. The sample restricted to the 542 villages where we have the alternative inequality measure.

TABLE M.3. Alternate Measure of Inequality: Empirical Results on Stochastic Dominance

	(1)	(2)	(3)	(4)
<i>Panel A: Consumption Metric</i>				
I fofd J	-0.0935*** (0.0193)	-0.136*** (0.0298)	-0.0707*** (0.0199)	-0.0975*** (0.0303)
J fofd I	0.0465** (0.0184)		0.0454** (0.0196)	
<i>Panel B: Self-Assessment Metric</i>				
I fofd J	-0.100*** (0.0177)	-0.170*** (0.0264)	-0.0715*** (0.0194)	-0.112*** (0.0288)
J fofd I	0.0730*** (0.0168)		0.0557*** (0.0188)	
Non-Comparable	Yes	No	Yes	No
Demographic Controls	No	No	Yes	Yes
Stratification Group FE	Yes	Yes	Yes	Yes

Notes: Same as Table 8. The sample restricted to the 542 villages where we have the alternative inequality measure.

TABLE M.4. Alternate Measure of Inequality: Empirical Results on Correlation between Hamlet Network Characteristics and Hamlet Level Error Rate

	(1)	(2)	(3)	(4)	(5)	(6)	(7)
<i>Panel A: Consumption Metric</i>							
Average Degree	-0.00704* (0.00369)						0.0142 (0.0138)
Average Clustering		-0.203*** (0.0728)					-0.296* (0.150)
Number of Households			0.000481 (0.000412)				0.000407 (0.000465)
First eigenvalue $\lambda_1(A)$				-0.00541* (0.00282)			-0.00932 (0.00813)
Fraction of Nodes in Giant Component					-0.141*** (0.0490)		-0.0815 (0.0814)
Link Density						-0.0419 (0.102)	0.280* (0.156)
R-squared	0.190	0.203	0.188	0.189	0.199	0.184	0.218
<i>Panel B: Self-Assessment Metric</i>							
Average Degree	-0.0116*** (0.00354)						0.00670 (0.0151)
Average Clustering		-0.283*** (0.0694)					-0.335** (0.141)
Number of Households			0.000901** (0.000422)				0.000560 (0.000482)
First eigenvalue $\lambda_1(A)$				-0.00613** (0.00265)			-0.00489 (0.00787)
Fraction of Nodes in Giant Component					-0.191*** (0.0492)		-0.0410 (0.0965)
Link Density						-0.187** (0.0835)	0.252 (0.160)
R-squared	0.261	0.278	0.259	0.254	0.271	0.255	0.278

Notes: Same as Table 9. The sample restricted to the 542 villages where we have the alternative inequality measure.