

Targeting High Ability Entrepreneurs Using Community Information: Mechanism Design In The Field

Reshmaan Hussam
Natalia Rigol
Benjamin N. Roth

Working Paper 20-082



Targeting High Ability Entrepreneurs Using Community Information: Mechanism Design In The Field

Reshmaan Hussam
Harvard Business School

Natalia Rigol
Harvard Business School

Benjamin N. Roth
Harvard Business School

Working Paper 20-082

Copyright © 2020 by Reshmaan Hussam, Natalia Rigol, and Benjamin N. Roth.

Working papers are in draft form. This working paper is distributed for purposes of comment and discussion only. It may not be reproduced without permission of the copyright holder. Copies of working papers are available from the author.

Funding for this research was provided in part by Harvard Business School.

Targeting High Ability Entrepreneurs Using Community Information: Mechanism Design In The Field

Reshmaan Hussam, Natalia Rigol and Benjamin N. Roth*

January 31, 2020

Abstract

Microentrepreneurs in low-income countries have widely varying marginal returns to capital, yet identifying those with the best opportunities remains a challenge due to a scarcity of verifiable information. With a cash grant experiment in India we demonstrate that community knowledge can help target high-growth microentrepreneurs; while the average marginal return to capital in our sample is 11% per month, microentrepreneurs reported to be in the top third of the community are estimated to have marginal return to capital between 23% and 35% per month. We cannot reject that microentrepreneurs ranked in the middle and bottom terciles of the community have a marginal return to capital of 0. Further we find evidence that community members distort their predictions when they can influence the distribution of resources. Finally we demonstrate that appropriately designed elicitation mechanisms can realign incentives for truthful reporting. These methods may be useful for using community information to target resources in other contexts, especially when targeting based on predicted treatment effects, or when community members may have incentives to distort their predictions.

*Hussam: Harvard Business School (email: rhussam@hbs.edu); Rigol: Harvard Business School (email: nrigol@hbs.edu); Roth: Harvard Business School (email: broth@hbs.edu). We would like to first thank our team for their tireless work on this project, especially our research manager Sitaram Mukherjee and our project assistants Prasenjit Samanta and Sayan Bhattacharjee. We are very grateful to Rohan Parakh for exceptional advice, research assistance, and field management. We are also grateful to Namita Tiwari, Suraj Jacob, and Meghana Mugikar for excellent assistance in the field. We thank Savannah Noray for meticulous technical assistance in Cambridge. We thank Shreya Chandra and Ana Paula Franco for help in preparing the paper for publication. This research was made possible with funding from the Asian Development Bank, Weiss Family Fund, PEDL, and the Schultz Fund at MIT. We received valuable feedback about this project from Rohini Pande, Benjamin Olken, Abhijit Banerjee, Pascaline Dupas, Esther Duflo, Peter Hull, Jonathan Roth, Christopher Woodruff, David McKenzie, Simone Schaner, Rema Hanna, Dean Karlan, and members of the development community more broadly. We are especially indebted to Arielle Bernhardt.

1 Introduction

Numerous experimental studies of microentrepreneurs in the developing world find widely heterogeneous returns to cash and credit (eg. Fafchamps et al. (2014); McKenzie et al. (2008); Banerjee et al. (2015)). Yet little is known about whether heterogeneity in returns to capital reflect idiosyncratic productivity shocks or persistent differences in entrepreneurial ability and opportunity. Moreover, to the extent that such persistent differences exist, lenders and foundations aiming to promote entrepreneurship often have little hard information with which to target high growth microentrepreneurs. We find that harnessing community information directly from a microentrepreneur’s peers may provide a viable approach to identifying high growth microentrepreneurs.

Our argument has three parts. First, we demonstrate that entrepreneurs in peri-urban Maharashtra have high quality information about one another along a variety of dimensions including marginal returns to capital. Their information is valuable for identifying high growth microentrepreneurs even after controlling for a wide range of demographic and business characteristics. Second we demonstrate that entrepreneurs manipulate their reports to favor themselves, their friends, and their family when the distribution of resources is at stake. Finally we identify several simple techniques motivated by mechanism design that effectively realign incentives for accuracy.

Specifically, we conducted a field experiment with 1,345 entrepreneurs from Amravati, a city in Maharashtra, India. We assigned respondents and their nearest neighbors to peer groups of 5 people. After collecting detailed baseline data from all respondents, we asked entrepreneurs to rank their peer group members on predicted marginal returns to capital, profits, and other firm, owner, and household characteristics. Once the community reports were complete, we randomly assigned USD 100 grants to one third of entrepreneurs in order to induce business growth and assess the accuracy of respondents’ predictions. We evaluate the accuracy of community information by comparing how well the rankings predict individuals’ true outcomes as reported at baseline or in subsequent follow-up surveys.

Our first main finding is that community members can identify high-return entrepreneurs. While the average marginal return to the grant was about 11% per month, our point estimates of the marginal returns to capital of entrepreneurs ranked in the top third range from 23% to 35%. Had we distributed our grants using community reports instead of random assignment, we would have more than tripled the total return on our investment.

To benchmark the value of community information, we compare its predictive accuracy against that of observable entrepreneur characteristics. We build a model to predict entrepreneurs’ marginal return to capital using a causal forest (a machine learning technique developed by Wager and Athey (2018) to predict heterogeneous treatment effects). We find that observable characteristics are indeed strong predictors of marginal return to capital. However, when we estimate marginal returns based on community information and control for the machine learning prediction, we still

find that those in the top tercile of the community prediction distribution earn 17% higher monthly returns than those in the bottom tercile. This finding suggests that community information is valuable above and beyond information that can be captured by observables.

Our second main finding is that strategic misreporting is a first-order concern when eliciting community information. By random assignment, half of respondents were told that their reports would be used only for research purposes (the “no-stakes” treatment) and the other half were told that their reports would be used to allocate USD 100 grants to members of their community (the “high-stakes” treatment). The correlation between community reports and true outcomes is on average 24% to 35% lower when allocation of resources is at stake, which significantly lowers the value of peer elicitation. We also identify who benefits from misreporting and by how much: we quantify the extent to which participants favor themselves, their family members, and their close friends (as identified by other group members).

Given the importance of strategic misreporting, we explore whether it is feasible to realign incentives to report truthfully. Alongside the high-stakes treatment, we cross-randomized treatments which varied respondents’ immediate benefit (or cost) for truthful responses. Respondents were assigned to report either in private or in a public setting, with their fellow neighbors observing their reports. Participants were also randomly assigned to receive monetary payments based on the truthfulness of their reports. Payments were calculated using the *Robust Bayesian Truth Serum* (RBTS), a peer prediction mechanism which determines participant scores as a function of the contemporaneous reports of other respondents.

Our third main finding is that methods grounded in mechanism design theory can be used to design a peer-elicitation environment in which truthtelling is incentive compatible. Monetary payments and public reporting do little to improve the accuracy of self-reports. But payments double the predictive power of reports that entrepreneurs make about other group members. We provide direct evidence that monetary payments reduce the likelihood that respondents favor their family members or their close friends. Finally, we find that public reporting doubles the predictive accuracy of reports about others when there are no stakes, but has no effect in a high-stakes setting. This nuanced finding may reflect a heterogeneous treatment effect, or a noisily estimated impact of observability on the quality of reports.

Beyond targeting cash grants to high-growth microentrepreneurs, the methods in this paper may prove useful in other contexts. We provide an experimental framework for predicting heterogeneous treatment effects *before* treatment implementation. Namely, by asking subjects of the experiment to predict their own and their peers’ treatment effects, researchers can leverage information embedded in their experimental contexts. This may serve as a complement to recently developed techniques to estimate heterogeneous treatment effects after the experiment is complete using observable characteristics (e.g. Wager and Athey, 2018). Dal Bó et al. (2018) employ a similar experimental design.

Our findings contribute to several literatures. The idea that social networks—friends, family, colleagues—are a rich source of information has deep roots in development economics. Of particular relevance are the studies that use community reports to inform policy, broadly construed. In the community targeting literature, Alatas et al. (2012) investigate whether villagers in Indonesia can select the village’s poorest residents to receive government transfers. They find that community targeting performs worse than a Proxy Means Test for assessing households’ level of consumption but better at capturing a household’s perception of their own poverty status. Basurto et al. (2019) find that village chiefs in rural Malawi are more likely to target fertilizer subsidies to households that self-report they would benefit from agricultural inputs than the standard PMT method. In the referrals literature, Beaman and Magruder (2012) find that high-quality workers in Kolkata, India can refer other high-quality laborers when incentivized to do so. In contrast, Bryan et al. (2015) find that borrowers in South Africa can do no better than the lending institution in selecting high-quality borrowers among their peers.¹ Giné and Karlan (2014) randomize joint versus individual liability contracts in microfinance and find little evidence that group members utilize local information to screen their partners. Lastly, Maitra et al. (2017) show that local traders in India can select microcredit borrowers for whom credit leads to larger increases in production and income than for borrowers selected by standard microcredit, with the caveat that both the selection method (traders’ screening versus self-selection into microfinance) and the contract type (individual versus joint liability loans) covary.

Our findings provide new insight into the depth and breadth of social knowledge contained in rural and peri-urban networks. The Alatas et al. (2012) study demonstrates that community members have reliable information regarding observable characteristics (wealth) of people across their social network. We show that community members can predict marginal returns to capital, a metric that is difficult to estimate even using rich observables or expert opinions. This is evidence that community members have accurate knowledge of one another that is much deeper than what has been previously shown.

Community knowledge—even if accurate—is only useful for allocative decision-making if those eliciting the information can be confident that they will gather *truthful* reports. And when allocation of resources is at stake, there is reason to be concerned that community members will lie. Yet strategic misreporting is not typically addressed in the design of programs which rely on community information to make decisions. For example, community-driven development projects, which leverage community information or community action to make decisions regarding public goods expenditures, are rarely designed to account for strategic behavior (Mansuri and Rao (2004); King (2013)).

We contribute to a young literature which addresses strategic misreporting in targeting programs. Alatas et al. (2019) examine whether elite capture poses a problem for community reporting,

¹All referred applicants had to also meet the bank’s eligibility criteria and, unlike in our setting, South Africa has a well-functioning credit bureau.

but incentives to manipulate the distribution of resources may extend beyond community elites. Though Alatas et al. (2019) conclude that elite capture is not a significant concern, we find that misreporting is common when community members are told that their reports will influence distribution of grants. Importantly, we find that community members distort their reports in favor of their family and friends, rather than toward community elites. Alatas et al. (2012) also elicit community reports in public in order to incentivize truth-telling. However, their experiment is not designed to evaluate the impact of public reporting on the accuracy of reports. Through random variation of the elicitation environment, we show that public reporting is not effective for realigning incentives with truth-telling when allocation of resources is at stake.

The rest of the paper proceeds as follows. Section 2 introduces the setting and study sample. Section 3 describes our conceptual approach to designing the elicitation environment, Section 4 describes our experimental setting and design, Section 5 describes the data and provides a brief discussion of the randomization, Section 6 discusses how well community members know one another, Section 7 provides our main results, and Section 8 concludes.

2 Study Sample and Context

Our study takes place in Amravati, a city of about 550,000 people in the state of Maharashtra, India. Households in our sample come from nine neighborhoods along the perimeter of Amravati; we selected these neighborhoods because they have a relatively high proportion of microentrepreneurs.² These are densely packed peri-urban slums; in each of these neighborhoods, there are roughly 900 household dwellings in a 500 by 700 ft. area. In September 2015, we conducted a complete door-to-door census of these neighborhoods, which encompassed 5,573 households. Based on households' responses to the census, we determined their eligibility for the study. In line with selection criteria of other recent "cash-drop" experiments (see e.g. de Mel et al. (2008)), all households in our sample have at least one enterprise with (i) USD 1,000 or less in total working and durable capital and (ii) no paid, permanent employees.³ Almost 30% of households in these neighborhoods owned at least one business and were eligible (1,576 households). Entrepreneurs in 1,345 of these households agreed to participate in our study so our sample population is reasonably representative of the universe of eligible enterprises in Amravati.

Characteristics of Microenterprise Owners. The modal entrepreneur in our sample is 40 years old and has roughly 8 years of formal education. Approximately 60% are male and almost all are married. Most entrepreneurs operate their business close to home, but they operate across a wide range of activities. About 30% of sample entrepreneurs work in manufacturing, typically as a tailor or stitcher. Another 30% work in services, mainly in food preparation and hair salons. A

²Our selection of neighborhoods was based on advice from local officials in the District Collector's Office. The nine neighborhoods are: Belpura, Vilash Nagar, Mahajan Pura, Akoli, New Saturna, Old Saturna, Wadali, and Pathan Chawk.

³Following de Mel et al. (2008)'s selection criteria, we excluded farmers and self-employed service people, such as domestic helpers and teachers. If there were multiple business owners in the household, we required that the household have at most USD 2000 in combined business capital.

further 30% work in retail, most commonly running a grocery shop. Outside of these three sectors, entrepreneurs are spread evenly across construction and livestock rearing. On average, sample entrepreneurs earn profits of Rs.4500 per month (USD 2.5 per day), which accounts for roughly half of their household income. Entrepreneurs also face a significant amount of risk: between the baseline and one year follow-up survey, about 10% of businesses in control group households were shut down. In over a third of these cases, the reason given for enterprise closure was illness of the business owner. Perhaps as a means of insuring against risk, households diversify across types of income-generating activities: in half of sample households, there is at least one fixed salary or daily wage worker and one fifth of households own more than one business.

Characteristics of Microentrepreneurs’ Peer Networks. In order to elicit entrepreneurs’ knowledge of one another, we assigned study participants to peer groups of roughly five people based on geographical proximity. Peer groups are the unit of information collection: entrepreneurs are asked to report on only themselves and their other group members, not on the entire community. Importantly, we find that peers know their group members well. On average, peers reported that they visited another group member on 22 occasions in the previous 30 days. Respondents were unable to identify another group member in less than 1% of cases. Two-thirds of respondents identify at least one other group member as a family member or close friend. In 70% of groups, at least two people operate a business in the same (broad) industry category. Entrepreneurs also actively maintain strong social ties within their group: over 50% of respondents reported that they regularly discuss private family and business matters with at least one other group member. And, entrepreneurs have at least some knowledge of almost every group member: 87% of respondents correctly identified for all other group members whether that person owned a motorcycle (half of respondents are motorcycle owners) and 80% correctly identified who among their peers had young children living in their home. In Section 6, we evaluate how well community members can predict household and enterprise characteristics.

3 Mechanisms to Incentivize Truthful Revelation

Community knowledge is only valuable for decision-making if it is incentive compatible for people to report truthfully. When the allocation of resources is at stake, strategic misreporting may be an important concern. Mechanism design offers an array of tools that make truth-telling incentive compatible in theory, and one of our goals is to understand which of these tools work to realign incentives in practice. In this section, we describe our conceptual approach for designing and evaluating the peer ranking elicitation environment.

Public Reporting. Fear of public reprisal is a powerful deterrent to socially undesirable behavior. This insight has been applied to incentivize costly actions across a number of settings (notable examples include using public notification of individuals’ voting record (Gerber et al., 2008) or electricity usage (Allcott and Rogers, 2014) to encourage behavioral change). Intuitively, conducting peer elicitation in public may reduce strategic misreporting because participants care about their

reputation for honesty. At the same time, publicity may exacerbate pressure to rank one’s family, friends, and influential members of the community more highly. To assess the relative strength of these competing effects, we randomly vary whether the peer elicitation exercise takes place in a private or public setting.

Paying for Truthfulness. Explicit monetary incentives for accuracy offer a promising deterrent to misreporting. One straightforward way to implement monetary incentives would be to pay respondents based on the closeness of their reports with an ex-post measure of accuracy. But often ex-post measures of accuracy are unavailable, or prohibitively costly to collect (such as in the case of estimating marginal returns to capital, which can never be confirmed for an individual entrepreneur). Further, even when signals of ex-post accuracy exist, using them necessitates a time-lag between the moment of elicitation and subsequent payment for reports. In settings with weak institutions, where trust in outsiders is minimal, respondents may demand to be paid contemporaneously with their reports. To circumvent these concerns we evaluate monetary incentives delivered via a peer prediction scheme, which rewards respondents based exclusively on their own reports and the contemporaneous reports of their peers. The particular payment rule we use is the *Robust Bayesian Truth Serum*, described in detail in Appendix A3.

Zero-sum Elicitation. During our peer elicitation exercise, entrepreneurs rank one another on metrics of business growth and profitability. Within each group of entrepreneurs, we evaluate two forms of community rankings: rankings relative to the particular members of the group, and reports placing each entrepreneur in quintiles relative to the community at large. The former has a zero-sum nature, in which promoting someone’s position necessitates diminishing another’s, and may therefore be more effective at inducing truthful reports (a respondent cannot merely place everyone in the highest position). However, if group members have correlated attributes, then these rankings may be less informative than rankings that assess each entrepreneur relative to the broader community. By examining both mechanisms we investigate which of these concerns dominates in practice.

Cross-Reporting. In the spirit of cross-reporting techniques which play a prominent role in mechanism design and implementation theory (see Maskin (1999)), we ask respondents to identify each group member’s closest peer in the group, with the intention of exploring whether group members identified as close peers distort their reports to favor one another. We also ask respondents to identify who in their peer group has the most accurate information regarding each ranking metric.

4 Experimental Design

4.1 Design of the Peer Elicitation Exercise

Recruitment. In October 2015, we visited the 1,576 eligible households and invited them to participate in our study. At the time of recruitment, households were told that a research team was

conducting a project to study entrepreneurship and business growth.⁴ In December 2015 - April 2016, we conducted baseline surveys of the 1,345 sample households. Separately, we also assigned respondents to groups of five based on geographic proximity, for a total of 274 groups across all neighborhoods.⁵ Once all baseline surveys in a given neighborhood were complete, surveyors returned to sample households to invite respondents to a meeting at the local town hall. Respondents were not given any information regarding the content of the meeting, or that they would be placed into groups with their peers. They were told, though, that to thank them for their participation in the study the research team would conduct a public lottery where some participants would be awarded a USD 100 grant.

Explanation of the Exercise to Respondents. Upon arrival at the town hall, respondents were each given 20 lottery tickets. They were told that, at the end of the activity, all people present would put their lottery tickets into an urn and grant winners would be selected by drawing lottery tickets. Participants were then separated and individually paired with a surveyor. Surveyors explained to participants that they would be asked to provide information about themselves and their neighbors. In order to ensure that participants were introduced to the elicitation exercise in a clear and consistent way, we created animated videos to introduce respondents to the concepts covered in the rankings questions and to guide them through the activity. When explaining the concept of marginal return to capital, we used examples to emphasize to respondents that an entrepreneur's projected marginal returns corresponds to their expected *change* in profits in response to the grant, and not their *level* of profits. After watching the videos, participants completed a series of quizzes to test their understanding of the activity and concepts. The introduction and subsequent ranking activity took place behind a privacy screen. The screen was there to ensure that coordination of responses would not be possible (as explained below, after collecting the rankings from respondents in the public reporting treatment, rankings were disclosed to their group members.) Surveyors also told participants which of their neighbors they would be ranking and gave them four to six placards, each with the name of a group member.

Questions Asked in the Ranking Exercise. First, we asked participants to rank themselves and their peers on predicted marginal returns to a USD 100 grant. We then asked respondents to rank themselves and their peers across several additional entrepreneur characteristics: educational attainment; average number of hours spent at work per week; performance in a digit span memory test; and, projected monthly profits 6 months post-grant disbursal, if the business owner were to receive a USD 100 grant. We also asked about a number of household-level characteristics: average monthly income over the past year, total value of assets; total medical expenses in the past 6 months; and, loan repayment trouble over the previous year.

To minimize respondent fatigue peer groups completed the ranking exercise only for a ran-

⁴No information regarding the community information nature of the project was disclosed to respondents at this time.

⁵We organized respondents into groups that would minimize the geographic distance between study households. The total number of respondents per neighborhood was not always a multiple of 5, so some groups had 4 or 6 clients.

domly assigned subset of these metrics (but all respondents completed the marginal returns ranking). For details on the assignment of ranking questions by treatment group, see the Appendix. And, participants completed both relative and quintile rankings for questions on marginal returns, business profits, and household income and assets, but only relative rankings for the remaining questions (this was also done to reduce fatigue). Finally, respondents were asked to cross-report on their peers: they identified one another’s closest peer in the group and, for each ranking question, respondents identified the group member they believed would have the information required to answer the question most accurately.

4.2 Description of Treatments

Respondents were cross-randomized (at the group level) to give their ranking reports under the following three treatment conditions, for a total of eight treatment cells: *No Stakes* vs. *High Stakes* (S_0 vs. S_1), *Private* vs. *Public* (P_0 vs. P_1), and *No Payments* vs. *Payments* (T_0 vs. T_1). We also randomly selected one-third of our sample to receive USD 100 grants. Grant randomization occurred at the individual level and was stratified by group. See Figure 2 for the randomization design.

High Stakes Environment (S_0 vs. S_1). For this treatment, participants were told that their responses in the ranking exercise would help determine the winner of the lottery that would occur at the completion of the activity. All participants across treatment groups were given twenty lottery tickets upon arrival at the town hall. Respondents in the high stakes treatment were told that, for each question, the peer ranked highest (on average) by group members would receive extra lottery tickets, and so would have a better chance of winning.⁶ In order to ensure that we would have sufficient power to evaluate the quality of predictions from the marginal returns rankings, all participants completed this ranking in a no-stakes setting (the marginal return ranking occurred prior to any mention of the high stakes treatment).⁷

Public Reporting (P_0 vs. P_1). Participants in both the Public and Private Reporting groups responded to each ranking question behind a privacy screen, in the presence of only their surveyor. But in the Public treatment, after completing each ranking question, peers came to the center of the room and sat in a circle with their response clipboard in front of them. Participants were told that they were doing this so that the survey coordinator could record their responses, but the primary purpose was to give them the opportunity to observe one another’s rankings.⁸ Crucially,

⁶We did not tell participants how many extra lottery tickets would be awarded to the person ranked highest; in order to keep the randomization as close to uniform as possible, we awarded only one extra lottery ticket per ranking. Respondents were in a high stakes setting for four ranking questions, and so a person in this treatment group could win at most four extra lottery tickets. Participants completed all rounds of ranking questions prior to the disbursal of the extra lottery tickets.

⁷Measures of profits among microentrepreneurs in settings like this one are notoriously noisy (see, for instance, de Mel et al. (2009)). Due to budget constraints, our experiment is just powered to detect how well marginal returns rankings predict realized marginal returns when accuracy of reports is not confounded by the incentive to lie present in a high-stakes setting.

⁸Surveyors report that respondents did in fact almost always look at their peers’ rankings.

participants understood ahead of doing the ranking exercise that their peers would see their responses. This was described to them in their introductory animation video and, to ensure that participants understood the set-up, groups performed several practice rounds. In the privacy treatment, respondents completed all ranking questions before interacting with peers and, even after the activity was completed, group members did not see each other’s individual responses.

Payments for Truthfulness (T_0 vs. T_1). The introductory video for participants in the monetary incentives group explained that they would be paid per ranking question, based on the truthfulness of their responses. As explained in Appendix A3 we did not explain the details of the RBTS scoring rule to participants. Instead, participants were told that people who reported what they truly believed would receive an extra Rs.100 on average (which is equivalent to 2/3 of the average daily wage). Payments were calibrated using the empirical distribution of reports from Rigol and Roth (2017) to maximize the strength of the incentive to tell the truth while adhering to a project budget constraint. Participants were shown an introductory video providing a basic overview of the payment rule and an explanation of the reporting requirements. Groups that were not in the monetary payments treatment were given a lump sum payment to compensate them for their time.

Enterprise Grant. Upon completion of the peer elicitation exercise, group members came to the center of the room and placed their lottery tickets into an urn. One respondent was blindfolded and then drew tickets to award USD 100 grants to one or two group members (the number of winners per peer group was determined by random assignment). Prior to grant randomization participants filled out worksheets specifying how they would invest the grant if they won. Participants were encouraged to invest grant money into their enterprise although this was not enforced. Grant money was distributed to winners via bank transfer.

4.3 Overview of Identification Strategy

In this section we provide an overview of our identification strategy; formal regression specifications are deferred to Sections 6 and 7.

Random assignment allows us to use the difference between post-period profits of grant winners and post-period profits of grant losers as an estimate of the average marginal return to the grant. We therefore identify the informational value of community members’ reports by testing the predictive power of respondents’ marginal return rankings against our estimates of true marginal returns.

Next, we assess whether community information extraction is susceptible to strategic misreporting when allocation of resources is on the line. We measure accuracy by comparing peer reports to self-reported values that participants provided at the time of the baseline survey.⁹ By comparing accuracy of peer reports for participants in the *No Stakes* and *High Stakes* groups (S_0 vs. S_1), we

⁹In order to ensure that we would have sufficient power to test predictions from the marginal returns rankings, all participants completed this ranking in a no-stakes setting.

identify the effect on strategic misreporting of shifting the elicitation environment to one in which reports can have consequences for allocation of grants.

Finally, we measure the efficacy of mechanisms to realign incentives for truthful reporting: a comparison of the accuracy of peer reports in the *Private* versus *Public* treatments (P_0 vs. P_1), or in the *No Payments* versus *Payments for Truthfulness* treatments (T_0 vs. T_1), identifies the effect each of these mechanisms has on respondents' truthfulness. Because we cross-randomize treatments, we can separately identify the strength of these mechanisms in the benchmark, *No Stakes* setting, and in the *High Stakes* setting, where respondents have a counteracting incentive to lie.

5 Data and Randomization Checks

Description of the Data. Our main analysis uses data from respondents' peer rankings during the elicitation exercise and from respondent surveys. Baseline surveys were conducted between December 2015 and April 2016, and three follow-up surveys were conducted between May 2016 and March 2017. For all survey rounds, each business owner in the household completed a detailed business module about her own enterprise and answered questions about her well-being. The business module included questions on enterprise costs; revenues; profits; seasonality; inventories; labor inputs; assets; and business history. At baseline, entrepreneurs also completed a digit span test and a set of psychometric questions.¹⁰ In each survey round, the study respondent also provided information regarding her household's finances. The household-level module included questions on income, health expenditures, credit history and loan repayment issues, and assets. For the asset section, the respondent indicated whether the household owned a particular type of asset and its current resale value. Surveyors were trained to visually verify that the household owned each of the assets about which they reported. At baseline, the respondent also completed a full household roster with education and labor history for each household member. For a complete timeline of the project and data, see Figure 3.

Randomization Checks. In Appendix Table 1, we present the randomization check of baseline characteristics by treatment. To check for balance we estimate the model,

$$Characteristic_{ij} = \tau_0 + \tau_1 Treatment_j + \epsilon_{ij}$$

where i indexes the individual and j indexes the group. $Treatment_j$ is a dummy for whether

¹⁰ Respondents answered each psychometric question in the module by providing their agreement with the given statement, where agreement was rated on a scale of one to five, with five indicating strong agreement and one indicating strong disagreement. A detailed description of the psychometric assessment module is in Appendix A4. The psychometric module questions are organized according to categories developed by industrial psychologists: polychronicity measures the willingness to juggle multiple tasks at the same time (Bluedorn et al. (1999)); impulsiveness is a measure of the speed at which a person makes decisions and savings attitudes (Barratt Impulsiveness Scale); tenacity measures a person's ability to overcome difficult circumstances (Baum and Locke (2004)); achievement is a measure of satisfaction in accomplishing a task well (McClelland (1985)); and locus of control measures a person's willingness to put themselves in situations outside of their control Rotter (1966).

the group was assigned to the *NoStakes* vs. *HighStakes* treatment (columns 1 and 2), the *NoPayments* vs. *Payments* treatment (columns 3 and 4), and the *Private* vs. *Public* treatment (columns 5 and 6).

The odd columns 1-7 show the average of each characteristic for the control group in each block. So column 1 shows the means of characteristics for groups that were assigned to *NoStakes*. The even columns show τ_1 for each treatment (the difference between treatment and control characteristics). The characteristics in *Panel A* are about the entrepreneur who was ranked during the ranking exercise and in *Panel B* are about her primary business. In *Panel C*, we show household level baseline measures. The variables “Value of Business Assets” and “Avg. Monthly Profits” are shown as aggregates over all household businesses. So if the ranked entrepreneur is the only business owner in the household, these reflect the values of only her businesses.

The majority of entrepreneur and household characteristics are balanced across treatment groups. Entrepreneurs assigned to *Payments* report lower household monthly income and entrepreneurs assigned to *Public* report lower value of household assets. At the bottom of the table, we present the F-test of whether the treatment group coefficients are jointly equal to zero. None of the joint tests of equality are rejected, suggesting that the randomization was effectively implemented.

6 Background Results - Entrepreneurs’ Community Knowledge

We begin our empirical analysis by investigating the depth of community members’ knowledge of one another. As discussed in Section 2, entrepreneurs have close social ties with peers in their neighborhood. Over half of respondents regularly discuss private family and business matters with at least one other group member and on average group members visit each other 22 times per month. In this section we show that community members have accurate knowledge about one another’s concurrent household finances and enterprise characteristics. In our main empirical analysis (Section 7), we will argue that community members also make accurate forward-looking predictions about entrepreneurs’ marginal returns.

During the ranking exercise, community members reported on their peers’ average monthly household income; predicted monthly profits if they were to receive a USD 100 grant; total value of household assets; household medical expenses over the previous six months; average weekly work hours; and, predicted performance on a working memory test.¹¹ At baseline, we asked each entrepreneur to self-report answers to these same questions (at the time of the baseline survey, respondents had no knowledge of the purpose of the study or of the peer ranking activity). To evaluate the accuracy of community reports, we estimate the relationship between entrepreneurs’ self reports and community members’ reports for that person. We use the following regression

¹¹We use a digit span test, which is a commonly used test for working memory. Respondents are shown flashcards with an increasing number of digits and asked to recall the numbers from memory. The surveyor records the total number of digits that the respondent correctly repeated back.

model:

$$Y_{ijq} = \beta_0 + \beta_1 \overline{Rank}_{ijq} + \gamma_n + \theta_m + \tau_s + \epsilon_{ijq}, \quad (1)$$

where $\overline{Rank}_{ijq} = \sum_{k=1}^n \frac{1}{n} * Rank_{ikjq}$, n is the total number of group members in group j and $Rank_{ikjq}$ is the rank that person k in group j assigns to person i (also in group j) on question q . So $\overline{Rank}_{ijq}$ is the average rank assigned to person i by the members of group j on question q . Y_{ijq} is the corresponding outcome (baseline survey self report) for question q of person i . To improve precision, we add neighborhood (γ_n), survey month (θ_m), and surveyor (τ_s) fixed effects. Standard errors are clustered at the group level.

In Table 1, we present the estimates of Specification 1.¹² To allow for comparability of estimates across questions, in *Panel A* we convert each outcome and the corresponding average rank for each question into percentiles. So, a 1 percentile increase in $\overline{Rank}_{ijq}$ is associated with a β_1 percentile increase in the outcome variable Y_{ijq} . In *Panel B*, we present results in levels of the outcome and the average rank, so that a 1 unit increase in $\overline{Rank}_{ijq}$ is associated with a β_1 increase in the value of the outcome variable Y_{ijq} .

Entrepreneurs have substantial knowledge of their peers' household and enterprise characteristics. For example, in column 3 of *Panel A*, a 1 percentile increase in the assets rank is associated with a 0.22 [SE=0.03] percentile increase in the distribution of actual household assets. They can also accurately assess even difficult to observe characteristics: for instance, a one unit increase in the average rank level is associated with a 0.62 [SE=0.10] extra digits recalled in the Digit Span Memory Test (column 5 of *Panel B*). So moving from the 5th percentile to the 95th percentile in average digitspan rank is associated with a doubling of the total number of digits an entrepreneur recalls. To contextualize the size of these estimates, we regress the business profits percentile on the percentile of the education of the entrepreneur and also the household assets percentile on the household income percentile: a 1 percentile increase in the education distribution is associated with a 0.12 percentile increase in the distribution of business profits, and a 1 percentile increase in the income distribution is associated with a 0.33 percentile increase in the assets distribution.

7 Main Results

7.1 Entrepreneurs' Average Marginal Returns to Capital

In the next section, we will investigate whether community members can accurately predict one another's returns to the grant. First, we assess the average impact of the intervention on entrepreneurs' profits. Following de Mel et al. (2008), we estimate average marginal returns to the grant with the primary specification,

$$Y_{ijt} = \alpha_0 + \alpha_1 Winner_{it} + \gamma_i + \sum_{t=1}^3 \delta_t + \theta_m + \tau_s + \epsilon_{ijt}. \quad (2)$$

¹²In Table 1, we pool across all treatment groups: *No Stakes* vs. *High Stakes* treatment, the *No Payments* vs. *Payments* treatment, and the *Private* vs. *Public*. In Sections 7.5 and 7.6, we break these estimates up by treatment.

where Y_{ijt} measures either total household business profits or household income of person i in survey round t .¹³ We measure business profits by asking entrepreneurs the following question: “Now that you have thought through your sales and your expenses from the past 30 days, I would like you to think about the profits of your business. By business profits, I mean taking the total income received from sales and subtracting all the cost of producing the items (raw material, wages to employees, fixed costs, etc). Can you tell me your business profits in the past 30 days?”¹⁴ Household income is also measured using a single question: “What is your total household income over the past 30 days from all income generating activities?” Like de Mel et al. (2008), we remove the outliers of the household income and total profits distributions (levels) by trimming the top 0.5% of both the absolute and percentage changes in profits measured from one period to the next. We also estimate regression Specification 2 for $\log(Y_{ijt} + 1)$ and the inverse hyperbolic sine of income and profits, using the untrimmed distributions.¹⁵ In the main specification, we utilize three rounds of follow-up surveys, so t ranges from 0 (baseline) to 3. $Winner_{it}$ is an indicator for whether person i won a grant at or before survey round t . Note that $Winner_{it}$ is 0 at period $t = 0$ for all people i . We also include the following fixed effects: person (γ_i), survey round (δ_t), survey month (θ_m), and surveyor (τ_s). Standard errors are clustered at the group level. The coefficient of interest in regression Specification 2 is α_1 , which measures average marginal return to the grant in the sample.

In the *No Stakes* treatment group, assignment of grant winners was uniformly random: all participants received twenty lottery tickets and each group member was equally likely to have their tickets drawn from the urn. But, as described in Section 4.2, respondents in the *High Stakes* group were eligible to receive up to four extra lottery tickets, based on whether their peers ranked them highest for the treatment questions.¹⁶ To account for this, we weigh all regressions by the *propensity score*—i.e. the probability of being assigned to the relevant treatment (Rosenbaum, 1987). In our setting, the probability of being assigned to treatment is fully determined by the number of lottery tickets that a subject receives, and the number of grants randomly allocated within each group. For instance, in a group with just one grant winner, the observation corresponding to respondent i who won the grant weighted by i ’s inverse probability of winning the grant lottery, $\frac{\text{Total Tickets}}{\text{Tickets Held by Subject } i}$. And the observation corresponding to a respondent i who did not win a grant is weighted by i ’s inverse probability of losing the lottery, $\frac{\text{Total Tickets}}{\text{Total Tickets} - \text{Tickets Held by Subject } i}$. In Appendix Figure A1, we plot the distribution of lottery tickets in the sample.

Table A2 presents results from estimating Specification 2. We find that the grant had a large

¹³Bernhardt et al. (2019) reanalyze data from several cash-drop experiments with microentrepreneurs and find that measures of returns to capital differ substantially when analyzed at the household versus enterprise level. We therefore aggregate profits of all household businesses, for all specifications.

¹⁴de Mel et al. (2009) find that asking one aggregate summary measure (rather than for the components) reduces noise in the estimation of profits.

¹⁵The results remain nearly identical whether we log-transform the trimmed or untrimmed income and profits distributions.

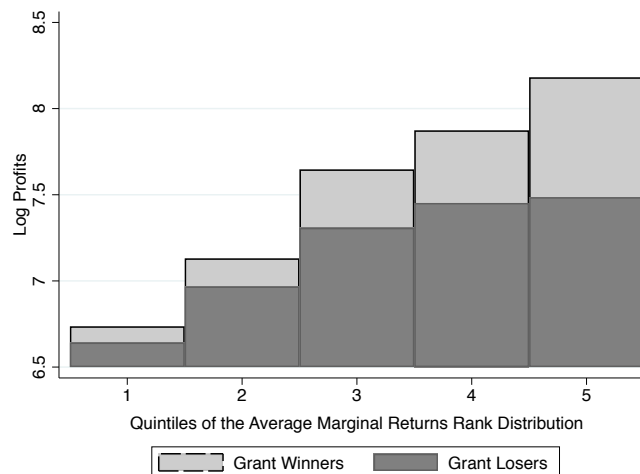
¹⁶For a more detailed description of the *High Stakes* treatment, please refer to Section 4.2.

positive effect on household income and total household profits. On average, households that win grants report an extra Rs.568.0 [SE=405.8] in household income and an extra Rs.684.8 [SE=319.1] in total household profits over households that were not awarded grants. These gains in household income and profits represent very high marginal returns to the grant: point estimates suggest that on average, households earn returns of 9.5% – 11.4% per month.¹⁷ These estimates are in line with average returns estimated from cash grants in other settings: de Mel et al. (2008) find marginal returns of 7.6% per month in response to a USD 100 grant and Fafchamps et al. (2014) find marginal returns of 9.7% per month in response to a USD 120 grant.

7.2 Can Communities Predict Entrepreneurs’ Marginal Returns To Capital?

Our measure of community knowledge is entrepreneurs’ average marginal returns rank. Respondents were asked, “Could you please rank your group members in order of who you think had the highest marginal returns to the Rs.6,000 grant? In other words, who would gain the most in monthly profits, or who would grow their business the most, from receiving a Rs.6,000 grant?” An entrepreneur’s average marginal returns rank is the mean of all the ranks assigned to her by her group members. We plot the distribution of average rank, which takes on values between one and five, in Appendix Figure A3. Since group members are in full agreement about an entrepreneur’s rank in fewer than 15% of cases, the distribution of average marginal return rank values is relatively smooth.

Figure 1: Marginal Returns to the Grant by Percentile of the Community Ranks Distribution



Notes: This figure plots the log of average post-grant profits (y-axis) by quintiles of the average marginal returns rank distribution (x-axis). The dark gray bars correspond to log profits of entrepreneurs who did not win grants and the light gray bars correspond to log profits of entrepreneurs who did win grants. Marginal returns rank is the rank assigned by a peer when asked “Could you please rank your group members in order of who you think had the highest marginal returns to the Rs.6,000 grant? In other words, who would gain the most in monthly profits, or who would grow their business the most, from receiving a Rs.6,000 grant?” Average marginal returns rank is the mean of the marginal returns ranks assigned to an entrepreneur by her peers and by herself.

¹⁷We arrive at this number by dividing the marginal increase in monthly income and profits by the size of the grant (Rs.6000).

In Figure 1, we plot the log of post-grant profits by grant treatment assignment and by quintile of average marginal returns rank. Each bar corresponds to average post-grant profits for entrepreneurs in a given quintile of the marginal returns rank distribution. Dark gray bars are profits of grant losers and light gray bars are profits of grant winners. We find that the gap in post-period profits between grant winners and grant losers—in other words, entrepreneurs’ marginal return to the grant—is increasing in the community’s rank report.

Figure 1 suggests both that there is significant heterogeneity in returns to the grant and that community members are able to identify accurately the ordering of their peers’ heterogeneous returns ex-ante. We further illustrate the accuracy of community members’ predictions in Figure 4. In that figure, we plot kernel-weighted local polynomial regressions (degree 1) of log profits at follow-up for grant winners and for grant losers on average marginal return rank percentile. We find that an entrepreneur’s marginal returns rank is strongly correlated with her increase in realized profits in response to the grant: below the 35th percentile of the ranks distribution, post-grant profits for winners and losers are statistically indistinguishable. But above the 35th percentile, the distance between treatment and control profits increases with marginal returns rank – this increasing distance is a measure of respondents’ prediction accuracy.

Our main specification is a difference-in-differences estimation of the relationship between community ranks and marginal returns to the grant. We extend the model from Specification 2 to incorporate peer ranks:

$$Y_{ijt} = \alpha_0 + \alpha_1 \text{Winner}_{it} + \alpha_2 \text{Winner}_{it} \times \overline{\text{Rank}}_{ij} + \gamma_i + \sum_{t=1}^3 \delta_t + \theta_m + \tau_s + \epsilon_{ijt}. \quad (3)$$

where $\overline{\text{Rank}}_{ij} = \sum_{k=1}^n \frac{1}{n} * \text{Rank}_{ikj}$, n is the total number of group members in group j and Rank_{ikj} is the rank that person k in group j assigns to person i (also in group j). So $\overline{\text{Rank}}_{ij}$ is the average marginal returns rank assigned to person i by the members of group j . The coefficient α_2 identifies the average additional marginal return to capital associated with a one unit increase in marginal return rank. The difference-in-differences specification estimates α_2 for a model in which marginal return increases linearly in the value of average rank. Motivated by the non-parametric estimates in Figure 4, we also estimate a non-linear model in which the ranks distribution is divided into terciles and rank tercile is interacted (as above) with Winner_{it} . In Table A3, we show that the sample is balanced across rank terciles and grant treatment groups at baseline. In Appendix Figure A4, we replicate Figure 4 with baseline profits and show that differences in marginal returns to the grant are not driven by baseline differences in profits.

Table 2 shows results of the difference-in-differences estimation of respondents’ ability to predict true marginal returns to capital. Outcome variables are household income and total household profits, in levels, logs, and the inverse hyperbolic sine (IHS). For the linear-in-rank version of the estimation (*Panel A*), the coefficient α_2 is large and positive for all six outcome variables. Coef-

ficients for the income specifications (columns 1-3) are significant at the 5% level; for profits, the coefficient on the level is significant at the 10% level while the coefficients for log and IHS profits are significant at the 5% level. An extra unit of average rank is associated with increases in profits and income of between Rs.466.4 [SE=276.1] and Rs.1,142.1 [SE=451.2] per month, respectively. These amounts translate to increases in monthly returns to the grant of between 7.8% and 19.0%. Average marginal return to capital in the sample is between 9.5% and 11.4% per month and an entrepreneur ranked one standard deviation above the mean has monthly marginal return to capital of 16.4% (the mean and standard deviation of the marginal return rank are 3.46 and 0.66, respectively). For an entrepreneur ranked two standard deviations above the mean, monthly returns to capital are 25.7%.

Panel B in Table 2 shows results from the non-linear, tercile rank version of the difference-in-differences estimation. Consistent with results from the local polynomial regressions in Figure 4, we cannot reject that the entrepreneurs in the bottom tercile of the marginal returns rank distribution have zero returns to the grant. For five of the six outcome variables (all but level of household profits), the coefficient on $Winner_{it}$ actually implies a negative return to the grant for entrepreneurs in the bottom tercile. Also consistent with Figure 4, the coefficients on log and IHS income and profits for the middle tercile are positive, but not significant, and the level effects are almost precisely zero.¹⁸ The strongest treatment effects of the grant are concentrated among entrepreneurs in the top tercile of the average rank distribution: depending on whether we use household income or profits, the coefficients on $Winner_{it} \times TopTercile_{ij}$ imply that monthly returns to the grant for the top tercile range from 23.3% to 35.2%. We can reject that the grant has the same effect for entrepreneurs in the middle and top tercile.

7.2.1 Robustness Checks

Evaluation of Community Information Using Cross-Sectional Variation Regression Specification 3 identifies the treatment effect of the grant off of the within-person differences in profits and income in the pre- and post- grant disbursal periods for grant winners and losers. As a robustness check, we also present results using an alternative specification in which the treatment effects are identified by comparing the cross-sectional differences between treatment and control groups in the post-grant disbursal periods, controlling for the baseline value of the outcome characteristic. The specification is:

$$Y_{ijt} = \beta_0 + \beta_1 Winner_{ijc} + \beta_2 Winner_{ijc} \times \overline{Rank}_{ijc} + \beta_3 \bar{Y}_{ijPRE} + \sigma_c + \theta_m + \tau_s + \epsilon_{ijt}, \quad (4)$$

where Y_{ijt} are post-treatment outcomes (so t ranges from 1 to 3 rather than 0 to 3 as in Specification 3) and $\bar{Y}_{ijPRE}$ is the pre-treatment (time 0) value of the outcomes. σ_c is a neighborhood cluster fixed effect. Standard errors are clustered at the group level. We present the analogue of Table A2

¹⁸Mechanically, since the middle tercile is fixed, the difference between the level and log results occurs because there are some extreme right-tail observations in the distribution of income and profits for the middle tercile ranks. The weight of these outliers in the regression is diminished when the distributions are log-transformed.

using Specification 4 in Table A4 and the analogue of Table 2 in Table A5. Results in the robustness specification are qualitatively similar in terms of the size and significance of coefficients.

Self Rank versus Community Ranks Throughout the analysis so far, our calculation of respondents’ average rank includes their self-rank. The impact of including respondents’ self-rank on community rank accuracy is ex-ante ambiguous. We might expect entrepreneurs to have better knowledge about themselves than they do about others. But it is also reasonable to assume that respondents will be more likely to strategically misreport in favor of themselves than when reporting about others. In Figure 5, we investigate the impact of self-rank on the community’s accuracy. We replicate the local polynomial regression of log profits at follow-up on marginal return rank percentile (as in Figure 4) with three specifications of the rank variable: (i) average rank including self-rank (*Panel 1*), (ii) average rank excluding self-rank (*Panel 2*), and (iii) only self-rank (*Panel 3*).¹⁹ Results shown in *Panel 1* and *Panel 2* are very similar, which indicates that entrepreneurs have strong knowledge of their peers and that community rank accuracy is not driven by the information contained in self-rank. We also see in *Panel 3* that entrepreneurs are able to predict their own marginal returns to the grant, though fewer entrepreneurs give themselves low rank values and so the correlation between the self-rank marginal returns prediction and actual marginal returns (the vertical distance between the profits of grant winners and losers) is weaker than it is with the average rank prediction. Finally, in Table A6, we replicate the results of Table 2 but exclude self-rank from the calculation of average rank. We find that results are nearly identical to those presented in Table 2, which again indicates that peers do indeed have important and valuable information about one another. In Sections 7.5 - 7.6, we further discuss the knowledge that entrepreneurs have about themselves and manipulation in self-reports.

Quintile versus Relative Ranks

We collect both zero-sum and quintile community ranks. Section 3 contains a more detailed discussion of the two ranking methods. All analysis in this section uses the (averaged) quintile community rankings. Results are qualitatively similar with both ranking methods but, because there is heterogeneity in peer groups’ average marginal return to capital, we find that quintile ranks are a more accurate assessment of an entrepreneur’s returns relative to the community. Results using the zero-sum rankings are in Table A7.

Individual versus Household-Level Profits

Bernhardt et al. (2019) shows that in households with multiple operating enterprises, grants and loans are not always invested in the targeted business: when women are the targeted recipient of a grant or loan and their husbands also operate a household business, resources are often invested in the husband’s rather than the wife’s business. We surveyed all household businesses and present all main results by aggregating across household enterprises. In Table A8, we present the results at

¹⁹Unlike average rank, which is the mean of 4-6 reports, the self-rank value is the result of a single report. As such, the self-rank variable only takes on integer values. For consistency across regressions in the three panels, we use rank value (rather than rank percentile as we did in Figure 4). As can be seen in Appendix Figure A3, there are few observations with a rank value below two. We therefore bottom code all three measures of rank.

the level of the client who was ranked by her peers.²⁰ Point estimates and standard errors remain nearly identical.

Demonetization

The month before we began our fifth (last) round of data collection, the Indian government removed from circulation two currency notes - the Rs.1,000 and Rs.500 bills - overnight. The result was a tremendous shock to the formal and informal economy. As Chodorow-Reich et al. (2020) report, traders experienced a 20% drop in sales due to demonetization. In fact, in the last round of surveying, over 50% of our sample reported being adversely affected by demonetization. For this reason, we exclude the post-demonetization wave of data from the analysis presented in the main tables. We replicate Table 2 with all five data rounds in Table A9. The results are qualitatively similar but marginally noisier in a few specifications.

Across specifications, we find that communities have deep knowledge of entrepreneurs' growth potential. Importantly, community members' predictions map to economically significant differences in returns to capital. Lending institutions would have good reason to target top-ranked entrepreneurs for credit: in 2016, the average yearly APR for microcredit in India was 24%. Entrepreneurs in the top tercile of community ranks earn *monthly* returns of 23.3% to 35.2%.²¹ But the point estimates in Table 2 also imply that entrepreneurs in the bottom tercile (and perhaps also middle tercile) may not have been able to cover the cost of a USD 100 loan without reducing their total household consumption (since these entrepreneurs do not increase their profits in response to the capital intervention).

7.3 Who are the Top-Ranked Entrepreneurs and How do They Invest Their Grant?

In this section, we explore whether differences in entrepreneurs' characteristics and investment decisions can help explain the large gaps in returns that we observe. For ease of exposition, the remaining main tables show only the rank tercile specification.

Entrepreneurs' Investment Decisions in Response to the Grant. In follow-up rounds of data collection, we asked grant winners to report on whether and how they had invested the grant money. Expenditures of the grant money were divided into business expenses (inventory, durable assets, labor, and other) and non-business expenses (loan repayment, giving out loans, household repairs, and other household expenses). We also asked respondents if they had supplemented the grant money with their own funds to make a business purchase. In Table A10 we examine the relationship between self-reported investment decisions and marginal returns rank. To do so we regress grant expenditures in each category (the sum of which is Rs.6,000) on whether

²⁰We stressed to clients that the person whose business the grant money would be invested in if they won it had to be the person who was ranked.

²¹There are many possible reasons why a loan might have induced different selection and investment patterns, but it is useful to benchmark entrepreneurs' returns against market rates. See (Fiala, 2018) for an experiment which randomly allocates loans or grants to entrepreneurs.

an entrepreneur is in the top or middle tercile of the marginal returns ranks distribution. The coefficients on the top and middle terciles indicate the difference in grant expenditures between entrepreneurs in those groups and entrepreneurs in the bottom tercile (the omitted group). Column 2 shows that business owners in the top tercile invest an extra Rs.903.1 [SE=276.9], or 15.1%, more of their grants in their enterprise than those in the bottom tercile. Both the top and middle terciles spend significantly less money on “Other Household Expenses”—medical expenses, education, food consumption, etc.—and are less likely to have saved their grant money for a future use.

Self-reports of grant expenditures suggest differences in investment behavior, but, since money is fungible, the observed effects might simply be due to mental accounting (see Karlan et al. (2016) for evidence and implications). To investigate whether grant investments translate to real increases in business inputs, we use regression Specification 3 to compare inventories, business assets, and labor outcomes of grant winners and losers. Results are shown in Table 3. We find that the grant induces top and middle ranked entrepreneurs to accumulate higher capital stocks: top tercile grant winners report an extra Rs.1,485.6 [SE=836.8] worth of inventory and an extra Rs.11,409.9 [SE=4,991.1] of durable assets. The treatment increases the capital stock (inventory plus durable assets) by approximately 215% of the grant amount. This treatment effect is within the confidence bound of increases in capital stock found in McKenzie et al. (2008).

The grant also induces increases in inputs complementary to capital: own, household, and non-household labor. In columns 3 and 4, we show that grant winners in the top tercile spend an extra 9.3 [SE=2.6] hours per week and an extra 2.2 [SE=1.0] days per month working when compared to their untreated counterparts. The treatment also has an impact on the amount of household and non-household labor. At baseline, 21% of enterprises in our sample employ household labor. Household workers in these enterprises contribute an average of 30 hours per week and almost none of them are officially paid a wage. Nine percent of households report using non-household labor in at least one of their businesses at baseline.²² Among these businesses, the average weekly wage bill at baseline is Rs.3,221. The grant induces top-ranked entrepreneurs to be 8.4% [SE=4.8%] more likely to have a household laborer and 8.0% [SE=3.8%] more likely to have a non-household laborer at follow-up when compared to their untreated counterparts (see columns 5 and 8).

We find that top ranked entrepreneurs’ investment behavior is markedly different from that of bottom ranked entrepreneurs: they invest a higher proportion of their grant into their business, turn those investments into higher business stock, and devote more time to working in their business.

Demographic Characteristics of Top-Ranked Entrepreneurs. In Table A11, we compare baseline characteristics of households and entrepreneurs in all three terciles of the marginal returns

²²For single-enterprise households, our eligibility criteria specified that businesses could not employ non-household labor at baseline. But households with multiple enterprises were eligible as long as there was at least one enterprise that met our eligibility criteria. Almost all households that report using non-household labor fall into this latter category. See Section 2 for a detailed explanation of eligibility criteria for households with multiple enterprises.

ranks distribution. In column 1, we present the mean of each characteristic for the bottom tercile group. We then estimate the following model:

$$Y_{ijc} = \beta_0 + \beta_1(MiddleTercile)_{ijc} + \beta_2(TopTercile)_{ijc} + \sigma_c + \theta_m + \tau_s + \epsilon_{ij} \quad (5)$$

In columns 2 and 3, we present the coefficients from regressions of each baseline characteristic on whether the respondent is ranked in the middle (β_1) or top (β_2) terciles, respectively. Coefficients can be interpreted as the difference in each characteristic associated with being in one of the upper terciles relative to being in the bottom tercile.

Top-ranked entrepreneurs are 8 percentage points more likely to be male; about 2 years younger; and, are less likely to be married than entrepreneurs in the bottom tercile. Entrepreneurs in the bottom and top terciles have roughly the same number of years of education, yet those who are top ranked remember an average of 0.57 digits more in the digit span memory test. Top-ranked business owners work an extra 5.0 hours per week and 1.1 days per month. We asked business owners how much a salaried job would have to pay per month in order for them to exit self-employment. Top ranked entrepreneurs report that they would require 17% higher monthly wages to leave their businesses. Top-ranked entrepreneurs are slightly more likely to be engaged in a food preparation business and less likely to engage in livestock than bottom-ranked entrepreneurs, but otherwise the industry distribution is similar across terciles.

Households with a top-ranked entrepreneur have the same total number of businesses as households in the lower terciles. But these households have enterprises that are 43.5% larger in terms of assets and earn 40.9% higher profits per month. They also earn 13.8% higher monthly income. Household labor composition is very similar across all three groups, but top and middle ranked households are slightly less likely to employ a household daily wage worker.

For the most part, entrepreneurs in the middle tercile have baseline characteristic means that lie between the means of the bottom and top ranked entrepreneurs. Two notable exceptions are that they have higher levels of education and business assets.

7.4 Benchmarking the Value of Community Information Against Observables

We showed in the previous section that top-ranked entrepreneurs differ from low-ranked entrepreneurs across several observable demographic characteristics. These findings raise the question: are community members simply using observable information to rank one another? In this section, we benchmark the value of community information against the value of observables. First, we investigate whether community information remains valuable for predicting high-return entrepreneurs even after controlling for baseline characteristics. Next, we compare the predictive power of each source of information. These questions are related but distinct: community members may use information that is orthogonal to information captured by observables, but the accuracy of community reports may still be lower than the accuracy of a selection mechanism based on observable characteristics.

Combining Community Information with Observable Characteristics. Is the value of peer reports diminished when we hold constant entrepreneurs’ baseline characteristics? We consider this question by estimating the value of community information while controlling for all the baseline characteristics presented in Table A1. These include demographic characteristics of the entrepreneur and her family, wealth and income measures of the household, and financial information about the business.

In Table 4, we present results of Specification 3 with the addition of the interaction of the baseline controls with $Winner_{it}$. Controlling for baseline characteristics only increases the size of the coefficients on community reports. This is because marginal returns rank is positively correlated with baseline profits and baseline profits are negatively correlated with marginal return to capital (implying that there are diminishing returns to capital). In Table A12, we include the same set of controls and estimate the value of community information using the robustness specification (Specification 4); results are qualitatively similar.

Finally, since psychological characteristics such as tenacity, polychronicity, and optimism, have also been shown to be predictive of credit worthiness and entrepreneurial aptitude (see Klinger et al. (2013)), we assess the value of entrepreneurs’ responses to a psychometric test. We find that the key estimates remain almost identical to the original results presented in Table 2 (results from this estimation are presented in Table A13).²³

Predicting High-Return Entrepreneurs Using Observable Characteristics. From our analysis in Section 7.2, we have an estimate of the value of community information. Next, we estimate the value of observable characteristics. This is a *prediction* problem, and not a parameter estimation problem: our goal in this section is not to understand the relationship between individual covariates (baseline characteristics) and entrepreneurs’ returns. Instead, we seek to combine the information contained in all covariates to produce a prediction of these returns. We apply the *Causal Forest* machine learning algorithm developed by Athey and Imbens (2016) and Wager and Athey (2018) to form a marginal returns ranking of entrepreneurs based on their baseline characteristics. We then compare the predictive power of this ranking to that of the community reports ranking.

Machine learning techniques typically require large datasets, separated into a subset used to train the predictive algorithm (the *training data*) and a subset used to evaluate the model (the *holdout data*). The models typically perform better on the training data than on the holdout data because they are in part predicting idiosyncratic features of the data that would not replicate in

²³Regressors are labeled according to the psychological trait for which they are meant to proxy (the specific wording of the statement is found in Appendix A4). There are two traits that are strongly predictive of marginal returns: optimism and achievement. We find that optimism negatively predicts marginal returns: business owners who are more likely to agree with the statements “In times of uncertainty I expect the best” and “I’m always optimistic about the future” and those who are more likely to disagree with “If something can go wrong with me, it will” have lower self-reported marginal returns. People who agree with the statement “Part of my enjoyment in doing things is improving my past performance” tend to have higher marginal returns.

a fresh sample (Hastie et al., 2008). Due to our limited sample size we both train and validate the model on the full dataset, which provides an upwardly biased estimate of the value of observables in predicting entrepreneurs’ marginal returns to capital. Therefore we are using conservative benchmark to evaluate community information.

Results.

In Table A14, we present the results of the machine learning exercise. In columns 1 and 2, we first replicate the results shown in column 4 of Table 2 and column 4 of Table 4, respectively.²⁴ The machine learning exercise produces a numerical prediction of the marginal returns of each individual in the sample based on their baseline characteristics. For comparison with our main specification for the community rankings estimates, we divide the predictions into terciles. In columns 3 and 4 we test how well the machine learning estimation predicts true marginal returns in our sample.

In column 3, the top tercile of entrepreneurs as identified by the causal forest algorithm earn an extra Rs.1,824.0 [SE=540.1] in marginal returns to the grant over the bottom tercile of entrepreneurs. This is comparable to the predictive value of community information reported in Section 7.2. In column 4, we add the community information prediction. First note that the coefficient on *Winner * TopTercile* is large and significant at the 5% level. Furthermore, the machine learning prediction estimate for the top tercile becomes a bit smaller, but remains a good prediction of the best entrepreneurs. The correlation between community rank and the machine learning prediction is 0.1. Taken together, these results indicate that community members are using additional information to rank beyond (detailed) covariates that are observable to the researcher and that peer information is very valuable in identifying high-ability entrepreneurs. Despite the fact that the model is overfit in-sample, community ranks continue to be predictive above and beyond the machine learning prediction.

7.5 Do Peers Distort Their Responses When There Are Real Stakes?

The analysis in the previous sections has shown that communities are well informed about members’ marginal returns to capital. But to be of practical use, community members need to report their opinions *truthfully*. In this section, we quantify whether and by how much community members distort their reports in high stakes settings.

The analyses in this section examine the relationship between community reports and entrepreneurs’ business characteristics; income, assets, and profits. As explained in Section 4.2, we did not randomize the *HighStakes* and *NoStakes* treatments until after the marginal returns ranking was completed due to power considerations. As a result, we cannot include predictions about marginal return to capital in our analyses of incentives.

In order to assess whether and how peers lie when there is incentive to strategically misreport

²⁴The estimates (and number of observations) differ slightly to ensure a comparable sample with the machine learning exercise. So in the replication of Table 4 column 4, we only control for the variables that we use in the machine learning exercise (a subset of the variables used in Table 4).

half of our sample was informed that their rankings would affect the probability that their peers (or themselves) would win the USD 100 grant (this is the *High Stakes* group). Respondents in the *No Stakes* group continued to believe that their ranking responses would only be used for research purposes. We assess strategic misreporting in Table 5 by amending Specification 1 to compare accuracy in the *High Stakes* and *No Stakes* groups:

$$Y_{ijq} = \alpha_0 + \alpha_1 \overline{Rank}_{ijq} + \alpha_2 Stakes_j + \alpha_3 Stakes_j \times \overline{Rank}_{ijq} + \gamma_n + \theta_m + \tau_s + \delta_q + \epsilon_{ijmq} \quad (6)$$

The model includes the following fixed effects: neighborhood (γ_n), survey month (θ_m), and surveyor (τ_s). Standard errors are clustered at the group level. α_1 captures the accuracy of the report in the control group (*No Stakes*). α_3 indicates the extent to which the rankings are differentially informative when respondents are told their reports will be used to help determine grant allocation.²⁵

²⁶ To increase power, we stack the percentilized outcomes and ranks across all 6 columns presented in *Panel A* of Table 1 and add a question fixed effect (δ_q) to the regression model.

Respondents may have idiosyncratic preferences for misreporting about certain peers in their group and may otherwise make idiosyncratic errors. One way to reduce noise is to average across all reports given about a particular group member.²⁷ So in columns 1-3 of Table 5, we show the regressions at the ranker-rankee level of observation ($Rank_{ijmq}$) and column 4-6 are the regressions with the average rank ($\overline{Rank}_{ijq}$). We observe that the average predictiveness of ranks in the (*No Stakes*) group increases significantly when reports are averaged: in column 1, a 1 percentile increase in the rank distribution is associated with a 0.16 [SE=0.02] shift in the outcome distribution in the individual regressions and a 0.25 [SE=0.02] shift in the average regression (column 4). Averaging reports nearly doubles the predictiveness of community reports.

Do respondents misreport in high stakes settings? We find that the coefficient on $Rank \times High\ Stakes$ is large, negative, and significant. We note that this was not ex-ante clear: the *High Stakes* treatment may have had a positive effect since introducing stakes may have caused respondents to focus or take the exercise more seriously. The regression implies that responses are significantly less accurate when respondents have an incentive to behave strategically: in the pooled individual regression in column 1, the responses become 34.8% less accurate in the *High Stakes* group.

Lastly, we asked respondents to rank their peers relative to others in the group (zero-sum ranking) and also relative to the community by reporting the quintile of the neighborhood distribution that they believe the peer to be in (quintile ranking). We hypothesized that quintile ranks could contain more valuable information about rankings because entrepreneurs are compared to

²⁵To reduce clutter in the regression tables, we have omitted the *High Stakes* coefficient from the regression report as it does not contain information relevant for the interpretation of results, but rather simply adjusts the constant.

²⁶In this section, we pool across the *Public* and *Payments* treatments.

²⁷In Table 1, all reports are averaged.

the community more broadly than only the group. But they could also be more susceptible to misreporting: unlike with zero-sum ranks, respondents could, for example, place all of their peers in the top quintile of the distribution indicating that everyone is equally excellent.

To compare these two elicitation methods, in columns 2-3 and 5-6, we show the results by separately stacking zero-sum and quintile rankings. In all four columns, the outcome variable is the same (percentile of Y_{ijq}). What changes is the method of reporting. In columns 2 and 5, the regressor is the percentile in the (individual or average) quintile rank distribution. In columns 3 and 6, the regressor is the percentile in the (individual or average) zero-sum rank distribution.²⁸

The coefficients on *Rank* in the individual (columns 2 and 3) and the average regressions (columns 5 and 6) are very similar, implying that in the absence of high-stakes, the value of information from relative and quintile ranks is very similar. While the coefficient on $Rank \times High\ Stakes$ in the quintile regressions is larger in magnitude both in the individual and average models, we cannot reject that respondents misreport by the same amount in either type reporting method.

Overall, we find that in the presence of real stakes, misreporting is an important problem.

7.6 Can Mechanism Design Tools Correct Incentives to Misreport?

Monetary Incentives and Public Reporting. Can we use tools from mechanism design to generate incentives for truthful reporting? And, are these tools effective even in high-stakes settings? We test the efficacy of two tools: payments for the accuracy of reports and reporting in public versus private.

In Table 6, we provide evidence of the *Public* and *Payments* treatments on the accuracy of reports. Again, following Specification 1 we estimate,

$$\begin{aligned} Y_{ijq} = & \eta_0 + \eta_1 \overline{Rank}_{ijq} + \eta_2 Public_j \times \overline{Rank}_{ijq} + \eta_3 Payments_j \times \overline{Rank}_{ijq} \\ & + \eta_4 Public_j \times Payments_j \times \overline{Rank}_{ijq} + \eta_5 Public_j + \eta_6 Payments_j \\ & + \eta_7 Public_j \times Payments_j + \gamma_n + \theta_m + \tau_s + \delta_q + \epsilon_{ijmq}. \end{aligned} \quad (7)$$

The coefficient η_1 identifies the accuracy of reports in groups in which respondents do not receive incentive payments and report in private. The coefficients on the first three interaction terms tell us the additional accuracy due to reporting (i) in public without monetary payments (η_2), (ii) in private with monetary payments (η_3), and (iii) in public with monetary payments (η_4).²⁹

²⁸In Table 1, we stacked the zero-sum and quintile ranks by question. So in column 1 of Table 1, the outcome variable is the household income and the regressors are the income quintile and zero-sum ranks, with a fixed effect for ranking type. Notice that the outcome variable is the same (household income) whether the regressor is a quintile or zero-sum ranking.

²⁹To reduce clutter in the regression tables, we have omitted the coefficients $Public_j \times Payments_j \times \overline{Rank}_{ijq}$, $Public_j$, $Payments_j$, $Public_j \times Payments_j$ from Table 6 as they do not contain information relevant for the interpretation of results, but rather simply adjust the intercept.

To determine how these tools perform in a high stakes setting, we split results by *No Stakes* (odd columns) and *High Stakes* (even columns). We also split the results by whether a respondent is reporting about herself (columns 1 and 2) or about her peers (columns 3 and 4).

We find that community members are both more accurate and less responsive to incentives for truthfulness when reporting about themselves. Putting respondents in a high-stakes setting decreases the accuracy of self-reports by 23.3%. Moreover, neither payments for truthfulness nor public reporting have any impact on the accuracy of self-reports. Note, though, that the accuracy of their self-reports (0.16 [SE=0.04] in column 2) in the high-stakes setting is approximately the same as the accuracy of reports about others in the group in the private and no payments treatment (0.14 [SE=0.05] in column 3).

When reporting about others, incentives for truthfulness can have a large impact on respondents' accuracy. First, in the *No Stakes* setting, the *Payments* and *Public* treatments both double the accuracy of reports (they each lead to increase in accuracy between 0.14 [SE=0.07] and 0.17 [SE=0.06]. The coefficient on the treatment in which respondents receive monetary incentives and report in public is large and negative ($Average\ Rank \times Payments \times Public$). But, we can reject at the 10% level that the accuracy of information in this group is the same as in the private reporting and no monetary incentives group. We therefore interpret the negative coefficient as an indication that monetary payments and public reporting are substitutes.

The monetary payments treatment is just as effective when allocation of resources is at stake: the *Payments* treatment still improves accuracy by 0.14 [SE=0.07], which is an increase in accuracy of over 100%. So, providing monetary payments corrects nearly all of the strategic misreporting that is induced by asking respondents to report in a high stakes setting.

In the *High Stakes* setting, we find that the *Public* treatment no longer has a significant impact on accuracy. As discussed in Section 3, the impact of public reporting on accuracy is ambiguous ex-ante. There may be pressure for respondents to up-rank their family members, but there may also be pressure from non-family members and other peers to be truthful. When we introduce stakes, both of these pressures are intensified: family members and close friends want the respondent to sway the grant allocation in their favor, but it may also be especially important to the community that members be truthful when there are high stakes. That we find different impacts of observability in the *High Stakes* and *No Stakes* treatment might reflect the differing intensities of these two competing forces, or it might reflect a lack of precision in our estimates.

How Do Respondents Distort Their Reports? So far we have established that respondents distort their reports when the distribution of resources is at stake, and that simple mechanisms can realign incentives for accuracy. Lastly we ask, for whom do respondents distort their reports to favor? At the start of the ranking exercise, we asked respondents to report their relationship with each peer in the group. We also asked each respondent to identify each other person's closest peer in the group. An entrepreneur's *cross-reported peer* is the peer that is most frequently reported as

their closest friend in the group.

To assess who respondents lie to favor, we analyze how the rankings themselves (not just accuracy) are affected by proximity between peers in Appendix Table A15. We see that respondents up-rank family members and cross-reported peers relative to other peers in the group in the absence of incentives and in private. But incentives and publicity reduce the average rank assigned to either of these groups.

Cross-Reporting. As shown in Table A15, community members are capable of successfully identifying people for whom a peer is likely to lie to favor. We also asked respondents to name the person who would be best able to predict who would provide the most accurate reports on average. In Table A16, we interact rank with whether a respondent has been selected by her group as the one who would provide the most accurate answers. In column 1, pooling across all questions, we see people who are selected provide information that is 50% more accurate than information provided by the standard respondent in the group.

8 Conclusion

We find that community members have information about their peers that is valuable for targeting even after controlling for a wide range of observable characteristics. Not only can community members identify characteristics of their peers' enterprises, they can also predict which of their peers have high returns to capital. But community information is also susceptible to strategic misreporting. In particular, we identify a tendency for respondents to favor their friends and family members in their reports. Moreover misreporting is exacerbated when respondents are told that their reports will influence the distribution of grants. If we assume that stakes would have reduced the accuracy of the marginal returns ranking by a third (the estimated average reduction in accuracy across the metrics evaluated in the high stakes treatment), then the marginal returns prediction for the top third of entrepreneurs would drop from 26% to 17.3% per month.

However we also find that a variety of techniques motivated by mechanism design theory are effective in realigning incentives for truthfulness. Relatively small monetary payments for accuracy and cross reporting techniques both substantially improve the accuracy of reports.

Is it worth it to invest in collecting community information and providing incentives to respondents? We calibrated the payment rule to pay, on average, Rs.25 per question per respondent. In total, we paid Rs.17,000 in incentives for the marginal returns question. If a lender were distributing 450 loans (as we did with grants), this would increase the cost on each loan by approximately Rs.40 per month. In Section 7.2, we estimated that the cost of interest that an MFI would charge per grant is Rs.570 per month. Adding the incentives costs (transferring it to the borrower) implies that the cost of the loan to each respondent per month would be Rs.610. Using the returns estimate from our preferred specification (Table 2 *Panel B*, column 4), borrowers would still earn a net return of Rs.785 per month if the full cost of the monetary incentives were passed on to the

borrowers.

Our hope is that the peer elicitation method identified in this paper can be useful for targeting in poorly developed financial markets in low-income countries, where information asymmetries are prevalent. Moreover, the tools developed in this paper may prove useful in other contexts in which researchers and policy makers aim to target resources using community information. This may be especially true when targeting is to be done based on *treatment effects* rather than observable characteristics, and in settings where the incentives of community members and policy makers may not be fully aligned.

References

- Alatas, V., A. Banerjee, R. Hanna, B. A. Olken, R. Purnamasari, and M. Wai-Poi (2019). Does elite capture matter? local elites and targeted welfare programs in indonesia. Technical report.
- Alatas, V., A. Banerjee, R. Hanna, B. A. Olken, and J. Tobias (2012). Targeting the poor: evidence from a field experiment in indonesia. *The American Economic Review* 102(4), 1206–1240.
- Allcott, H. and T. Rogers (2014). The short-run and long-run effects of behavioral interventions: Experimental evidence from energy conservation. *The American Economic Review* 104(10), 3003–3037.
- Athey, S. and G. W. Imbens (2016). Recursive partitioning for heterogeneous causal effects. *Proceedings of the National Academy of Sciences of the U.S.A.* 113(27), 7353–60.
- Banerjee, A., D. Karlan, and J. Zinman (2015). Six randomized evaluations of microcredit: Introduction and further steps. *American Economic Journal: Applied Economics* 7(1), 1–21.
- Basurto, M. P., P. Dupas, and J. Robinson (2019). Decentralization and efficiency of subsidy targeting: Evidence from chiefs in rural malawi. *Journal of Public Economics Forthcoming*.
- Baum, J. R. and E. A. Locke (2004). The relationship of entrepreneurial traits, skill, and motivation to subsequent venture growth. *Journal of Applied Psychology* 89(4), 587.
- Beaman, L. and J. Magruder (2012). Who gets the job referral? evidence from a social networks experiment. *The American Economic Review* 102(7), 3574–3593.
- Bernhardt, A., E. Field, R. Pande, and N. Rigol (2019). Household matters: Revisiting the returns to capital among female microentrepreneurs. *American Economic Review: Insights* 1(2), 141–60.
- Bluedorn, A. C., T. J. Kalliath, M. J. Strube, and G. D. Martin (1999). Polychronicity and the inventory of polychronic values (ipv) the development of an instrument to measure a fundamental dimension of organizational culture. *Journal of Managerial Psychology* 14(3/4), 205–231.
- Bryan, G., D. Karlan, and J. Zinman (2015). Referrals: Peer screening and enforcement in a consumer credit field experiment. *American Economic Journal: Microeconomics* 7(3), 174–204.

- Chodorow-Reich, G., G. Gopinath, P. Mishra, and A. Narayanan (2020). Cash and the economy: Evidence from india’s demonetization. *The Quarterly Journal of Economics* 135, 57–103.
- Dal Bó, E., F. Finan, N. Y. Li, and L. Schechter (2018). Government decentralization under changing state capacity: Experimental evidence from paraguay. *Working Paper*.
- de Mel, S., D. McKenzie, and C. Woodruff (2008). Returns to capital in microenterprises: Evidence from a field experiment. *The Quarterly Journal of Economics* 123(4), 1329–1372.
- de Mel, S., D. J. McKenzie, and C. Woodruff (2009). Measuring microenterprise profits: Must we ask how the sausage is made? *Journal of Development Economics* 88(1), 19–31.
- Fafchamps, M., D. McKenzie, S. Quinn, and C. Woodruff (2014). Microenterprise growth and the flypaper effect: Evidence from a randomized experiment in ghana. *Journal of Development Economics* 106, 211–226.
- Fiala, N. (2018). Returns to microcredit, cash grants and training for male and female microentrepreneurs in uganda. *World Development* 105, 189–200.
- Gerber, A. S., D. P. Green, and C. W. Larimer (2008). Social pressure and voter turnout: Evidence from a large-scale field experiment. *American Political Science Review* 102(1), 33–48.
- Giné, X. and D. Karlan (2014). Group versus individual liability: Short and long term evidence from philippine microcredit lending groups. *Journal of Development Economics* 107, 65–83.
- Hastie, T., R. Tibshirani, and J. Friedman (2008). *The Elements of Statistical Learning: Data mining, Inference, and Prediction*. Springer.
- Karlan, D., A. Osman, and J. Zinman (2016). Follow the money not the cash: Comparing methods for identifying consumption and investment responses to a liquidity shock. *Journal of Development Economics* 121, 11–23.
- King, E. (2013). A critical review of community-driven development programmes in conflict-affected contexts. *London: DFID / International Rescue Committee*.
- Klinger, B., A. I. Khwaja, and C. del Carpio (2013). *Enterprising Psychometrics and Poverty Reduction*. Springer-Verlag New York.
- Maitra, P., S. Mitra, D. Mookherjee, A. Motta, and S. Visaria (2017). Financing smallholder agriculture: An experiment with agent-intermediated microloans in india. *Journal of Development Economics* 127, 306–337.
- Mansuri, G. and V. Rao (2004). Community-based and driven development: A critical review. *The World Bank Research Observer* 19(1), 1–39.
- Maskin, E. (1999). Nash equilibrium and welfare optimality. *The Review of Economic Studies* 66(1), 23–38.

- McClelland, D. C. (1985). How motives, skills, and values determine what people do. *American Psychologist* 40(7), 812.
- McKenzie, D., S. de Mel, and C. Woodruff (2008). Returns to capital: Results from a randomized experiment. *Quarterly Journal of Economics* 123(4), 1329–72.
- Prelec, D. (2004). A bayesian truth serum for subjective data. *Science* 306(5695), 462–466.
- Rigol, N. and B. Roth (2017). Paying for the truth: The efficacy of peer prediction in the field. *Working Paper*.
- Rosenbaum, P. R. (1987). Model-based direct adjustment. *Journal of the American Statistical Association* 82(398), 387–394.
- Rotter, J. B. (1966). Generalized expectancies for internal versus external control of reinforcement. *Psychological Monographs: General and Applied* 80, 1–28.
- Selten, R. (1998). Axiomatic characterization of the quadratic scoring rule. *Experimental Economics* 1, 43–62.
- Wager, S. and S. Athey (2018). Estimation and inference of heterogeneous treatment effects using random forests. *Journal of the American Statistical Association* 113(523).
- Witkowski, J. and D. Parkes (2012). A robust bayesian truth serum for small populations. *Proceedings of the 26th AAAI Conference on Artificial Intelligence (AAAI’12)*.

Figure 2: Randomization Design

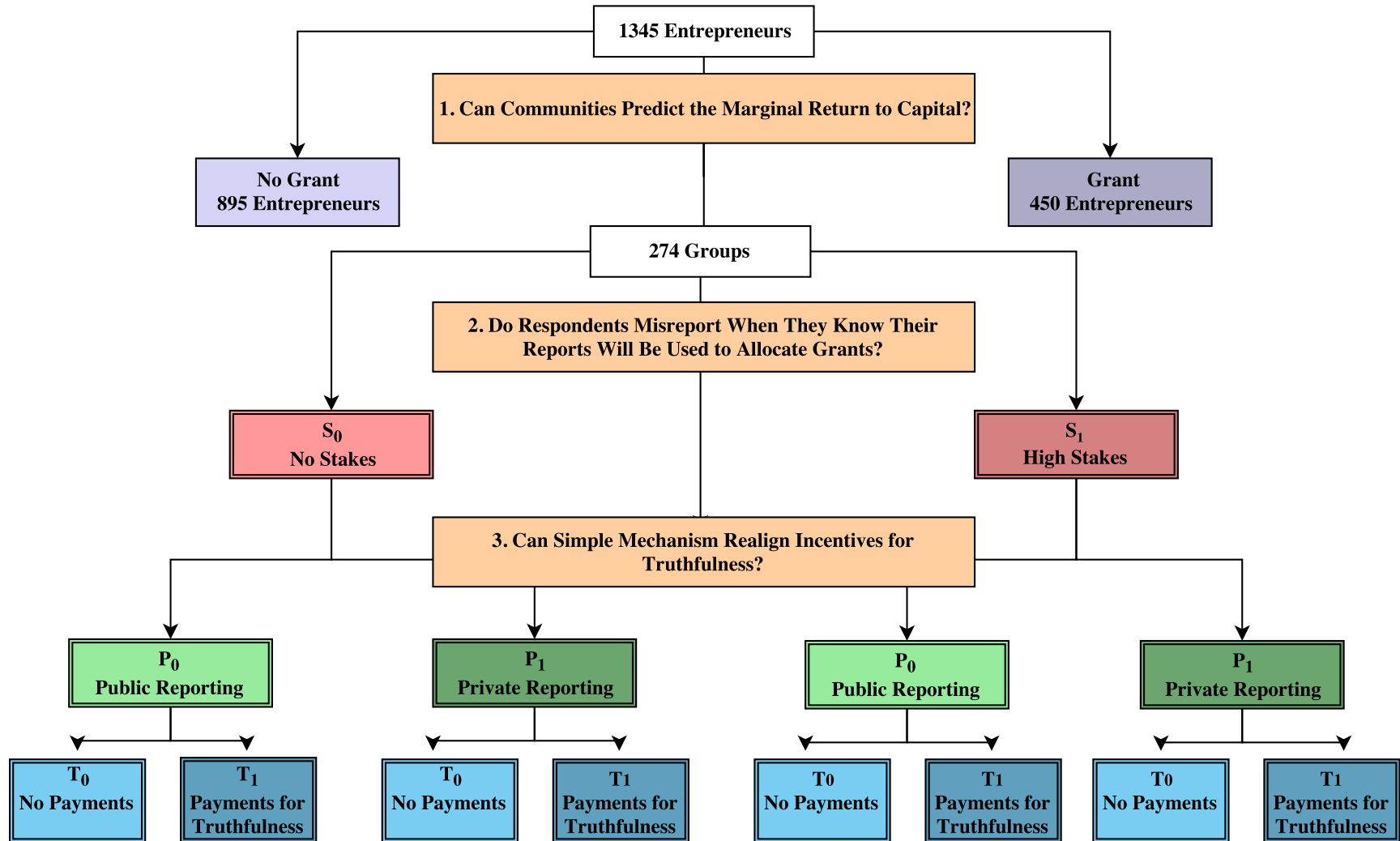


Figure 3: Timeline

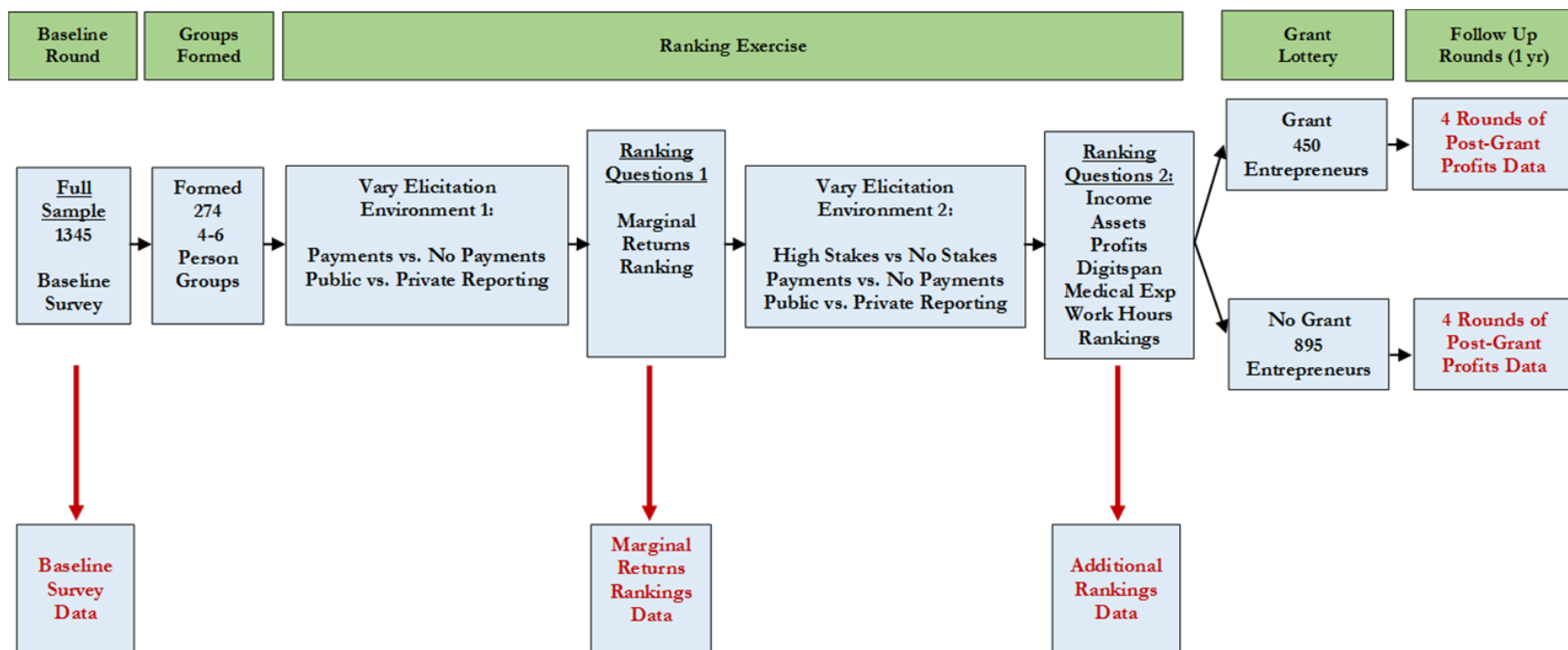
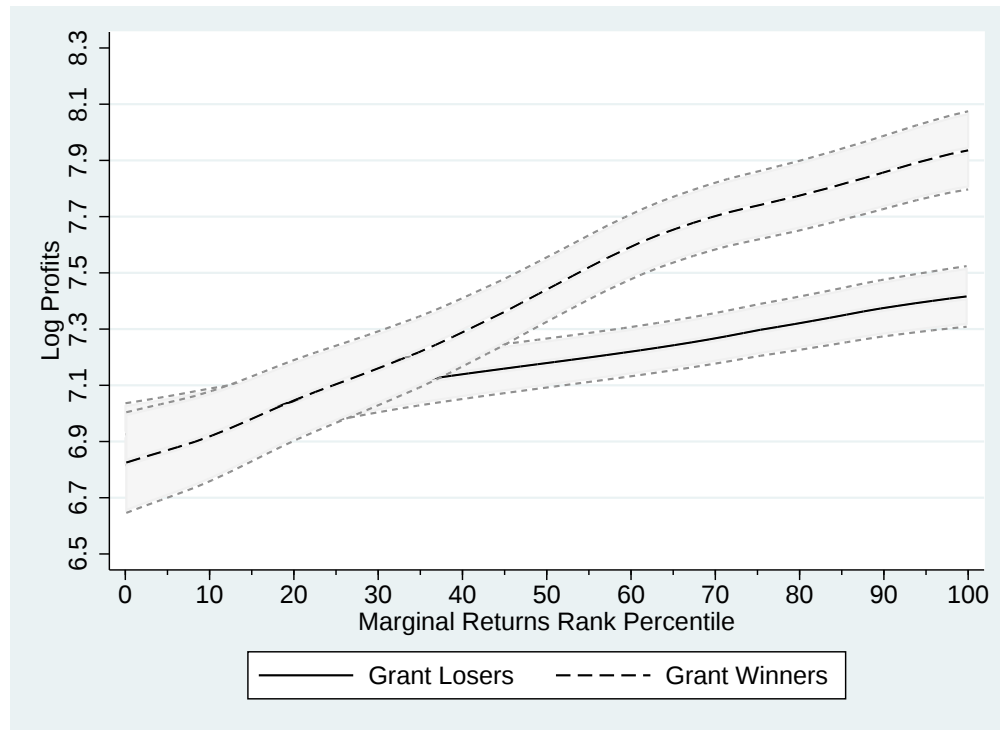
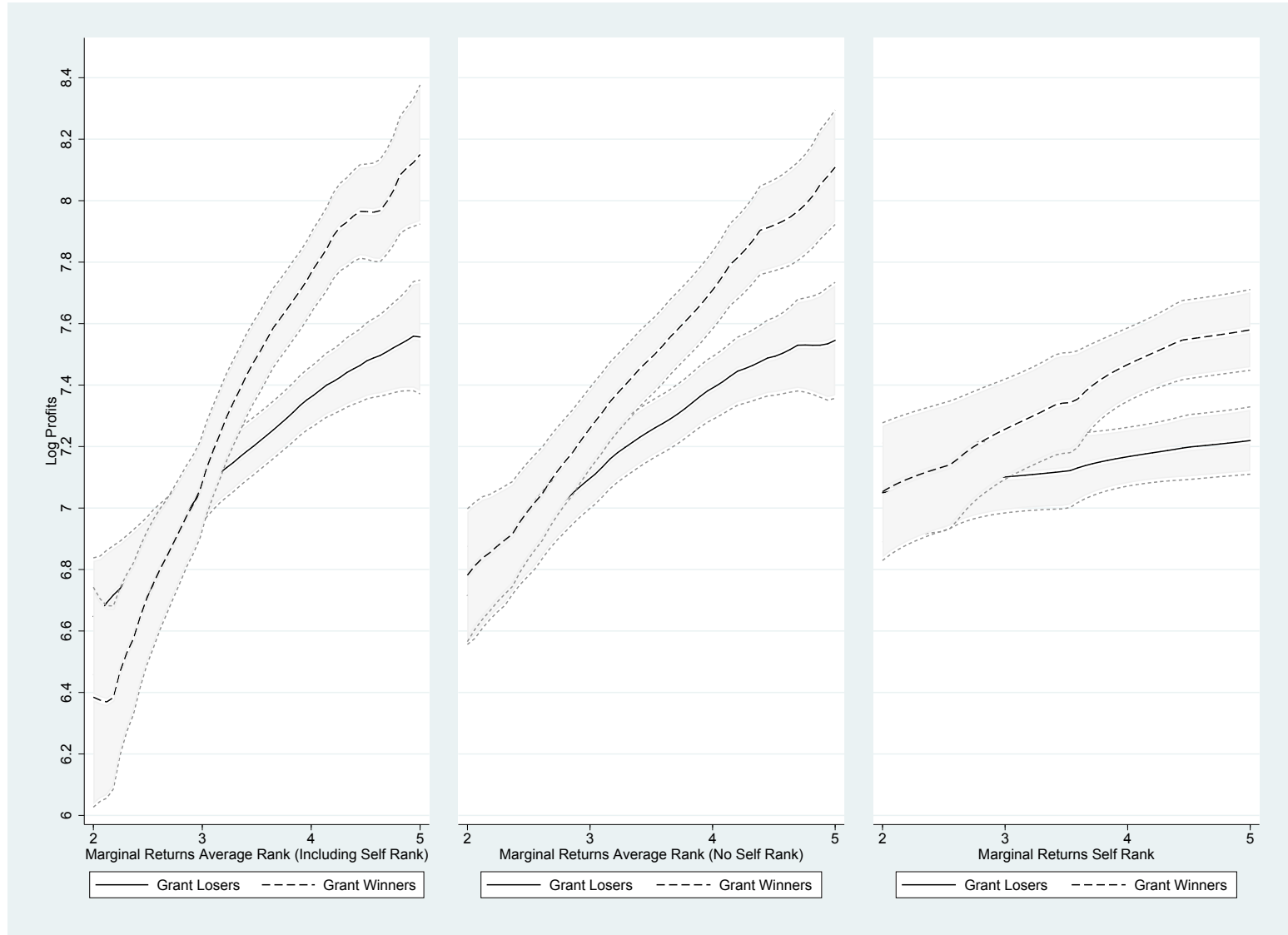


Figure 4: Marginal Returns to the Grant by Percentile of the Average Community Ranks Distribution



Notes: This figure plots two kernel-weighted local polynomial regressions of log profits on the marginal returns rank percentile, estimated separately for respondents who won and respondents who did not win grants. Log profits is the log value of average profits in the post grant disbursement periods. The marginal returns rank percentile is the percentile of the average rank assigned to person i by all of her peers in her group. 90% confidence bands are shown.

Figure 5: Marginal Returns to the Grant by Average Community Ranks



Notes: This figure plots two kernel-weighted local polynomial regressions of log profits on (1) the marginal returns rank that includes self rank (Panel 1), or (2) the marginal returns average rank that excludes self rank (Panel 2), or (3) the marginal returns self rank (Panel 3), estimated separately for respondents who won and respondents who did not win grants. Log profits is the log value of average profits in the post grant disbursement periods. 90% confidence bands are shown.

Table 1: What Do Respondents Know About One Another?

	(1) Income	(2) Profits	(3) Assets	(4) Medical Exp.	(5) Digitspan	(6) Work Hours
<i>Panel A: Average Rank Percentile</i>						
Average Rank	0.225*** (0.027)	0.224*** (0.028)	0.224*** (0.028)	0.219*** (0.061)	0.280*** (0.041)	0.089 (0.063)
N	1910	1968	1834	259	277	276
<i>Panel B: Average Rank Level</i>						
Avg. Rank Level	1897.179*** (259.642)	1543.777*** (216.584)	1.24e+05*** (23228.214)	1367.119*** (469.760)	0.621*** (0.097)	3.500* (1.791)
Mean of Outcome	8851.27 [6863.01]	6914.69 [5993.13]	474730.62 [718455.42]	2886.46 [5428.52]	5.19 [1.70]	61.32 [22.91]
N	1910	1968	1834	259	277	276
No. HHs	1021	1032	990	259	277	276

* $p \leq 0.10$, ** $p \leq 0.05$, *** $p \leq 0.01$. Notes: Data in this table come from round 1 (baseline) of data collection. Robust standard errors clustered at group level in parentheses. The model includes randomization cluster, surveyor, and date of survey fixed effects. In Panel A, the outcome variable is the percentile of the outcome in the column title and the regressor is the percentile of the average rank given to a respondent, computed by question. In Panel B, the outcome variable is the level of the outcome in the column title and the regressor is the average rank level for that particular question. The level of observation is the rankee. The number of observations varies across questions because each respondent answered only a subset of the questions. See the Implementation Appendix for details.

Table 2: Do Peer Reports Predict True Marginal Returns to the Grant?

	(1) Income	(2) Log Income	(3) IHS Income	(4) Profits	(5) Log Profits	(6) IHS Profits
<i>Panel A: Average Rank Value</i>						
Winner*Rank	1142.055** (451.162)	0.214** (0.103)	0.227** (0.110)	466.416* (276.063)	0.397** (0.177)	0.432** (0.191)
Winner	-3399.676** (1650.499)	-0.606* (0.359)	-0.640* (0.384)	-935.557 (926.462)	-1.043* (0.630)	-1.141* (0.680)
<i>Panel B: Average Rank Tercile</i>						
Winner*Top Tercile Rank	2111.346*** (760.202)	0.477** (0.203)	0.503** (0.217)	1395.207*** (531.256)	0.842*** (0.295)	0.909*** (0.319)
Winner*Middle Tercile Rank	222.587 (779.089)	0.068 (0.164)	0.069 (0.175)	-21.409 (392.529)	0.055 (0.281)	0.060 (0.303)
Winner	-298.654 (569.847)	-0.062 (0.147)	-0.065 (0.157)	166.726 (347.214)	-0.000 (0.224)	-0.004 (0.242)
<i>P-value from F-Test</i>						
Winner*Top Tercile Rank= Winner*Middle Tercile Rank	0.022**	0.020**	0.020**	0.009***	0.004***	0.004***
Mean of Outcome for Grant Losers	8197.92 [6412.96]	8.62 [1.35]	9.30 [1.42]	4552.74 [5160.11]	7.33 [2.55]	7.95 [2.74]
N	5324	5342	5342	5319	5337	5337
No. HHs	1336	1336	1336	1336	1336	1336

* $p \leq 0.10$, ** $p \leq 0.05$, *** $p \leq 0.01$. Notes: Rank indicates the average ranking the entrepreneur was given by her peers for the marginal returns to grant ranking question. Top (Middle) Tercile Rank is a dummy for whether the entrepreneur is in the top (middle) tercile of the average marginal return rank distribution. Winner indicates that the household is a grant recipient after baseline (after round 1 of data collection). The unit of observation is the household. Robust standard errors clustered at the group level in parentheses. All regressions include household, survey month, survey round, and surveyor fixed effects. All regressions are weighed by the inverse propensity score described in Section 7. Data in this table comes from rounds 1-4 of data collection.

Table 3: Impact of Grant on Business Inputs

	Business Assets		Owner Labor		Household and Non-Household Labor					
	(1)	(2)	(3)	(4)	(5)	(6)	(7)	(8)	(9)	(10)
	Business Inventory	Durable Business Assets	Total Hours Worked Past Week	Total Days Worked Past Month	Uses Household Labor	Household Labor Hours Past Week	HH Labor Wage Bill Past Week	Uses Non-HH Labor	HH Labor Hours Past Week	Non-HH Labor Wage Bill Past Week
Winner*Top Tercile Rank	1485.575*	11409.913**	9.313***	2.243**	0.084*	4.661*	66.495	0.080**	7.240*	201.630
	(836.752)	(4991.090)	(2.641)	(1.019)	(0.048)	(2.565)	(51.373)	(0.038)	(3.766)	(181.499)
Winner*Middle Tercile Rank	1132.390*	2513.748	2.217	1.469	0.044	3.849*	65.545	0.002	2.422	195.809
	(632.570)	(2837.986)	(2.727)	(0.972)	(0.050)	(2.105)	(51.179)	(0.038)	(3.500)	(251.063)
Winner	-504.699	-2740.186	-3.647	-0.837	-0.018	-3.784**	-45.141	-0.002	-2.174	20.164
	(477.063)	(2235.107)	(2.276)	(0.793)	(0.034)	(1.837)	(51.126)	(0.028)	(2.444)	(155.188)
<i>P-value from F-Test</i>										
Winner*Top Tercile Rank= Winner*Middle Tercile Rank	0.708	0.088*	0.004***	0.413	0.462	0.706	0.891	0.052*	0.278	0.984
Mean of Outcome for Grant Losers	4764.69 [12317.31]	40241.38 [91800.55]	37.18 [25.68]	21.27 [8.48]	0.18 [0.39]	5.18 [16.18]	12.59 [250.64]	0.09 [0.28]	6.96 [36.98]	269.42 [1706.11]
N	5302	5299	5229	5193	2672	2672	2672	2672	2672	2672
No. HHs	1335	1332	1336	1336	1336	1336	1336	1336	1336	1336

* $p \leq 0.10$, ** $p \leq 0.05$, *** $p \leq 0.01$. Notes: Rank indicates the average ranking the entrepreneur was given by her peers for the marginal returns to grant ranking question. Top (Middle) Tercile Rank is a dummy for whether the entrepreneur is in the top (middle) tercile of the average marginal return rank distribution. Winner indicates that the household is a grant recipient after baseline (after round 1 of data collection). The unit of observation is the household. Robust standard errors clustered at the group level in parentheses. All regressions include household, survey month, survey round, and surveyor fixed effects. All regressions are weighed by the inverse propensity score described in Section 7. The number of observations in columns 1-4 varies due to missing data across the rounds. Data for these columns comes from rounds 1-4 of data collection. Variables reported in columns 5-10 were only collected at baseline and in round 4.

Table 4: Returns with Baseline Controls

	(1) Income	(2) Log Income	(3) IHS Income	(4) Profits	(5) Log Profits	(6) IHS Profits
<i>Panel A: Average Rank Value</i>						
Winner*Rank	1328.764*** (354.992)	0.216** (0.098)	0.226** (0.104)	837.261*** (242.819)	0.491*** (0.154)	0.528*** (0.166)
Winner	3272.715 (2786.414)	0.449 (0.679)	0.425 (0.721)	-693.861 (2106.249)	-0.933 (1.338)	-1.081 (1.451)
<i>Panel B: Average Rank Tercile</i>						
Winner*Top Tercile Rank	2596.110*** (568.905)	0.479** (0.189)	0.502** (0.203)	1886.685*** (388.165)	0.928*** (0.256)	0.990*** (0.277)
Winner*Middle Tercile Rank	876.071 (551.498)	0.069 (0.168)	0.066 (0.180)	501.861 (336.401)	0.099 (0.264)	0.099 (0.286)
<i>P-value from F-Test</i>						
Winner*Top Tercile Rank= Winner*Middle Tercile Rank	0.006***	0.014**	0.014**	0.002***	0.002***	0.002***
Mean of Outcome for Grant Losers	8197.92 [6412.96]	8.62 [1.35]	9.30 [1.42]	4552.74 [5160.11]	7.33 [2.55]	7.95 [2.74]
N	5249	5267	5267	5243	5261	5261
No. HHs	1336	1336	1336	1336	1336	1336

* $p \leq 0.10$, ** $p \leq 0.05$, *** $p \leq 0.01$. Notes: Rank indicates the average ranking the entrepreneur was given by her peers for the marginal returns to grant ranking question. Top (Middle) Tercile Rank is a dummy for whether the entrepreneur is in the top (middle) tercile of the average marginal return rank distribution. Winner indicates that the household is a grant recipient after baseline (after round 1 of data collection). The unit of observation is the household. Robust standard errors clustered at the group level in parentheses. All regressions include household, survey month, survey round, and surveyor fixed effects. All regressions are weighed by the inverse propensity score described in Section 7. Regressions in the odd columns include Winner interacted with the following controls: gender, education, married, age, digitspan, household size, household demographics, number of fixed salary, daily wage, and self-employed workers, and business type. The regressions in the even columns include all controls in Table 4. Data in this table comes from rounds 1-4 of data collection.

Table 5: Do Respondents Distort Responses?

	(1)	(2)	(3)	(4)	(5)	(6)
	All Questions Pooled	Quintile Questions	Relative Questions	All Questions Pooled	Quintile Questions	Relative Questions
Rank	0.161*** (0.016)	0.159*** (0.017)	0.163*** (0.018)			
Rank*Stakes	-0.056*** (0.021)	-0.066*** (0.024)	-0.049** (0.023)			
Average Rank				0.251*** (0.024)	0.264*** (0.029)	0.243*** (0.025)
Average Rank*Stakes				-0.060* (0.034)	-0.093** (0.042)	-0.038 (0.036)
Reports	Individual	Individual	Individual	Average	Average	Average
N	32009	13101	18908	6524	2669	3855
No. Obs	1336	1336	1336	1336	1336	1336

* $p \leq 0.10$, ** $p \leq 0.05$, *** $p \leq 0.01$. Notes: Data in this table comes from Round 1 (Baseline) of data collection. The characteristics in Panels A and B are of the entrepreneur that was ranked in the elicitation exercise. Standard errors are clustered at group level. The model includes neighborhood cluster, surveyor, and date of survey fixed effects. The left hand side variable is the percentile of the outcome in question. The regressor is the percentile of the average rank given to a respondent, computed by question. The level of observation is the ranker-rankee in Columns 1-3 and the rankee in Columns 4-6.

Table 6: How Do Incentives and Public Reporting Affect Responses?

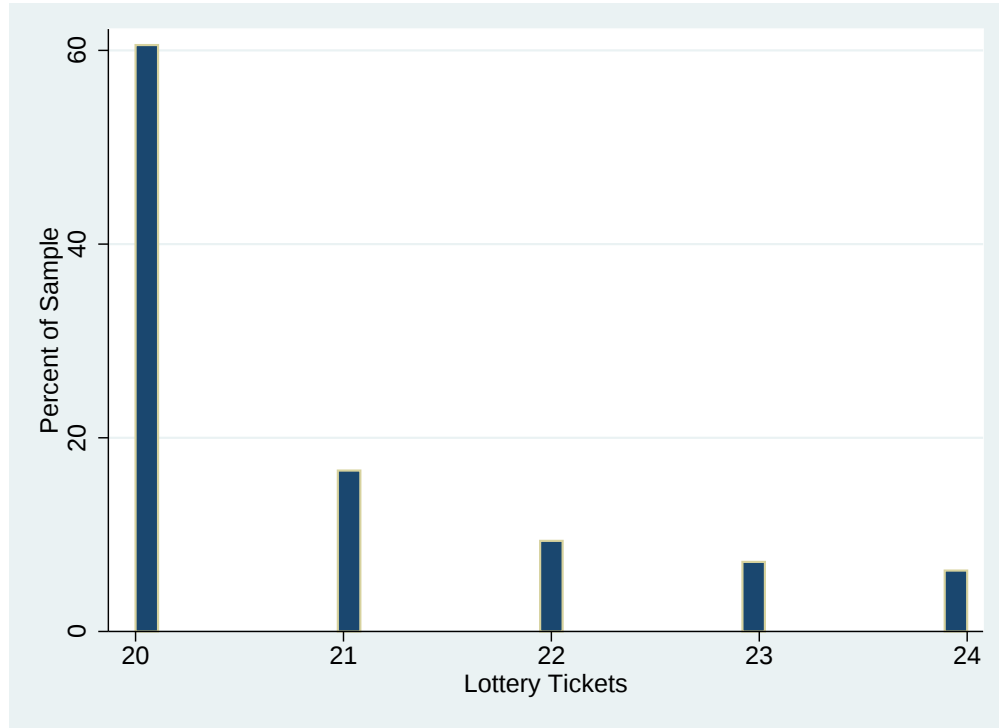
	(1)	(2)	(3)	(4)
	All Questions	All Questions	All Questions	All Questions
	Pooled	Pooled	Pooled	Pooled
Average Rank	0.210*** (0.035)	0.161*** (0.042)	0.136*** (0.046)	0.112** (0.048)
Average Rank*Public	-0.002 (0.053)	0.000 (0.061)	0.165** (0.065)	0.030 (0.060)
Average Rank*Incentives	-0.019 (0.061)	-0.086 (0.065)	0.148** (0.067)	0.143** (0.072)
Average Rank*Incentives*Public	-0.020 (0.092)	0.050 (0.098)	-0.228** (0.095)	-0.119 (0.099)
Who is Ranked?	Self	Self	Not Self	Not Self
Treatment	[No Stakes]	[Stakes]	[No Stakes]	[Stakes]
N	3218	3276	3231	3289
No. Obs	1330	1330	1336	1336

* $p \leq 0.10$, ** $p \leq 0.05$, *** $p \leq 0.01$. Notes: Data in this table comes from Round 1 (Baseline) of data collection. The characteristics in Panels A and B are of the entrepreneur that was ranked in the elicitation exercise. Standard errors are clustered at group level. The model includes neighborhood cluster, surveyor, and date of survey fixed effects. The left hand side variable is the percentile of the outcome in question. The regressor is the percentile of the average rank given to a respondent, computed by question. The level of observation is the rankee.

Appendix

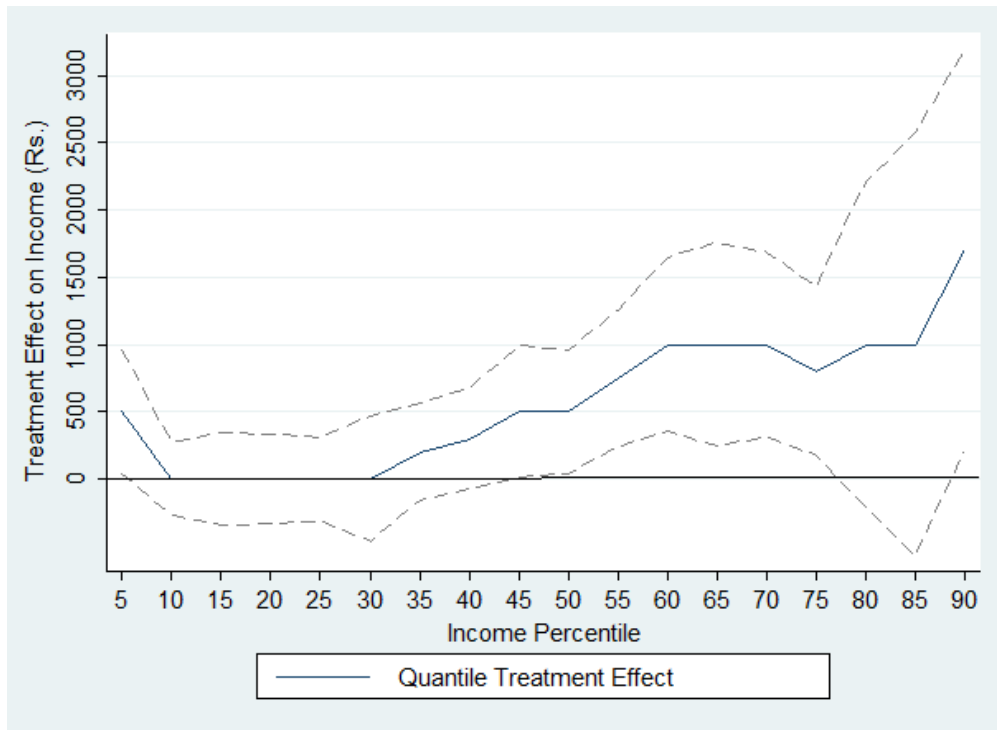
A1 Appendix Figures

Figure A1: Distribution of Lottery Tickets



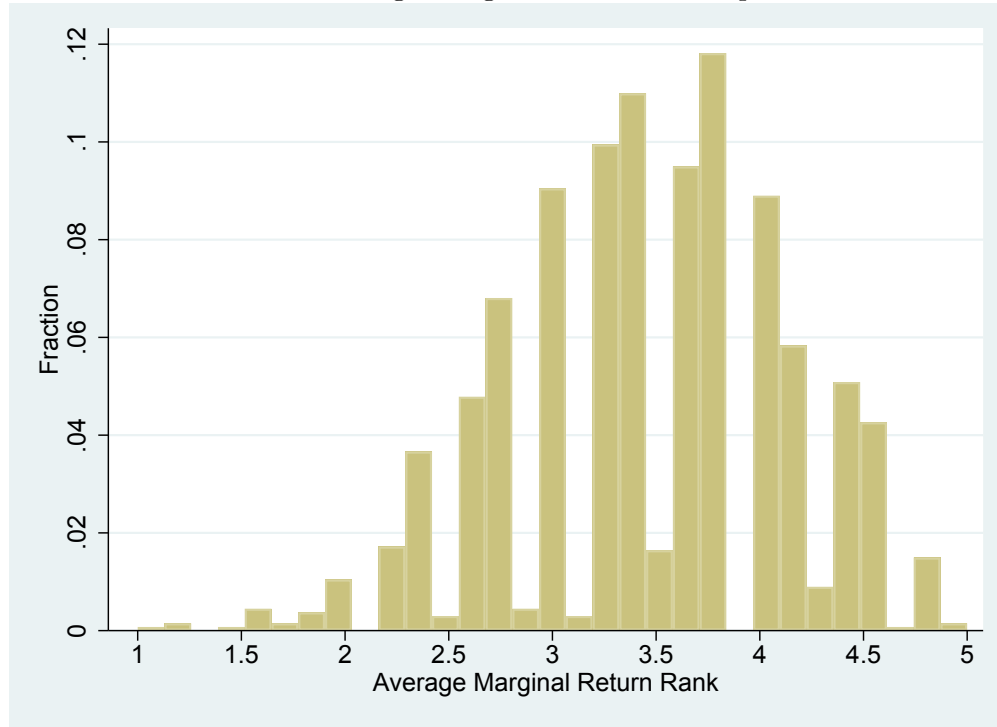
Notes: This figure plots the distribution of lottery tickets in the sample. Lottery tickets were used to select the grant winners. In the *NoStakes* treatment group, participants received 20 lottery tickets and each group member was equally likely to have their tickets drawn from the urn. In the *HighStakes* group, participants were eligible to receive up to 4 extra lottery tickets, based on whether their peers ranked them highest for the treatment questions.

Figure A2: Quantile Treatment Effects



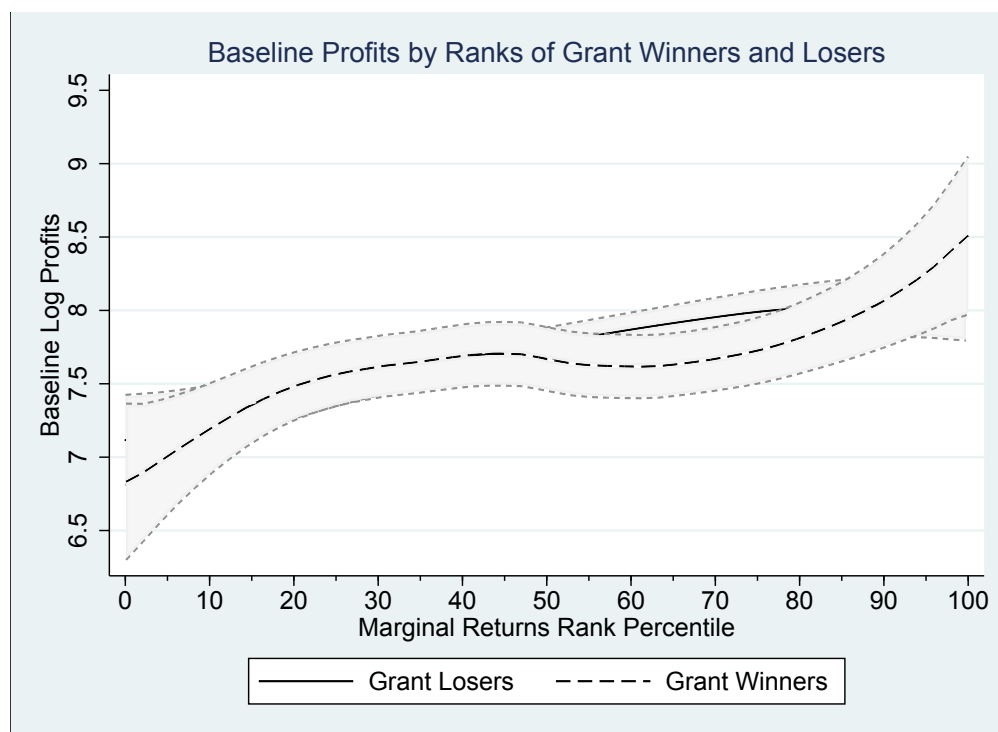
Notes: This figure plots the quantile treatment effects (blue line) obtained from quantile regressions from the 5th to the 95th quantile. The regressions include surveyor, survey month, and randomization strata fixed effects. Standard errors are clustered at the group level. The 90% confidence bands are represented by the dotted lines.

Figure A3: Distribution of the Average Marginal Returns Rank
of the Average Marginal Returns Rank.pdf



Notes: This figure plots the distribution of the average marginal return rank. The average marginal return rank is the mean of every rank assigned to person i by all of her peers in her group. As groups consist of 4-6 members, the average rank ranges between 1 and 5.

Figure A4: Baseline Profits by Percentile of the Average Community Ranks Distribution



Notes: This figure plots two kernel-weighted local polynomial regressions of baseline log profits on the marginal returns rank percentile, estimated separately for respondents who won and respondents who did not win grants. Baseline log profits is the log value of average profits in the pre grant disbursement period. The marginal returns rank percentile is the percentile of the average rank assigned to person i by all of her peers in her group. 90% confidence bands are shown.

A2 Appendix Tables

Table A1: Balance Check

	(1) No Stakes Mean	(2) Stakes Difference	(3) No Incentives Mean	(4) Incentive Difference	(5) Private Mean	(6) Public Difference	(7) Grant Loser Mean	(8) Grant Winner Difference	(9) N
<i>Panel A: Individual Characteristics of Ranked Entrepreneur</i>									
Male	0.605	0.035	0.640	-0.036	0.620	0.010	0.627	0.001	1334
Education	7.346	-1.468	7.360	-1.443	5.859	1.727	7.205	-1.897	1335
Married	1.269	-0.015	1.243	0.053	1.257	0.014	1.262	0.008	1335
Age	40.517	1.147	41.045	0.220	40.847	0.653	41.067	0.213	1334
Digitspan	5.275	-0.095	5.257	0.004	5.244	0.009	5.224	0.054	1339
Wage Exit Self-Employment	13384.501	-880.867	13197.333	-575.457	13272.793	-858.485	13122.148	-341.264	1344
Total Hours Worked Past Week	40.594	2.048	36.662	-1.183	36.616	-0.417	36.664	-0.365	1336
Total Days Worked Past Month	22.659	0.148	21.178	-0.331	21.039	0.495	21.115	-0.263	1344
<i>Panel B: Characteristics of Household Businesses</i>									
Business Type- Manufacturing	0.256	-0.026	0.244	-0.004	0.237	0.004	0.238	0.004	1344
Business Type- Retail	0.323	0.004	0.316	0.021	0.334	-0.012	0.331	-0.016	1344
Business Type- Service	0.219	-0.016	0.215	0.007	0.224	-0.016	0.211	0.011	1344
Business Type- Piecerate	0.079	-0.003	0.064	0.024	0.064	0.023	0.084	-0.017	1344
Business Type- Livestock	0.031	0.022**	0.047	-0.015	0.043	-0.004	0.045	-0.009	1344
Business Type- Food Preparation	0.058	0.028*	0.083	-0.027*	0.071	-0.000	0.056	0.044***	1344
Business Type- Construction	0.022	-0.004	0.024	-0.008	0.022	-0.002	0.022	-0.007	1344
Business Type- Agricultural	0.001	-0.000	0.000	0.003	0.000	0.004	0.002	-0.002	1344
<i>Panel C: Household Characteristics</i>									
Household Size	3.805	-0.038	3.747	0.060	3.812	-0.071	3.800	-0.054	1334
No. Children 0-5	0.458	-0.088**	0.381	0.058	0.418	-0.023	0.428	-0.051	1344
No. Children 6-12	0.490	0.085*	0.566	-0.059	0.541	-0.024	0.569	-0.106**	1344
No. Salaried HH Members	0.471	-0.061	0.424	0.053	0.443	-0.001	0.450	-0.012	1344
No. Daily Wage HH Members	0.283	-0.008	0.264	0.005	0.288	-0.031	0.283	-0.047	1344
Baseline Assets	41028.608	25085.131	43906.256	21000.920	74981.796	-37743.733	59977.609	-15687.780	1344
Value of HH Assets	467921.870	49096.021	508773.834	-39818.455	533669.683	-85022.800*	509037.625	-48553.286	1344
Avg. Monthly Profits	4983.018	140.667	5125.222	-122.815	5176.845	-184.240	5075.006	-5.231	1344
Avg. Monthly Income	9026.006	-450.250	9041.556	-690.971*	8751.447	-28.446	8542.329	701.001*	1344
<i>P-Value Joint F-Test</i>		.26		.43		.67		.20	

* p ≤ 0.10, ** p ≤ 0.05, *** p ≤ 0.01. Notes: Data in this table comes from Round 1 (Baseline) of data collection. The characteristics in Panels A and B are of the entrepreneur that was ranked in the elicitation exercise. Standard errors are clustered at group level. The model includes neighborhood cluster, surveyor, and date of survey fixed effects.

Table A2: Average Monthly Return to the Grant

	(1)	(2)	(3)	(4)	(5)	(6)
	Income	Log Income	IHS Income	Profits	Log Profits	IHS Profits
Winner	567.994 (405.829)	0.139 (0.093)	0.147 (0.099)	684.814** (319.068)	0.335** (0.139)	0.358** (0.150)
Mean of Outcome for Grant Losers	8311.12 [6609.07]	8.62 [1.39]	9.31 [1.46]	4588.03 [5173.47]	7.35 [2.53]	7.98 [2.72]
N	5324	5342	5342	5319	5337	5337
No. Obs	1336	1336	1336	1336	1336	1336

* $p \leq 0.10$, ** $p \leq 0.05$, *** $p \leq 0.01$. Notes: The unit of observation is the household. Winner indicates that the household is a grant recipient after baseline (round 1 of data collection). Robust standard errors clustered at the group level in parentheses. All regressions include household, survey month, survey round, and surveyor fixed effects. All regressions are weighed by the inverse propensity score described in Section 7. Data in this table comes from rounds 1-4 of data collection.

Table A3: Balance Check by Tercile of Marginal Return Rank

	Top Tercile		Middle Tercile		Bottom Tercile	
	(1)	(2)	(3)	(4)	(5)	(6)
	Grant Loser	Grant Winner	Grant Loser	Grant Winner	Grant Loser	Grant Winner
	Mean	Difference	Mean	Difference	Mean	Difference
<i>Panel A: Individual Characteristics of Ranked Entrepreneur</i>						
Male	0.647	0.072*	0.625	-0.075	0.602	0.000
Education	8.202	-4.422	7.118	-0.101	6.038	0.132
Married	1.242	-0.071	1.253	0.038	1.295	0.125
Age	39.905	0.955	40.740	1.504	42.920	-2.310
Digitspan	5.596	0.058	5.064	0.042	4.958	0.096
Wage Exit Self-Employment	13794.817	-44.102	13521.667	-872.300	11381.132	184.822
Total Hours Worked Past Week	49.439	-3.481	44.547	5.129	41.917	-0.402
Total Days Worked Past Month	25.832	-0.355	25.550	-0.286	23.955	0.809
<i>Panel B: Characteristics of Household Businesses</i>						
Business Type- Manufacturing	0.238	0.011	0.247	0.020	0.230	-0.029
Business Type- Retail	0.354	-0.022	0.337	-0.026	0.294	0.008
Business Type- Service	0.198	0.011	0.203	0.022	0.245	-0.012
Business Type- Piecerate	0.085	-0.038*	0.080	-0.013	0.087	0.006
Business Type- Livestock	0.021	0.018	0.040	-0.016	0.079	-0.026
Business Type- Food Preparation	0.070	0.047*	0.053	0.039	0.042	0.032
Business Type- Construction	0.024	-0.019*	0.023	-0.018	0.019	0.026
Business Type- Agricultural	0.003	-0.003	0.003	-0.003	0.000	0.000***
<i>Panel C: Household Characteristics</i>						
Household Size	3.788	-0.001	3.851	-0.049	3.750	-0.176
No. Children 0-5	0.390	0.003	0.453	-0.117	0.442	-0.087
No. Salaried HH Members	0.454	-0.088	0.433	0.007	0.468	0.053
No. Daily Wage HH Members	0.186	-0.033	0.273	-0.025	0.415	-0.127*
Baseline Assets	105005.116	-63757.613	50203.457	4919.489	15306.755	16734.989*
Value of HH Assets	620032.765	-1.545e+05	549928.290	-36274.015	324417.038	68597.208
Avg Monthly Profits	6104.413	-155.569	4918.372	103.011	4027.783	56.697
Avg Monthly Income	9300.610	496.290	7849.667	1801.657**	8387.925	-110.264

* p≤ 0.10, ** p≤ 0.05, *** p≤ 0.01. Notes: Data in this table comes from Round 1 (Baseline) of data collection. Robust standard errors clustered at group level in parentheses. The model includes randomization cluster, surveyor, and date of survey fixed effects.

Table A4: ANCOVA Average Monthly Returns to the Grant

	(1)	(2)	(3)	(4)	(5)	(6)
	Income	Log Income	IHS Income	Profits	Log Profits	IHS Profits
Winner	714.294** (279.068)	0.080 (0.057)	0.080 (0.060)	406.313* (215.107)	0.276** (0.109)	0.296** (0.117)
Mean of Outcome	8311.12 [6609.07]	8.62 [1.39]	9.31 [1.46]	4588.03 [5173.47]	7.35 [2.53]	7.98 [2.72]
N	3988	4006	4006	3981	3999	3999
No. Obs	1336	1336	1336	1336	1336	1336

* $p \leq 0.10$, ** $p \leq 0.05$, *** $p \leq 0.01$. Notes: Rank indicates the average ranking the entrepreneur was given by her peers for the marginal returns to grant ranking question. Winner indicates that the household is a grant recipient after baseline (after round 1 of data collection). The unit of observation is the household. Robust standard errors clustered at the group level in parentheses. All regressions include household, survey month, survey round, and surveyor fixed effects. All regressions are weighed by the inverse propensity score described in Section 7. Data in this table comes from rounds 1-4 of data collection.

Table A5: ANCOVA Monthly Returns by MR Rank

	(1)	(2)	(3)	(4)	(5)	(6)	(7)	(8)	(9)	(10)	(11)	(12)
	Income	Income	Log Income	Log Income	IHS Income	IHS Income	Profits	Profits	Log Profits	Log Profits	IHS Profits	IHS Profits
Winner*Rank	465.187 (369.014)		0.138 (0.084)		0.148 (0.089)		563.291** (277.026)		0.302* (0.158)		0.323* (0.170)	
Winner*Top Tercile Rank		1676.339** (674.741)		0.363** (0.142)		0.379** (0.150)		1549.716*** (536.544)		0.658** (0.264)		0.695** (0.284)
Winner*Middle Tercile Rank		1309.838** (520.526)		0.320** (0.142)		0.329** (0.151)		663.525* (382.859)		0.257 (0.270)		0.264 (0.292)
Winner	-996.348 (1218.452)	-427.236 (397.861)	-0.426 (0.304)	-0.183 (0.112)	-0.458 (0.323)	-0.193 (0.119)	-1649.447* (890.908)	-479.401* (267.627)	-0.821 (0.571)	-0.093 (0.213)	-0.879 (0.617)	-0.092 (0.231)
<i>P-value from F-Test</i>												
Winner*Top Tercile Rank= Winner*Middle Tercile Rank		0.551		0.728		0.705		0.073*		0.081*		0.081*
Mean of	8311.12	8311.12	8.62	8.62	9.31	9.31	4588.03	4588.03	7.35	7.35	7.98	7.98
Outcome	6609.07	6609.07	1.39	1.39	1.46	1.46	5173.47	5173.47	2.53	2.53	2.72	2.72
N	3988	3988	4006	4006	4006	4006	3981	3981	3999	3999	3999	3999
No. Obs	1336	1336	1336	1336	1336	1336	1336	1336	1336	1336	1336	1336

* $p \leq 0.10$, ** $p \leq 0.05$, *** $p \leq 0.01$. Notes: Rank indicates the average ranking the entrepreneur was given by her peers for the marginal returns to grant ranking question. Top (Middle) Tercile Rank is a dummy for whether the entrepreneur is in the top (middle) tercile of the average marginal return rank distribution. Winner indicates that the household is a grant recipient after baseline (after round 1 of data collection). The unit of observation is the household. Robust standard errors clustered at the group level in parentheses. All regressions include household, survey month, survey round, and surveyor fixed effects. All regressions are weighed by the inverse ropensity score described in Section 7. Data in this table comes from rounds 1-4 of data collection.

Table A6: Returns without Controls—Excluding Self Rank

	(1)	(2)	(3)	(4)	(5)	(6)
	Income	Log Income	IHS Income	Profits	Log Profits	IHS Profits
<i>Panel A: Average Rank Value</i>						
Winner*Rank (No Self)	993.762*** (375.339)	0.203** (0.080)	0.216** (0.085)	390.117 (249.278)	0.391*** (0.133)	0.426*** (0.143)
Winner	-2828.980** (1272.213)	-0.604** (0.266)	-0.646** (0.283)	-856.053 (783.712)	-1.053** (0.460)	-1.152** (0.496)
<i>Panel B: Average Rank Tercile</i>						
Winner*Top Tercile Rank (No Self)	1825.550*** (697.061)	0.318* (0.175)	0.337* (0.186)	906.772* (489.602)	0.700** (0.282)	0.758** (0.304)
Winner*Middle Tercile Rank (No Self)	408.567 (633.808)	0.052 (0.152)	0.056 (0.162)	-47.407 (375.334)	0.173 (0.274)	0.192 (0.296)
Winner	-331.269 (531.776)	-0.062 (0.130)	-0.067 (0.138)	122.331 (333.397)	-0.066 (0.225)	-0.077 (0.243)
<i>P-value from F-Test</i>						
Winner*Top Tercile Rank= Winner*Middle Tercile Rank	0.034**	0.097*	0.100*	0.042**	0.034**	0.036**
Mean of Outcome for Grant Losers	8197.92 [6412.96]	8.62 [1.35]	9.30 [1.42]	4552.74 [5160.11]	7.33 [2.55]	7.95 [2.74]
N	5320	5338	5338	5315	5333	5333
No. HHs	1336	1336	1336	1336	1336	1336

* $p \leq 0.10$, ** $p \leq 0.05$, *** $p \leq 0.01$. Notes: Rank indicates the average ranking the entrepreneur was given by her peers for the marginal returns to grant ranking question. Top (Middle) Tercile Rank is a dummy for whether the entrepreneur is in the top (middle) tercile of the average marginal return rank distribution. Winner indicates that the household is a grant recipient after baseline (after round 1 of data collection). The unit of observation is the household. Robust standard errors clustered at the group level in parentheses. All regressions include household, survey month, survey round, and surveyor fixed effects. All regressions are weighed by the inverse propensity score described in Section 7. Data in this table comes from rounds 1-4 of data collection.

Table A7: Returns without Controls—MR Relative Ranking

	(1) Income	(2) Log Income	(3) IHS Income	(4) Profits	(5) Log Profits	(6) IHS Profits
<i>Panel A: Average Rank Value</i>						
Winner*Relative Rank	586.226 (362.726)	0.090 (0.087)	0.097 (0.092)	141.020 (266.155)	0.226* (0.126)	0.250* (0.136)
Winner	-1239.137 (1047.211)	-0.190 (0.260)	-0.207 (0.277)	37.158 (756.613)	-0.412 (0.397)	-0.467 (0.428)
<i>Panel B: Average Rank Tercile</i>						
Winner*Top Tercile Relative Rank	1231.924* (686.802)	0.083 (0.185)	0.087 (0.197)	273.555 (483.169)	0.472* (0.275)	0.523* (0.297)
Winner*Middle Tercile Relative Rank	-27.193 (599.114)	-0.117 (0.160)	-0.120 (0.171)	221.704 (378.066)	0.512** (0.254)	0.572** (0.274)
Winner	93.164 (480.603)	0.093 (0.140)	0.097 (0.150)	288.020 (336.017)	-0.077 (0.216)	-0.098 (0.234)
<i>P-value from F-Test</i>						
Winner*Top Tercile Rank= Winner*Middle Tercile Rank	0.062*	0.204	0.215	0.917	0.866	0.849
Mean of Outcome for Grant Losers	8197.92 [6412.96]	8.62 [1.35]	9.30 [1.42]	4552.74 [5160.11]	7.33 [2.55]	7.95 [2.74]
N	5324	5342	5342	5319	5337	5337
No. HHs	1336	1336	1336	1336	1336	1336

* $p \leq 0.10$, ** $p \leq 0.05$, *** $p \leq 0.01$. Notes: Rank indicates the average ranking the entrepreneur was given by her peers for the marginal returns to grant ranking question. Top (Middle) Tercile Rank is a dummy for whether the entrepreneur is in the top (middle) tercile of the average marginal return rank distribution. Winner indicates that the household is a grant recipient after baseline (after round 1 of data collection). The unit of observation is the household. Robust standard errors clustered at the group level in parentheses. All regressions include household, survey month, survey round, and surveyor fixed effects. All regressions are weighed by the inverse propensity score described in Section 7. Data in this table comes from rounds 1-4 of data collection.

Table A8: Do Peer Reports Predict True Marginal Returns to the Grant? Client Level Regressions

	(1)	(2)	(3)	(4)	(5)	(6)
	Income	Log Income	IHS Income	Profits	Log Profits	IHS Profits
<i>Panel A: Average Rank Value</i>						
Winner*Rank	845.889** (387.166)	0.198** (0.087)	0.211** (0.093)	263.664 (240.365)	0.371** (0.146)	0.406** (0.158)
Winner	-2425.919* (1373.998)	-0.609** (0.304)	-0.655** (0.324)	-396.647 (805.601)	-1.078** (0.528)	-1.189** (0.570)
<i>Panel B: Average Rank Tercile</i>						
Winner*Top Tercile Rank	1641.495** (660.235)	0.462*** (0.174)	0.492*** (0.186)	937.607** (443.212)	0.802*** (0.261)	0.871*** (0.282)
Winner*Middle Tercile Rank	204.056 (619.962)	0.127 (0.141)	0.134 (0.150)	88.056 (350.157)	0.139 (0.272)	0.150 (0.294)
Winner	-172.311 (488.504)	-0.140 (0.123)	-0.151 (0.131)	126.644 (292.720)	-0.140 (0.210)	-0.156 (0.227)
<i>P-value from F-Test</i>						
Winner*Top Tercile Rank= Winner*Middle Tercile Rank	0.035**	0.036**	0.035**	0.068*	0.007***	0.006***
Mean of Outcome for Grant Losers	8192.90 [6411.52]	8.62 [1.35]	9.30 [1.43]	4175.47 [4896.27]	7.16 [2.65]	7.78 [2.84]
N	5316	5333	5333	5311	5328	5328
No. HHs	1334	1334	1334	1334	1334	1334

* $p \leq 0.10$, ** $p \leq 0.05$, *** $p \leq 0.01$. Notes: Rank indicates the average ranking the entrepreneur was given by her peers for the marginal returns to grant ranking question. Top (Middle) Tercile Rank is a dummy for whether the entrepreneur is in the top (middle) tercile of the average marginal return rank distribution. Winner indicates that the household is a grant recipient after baseline (after round 1 of data collection). The unit of observation is the household. Robust standard errors clustered at the group level in parentheses. All regressions include household, survey month, survey round, and surveyor fixed effects. All regressions are weighed by the inverse propensity score described in Section 7. Data in this table comes from rounds 1-4 of data collection.

Table A9: Do Peer Reports Predict True Marginal Returns to the Grant? Includes Demonitization Survey Wave

	(1) Income	(2) Log Income	(3) IHS Income	(4) Profits	(5) Log Profits	(6) IHS Profits
<i>Panel A: Average Rank Value</i>						
Winner*Rank	585.892*	0.148*	0.160*	106.766	0.333**	0.368**
	(353.974)	(0.086)	(0.091)	(232.859)	(0.143)	(0.154)
Winner	-1568.151	-0.475	-0.515	19.416	-0.972*	-1.079*
	(1279.471)	(0.298)	(0.317)	(796.241)	(0.518)	(0.559)
<i>Panel B: Average Rank Tercile</i>						
Winner*Top Tercile Rank	1259.024**	0.370**	0.396**	724.050*	0.725***	0.789***
	(580.769)	(0.164)	(0.176)	(410.899)	(0.253)	(0.273)
Winner*Middle Tercile Rank	299.392	0.120	0.129	-45.193	0.276	0.305
	(612.973)	(0.139)	(0.148)	(385.429)	(0.261)	(0.282)
Winner	-106.904	-0.141	-0.153	118.307	-0.179	-0.199
	(457.136)	(0.118)	(0.126)	(309.941)	(0.205)	(0.221)
<i>P-value from F-Test</i>						
Winner*Top Tercile Rank= Winner*Middle Tercile Rank	0.152	0.096*	0.095*	0.086*	0.050**	0.050*
Mean of Outcome for Grant Losers	8164.55	8.60	9.28	4473.72	7.23	7.85
	[6433.29]	[1.42]	[1.50]	[4990.59]	[2.68]	[2.88]
N	6654	6677	6677	6649	6672	6672
No. HHs	1336	1336	1336	1336	1336	1336

* $p \leq 0.10$, ** $p \leq 0.05$, *** $p \leq 0.01$. Notes: Rank indicates the average ranking the entrepreneur was given by her peers for the marginal returns to grant ranking question. Top (Middle) Tercile Rank is a dummy for whether the entrepreneur is in the top (middle) tercile of the average marginal return rank distribution. Winner indicates that the household is a grant recipient after baseline (after round 1 of data collection). The unit of observation is the household. Robust standard errors clustered at the group level in parentheses. All regressions include household, survey month, survey round, and surveyor fixed effects. All regressions are weighed by the inverse propensity score described in Section 7. Data in this table comes from rounds 1-4 of data collection.

Table A10: Do Marginal Returns Ranks Predict Grant Usage?

	(1)	(2)	(3)	(4)	(5)	(6)	(7)	(8)	(9)	(10)	(11)
	Rs. Added to Grant Amount	Business Expenditures	Inventory	Equipment	Labor	Other Business Expenditures	Household Expenditures	Loan Repayment	Household Repairs	Other Household Expenditures	Amt of Grant Saved
Top Tercile Rank	531.579 (411.113)	903.051*** (276.894)	692.839** (308.220)	183.343 (299.405)	-4.119 (19.505)	30.987 (72.512)	-557.747** (217.506)	-2.977 (88.292)	61.071 (38.368)	-615.840*** (195.208)	-357.152* (199.413)
Middle Tercile Rank	163.772 (240.717)	517.686* (305.849)	364.869 (333.065)	0.097 (294.817)	-7.452 (14.342)	160.172* (92.349)	-573.271** (222.125)	-60.075 (89.715)	-1.019 (12.014)	-512.177** (203.649)	47.548 (241.655)
<i>P-value from F-Test</i>											
Winner*Top Tercile Rank=	0.431	0.137	0.322	0.551	0.890	0.216	0.932	0.411	0.140	0.527	0.050*
Winner*Middle Tercile Rank											
Mean of	852.13	4538.00	2596.08	1776.91	14.13	150.90	737.02	82.39	27.09	627.54	729.93
Outcome	3111.12	2255.73	2606.43	2497.10	160.76	722.13	1633.86	624.25	348.48	1507.78	1735.95
N	445	445	445	445	445	445	443	443	443	443	445
No. HHs	446	446	446	446	446	446	443	443	443	443	446

* p≤ 0.10, ** p≤ 0.05, *** p≤ 0.01. Notes: Rank indicates the average ranking the entrepreneur was given by her peers for the marginal returns to grant ranking question. Top (Middle) Tercile Rank is a dummy for whether the entrepreneur is in the top (middle) tercile of the average marginal return rank distribution. The unit of observation is the household. Robust standard errors clustered at the group level in parentheses. All regressions include household, survey month, survey round, and surveyor fixed effects. Data in this table comes from rounds 1-4 of data collection for grant winners.

Table A11: Baseline Differences Between Top, Middle, and Bottom-Ranked Entrepreneurs

	(1) Bottom Tercile Rank Mean	(2) Middle Tercile Rank Difference	(3) Top Tercile Rank Difference
<i>Panel A: Entrepreneur Characteristics</i>			
Male	0.605	-0.004	0.083***
Education	6.061	0.923***	0.236
Married	1.319	-0.053	-0.104**
Age	42.219	-1.060	-2.032***
Digitspan	4.964	0.158	0.570***
Wage Exit Self-Employment	11700.508	983.468	1984.757***
Owner Hours Worked Past Week	39.094	2.610	5.033***
OwnerDays Worked Past Month	22.008	0.917*	1.117**
Business Employed in 5 Yrs	0.816	0.033	0.021
Monthly Sales Change 2014	345.951	182.353	395.467***
<i>Panel B: Business Type</i>			
Business Type- Manufacturing	0.221	0.032	0.017
Business Type- Retail	0.297	0.030	0.050
Business Type- Service	0.241	-0.030	-0.029
Business Type- Piecerate	0.091	-0.020	-0.030
Business Type- Livestock	0.066	-0.027*	-0.036***
Business Type- Food Preparation	0.053	0.015	0.030*
Business Type- Construction	0.025	-0.004	-0.005
Business Type- Agricultural	0.000	0.002	0.002
Uses Household Labor	0.188	0.034	0.032
Uses Non-HH Labor	0.063	0.030	0.042**
<i>Panel C: Household Characteristics</i>			
Household Size	3.699	0.126	0.108
No. Children 0-5	0.420	0.004	-0.037
No. Children 6-12	0.519	0.004	0.045
Total No. HH Businesses	1.117	0.044*	0.011
No. Salaried HH Members	0.496	-0.066	-0.074
No. Daily Wage HH Members	0.366	-0.092**	-0.180***
Capital	31470.596	37032.845**	18424.085**
Value of HH Assets	352950.361	166833.467***	153664.465***
Avg. Monthly Profits	4017.595	916.613***	1646.507***
Avg. Monthly Income	8270.738	146.134	1141.693***

* $p \leq 0.10$, ** $p \leq 0.05$, *** $p \leq 0.01$. Notes: Data in this table comes from Round 1 (Baseline) of data collection. The characteristics in Panels A and B are of the entrepreneur that was ranked in the elicitation exercise. Standard errors are clustered at group level. The model includes neighborhood cluster, surveyor, and date of survey fixed effects.

Table A12: ANCOVA Returns by MR Rank with Controls

	(1)	(2)	(3)	(4)	(5)	(6)	(7)	(8)	(9)	(10)	(11)	(12)
	Income	Income	Log Income	Log Income	IHS Income	IHS Income	Profits	Profits	Log Profits	Log Profits	IHS Profits	IHS Profits
Winner*Rank	384.349 (373.532)		0.119 (0.087)		0.127 (0.092)		321.437 (240.588)		0.314** (0.151)		0.340** (0.163)	
Winner*Top Tercile Rank		1450.852** (658.884)		0.318** (0.142)		0.331** (0.151)		1262.707*** (481.373)		0.619** (0.254)		0.659** (0.274)
Winner*Middle Tercile Rank		1019.479* (535.477)		0.243* (0.135)		0.247* (0.144)		725.124* (373.890)		0.221 (0.267)		0.224 (0.289)
<i>P-value from F-Test</i>												
Winner*Top Tercile Rank= Winner*Middle Tercile Rank		0.451		0.540		0.518		0.222		0.068*		0.065*
Mean of	8311.12	8311.12	8.62	8.62	9.31	9.31	4588.03	4588.03	7.35	7.35	7.98	7.98
Outcome	6609.07	6609.07	1.39	1.39	1.46	1.46	5173.47	5173.47	2.53	2.53	2.72	2.72
N	3956	3956	3974	3974	3974	3974	3948	3948	3966	3966	3966	3966
No. Obs	1336	1336	1336	1336	1336	1336	1336	1336	1336	1336	1336	1336

* $p \leq 0.10$, ** $p \leq 0.05$, *** $p \leq 0.01$. Notes: Rank indicates the average ranking the entrepreneur was given by her peers for the marginal returns to grant ranking question. Top (Middle) Tercile Rank is a dummy for whether the entrepreneur is in the top (middle) tercile of the average marginal return rank distribution. Winner indicates that the household is a grant recipient after baseline (after round 1 of data collection). The unit of observation is the household. Robust standard errors clustered at the group level in parentheses. All regressions include household, survey month, survey round, and surveyor fixed effects. All regressions are weighed by the inverse propensity score described in Section 7. This regression also includes business sector interacted with winner fixed effects. Data in this table comes from rounds 1-4 of data collection.

Table A13: Returns with Psychometric Controls

	(1)	(2)	(3)	(4)	(5)	(6)	(7)	(8)
	Income	Income	Log Income	Log Income	Profits	Profits	Log Profits	Log Profits
Winner*Rank	861.000** (388.615)		0.193** (0.088)		337.492 (253.744)		0.502*** (0.140)	
Winner*Top Tercile Rank		1585.154** (669.373)		0.426** (0.166)		1131.645** (454.037)		1.007*** (0.248)
Winner*Middle Tercile Rank		341.732 (638.896)		0.147 (0.155)		133.873 (406.035)		0.366 (0.255)
Winner*Impulsiveness I	-912.830 (706.038)	-886.663 (695.462)	-0.126 (0.185)	-0.123 (0.183)	-998.325* (538.301)	-984.227* (531.670)	-0.108 (0.262)	-0.099 (0.262)
Winner*Impulsiveness II	-240.966 (400.889)	-257.499 (399.597)	0.027 (0.080)	0.028 (0.079)	66.378 (259.421)	58.722 (258.634)	0.015 (0.148)	0.017 (0.147)
Winner*Impulsiveness III	-147.462 (409.497)	-143.547 (407.160)	0.077 (0.101)	0.076 (0.101)	-7.090 (276.568)	-22.201 (272.916)	0.024 (0.183)	0.024 (0.183)
Winner*Optimism I	6.008 (450.948)	34.995 (452.763)	0.108 (0.119)	0.114 (0.121)	303.933 (322.613)	326.840 (329.188)	0.041 (0.207)	0.055 (0.206)
Winner*Optimism II	-457.783 (416.377)	-486.347 (416.449)	-0.025 (0.102)	-0.030 (0.102)	-364.490 (326.743)	-380.792 (326.790)	-0.040 (0.124)	-0.054 (0.122)
Winner*Optimism II	516.584 (322.019)	508.376 (323.551)	0.070 (0.060)	0.068 (0.059)	584.927** (236.867)	568.577** (237.053)	0.240** (0.100)	0.235** (0.100)
Winner*Optimism IV	-49.522 (719.360)	-70.763 (723.757)	-0.066 (0.150)	-0.070 (0.154)	519.642 (461.211)	539.055 (468.534)	0.218 (0.270)	0.202 (0.271)
Winner*Tenacity I	1282.898 (791.424)	1206.592 (778.730)	0.194 (0.136)	0.178 (0.135)	305.558 (316.589)	242.946 (315.445)	0.241 (0.210)	0.206 (0.212)
Winner*Tenacity I	56.061 (369.500)	69.123 (376.304)	0.063 (0.091)	0.065 (0.093)	164.188 (237.048)	190.786 (239.809)	0.117 (0.145)	0.118 (0.147)
Winner*Polychronicity I	-384.990 (239.928)	-362.155 (240.609)	-0.064 (0.070)	-0.060 (0.069)	-202.098 (180.933)	-191.365 (177.694)	0.114 (0.102)	0.126 (0.101)
Winner*Polychronicity II	-312.245 (442.865)	-327.270 (443.888)	-0.259** (0.126)	-0.262** (0.127)	-638.545* (326.172)	-652.558** (327.717)	-0.181 (0.138)	-0.187 (0.137)
Winner*Polychronicity III	-421.270 (464.030)	-393.700 (453.574)	-0.029 (0.089)	-0.026 (0.088)	-105.941 (326.695)	-112.911 (325.616)	-0.107 (0.313)	-0.093 (0.311)
Winner*Locus of Control I	374.551 (584.160)	408.898 (592.791)	0.142 (0.187)	0.158 (0.186)	712.464 (450.445)	769.577* (451.903)	0.501 (0.315)	0.531* (0.312)
Winner*Locus of Control II	-397.290 (313.866)	-395.501 (316.193)	0.006 (0.077)	0.008 (0.077)	-311.272 (210.586)	-288.680 (213.036)	0.072 (0.107)	0.075 (0.107)
Winner*Achievement I	51.619 (420.963)	63.844 (416.943)	0.104 (0.092)	0.107 (0.093)	-145.959 (266.704)	-151.602 (268.939)	-0.073 (0.151)	-0.064 (0.151)
Winner*Achievement II	746.574 (741.685)	729.900 (736.380)	0.022 (0.132)	0.016 (0.132)	-91.287 (381.756)	-108.567 (385.683)	-0.422 (0.298)	-0.435 (0.297)
Winner*Organization	-409.176 (697.651)	-363.428 (698.562)	-0.360** (0.166)	-0.356** (0.164)	-186.069 (502.509)	-187.588 (494.310)	-0.198 (0.299)	-0.183 (0.294)
Winner	837.866 (5206.953)	2984.549 (4976.832)	0.051 (1.218)	0.479 (1.156)	113.577 (3907.813)	798.057 (3820.777)	-3.626* (1.946)	-2.476 (1.791)
<i>P-value from F-Test</i>								
Winner*Top Tercile Rank=		0.056*		0.060*		0.043**		0.005***
Winner*Middle Tercile Rank								
Mean of	8311.12	8311.12	8.62	8.62	4588.03	4588.03	7.35	7.35
Outcome	6609.07	6609.07	1.39	1.39	5173.47	5173.47	2.53	2.53
N	5292	5292	5310	5310	5287	5287	5305	5305
No. HHs	1336	1336	1336	1336	1336	1336	1336	1336

* $p \leq 0.10$, ** $p \leq 0.05$, *** $p \leq 0.01$. Notes: Data in this table comes from Round 1 (Baseline) of data collection. The characteristics in Panels A and B are of the entrepreneur that was ranked in the elicitation exercise. Standard errors are clustered at group level. The model includes neighborhood cluster, surveyor, and date of survey fixed effects. All regressions are weighed by the inverse propensity score described in Section 7. Data in this table comes from rounds 1-4 of data collection.

Table A14: Marginal Returns Predictions Using Machine Learning versus Community Information

	(1) Profits	(2) Profits	(3) Profits	(4) Profits
Winner*Rank	468.626* (276.397)			
Winner*Top Tercile Rank		1394.302*** (531.557)		1240.201** (498.735)
Winner*Middle Tercile Rank		-10.486 (388.844)		61.916 (396.501)
Winner	-905.308 (927.676)	201.490 (348.990)	-289.479 (348.585)	-731.947* (433.784)
Winner*ML Top Tercile Rank (In)			2106.883*** (618.226)	1934.187*** (583.170)
Winner*ML Middle Tercile Rank (In)			993.840*** (377.242)	1033.727*** (382.175)
<i>P-value from F-Test</i>				
Winner*Top Tercile Rank= Winner*Middle Tercile Rank		0.009***		
Mean of	4591.42	4591.42	4591.42	4591.42
Outcome	5180.71	5180.71	5180.71	5180.71
N	5326	5326	5326	5326
No. HHs	1336	1336	1336	1336

The first column replicates the main regression in Column 4 of Table 2. The second column replicates Column 4 of Table . The ML Top Tercile Rank (In) and ML Middle Tercile Rank (In) are dummy variables for the top and middle tercile ranks of a marginal returns prediction generated by a generalized method of forests algorithm. The model is trained using data from the India experiment (therefore this is an in-sample estimate). Cross-validation yields an optimal minimum node size of 150 and the model is produced by growing 10000 trees. All models include surveyor, and date of survey fixed effects. All regressions are weighed by the inverse propensity score described in Section 7.1. Standard errors are clustered at the group level.

Table A15: How do Respondents Lie? Individual Regressions

	(1)	(2)	(3)	(4)	(5)	(6)
	Rank	Rank	Rank	Rank	Rank	Rank
Characteristic	0.373*** (0.083)	0.316*** (0.119)	0.429*** (0.113)	0.215*** (0.050)	0.063 (0.070)	0.351*** (0.062)
Characteristic*Public	-0.245** (0.122)	-0.075 (0.181)	-0.403** (0.157)	-0.056 (0.079)	0.040 (0.110)	-0.127 (0.107)
Characteristic*Incentives	-0.117 (0.123)	-0.064 (0.173)	-0.128 (0.177)	-0.165* (0.084)	-0.034 (0.114)	-0.275** (0.120)
Characteristic*Public*Incentives	0.150 (0.176)	0.019 (0.241)	0.280 (0.254)	0.245** (0.124)	0.116 (0.169)	0.347* (0.179)
Mean of Outcome	3.15 [1.37]	3.15 [1.37]	3.15 [1.37]	3.15 [1.37]	3.15 [1.37]	3.15 [1.37]
Characteristic Treatment	Family [Pooled]	Family [Stakes]	Family [No Stakes]	Peer(CR) [Pooled]	Peer(CR) [Stakes]	Peer(CR) [No Stakes]
N	25491	12911	12580	32009	16187	15822
No. HHs	1336	1336	1336	1336	1336	1336

* $p \leq 0.10$, ** $p \leq 0.05$, *** $p \leq 0.01$. Notes: Data in this table comes from Round 1 (Baseline) of data collection. The characteristics in Panels A and B are of the entrepreneur that was ranked in the elicitation exercise. Standard errors are clustered at group level. The model includes neighborhood cluster, surveyor, and date of survey fixed effects. The outcome variable is the rank pooled across the three treatment questions - income, assets, and profits - for quintile and relative ranks. The regressor is an interaction of the characteristics described at the bottom of each column and the treatment status. Self is a dummy for the respondent is ranking herself. Family is dummy for whether the respondent and the rankee are family members. Peer (SR) is a self-report of the respondent's closest peer in the group. Peer (CR) is a cross-report of the respondent's closest peer in the group given by other group members. It is a dummy for whether at least two group members agreed that the same person is the closest peer of the respondent. In columns 5-8, we drop the respondent's ranking of herself and re-rank the group members (maintaining the same original order of the relative rank). Note that for quintiles, the evaluation remains exactly the same.

Table A16: Cross Report: Can Respondents Identify Who Has the Best Information?

	(1) Questions [Pooled]	(2) Questions [Quintile]	(3) Questions [Zero-Sum]	(4) Income [Quintile]	(5) Income [Zero-Sum]	(6) Profits [Quintile]	(7) Profits [Zero-Sum]	(8) Assets [Quintile]	(9) Assets [Zero-Sum]
Rank*Cross Report	0.074* (0.042)	0.065 (0.056)	0.083 (0.058)	0.282*** (0.103)	0.059 (0.075)	0.053 (0.078)	0.090 (0.104)	-0.036 (0.080)	0.105 (0.120)
Rank	0.132*** (0.012)	0.121*** (0.012)	0.144*** (0.013)	0.127*** (0.018)	0.136*** (0.019)	0.099*** (0.017)	0.144*** (0.018)	0.135*** (0.019)	0.151*** (0.021)
Cross Report	-0.054* (0.031)	-0.047 (0.047)	-0.062 (0.039)	-0.214** (0.108)	-0.019 (0.062)	-0.028 (0.061)	-0.176*** (0.058)	0.026 (0.062)	-0.011 (0.070)
Mean of Outcome	0.51 [0.29]	0.51 [0.29]	0.51 [0.29]	8870.25 [6868.28]	8802.26 [6826.21]	6872.76 [6066.39]	6951.65 [5965.52]	473101.96 [729425.49]	477499.21 [711400.77]
N	28233	13179	15054	4375	5051	4651	5116	4153	4887
No. HHs	1344	1344	1344	895	1029	942	1038	848	996

Notes: Robust standard errors clustered at group level in parentheses. The model includes randomization cluster, surveyor, and date of survey fixed effects. The outcome variable is the percentile of the outcome in the column header. The regressor is the percentile of the average rank given to a respondent, computed by question. * $p \leq 0.10$, ** $p \leq 0.05$, *** $p \leq 0.01$. The outcome variable is the level of the outcome in the column header. The regressor is the percentile of the rank given to a respondent by each group member, computed by question. The level of observation is the ranker-rankee pair for each question.

A3 The Robust Bayesian Truth Serum

This discussion is based on Rigol and Roth (2017).

Peer prediction mechanisms, including Witkowski and Parkes (2012) *Robust Bayesian Truth Serum* (RBTS), incentivize truthful reporting of beliefs without reference to ex-post measures of accuracy.³⁰ Instead, these mechanisms determine payments as a function of the contemporaneous reports of several respondents.

We implemented a variant of RBTS, which requires elicitation of agents' first order beliefs (the ranking that an agent assigns to each of his peers) and second order beliefs (the probability distribution the agent assigns to each possible ranking his peers may give one another). RBTS rewards an agent's second order beliefs based on their proximity to the empirical distribution of stated first order beliefs. First order beliefs are evaluated based on how "surprisingly common" they are relative to other agents' stated second order beliefs. That is, agents are compensated for first order beliefs that have empirical frequencies higher than predicted by other agents' stated second order beliefs. Witkowski and Parkes (2012) show that under the assumption of a common and admissible prior, truthful reporting is a Bayesian Nash Equilibrium. Details on the mechanics of the payment rule are deferred to the following section.

Implementation of the Robust Bayesian Truth Serum. Peer prediction methods are attractive because they make truthtelling incentive compatible and circumvent the need for ex-post verification of outcomes. The principal challenge to implementation of RBTS is its complexity. It is infeasible to describe RBTS (and its incentive compatibility) to respondents in our setting who are largely innumerate. This is a challenge shared by many mechanisms implemented in practice (most notably, two-sided matching algorithms, versions of which are commonly used in education and entry-level labor markets). A common tactic, which we take in this study, is simply to assert to respondents that they can do no better than to tell the truth.³¹

In Rigol and Roth (2017) we provide evidence that this is a reasonable tactic. We report on an experiment among a sample drawn from a very similar population to that of our current study, in which compare the accuracy of peer reports when paying agents for truthfulness using a straightforward payment rule based on ex-post accuracy and when paying agents using peer prediction mechanisms. Surveyors carefully and completely explained the ex-post payment rule to respondents. For the peer prediction method, surveyors simply asserted to respondents that they would maximize their incentive payments by telling the truth. We elicit information regarding

³⁰See Prelec (2004) for a seminal contribution to this literature.

³¹The National Resident Matching Program, which matches new physicians to residency spots in the United States, has a video explanation of the steps involved in the mechanism and advises physicians that "To make the matching algorithm work best for you, create your rank order list in order of your true preferences, not how you think you will match." The video explanation and accompanying instructions do not attempt to explain why truthtelling is a dominant strategy. The website is: staging-nrmp.kinstanet.com/matching-algorithm. For the Boston Public Schools matching system, parents are told "List a number of choices (BPS recommends at least five) and order them in the true order of preference to increase the chances of getting the school that you want."

borrower reliability and entrepreneurial ability and we find that the additional accuracy induced by the simple ex-post incentive is statistically and economically indistinguishable from that induced by the peer prediction method. Both payment methods led to significantly more accurate reports than elicitation without monetary payments.

That respondents believe our assertion that they should tell the truth is reassuring, but it may nevertheless be desirable to verify that RBTS’s theoretical properties hold in practice. While RBTS is incentive compatible in theory, it may be that given the empirical distribution of beliefs, respondents can indeed increase their payoff with deceptive reports. In Rigol and Roth (2017), we verify that the payment method is incentive compatible in practice. To do so, we estimate the higher order beliefs of respondents in the sample and used these beliefs to determine respondents’ subjective expected payments from RBTS.

That RBTS is incentive compatible in practice is encouraging for several reasons. First, we do not want to deceive respondents when we tell them they can do no better than to tell the truth. Second, that assertion will only be reinforced with repeated use — because RBTS is incentive compatible, agents will receive experiential feedback over time that truth-telling is the highest paying strategy.

Details: Theory and Intuition

In this appendix section we discuss the details of the Robust Bayesian Truth Serum, an intuition for the underlying incentive properties, and our implementation of the payment rule in the field. The following discussion of the model is based on Witkowski and Parkes (2012).

Suppose there is a binary state of the world $t \in (h, l)$ (high, low) representing the entrepreneurial quality of a community member. Agents get a binary signal which is informative of the state of the world. That is each agent receives a signal $s \in \{h, l\}$ which may represent what they observe about their peer (e.g. they appear responsible, smart etc). Suppose further that all agents share a common prior about the state of the world such that they all agree on the prior probability of a high state, and they all agree on the distribution of signals conditional on the state. Let $p_h = P(s_j = h | s_i = h)$ be the probability an agent assigns to one of his peers receiving a high signal conditional on himself receiving a high signal, and analogously let $p_l = P(s_j = h | s_i = l)$. We say the common prior is *admissible* if $p_h > p_l$, which in English implies that the probability that one’s peer receives a high signal is higher if the agent himself receives a high signal. Many natural distributions satisfy this weak requirement.

In order to define the RBTS we must first define the quadratic scoring rule. Let

$$R_q(y, \omega) = \begin{cases} 2y - y^2 & \text{if } \omega = 1 \\ 1 - y^2 & \text{if } \omega = 0 \end{cases}$$

Imagine an agent trying to predict whether some true state ω is 1 or 0. The quadratic scoring

rule has the property that his expected score is maximized by reporting his true belief about the probability the state ω is 1 (see e.g. Selten (1998)).

The RBTS is implemented as follows. Every agent states their first order belief (their signal), in a report $x_i \in \{0, 1\}$ (imagine $x_i = 1$ corresponding to $s_i = h$). Further they report their second order belief $y_i \in [0, 1]$ (this is the fraction of the population they believe will report a high signal, $x_k = 1$). For each agent i , assign them a peer agent j , and a reference agent k , and calculate

$$y'_i = \begin{cases} y_j + \delta & \text{if } x_i = 1 \\ y_j - \delta & \text{if } x_i = 0 \end{cases}$$

for arbitrary δ . The RBTS payment for agent i is

$$u_i = R_q(y'_i, x_k) + R_q(y_i, x_k)$$

The main theorem of Witkowski and Parkes (2012) is that under the assumption of an admissible prior and risk neutral agents, there is a Bayes' Nash Equilibrium in which all agents report their first and second order beliefs truthfully.

The intuition behind the payment rule is fairly straightforward. The payment rule has two components. The second component incentivizes the agent to be truthful about his second order beliefs. That is, the agent is paid via the quadratic scoring rule to predict what some reference agent k will announce as his signal. And by the discussion above, agent i maximizes his expected payment from this component of the scoring rule by truthfully announcing his belief y_i about the likelihood agent k will announce a high signal. In simpler terms, the payment rule rewards agent i for choosing a second order belief as close as possible to the truth (the realized distribution of first order beliefs).

The first component of the payment rule incentivizes the agent to be truthful about his first order beliefs. The term y'_i takes an arbitrary person j 's second order belief y_j and either raises or lowers it depending on i 's report x_i . RBTS pays agent i $R_q(y'_i, x_k)$, and so i wants y'_i to be as near as possible to the true distribution of responses in the population. The admissibility assumption guarantees that if person j were to know that person i 's signal were high, then person j would increase his assessment as to the number of people in the group who received high signals. Likewise, if j were to learn that i 's signal were low, j would lower his assessment about the number of people in the group who received high signals. In effect the mechanism raises or lowers j 's assessment based on i 's report, and then pays i based on the closeness of this modified report to the truth. Thus i can do no better than to tell the truth.

Practical Implementation

We used this payment rule in the field to incentivize rank order responses about members of each group. The model and payment rule, however, were designed for binary responses. Thus while

responses contain a rank ordering of 5 people, we treat each ranking as a composite response to 25 yes/no questions of the form “Is person i the highest ranking individual in the group?”, “Is he the second highest?” and so on. We elicited second order beliefs of the form “How many people will say person i is the highest ranking individual in the group?” “How many will say he is the second highest?” and so on. From there we directly applied the payment rule, calibrated so that the expected difference between payments arising from truthful and deceptive answers was large. Note that the accuracy of responses across various questions in a single ranking were correlated, but under the assumption of risk neutrality (which is maintained throughout the peer prediction literature and may be empirically reasonable with respect to moderate sums of money), these correlations are irrelevant.

A4 Entrepreneurial Psychology

Impulsiveness:

- I plan tasks carefully.
- I make up my mind quickly
- I save regularly.

Optimism:

- In uncertain times I usually expect the best.
- If something can go wrong for me, it will.
- I’m always optimistic about my future.
- Generally speaking, most people in this community are honest and can be trusted

Locus of Control

- A person can get rich by taking risks.
- I only try things that I am sure of.

Tenacity

- I can think of many times when I persisted with work when others quit
- I continue to work on hard projects even when others oppose me.

Polychronicity:

- I like to juggle several activities at the same time
- I would rather complete an entire project every day than complete parts of several projects.
- I believe it is best to complete one task before beginning another.

Achievement

- Part of my enjoyment in doing things is improving my past performance
- If given the chance, I would make a good leader of people.

Organized person:

- My family and friends would say I am a very organized person