

Targeting, Personalization, and Engagement in an Agricultural Advisory Service*

Susan Athey
Stanford GSB

Shawn Cole
Harvard University

Shanjukta Nath
University of Georgia

Jessica Zhu
PxD

August 9, 2023

Abstract

ICT is increasingly used to deliver customized information in developing countries. We examine whether individually targeting the timing of automated voice calls meaningfully increases engagement in an agricultural advisory service. We define, estimate, and evaluate a novel recommendation system that customizes contact times to individual characteristics. This system generates significant gains, up to an 8% increase over the baseline pickup rate of 0.31. Our approach, delivered at scale, is well-suited for developing country settings. We show how to optimize around resource constraints, measure equity-efficiency trade-offs when targeting vulnerable groups, and evaluate the robustness of recommendations to technology or preference shocks.

Keywords: Recommendation Systems, Agricultural Extension Services, Agriculture, Causal Inference, Randomized Controlled Trials.

JEL Classifications: O13, C90, Q16.

*Authors are listed in alphabetical order. The Golub Capital Social Impact Lab at Stanford Graduate School of Business provided funding for this research. The authors would like to thank Sam van Herwaarden, Prasanna Kumar, Otini Mpinganjira, Tarun Pokiya, Aparna Priyadarshi, Theresa Solenski, and Hannah Timmis at PxD and Ayush Kanodia at Stanford University for their collaboration and support throughout the project. Usual disclaimers apply.

1 Introduction

Technology-enabled interventions have the potential to improve people’s access to useful information, thereby increasing welfare and reducing inequality. The benefits of automated delivery of information have been shown in various sectors, such as public health, education, labor, and agriculture.¹ Information delivery through mobile phones may hold particular promise in developing-country settings, where access to the internet may be limited or nonexistent, and information barriers can be severe. Automated delivery further enables the customization of both the content and the schedule for delivery.

In this paper, we study a particular form of customization — the scheduling of delivery of information at a convenient time for the recipient — in the context of an agricultural advisory service in India that serves over one million farmers.² The service delivers weekly information to farmers via prerecorded telephone calls; the information is customized to the farmer’s activities (e.g., geography, land type, seed variety, planting method) at a specific point in the growing season. This project develops, implements, and evaluates a program that targets call times to farmers based on their observable characteristics and history of interaction with the service. We find substantial scope for gains from customization, with our preferred estimate suggesting an 8% improvement in user engagement from a baseline of 31 percentage points (pp) relative to a non-optimized policy. If we were to implement the policy on the entire user base of farmers, this would translate to additional engagement for approximately 30,000 farmers (over the baseline non-optimized policy) with the service each week. Those increased engagements are the first crucial step in farmers’ learning and adoption of improved agriculture practices. We further identify efficiency-equity trade-offs in the implementation of targeted policies, and we show that a small sacrifice in overall

¹The benefits of automated delivery are demonstrated in public health (Araya et al., 2021; Berman and Fenaughty, 2005; Ekeland et al., 2010; Knight et al., 2021; Lee et al., 2021; Voss et al., 2019), education (Agrawal et al., 2022; Hoxby, 2014; Rodriguez-Segura, 2022), labor (Dammert et al., 2013, 2015), and agriculture (Cole and Fernando, 2021; Fabregas et al., 2019).

²The advisory service served 1.3 million farmers at the time of this study, while its user base has grown over time. As of April 2023, approximately 5.3 million farmers are registered with this service.

expected outcomes can yield a significant increase in gender equity.

Prior to our engagement with the NGO, calls were spread roughly evenly across the week to manage bandwidth constraints. During that time, no efforts were made to match call times to user preferences or availability. In collaboration with the NGO, we changed the “default” policy from a purely ad hoc one to a purely random “uniform policy” in which users were assigned with equal probability to each of the 91 potential calling times within the week.

As an initial step, we used data from the outcome of this “uniform policy” to train a machine learning model, with the goal to maximize user engagement while obeying bandwidth constraints, resulting in an “estimated optimal policy.” Next, we conducted a prospective randomized evaluation, which compares the estimated optimal policy against the uniform policy by randomly assigning farmers to either their predicted optimal time (the targeted policy) or a time at random (the uniform policy). Using the data collected from these two policies allowed us to conduct two distinct types of evaluations. First, we conducted “on-policy” evaluations: since farmers were randomly assigned (as part of real world service operations) to two groups, where one group received the targeted policy and the other the uniform policy, we were able to estimate the treatment effect of the targeted policy by comparing the sample mean outcomes from the two groups of users. Second, we conducted “off-policy” evaluations, where we estimated the counterfactual benefits of alternative policies (policies mapping farmer characteristics to call times) that were not used in practice. Our off-policy evaluations took advantage of the fact that millions of calls were randomly assigned to call times under the uniform policy. We repeated the process of estimating optimal policies based on previous data and then implementing them in a randomized evaluation several times, each time updating the estimated optimal policy using the most recent data. We also provide new evidence on the reliability of off-policy estimates by comparing the gains estimated from off-policy estimates obtained at the point of policy design (using data until time t) to actual gains that resulted from on-policy evaluation

after time t . This analysis highlights the extent to which changes in user behavior (driven, for example, by changes in preferences, circumstance, or other shocks) created differences between on- and off-policy estimates. On-policy estimates also show the improvements in statistical power from implementing policies in practice in a setting where there are many possible call times to choose from.

The developing-country setting where we implemented our interventions introduced significant challenges, requiring us to use novel approaches to develop and evaluate recommendation systems.

First, the mission of the NGO with which we partnered is to improve the lives of the very poor. All of their communication occurs via telephone, both because many farmers have limited literacy and because few have access to the internet. The NGO has attracted a user base of over one million farmers receiving weekly calls. Although the user base is large relative to many economic field experiments, it is small relative to settings where recommendation systems have been employed in practice, such as large online shopping platforms or popular apps. Relative to settings where policy estimation and evaluation have typically been rigorously analyzed, which usually involve a handful of treatment arms, our setting with many treatment arms (possible call times) introduces additional challenges. We demonstrate the benefit of deploying policies rather than using historical data for off-policy evaluation; even with a million users, on-policy evaluations have substantially more statistical power than off-policy evaluations, due to the large number of treatment arms.

Second, because the agricultural advisory service, which is run in partnership between a nonprofit organization and the government, has significant bandwidth constraints, the development of targeted treatment assignment policies requires complex optimization. It further introduces the possibility of trade-offs between equity and efficiency when determining whether to allocate scarce time slots to vulnerable groups. Our approach allows us to quantify the magnitude of these trade-offs for vulnerable groups such as female farmers or non-smartphone users. For instance, we estimate counterfactually that prioritizing female

farmers can reduce the gender gap in farmer engagement with the service from 3 pp to 1.3 pp, with only a very modest reduction in total farmers reached.

Third, agriculture is a seasonal undertaking, and farmer behavior (or technology) may be subject to time-varying shocks; this can degrade the performance of recommendation systems. We suggest two approaches to mitigate their cost using off-policy evaluations. First, we show that placing greater weight on more recent data than equally weighing training samples from a longer time series yields substantial gains. Second, the call center did two follow-up calls if the farmer did not pick up the call on the first attempt. When we implemented the estimated optimal policies in practice, we only customized the time of the first attempt. In the presence of shocks, the first call time might fail, but we show using off-policy evaluation that if the advisory delivery system can potentially customize the time of the follow-up calls, further gains are possible.

The paper contributes to several literatures. First, while recommendation systems have rapidly grown in popularity and been adopted in the e-commerce, social network, and tech sectors in developed countries ([Portugal et al., 2018](#)), they have been evaluated less for social impact applications in the developing world ([Agrawal et al., 2022](#)).³ To this limited empirical literature, we contribute findings for a technology (automated voice calls) that can serve the poorest billion individuals who lack access to the internet in a developing country.

A second contribution is our approach to optimization. Historically, practitioners and academics have considered the problem of prioritizing a costly treatment to a subset of users (e.g., [Ascarza \(2018\)](#)). In contrast, we solve for the optimal allocation of scarce slots across the entire population. Resources are often constrained in developing-country settings; this is one of very few research papers that augments recommendation systems (typically developed in settings without constraints) ([Aher and Lobo, 2013](#); [Samin and Azim, 2019](#)) by incorporating constrained optimization in a real-world application and

³[Agrawal et al. \(2022\)](#) show the gains from personalizing content in an educational app in India.

directly measuring its impact.⁴

As a third improvement, we contribute to a recent, growing literature that compares the relative merits of on-policy evaluation to off-policy evaluation (Hitsch and Misra, 2018; Yang et al., 2020). The control group in our experiment received a uniform random treatment, which allowed us to evaluate as many new targeting policies as we liked without any further implementation cost (off-policy). Collecting data using uniform randomization of a set of actions and then conducting off-policy evaluation to assess targeted policies is an approach recently used in several field experiments in marketing (Gabel and Timoshenko, 2022; Hitsch and Misra, 2018; Simester et al., 2020). But, importantly, in addition to off-policy evaluations, we tested our designed targeted policy at scale (on-policy). On-policy evaluation provides a real-world impact evaluation with greater statistical power and allows us to compare the performance of off-policy and on-policy evaluations. Importantly, using on-policy evaluations, we find that preference or technology shocks may be large enough that off-policy evaluations do not reliably predict real-world gains (Hitsch and Misra, 2018; Kuang et al., 2020; Rabanser et al., 2019). We then further examine methods that can be used to mitigate shocks to preferences, a common challenge in practice (Hitsch and Misra, 2018). If the shocks are related to the agricultural season, then giving more weight to recent engagement data may be more useful in training the recommendation system. This approach produced better outcomes in off-policy evaluations using subsequent weeks of data. Second, we show that evaluating the targeted policy using the overall engagement in the presence of shocks, such as counting pickup rates of both the first and second call attempts, can lead to additional gains compared to evaluating only using the first call pickup rate, which is the focus of the design of the targeted policy. Moreover, personalizing the second call can result in additional gains in engagement.

⁴A different problem involving constraints that has received attention in the machine learning literature is “calibration,” where recommendation systems are constrained to give users a consistent mix of different types of content (Seymen et al., 2021; Steck, 2018). The empirical components of these papers make use of historical data and do not deploy the policies they developed in practice.

Last but not least, we studied the equity-efficiency trade-off and developed policies that can be used to reduce the engagement gap between male and female farmers. We provide a real-world empirical application that quantifies the trade-off, thus contributing to the fairness discussion in machine learning ([Athey et al., 2022](#); [Beretta et al., 2019](#); [Rambachan et al., 2020](#); [Yang and Stoyanovich, 2016](#); [Zehlike et al., 2017](#)). Our work also relates to the machine learning literature on multi-objective optimization and fairness, which attempts to find Pareto-optimal algorithms within a family of recommendation systems ([Jambor and Wang, 2010](#); [Xiao et al., 2017](#)).

Despite the fact that recommendation systems are widely deployed in industry, they are rarely evaluated externally or subject to outside scrutiny. Randomized experiments are also commonly used (but rarely published) in the industry to evaluate improvements to algorithms, but it is much less common to evaluate the introduction of recommendation systems relative to a baseline using a randomized experiment. This project may be considered an example of the “economist as plumber” approach ([Duflo, 2017](#)), in which we develop and implement improved policies in a production system, testing their efficacy in the real world. We show how to confront and address a variety of issues that arise in practice and further show how to apply the causal inference framework familiar to economists in a recommendation systems context.

The remainder of this paper is organized as follows. Section 2 provides details on the agricultural extension service and the type of crop advisory messages provided to farmers through the service. Section 3 lays out the setup for the experiment. Section 4 illustrates the findings based on data collected from users randomly assigned to the uniform policy, who were subject to randomization between 91 call times. Section 5 provides details on the methodology used to estimate, evaluate, and deploy estimated policies in this experiment. Section 6 provides details on the off-policy evaluation, equity-efficiency trade-offs, gaps between off- and on-policy results, and the ways to mitigate shocks in future weeks. We conclude in Section 7.

2 Context

Improving agricultural practices is a critical part of strategies to reduce poverty, promote food security, and address environmental concerns. While many developing countries have invested heavily in agricultural extension services, their reach is often limited, and the empirical evidence on their efficacy is quite mixed ([Anderson and Feder, 2004](#)). India faces such challenges. The rural population comprises 64% of the total population in India, and the majority (about 90%) of the poor reside in rural areas.⁵ The Government of India operates a pluralistic extension system with 90,000 extension agents ([Swanson and Davis, 2014](#)). However, less than 6% of farmers report having received extension services ([Cole and Fernando, 2021](#)).

A critical development in the past decade affecting the agricultural sector in developing countries has been the widespread availability of low-cost telephone services. Technology has opened up new opportunities for sharing information with farmers ([Aker et al., 2016](#)). The development of machine-learning techniques offers an even greater opportunity, potentially enabling agricultural advice to be customized, in an automated fashion, for millions of farmers. This paper represents the first attempt that we are aware of to bring such techniques to large populations in a developing country.

To investigate the impact of service customization and targeting, we conducted a multistage experiment with a phone-based agricultural extension service in India. This digital extension service, launched in 2018, has been developed and implemented by NGOs in collaboration with an Indian state government.⁶ By the end of 2021, it served 1.3 million smallholder farmers throughout the state with a two-way, mobile-phone-based platform and a live call center.

⁵The percent of the population in rural areas in India is for 2022 from the World Bank's DataBank, and the fraction of poor corresponds to the 2022 Multidimensional Poverty Index (MPI) report by the United Nations Development Program (UNDP).

⁶Specifically, this digital extension service has been developed and implemented by Precision Development and the Abdul Latif Jameel Poverty Action Lab in partnership with an Indian state government, with support from the Bill & Melinda Gates Foundation.

The service provides customized advisories on 21 crops, livestock, and fisheries using farmer covariate information (i.e., language, location, crop, water management) and agricultural data (i.e., weather forecast, market information, pest/disease outbreaks). Users of this service receive agricultural information through three channels: 1) weekly interactive voice response (IVR) calls that provide farmers customized farming advisories timed to the crop calendar (Outbound Calls); 2) an IVR platform that farmers can call in to listen to content from an advisory library and record their questions; and 3) a call center where farmers can call in and ask agricultural-related questions. Questions are answered by local agronomists, who send recorded answers within 48 hours.

The experiment was conducted over six weeks in October and November 2021. The weeks of this multistage experiment are defined in Table 1. Every week, we sent out an agricultural advisory call to nearly 1 million farmers, prioritizing advisories on rice, one of the most important staple crops in this Indian state. Appendix A shows example scripts of agricultural advisory messages sent to farmers.⁷

3 Setup

In this section, we first provide a high-level overview of the different data collection methods used during the six weeks of the experiments. Next, we introduce notation that will be needed to describe our experiments and analysis more formally.

As described in the introduction, two data collection methods were used for this experiment. The data collection method determines the probability (μ_{ij}) that farmer $i \in \{1, 2, \dots, N_t\}$ is called in day-hour block $j \in \mathcal{J} = 1, \dots, 91$. The first data collection method, which we refer to as “uniform randomization,” assigns each farmer i to each of 91 call times, denoted j , with equal probability, so that $\mu_{ij} = 1/91$ for all i and j . The second data col-

⁷We provide example advisory messages for two types of messages in Appendix A. One is an advisory message on pest management, and the other is an advisory message on basal fertilizer application for transplanting.

Table 1: Experiment Weeks in October and November 2021

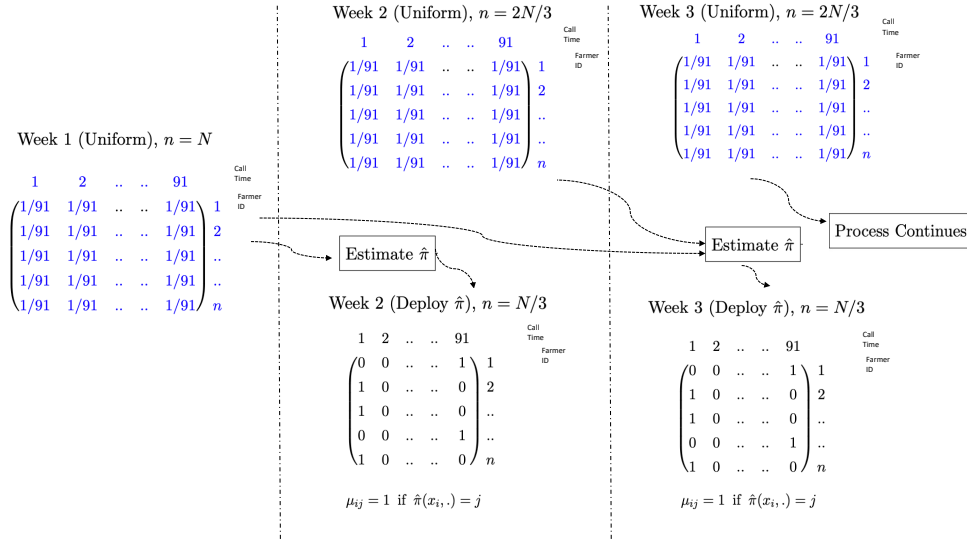
Week No.	Date	Uniform Randomization		Targeted Policy	
		Sample Name	Sample Size	Sample Name	Sample Size
1	Oct 5-10*	$N_{1,\mathcal{U}}$	881,891	—	—
2	Oct 18-24	$N_{2,\mathcal{U}}$	616,656	$N_{2,\hat{\pi}_A}$	265,188
3	Oct 25-31	$N_{3,\mathcal{U}}$	879,109	—	—
4	Nov 1, Nov 4-7*	$N_{4,\mathcal{U}}$	707,644	$N_{4,\hat{\pi}_B}$	167,995
5	Nov 8-14	$N_{5,\mathcal{U}}$	624,801	$N_{5,\hat{\pi}_C}$	234,276
6	Nov 17-23	$N_{6,\mathcal{U}}$	587,608	$N_{6,\hat{\pi}_D}$	227,386

Notes: *In these weeks, data was collected for only the indicated days of the week and not the entire week. $N_{t,\mathcal{U}}$ is used to denote the sample that is allocated to uniform randomization in week t . $N_{t,\hat{\pi}_e}$ is the sample that receives targeted policy $\hat{\pi}_e$ where $e \in \{A, B, C, D\}$ and t denotes the week. Our sample includes 880,000 farmers for whom we observe a complete set of covariates. The sample varies from around 800,000-880,000, depending on which farmers are designated to receive messages each week.

lection method, which we call a “targeted policy,” is deterministic and depends on the observable characteristics of a farmer x_i . The targeted policy π is a mapping from farmer covariates x_i to treatment arms j . Thus, the probability that farmer i is allocated to call time j is 1 if $\pi(x_i, \cdot) = j$. During the course of this study, we implemented several different targeted policies. In addition, in selected weeks, we implemented a higher-level experiment, where eligible farmers were randomized (with constant probability) into either the uniform randomization data collection or a targeted policy. Different targeted policies were used in different weeks. Data collected using the uniform randomization method from previous weeks was used to estimate a model and develop a targeted policy for the subsequent week. Figure 1 shows the data collection process for every week of the experiment.

Formally, the experiment was conducted over six weeks $t \in \mathcal{T} \equiv \{1, 2, 3, \dots, 6\}$. Each week of data is a sample N_t of farmers drawn from the population. N_t consists of those farmers who have a complete set of covariates and have signed up to receive messages assigned for week t . The treatment is call time, which is a combination of hour and day of

Figure 1: Two Data Collection Methods



Notes: The figure illustrates the two different data collection methods. We start with uniform randomization between 91 treatment arms. This data is used to estimate the targeted policies. In subsequent weeks the targeted policies are deployed (our second method of collecting data).

the week. The call center operates from 8 AM to 9 PM, seven days a week, so there are 91 hour-day call times. The experiment is conducted on a weekly basis as new crop advisories are provided to the farmers every week through the extension service. Every farmer $i \in \{1, 2, \dots, N_t\}$ can receive only one of the individual treatments $j_{it} \in \mathcal{J} \equiv \{1, 2, \dots, 91\}$ in each week. The treatment J is a random variable with support in $\mathcal{J}$ and $J \sim F_J$. Note that j is a realization of J .

The outcome is binary pickup dummy Y_{it}^{obs} . Let $Y_{it}(j)$ be the potential outcome for farmer i in week t and call time j . The observed outcome can therefore be written in terms of potential outcomes as $Y_{it}^{\text{obs}} = \sum_{j=1}^{|\mathcal{J}|} \mathbb{1}(J_{it} = j) Y_{it}(j)$. Also, Y_{it}^{obs} is a random variable with support in $\mathcal{Y} \equiv \{0, 1\}$. $x_i \in \mathcal{X}$ is the vector of observed covariates for individual i . This vector does not vary with time, so we do not have the t subscript. We let F_X denote the distribution of X with support in $\mathcal{X}$. Each week, the observed covariates are determined before an intervention is assigned. Table A1 shows the set of observed covariates. It includes gender, access to irrigation, land size, smartphone ownership, district of residence,

and historical engagement with the service prior to the start of our experiment.

The potential outcome for farmer i in call time j for week t , $Y_{it}(j)$, follows a Bernoulli distribution with probability of pickup given by $\mu(x, j, t)$. Therefore, $\mu(x, j, t) = \mathbb{E}_Y[Y_{it}(j)|x]$.

In order to construct an estimator for μ , we need to specify a prediction model m (e.g., a LASSO regression model with a given specification), where we let $\mathcal{M}$ denote the set of possible models, and a training dataset $\mathcal{S} \in [\mathcal{Y}, \mathcal{X}]^n$. Then, we denote an estimator by $\hat{\mu} : \mathcal{X} \times \mathcal{J} \times \mathcal{T} \times \mathcal{M} \times [\mathcal{Y}, \mathcal{X}]^n \rightarrow [0, 1]$. Given a ML model $m \in \mathcal{M}$ and a training dataset $\mathcal{S}$, $\hat{\mu}(x, j, t; m, \mathcal{S})$ is the estimate of $\mu(x, j, t)$.

3.1 Targeted and Uniform Policies

Several policies $\pi : \mathcal{X} \rightarrow \mathcal{J}$ were estimated, deployed, and evaluated in this project. In most scenarios in this project, we value reaching each farmer equally and do not put welfare weights on the outcome. The population object of interest is therefore defined as:

$$V(\pi, \mathcal{S}^{\text{eval}}) = \mathbb{E}_{X \sim F_X}[\mathbb{E}_Y Y_t[(\pi(X))]]$$

where $\mathcal{S}^{\text{eval}}$ is the evaluation sample. Class Π is to denote the constraints used in this experiment. For this paper, it encodes budget constraints.⁸ Therefore, the optimal policy π^* is $\pi^*(\Pi) = \arg \max\{V(\pi, \mathcal{S}^{\text{eval}}) : \pi \in \Pi\}$. This paper uses machine learning models to estimate targeted policies. The goal is to learn about the best call times for every farmer in the sample using predicted pickup rates. The estimator of π^* is denoted as $\hat{\pi}(X, t, \Pi; m, \mathcal{S})$. This estimator is a function of covariates x and is parameterized by the week t , class Π of policies, a ML model m , and the training data $\mathcal{S}$ used to estimate the parameters of the model. It returns the best call time based on $\hat{\mu}$. Hence $\hat{\pi} : \mathcal{X} \times \mathcal{T} \times \mathcal{P} \times \mathcal{M} \times [\mathcal{Y}, \mathcal{X}]^n \rightarrow \mathcal{J}$

⁸In Section 6.4, we incorporate weights into this population object of interest to discuss the equity efficiency trade-off.

where,

$$\hat{\pi}(x; t, \Pi; m, \mathcal{S}) = \arg \max_j \hat{\mu}(x; j, t; m, \mathcal{S}), \quad \text{s.t. } \pi \in \Pi$$

We call the estimated policies $\hat{\pi}$ “targeted policies” throughout this paper. The population object of interest for the targeted policy is the value of the targeted policy evaluated on an evaluation sample $\mathcal{S}^{\text{eval}}$ ($V(\hat{\pi}, \mathcal{S}^{\text{eval}})$). Therefore $V : \Pi \times \{\mathcal{Y}, \mathcal{X}\}^n \rightarrow \mathcal{R}$ and

$$V(\hat{\pi}, \mathcal{S}^{\text{eval}}) = \mathbb{E}_{X \sim F_X} [\mathbb{E}_Y Y_t[(\hat{\pi}(X, \cdot))]]$$

As illustrated in Figure 1, dividing the population into a targeted policy and a group whose call time is uniformly randomized serves to evaluate the targeted policy against a “control group” (uniform randomization), but the control group itself generates data that can be used to evaluate counterfactual policies and help design subsequent targeted policies.

In order to motivate the definition of the uniform policy, we first define a “fixed time policy” π^j where every farmer is called during call time j . Formally, $\pi^j(x) = j \quad \forall x$. The uniform call time policy is explained using the following steps. For farmer i , step 1 is to draw a random variable J_{it} from a discrete uniform distribution $\mathcal{U}\{1, 91\}$. Step 2 is if $J_{it} = j$, then call farmer i in call time j . The population object of interest for the fixed call time policy is

$$\begin{aligned} V(\pi^j, \mathcal{S}^{\text{eval}}) &= \mathbb{E}_{X \sim F_X} \mathbb{E}_Y (Y_t(\pi^j(X))) \\ &= \mathbb{E}_{X \sim F_X} (\mu(X, j, t)). \end{aligned}$$

Consequently, the population object of interest for the uniform policy is

$$\begin{aligned}
\bar{V}(\mathcal{U}, \mathcal{S}^{\mathcal{U}}) &= \mathbb{E}_{X \sim F_X} [\mathbb{E}_Y [\mathbb{E}_{J \sim \mathcal{U}} [\sum_j (Y_t(\pi^j) \mathbb{1}(J = j))]]] \\
&= \sum_j \mathbb{E}_{X \sim F_X} [\mathbb{E}_Y [Y_t(\pi^j) \mathbb{P}(J = j)]] \\
&= \sum_j \mathbb{E}_{X \sim F_X} [\mu(X, j, t) \mathbb{P}(J = j)] \\
&= \frac{1}{91} \sum_{j=1}^{91} V(\pi^j)
\end{aligned}$$

where $\mathcal{S}^{\mathcal{U}}$ is the evaluation sample for the uniform policy.⁹

4 The Effect of Call Times on Outcomes: Evidence from Uniform Randomization

As described in Figure 1, in this project, uniform randomization was used as a data collection method for a randomly selected set of farmers. Although the project's ultimate goals include estimating and evaluating alternative targeted policies, we motivate this work by reporting findings from the uniform randomization sample about the variation in engagement across 91 call times. These findings showcase the importance of call time as well as the heterogeneity of preferences across demographic groups.

Figure A1 shows farmers' engagement with the service. The outcome variable is a binary outcome representing whether a farmer picked up the first call.¹⁰ The average pickup rate is around 31%, but there is substantial variation in pickup rates at different points of time during the week. They suggest that evening and morning hours have relatively higher

⁹In this paper, $\bar{V}$ is used to denote the value of a random policy like uniform policy, and V is used for the value of a deterministic policy like the targeted policies estimated in this paper.

¹⁰In most analyses in this paper, we focus on the pickup of the first attempt of the call, while each call is attempted up to three times. In Section 6.3, we discuss the follow-up calls.

pickup than afternoon hours. We demonstrate the magnitude and statistical significance of these differences in an abbreviated manner by showing the mean pickup rates for morning, afternoon, and evening times and comparing afternoon pickup rates with morning and evening pickup rates separately. Table 2 shows the average pickup in the afternoon is 1.3 (1.8) percentage points lower than the morning (evening).

The uniform randomization data helps us identify several hours during the weekends that are high engagement hours, especially the evening hours. Historically, the agricultural advisory service did not schedule many calls during the weekends even though it operated all seven days of the week. Hence, it had inadvertently failed to take advantage of very high engagement hours.

Table 2: Variation in Engagement by Call Times

	Morning	Afternoon	Difference
Mean	0.320	0.307	0.013
Std. Error	[0.00051]	[0.00046]	[0.00069]
	Evening	Afternoon	Difference
Mean	0.325	0.307	0.018
Std. Error	[0.00052]	[0.00046]	[0.00069]

Notes: We pool the data collected using uniform randomization in weeks 1, 2, and 3 of the experiment (see Table 1 for sample size). Morning hours are between 8 AM to 12 PM, afternoon hours are between 12 PM to 5 PM, and evening hours are between 5 PM to 9 PM.

Next, we present the socio-demographic characteristics of the sample and variation in engagement for different farmer subgroups. Table A1 provides summary statistics on farmer characteristics. About 18% of the sample are female farmers. Around 37.4% of the farmers are smartphone users, and a little over 44% use irrigation systems on their land. Additionally, the residential districts for the farmers are provided in the covariates data.

Previous research has shown substantial gender gaps in access to technology in agriculture (Owusu et al., 2018; Quisumbing and Pandolfelli, 2010). Consistent with the literature, we observe a persistent but time-varying gap in engagement rates for female farmers (0.29)

relative to male farmers (0.32), as shown in Panel (a) in Figure 2. The graph also suggests that the high pickup times often overlap for male and female farmers. The overlap hours have an important implication for policy design, as simply maximizing overall pickup rates (equal weighting) could result in giving the most valuable time slots to men; alternatively, placing a higher weight on reaching women could preserve some of the higher-value slots for women. This could help reduce the gender gap in engagement with the service.

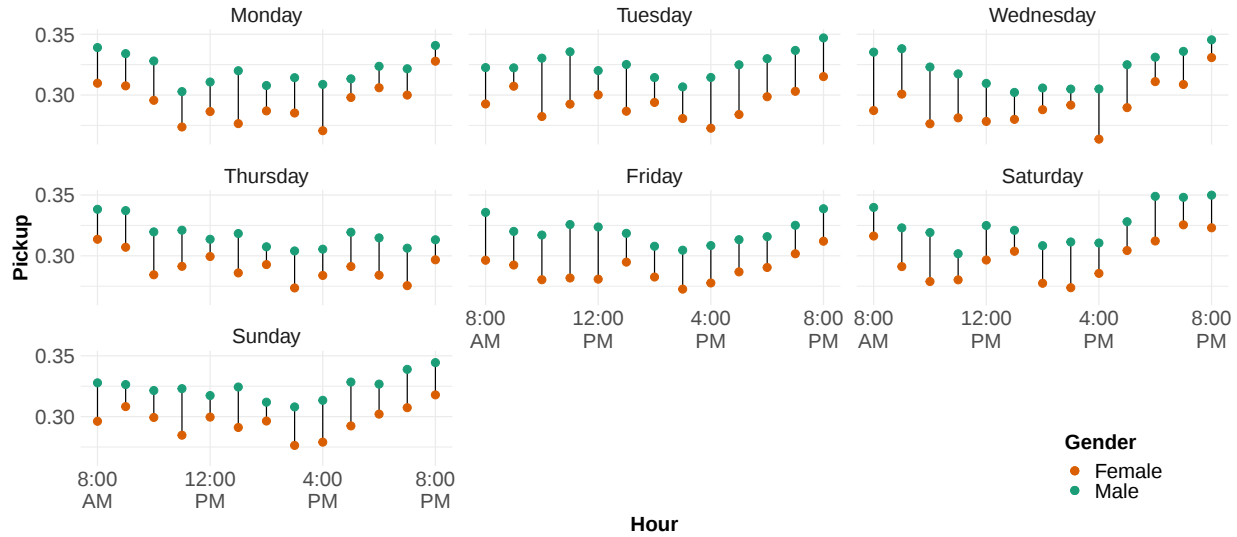
We further explore the heterogeneity in farmers' engagement by their access to information and wealth levels. Our first result is that the average pickup rate among smartphone users is 4.1 percentage points lower than among non-smartphone users. Figure A2 shows this comparison of engagement by smartphone ownership.

We then examine the variation in pickup rates by the distribution of land size, which serves as a proxy for wealth. We divide the sample of farmers into deciles based on their land size. Panel (b) in Figure A2 shows that farmers at the high end of the land size distribution have lower engagement than the farmers with smaller land sizes. These figures together suggest that non-smartphone users and poor farmers are more engaged with the service, consistent with the hypothesis that they have limited outside options to access farming-related information.

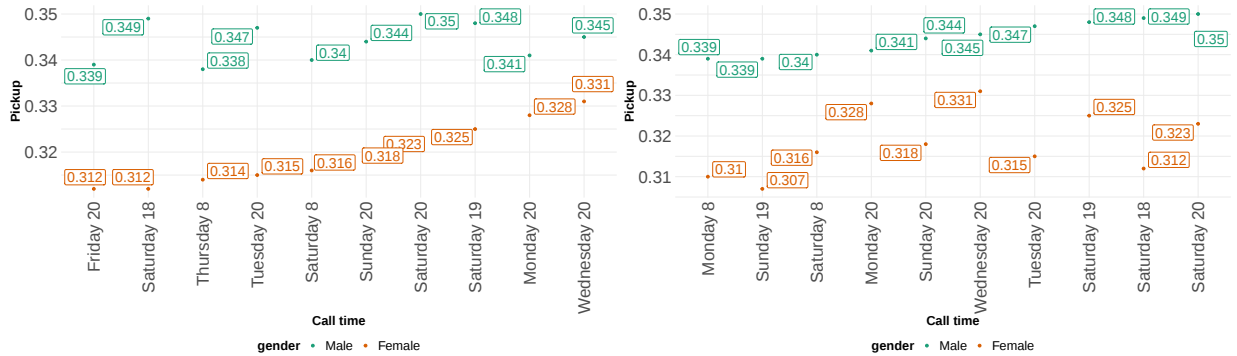
We additionally discuss some crucial features of the farmers based on their historical engagement behavior. We have data on farmers' historical engagement with the agricultural advisory service from the initiation of the service (July 2018) till right before our experiment (September 2021). We also use historical engagement data to estimate targeted policies. Figure A3 shows the distribution of pickup across the 91 call times in the historical data.

The farmers in the experiment sample registered with the agricultural advisory service at different points in time. Panel (a) in Figure A4 shows the distribution of the length of the total time with the service. The mean farmer duration with the service is 516 days. Panel (b) shows the variation in the month when they started receiving calls from the extension service. High enrollments are observed for the months of April, September, and

Figure 2: (a) Pickup by time and gender, (b) Highest Pickup arms for female farmers and (c) male farmers in the dataset.



a. Variation in Pickup by Gender for 91 Treatment Arms



b. Top 10 Buckets for Females

c. Top 10 Buckets for Males

Notes: We pool the data collected using $\mathcal{U}$ in weeks 1, 2, and 3 of the experiment. Panel (a) illustrates the gender gap in engagement over the 91 call times. Panel (b) shows the most popular call times for female farmers arranged in increasing order of popularity with the pickup for male farmers for comparison. Panel (c) shows the most popular call times for male farmers arranged in increasing order of popularity with the pickup for female farmers.

December. Since the experiment was conducted in the months of October and November 2021, it is possible that farmers who signed up with the service at different points in time value the information sent out during the months of the experiment differently. Finally, we show the variation in engagement during the weeks of the experiment by their duration with the service. Figure A4 suggests the possibility of information fatigue directly related to the time the farmers have been associated with the service. We define a new farmer dummy, which takes a value of 1 if the farmer spent less than the mean duration with the service and 0 otherwise. The engagement of new farmers is consistently higher than the engagement of the old farmers (Panel (c) in Figure A4).

Finally, the assignment of calls to hours in the historical data followed a changing set of ad hoc rules, resulting in a nonuniform distribution of calls across the 91 treatment times. Figure A5 illustrates this separately for each historical year. In 2018, no calls were made on Tuesdays and very few calls were made on Sundays. In 2020 and 2021, very few calls were made on Sundays especially during evening. This highlights the benefits of uniform randomization of call times to ensure that there is sufficient data to learn about the relationship between farmer characteristics and engagement across all call times.

5 Method for Estimating and Evaluating a Targeted Policy

As Figure 1 outlined, this project had multiple stages. In week 1, training data was collected using $\mathcal{U}$. In week 2, the uniform randomization data from week 1 was used to construct and deploy a policy in a randomized controlled trial, comparing the targeted policy to the uniform policy. In each subsequent week, data from the previous weeks' uniform policies were used to design a new targeted policy, which was again evaluated against the uniform policy.

5.1 Estimate $\hat{\pi}$ Using Uniform Randomized Data

This section describes how targeted policies were estimated. First, we selected an estimator for the outcome model. With sufficient data, it would be possible to separately estimate $\mu(\cdot, j, t)$ for each treatment arm j and week t . However, even though we have approximately 880,000 farmers in the sample with 91 treatment arms and many covariates, it was efficient to pool the data into a single model and to use data-driven model selection. Among many reasonable machine learning models (e.g., LASSO, random forest) or matrix factorization methods (e.g., recommendation systems), we chose LASSO regularized regression due to a large number of potential treatments and covariates. Our regressors include treatment indicators for each call time j , covariates x_i , and interactions of those covariates with treatment indicators.

Because the outcome model will be used to select treatments, it is important that our estimator $\hat{\mu}(X, J, t; m, \mathcal{S})$ yields useful estimates of the difference in expected outcomes across treatment arms. Since our training data $\mathcal{S}$ has uniform randomization, we do not need to be concerned that farmers assigned to different treatment arms are different in some unobserved way; the randomization ensures independence ($\{Y_i(j)\}_{j=1}^J \perp\!\!\!\perp J$) and thus unconfoundedness. Formally, the estimator takes the form:

$$\text{logit}(\mu_{ij}) = X_i\beta + \sum_{j=1}^{91} \delta_j J_i + \sum_{j=1}^{91} \gamma_j X_i J_i \quad (1)$$

where μ_{ij} is the probability of call pickup, X denotes the covariate matrix, and J denotes the treatment dummies. The objective of LASSO is to minimize the following.

$$-l(\theta, X, T) + \lambda\|\beta\| + \lambda\|\gamma\|, \quad (2)$$

where $\theta = (\beta, \delta, \gamma)$ and l denotes the log-likelihood of the outcome. The selection of the regularization parameter λ is done using cross-validation. We consider the full training

data and split it into k folds. The optimal λ_0 minimizes the average mean squared error across the k folds. The mean squared error is computed as the squared difference between the observed and predicted outcome.

The model incorporates the main effects and the covariates' interaction with the treatment dummies. Lastly, we do not penalize the coefficients on the treatment dummies. We also estimate a few variations of the above specification. For instance, a model with a polynomial function for the treatment dummies is estimated in Appendix F.1. We also estimate an additional specification where the penalty on the regularization parameters varies based on the order of the interaction terms. We call this specification the hierarchical LASSO. The interaction effects are likely smaller in models with two-way and higher-order interactions than the main effects. Hence, the penalties are allowed to increase with the degree of the interactions. Appendix F.2 discusses this variation and result.

Below, we evaluate the expected benefit from assigning farmers using a targeted policy derived from $\hat{\mu}$. If the same data is used to generate the targeted policy and evaluate its effectiveness, the result will likely overstate the benefits of the policy. We use cross-fitting to address this concern, dividing the data into K equal groups (folds). Denoting $k(i)$ the fold which contains farmer i , for each fold k , we estimate the model on all folds except k , and then use the model parameters to predict pickup for farmers in the k^{th} fold. We repeat this k times to generate $\hat{\mu}_{-k(i)}$ for all farmers in the sample (Figure A6).¹¹

The estimate of $\hat{\pi}_{-k}$ also depends on the constraints related to technological limits on the number of farmers that can be called in an hour-day combination. We account for this by using a two-step process. Step 1 uses LASSO and data collected using uniform randomization to predict $\hat{\mu}$ with cross-fitting. Step 2 uses $\hat{\mu}$ and the budget constraints to

¹¹The concept of cross-fitting is similar but distinct from cross-validation. Cross-validation is about model selection and selecting tuning parameters where the ultimate goal is to estimate a single model on the entire training data (the one that minimizes the cross-validation error), while with cross-fitting, we retain the models estimated on the different folds as an input to a subsequent step of the analysis. The subsequent step in our context is estimating the targeted policy.

allocate farmers to their optimal or near-optimal call times as shown below.

$$\begin{aligned} \max_z \quad & \sum_{i=1}^n \sum_{j=1}^J z_{ij} \hat{\mu}_{-k(i)}(x_i, j, \cdot) \\ \text{s.t.} \quad & \sum_j z_{ij} = 1 \quad \forall i, \quad \sum_i z_{ij} \leq b_j \quad \forall j \quad \text{and} \quad z_{ij} \in \{0, 1\} \end{aligned} \quad (3)$$

where b_j is the capacity limit on call time j . The decision variables are z_{ij} , which takes a value of 1 if farmer i is allocated to call time j and 0 otherwise. The decision variables can be represented in the form of a $N \times J$ matrix z , where each row corresponds to a farmer and each column represents a treatment arm. The set of constraints is combined to generate a matrix A that has dimension $(N + J) \times (N * J)$. The first N rows of A ensure that every farmer i can be allocated to only one call time j . The remaining J rows ensure that the sum of allocation in each call time cannot exceed the total capacity, and, therefore, they should add to be less than or equal to b_j . The resulting set of equations for constraints is $A \times \text{vec } z \leq [1, \dots, 1, b_1, \dots, b_J]$. Since the decision variables are discrete, this can be set up as an integer programming problem.

We use the Gurobi Parallel Mixed Integer Programming solver. This solver uses a linear-programming-based branch-and-bound algorithm for mixed integer programming problems. All the steps are summarized in the targeting algorithm below for a training dataset $\mathcal{S}$.

Algorithm 1 Algorithm for Estimating $\hat{\pi}$

Input: Capacity limits b ; Training data $\mathcal{S}$

Result: Model parameters $(\hat{\delta}, \hat{\beta}, \hat{\gamma})$, Allocation z

- 1: Partition data into K mutually exclusive folds
 - 2: Estimate LASSO using all folds except k (Equation 2): Obtain $\hat{\delta}, \hat{\beta}, \hat{\gamma}$
 - 3: Predict for each i in k : $\hat{\mu}_{-k}(x_i, J_i, \cdot)$
 - 4: Repeat 1, 2, and 3 for all folds: Append $\hat{\mu}_{(-k)} \forall k \in K$.
 - 5: Constrained Optimization: $\max_z \sum_{i=1}^n \sum_{j=1}^J z_{ij} \hat{\mu}_{-k(i)}(x_i, j, \cdot) \quad \text{s.t.} \quad \sum_j z_{ij} = 1 \forall i \quad \text{and} \quad \sum_i z_{ij} \leq b_j \forall j$ (Equation 3)
 - 6: Use $\hat{\mu}$, capacity limit b and solve for z
 - 7: **return** $[\hat{\delta}, \hat{\beta}, \hat{\gamma}, z]$
-

Once we have estimated $\hat{\pi}$, we can evaluate it in two ways. In on-policy evaluations, we implement the targeted policy in practice and compare it to the alternative uniform policy. Alternatively, we can employ off-policy or counterfactual evaluation using the data collected via $\mathcal{U}$ to evaluate the counterfactual outcome from implementing any targeted policy. Below we describe the two concepts used to evaluate targeted policies in this project.¹²

5.2 Off-Policy Evaluation

Off-policy evaluation is used to compute the value of targeted policies from a dataset that is collected using a different policy (Athey and Wager, 2021; Zhou et al., 2022). For instance, for this experiment, data is collected using uniform randomization for weeks $t \in \{1, 2, \dots, 6\}$. This uniform randomized data can be used to estimate and evaluate $\hat{\pi}$.

Appendix D provides details on off-policy evaluation for a simplified setting with two covariates and four treatment arms. The key idea is that under uniform randomization, $\frac{N}{91}$ farmers receive a treatment under $\mathcal{U}$ that matches their assignment under $\hat{\pi}$. The outcome of this subset of farmers can be used to estimate the value of targeted policy counterfactually.

Estimator of Policy Value: The estimators for the population objects of interest (defined in Section 3) are obtained using data collected from the uniform random policy. All targeted policies are based on an underlying outcome model estimated via cross-fitting. The estimator for the value of $\hat{\pi}$ under the off-policy evaluation is denoted by $\hat{V}^{\text{Off}}(\hat{\pi}, \mathcal{S}^{\text{eval}})$.

$$\hat{V}^{\text{Off}}(\hat{\pi}, \mathcal{S}^{\text{eval}}) = \frac{\sum_{i \in \mathcal{S}^{\text{eval}}} Y_{it}^{\text{obs}} \mathbb{1}(\hat{\pi}(x_i, \cdot) = J_{it})}{\sum_{i \in \mathcal{S}^{\text{eval}}} \mathbb{1}(\hat{\pi}(x_i, \cdot) = J_{it})}$$

This is an unbiased estimator of the value of the targeted policy given randomized data and the use of cross-fitting in constructing the policy. Note that after we estimate

¹²For comparing the value of two policies in the same data-set, we cannot analyze the two components of the difference separately since the same observations appear in both terms in regions of overlap. Hence, we sum over observations and for each observation take the difference.

$\hat{\pi}$ using the uniform randomization data, there is a further important step of conducting the off-policy evaluation, as we want to counterfactually estimate the value of $\hat{\pi}$ prior to deployment. In a production setting, we would only deploy $\hat{\pi}$ if, counterfactually, we see significant gains of deploying $\hat{\pi}$ over our baseline uniform policy.

5.3 On-Policy Evaluation

Once the off-policy evaluation showed that there were significant gains from deploying $\hat{\pi}$ over the uniform policy, the next step was to deploy the targeted policy $\hat{\pi}$. In the subsequent week, we randomly divided the farmers into two groups. Group 1 received calls according to $\hat{\pi}$. Group 2 received calls according to the uniform policy. At the end of the week, engagement data was collected for both the groups, as illustrated in Figure 1. This data allows us to do on-policy evaluation.

The population object of interest for the on-policy evaluation is defined as $\delta(\hat{\pi}, \mathcal{U}) = V(\hat{\pi}, \mathcal{S}^{\text{eval}}) - \bar{V}(\mathcal{U}, \mathcal{S}^{\mathcal{U}})$, where $\mathcal{S}^{\text{eval}}$ is the evaluation data for the targeted policy and $\mathcal{S}^{\mathcal{U}}$ is the evaluation data for the uniform policy group. In addition to the overall difference between the two policies, these differences can also be computed for the popular hours. In order to do this comparison for call time j , the covariate space is split in a way such that, for the subset of covariates, the best call time is j . This subset and the population objects of interest are defined below. Appendix D explains on-policy evaluation for the simplified setting. We use $R = \{x \in \mathcal{X} : \hat{\pi}(x, \cdot) = j\}$ to denote the subset. Moreover, difference in the value of two policies (δ) is defined below.

$$\delta(\hat{\pi}, \mathcal{U})^j = \mathbb{E}_{x \in R}[\mathbb{E}_Y[(Y_t(\hat{\pi}(\cdot)))] - \mathbb{E}_{x \in R}[\mathbb{E}_Y[\mathbb{E}_{J \sim \mathcal{U}}[\sum_j (Y_t(\pi^j) \mathbb{1}(J = j))]]]$$

Note that on-policy estimation is straightforward: It simply entails taking a sample mean on a dataset where the relevant policy was applied. For reference, we define $\hat{V}^{\text{On}}(\mathcal{S}) = \frac{\sum_{i \in \mathcal{S}} Y_{it}^{\text{obs}}}{|\mathcal{S}|}$.

6 Results: Estimating, Evaluating, and Deploying Targeted Policies

In this section, we present our main results employing on- and off-policy evaluation. Our first step was to collect the data using uniform randomization. This uniform randomization data was then used to estimate $\hat{\pi}$. We present the $\hat{\pi}$ that was estimated using the uniform randomization data from weeks 1, 2, and 3 and call it $\hat{\pi}_D$.

The machine learning model used for this analysis is the LASSO model (Equation 2). This model along with the datasets were used to predict the probability of pickup for every farmer for each of the 91 call times ($\hat{\mu}$). We do not penalize the coefficients on the treatment dummies. The optimal penalty (λ_0) for all other coefficients is chosen to minimize cross-validation error. Figure 3 (a) shows the Mean Squared Error (MSE) corresponding to different values of λ .¹³

To assess the out-of-sample prediction performance of the LASSO model, the data can be divided K folds for cross-fitting. The cross-fitting process is discussed in section 5. Panel (b) in Figure 3 shows the Receiver Operating Characteristics (ROC) plot for the binary pickup (out of sample cross-fitted data). The Area Under the Curve (AUC) is 0.683. This suggests that the LASSO model would be able to correctly predict whether a user will answer the call or not in about 68.3% of the cases. Panel (c) shows the distribution of predicted responses by the class of the binary pickup variable. The predicted response for farmers with actual pickup as 1 is higher when compared to those who did not pick up the calls. Lastly, Panel (d) in Figure 3 shows the calibration plot between the predicted and true pickup. The calibration plots done separately by farmer covariates are shown in Table

¹³Panel (a) in Figure 3 shows the output of the cross-validation exercise used to choose the optimal λ . It displays two special values of λ with vertical bars in the graph. The λ corresponding to the vertical bar on the left minimizes the cross-validated error. The λ corresponding to the vertical bar on the right provides the highest penalized model such that the cross-validated MSE is within one standard error of the lowest λ . We use the λ that minimizes the cross-validated error. λ_{1SE} is also reported for extreme cases when too many variables get dropped from the specification. However, we do not have such an extreme situation for our estimation, and we use $\lambda_{\min}$.

A7. We provide the plots by farmer gender and by farmer land size in this appendix.

Following the steps in Algorithm 1 on the uniform randomized data for weeks 1, 2, and 3, off-policy evaluation can be used to estimate the value of $\hat{\pi}_D$ counterfactually. The means and standard errors corresponding to off-policy estimator $\hat{V}^{\text{Off}}(\hat{\pi}_D, \mathcal{S}^{\text{eval}})$ defined in Section 5.2 can be used to estimate the value of targeted policy counterfactually. Moreover, the means and standard errors of the estimator defined in Section 5.3 can be used to estimate the value of the uniform policy. Figure 4 shows that there are substantial gains (2.6 pp or 8%) to calling farmers according to $\hat{\pi}_D$ over the baseline of uniform policy. To provide a sense of the gains for different optimal call times, Figure 4 also reports the off-policy gains for the 12 most popular call times (which comprise 91% of the targeted policy sample.)

Next, we explore whether there are gains from adding additional covariates to the model, incorporating farmer duration, start month, and mean past engagement as covariates.¹⁴ To examine the prediction accuracy of this model, we repeat the same steps as above and compute the ROC curve and AUC. The AUC is 0.681 when we incorporate these additional measures from the historical data. The targeted policy using this model produces gains close to 8%, which is comparable to the gains observed using $\hat{\pi}_D$.

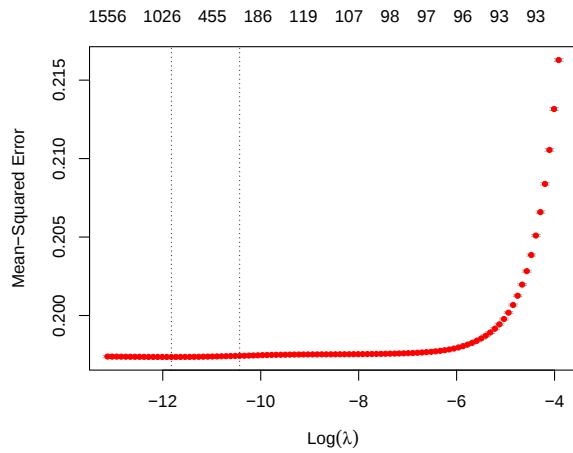
6.1 Deploying and Evaluating $\hat{\pi}$

As part of the experiment, $\hat{\pi}$ was estimated, deployed, and evaluated; we iteratively updated $\hat{\pi}$ as we collected additional uniform data. We also refined some details of our approach. For example, in week 2, we used 21 time slots (morning, afternoon, and evening) over seven days, subsequently shifting to 91 slots (13 hour-slot per day), and, in other weeks, holidays or other real-world constraints affected the set of days in which we could implement policies. These details are described in Appendix E.

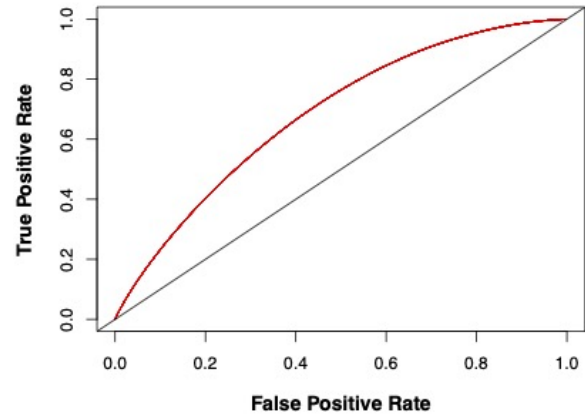
We now present results from the final week of our study, in which we conduct an

¹⁴The details on the matrix completion exercise are presented in Appendix F.3.

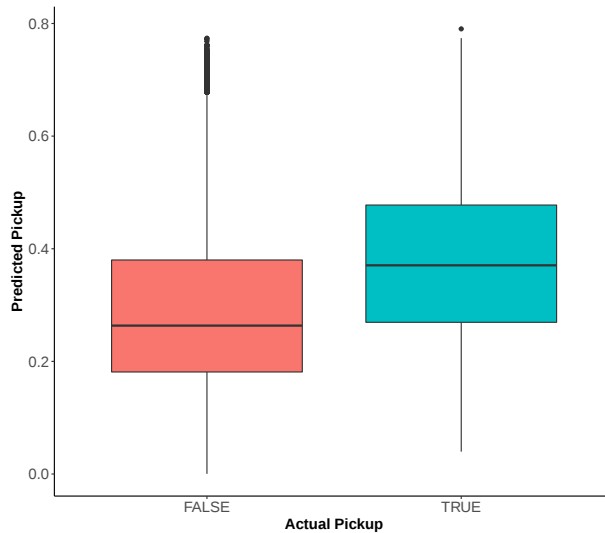
Figure 3: Out-of-Sample Predictive Performance



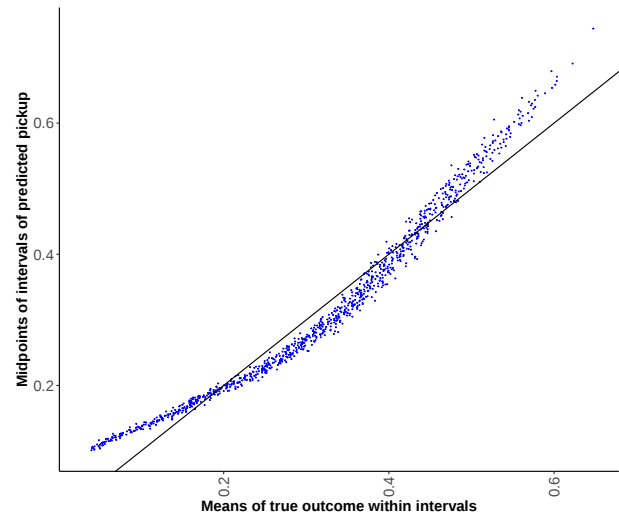
a. Cross-Validation Error



b. ROC curve



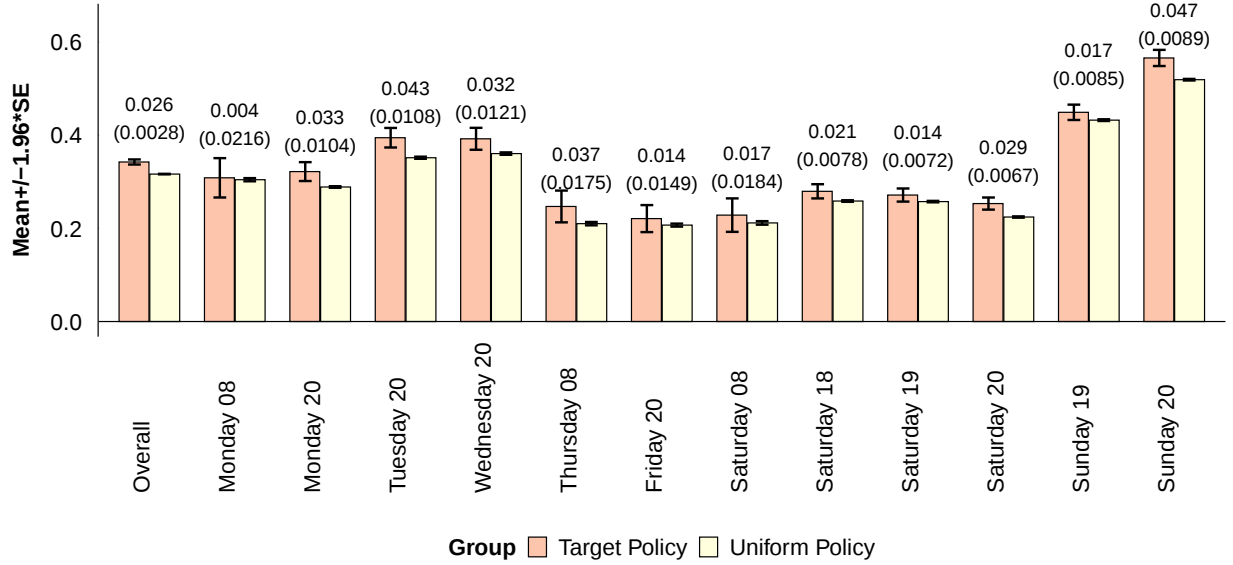
c. Predicted Response



d. Calibration Plot

Notes: (a) shows the MSE corresponding to the different regularization parameters (λ) used in our analysis. (b) shows the ROC curve for the out-of-sample prediction using cross-fitted data. The AUC is 0.683. (c) shows the distribution of predicted response by the binary true pickup in the data. (d) shows the calibration plot for true and predicted pickup using cross-fitted data.

Figure 4: Off-Policy Evaluation: $\hat{\pi}_D$ and Uniform



Notes: Sample includes data from uniform randomization in weeks 1, 2, and 3. Targeted policy estimates are computed using $\hat{V}(\hat{\pi}_D, \mathcal{S}^{\text{eval}})$ (Section 5.2). Uniform policy estimates are computed using the sample mean (Section 5.3). The estimate for the differences between $\hat{V}(\hat{\pi}_D, \mathcal{S}^{\text{eval}})$ and $V(\mathcal{U}, \cdot)$ is displayed at the top of the bars and the standard error of this difference is in parentheses.

on-policy evaluation of $\hat{\pi}_D$.¹⁵ In the final week, we randomly assigned farmers into two groups. The first group was assigned targeted policy $\hat{\pi}_D$, and the second group was called according to the uniform randomization. We compare the average pickup for the two groups in Table 3. Table 3 shows the difference between the sample mean of farmers who received calls according to $\hat{\pi}_D$ and those that got called according to the uniform policy. The differences are shown for the overall sample as well for the sample of female and male farmers. Surprisingly, there is only a .4 percentage point increase in overall pickup relative to the control group (30.9 pp); this gain is much smaller than the 2.6 percentage point gains predicted by off-policy evaluation (Figure 4).

Our finding of lower-than-expected performance is in line with the recent work discussing distribution or dataset shifts in the test data (Kuang et al., 2020; Rabanser et al., 2019). In the next two subsections, we seek to understand better why the on-policy per-

¹⁵Appendix C provides details on the intermediate policies deployed during the experiment in weeks 2, 4, and 5, respectively.

formance was lower than expected and propose approaches to improve the robustness of estimated policies.

Table 3: On-Policy Evaluation in Week 6

Data Collection-policy	$\hat{\pi}_D$	Uniform	Difference
Outcome Variable			
	All Farmers		
First Call	0.313 [0.001]	0.309 [0.001]	0.004 [0.001]
	Female Farmers		
First Call	0.2971 [0.0023]	0.2901 [0.0014]	0.0069 [0.0027]
	Male Farmers		
First Call	0.3167 [0.0011]	0.3135 [0.0007]	0.0032 [0.0013]
N	227,386	587,608	

Notes: On-policy evaluation estimates use sample means for each of the two data collection mechanisms from Week 6.

6.2 Transportability of Targeted Policy to Future Weeks

A robust, targeted policy should work not just out-of-sample for the time period it is estimated but also out-of-sample in subsequent weeks when it is prospectively deployed. However, if the distribution of the outcome variable is subject to systematic variation over time, the policy designer may face a trade-off: Placing greater weight on more recent data may focus on behavior that is closer to what will be observed in the near future but at a cost of using less of the information available from earlier weeks.

In this section, we take advantage of the fact that we have a number of weeks of random uniform data, which allows us to vary the number of historic weeks used when estimating a policy and explore the robustness of the policy in subsequent weeks of data (e.g., out-of-sample). This relates to the ideas of stability and transportability of targeted policies (Hitsch and Misra, 2018). Figure 5 provides an overview of our approach: We consider four scenarios in which older data is down-weighted in our policy estimation. This is possible in

a setting such as agriculture where seasonality considerations may affect farmers' workload and task distribution (Gill et al., 1991; Vemireddy and Pingali, 2021). The real world may also be subject to other events, such as festivals or cricket matches, which alter time-use preferences and hence affect actual gains of implemented policies.

Using our notation for training data $\mathcal{S}$ and evaluation data $\mathcal{S}^{\text{eval}}$, we see that the gains from off-policy evaluations do indeed depend on the degree to which the estimation weighs more recent vs. less recent data. In Figure 5, "Train" indicates the weeks used to train the policy, and "Test" indicates the week used in the off-policy evaluation. The weights refer to the relative weighting in the LASSO model.¹⁶ The bars give the pickup rate for the targeted policy and the uniform policy. The estimate for the differences in these values is provided at the top of the bars. We find substantial gains (5-8%) for the value of the targeted policy counterfactually estimated over the uniform policy for scenarios where the training data weeks closest to the test week are weighed the highest in the machine learning model. This is likely because the most recent weeks of data capture the technology and preference shocks much better than the older weeks (Figure 5).

6.3 Robustness to Shocks

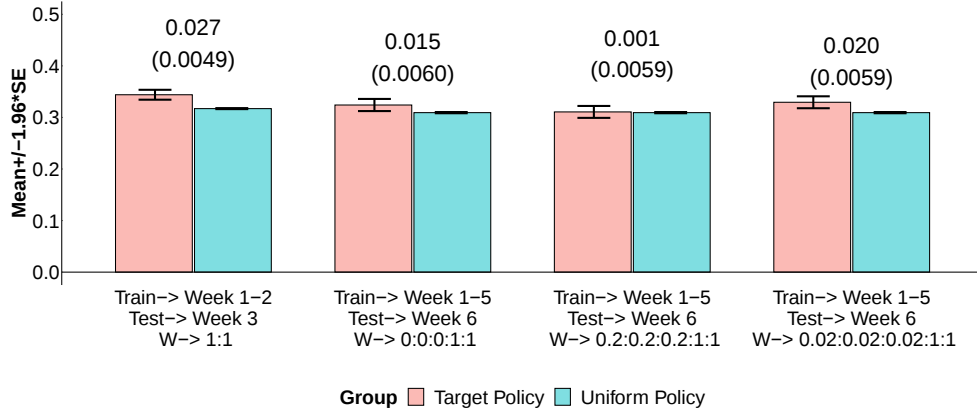
The overall objective of this project is to maximize farmer engagement with the service. The policy was estimated and deployed for the timing of the first call attempt each week to reach a farmer. In fact, if a farmer does not answer the phone, the dialing software automatically attempts a second call 24 hours later, and if that is not answered, a third and final attempt is made 24 hours later (Figure A8). While, in principle, the second and third calls could have been optimized, the nonprofit did not want to take on such

¹⁶The objective function for the weighted LASSO is provided below

$$\sum_t -W_t l(\theta, X_t, J) + \lambda \|\beta\| + \lambda \|\gamma\|$$

where W_t is the weight associated with the training sample in week t and $t \in \{1, 2, \dots, 5\}$.

Figure 5: Off-Policy Evaluation of Targeted Policies Using Subsequent Weeks' Evaluation Data



Notes: We use varying subsets of uniform randomization data over the six weeks of the experiment for the estimation (“Train”) and evaluation (“Test”) of targeted policies in this figure counterfactually. In scenario 1, the targeted policy is estimated with data from weeks 1 and 2, while the evaluation data comes from week 3. The entries for targeted policy in the paper are means and standard errors of $\hat{V}^{\text{Off}}(\hat{\pi}, \mathcal{S}_{\mathcal{U},3})$, as described in Section 5.2. The entries for uniform policy are on-policy sample averages of outcomes for units assigned to uniform policy in the Test week; for scenario 1, this is week 3. Similarly, in scenarios 2, 3, and 4, the targeted policy is estimated using data from weeks 1, 2., and 5 with varying weights.

operational complexity. We can, however, use these follow-up calls to provide a sense of how follow-up calls can help mitigate some of the shocks to the first call.

This section explores the efficacy of the rule of thumb of calling farmers 24 hours later and proposes a method for counterfactual evaluation of all three calls. This section illustrates how the targeted policy can improve the overall engagement of farmers relative to the uniform policy. We provide the first evidence for this mechanism using the overall pickup over the three calls for our on-policy evaluation of $\hat{\pi}_D$ (Table 4). We find the impact of targeted policy is higher when examining overall engagement by adding the pickup over call attempts 1, 2, and 3 than only examining the engagement over the first call. If there are shocks to the first call, follow-up calls can mitigate some of the cost.

Next, we provide evidence of the impact of the first-call targeted policy on overall engagement using off-policy evaluations. Figure A9 provides evidence that the benefit of targeting the first calls has benefits for overall farmer engagement. For this exercise, $\hat{\pi}$ is estimated on uniformly randomized data for the first calls in week 1 and week 2. We

Table 4: Incorporating Follow-Up Calls to Mitigate Shocks: On-Policy Evaluation

Data Collection-policy	$\hat{\pi}_D$	Uniform	Difference
Outcome Variable			
	All Farmers		
Call 1,2,3	0.552 [0.001]	0.544 [0.001]	0.008 [0.001]
	Female Farmers		
Call 1,2,3	0.5249 [0.0025]	0.5133 [0.0016]	0.0116 [0.0029]
	Male Farmers		
Call 1,2,3	0.5577 [0.0011]	0.5500 [0.0007]	0.0077 [0.0014]
N	227,386	587,608	

Notes: This table shows the on-policy evaluation results for targeted policy $\hat{\pi}_D$ in week 6. The farmers are randomized between two groups. Group A gets called according to $\hat{\pi}_D$, and Group B gets called according to uniform randomization. Sample means are used to estimate the value of $V(\hat{\pi}_D, \mathcal{S}^{\text{eval}})$ and $V(\bar{\mathcal{U}})$.

evaluate the first-call policy not just on the first call but also on the follow-up calls. The sample considered for evaluation for the targeted policy corresponds to a subset of people in the evaluation set whose actual assignment matched the targeted policy according to the first call targeted policy. The first two bars in Figure A9 show the gains of targeting the first call on the first call. Next, we add the pickup over the first and second calls for the evaluation sample. We continue to see that the pickup over the first two calls for the targeted policy group is higher than the pickup for the first and second calls for the uniform group (2 pp (SE=0.36)). Next, we add pickup over calls 1, 2, and 3. We see that the targeted policy engagement is higher than the uniform policy engagement (1.6 pp (SE=0.35)).

Furthermore, we provide some evidence of the benefit of targeting second calls. We estimate a targeted policy for the second call using data on whether farmers answered a second call. Among the farmers who did not pick up the first call in week 1 and week 2, we

estimate a second call targeted policy. We evaluate the second call targeted policy on the second call data but adjust for the propensity of the pickup in the first call. Here are the steps to this evaluation: We predict pickup for first calls under the first call counterfactual policy (first call targeted policy). Next, we reweigh all of the second call data in the evaluation step according to the inverse propensity weights. This adjusts for the fact that targeting the first call changes the set of people left over for the second call. The results are provided in Table 5. We observe a 4.3 pp gain from targeting the second calls relative to the baseline mean of the uniform group for the inverse propensity-weighted pickup of 29.8%.

Table 5: Incorporating Follow-Up Calls to Mitigate Shocks

	Targeted Policy	Uniform	Difference
Impact of Policy Targeting Second Call on Pick-up of Second Call			
Second Call	0.341 [0.0061]	0.298 [0.0006]	0.043 [0.0061]

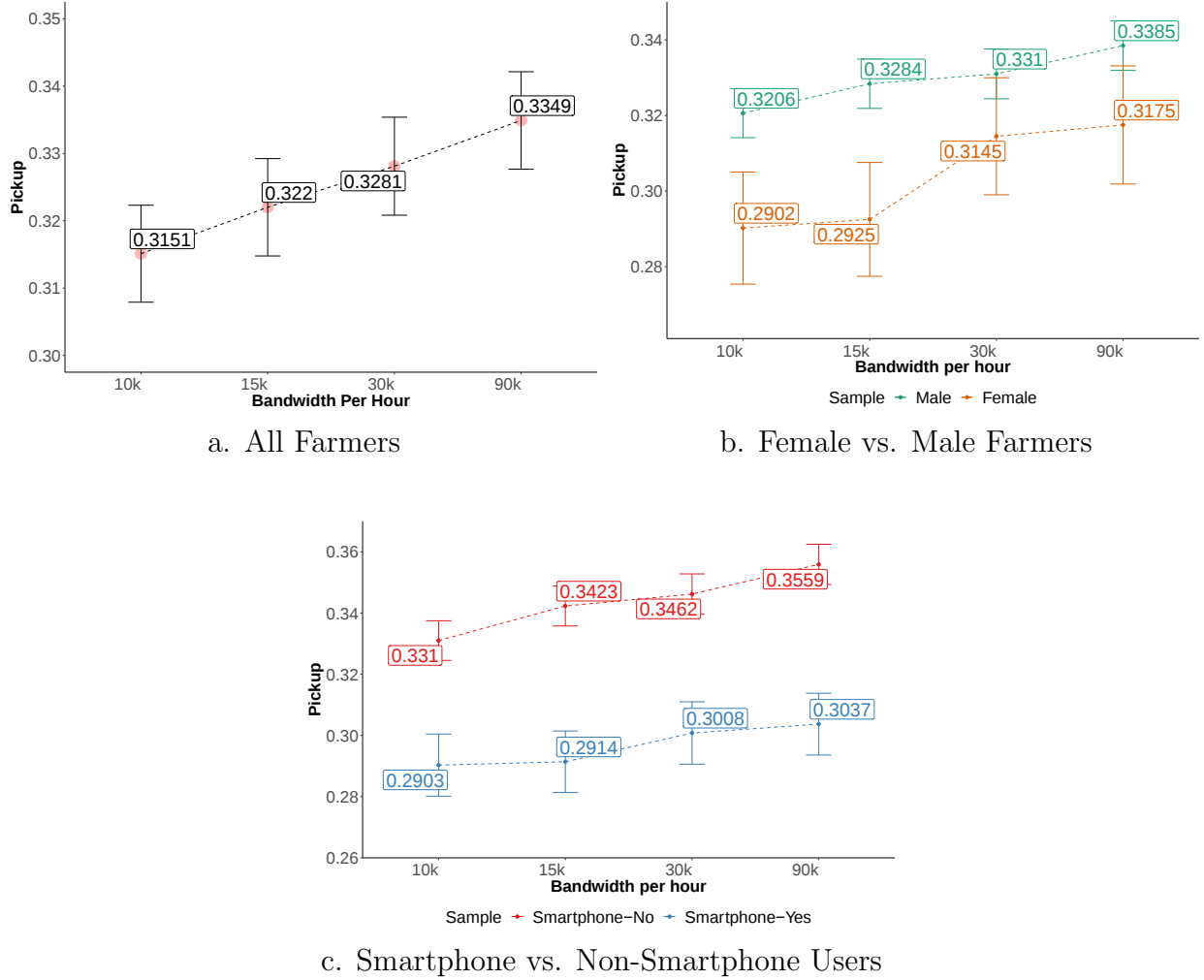
Notes: The training and evaluation data for estimating the value of targeted policy counterfactually and the uniform policy uses the uniform randomization data for week 1 and week 2. The training data for estimating the second call target policy consists of $N_{1,\mathcal{U}}^{\text{second}} = 630,383$, $N_{2,\mathcal{U}}^{\text{second}} = 431,290$. Note the estimate for the $\hat{V}^{\text{Off}}(\hat{\pi}, S_{\mathcal{U},1,2}^{\text{second}})$ is done for the inverse propensity weighted second call pick-up instead of the non-adjusted pick-up.

6.4 Bandwidth Constraints and Equity-Efficiency Trade-Off

We begin this section by first estimating the benefits of expanding the bandwidth of the agricultural advisory service. We do this computation counterfactually using the uniform randomization data for weeks 1, 2, and 3 of the experiment. We estimate policies by modifying Algorithm 1, in particular by altering the bandwidth limits in the constrained optimization step. We start with a bandwidth limit of 10,000 farmers per treatment arm and then gradually expand it to 90,000 farmers. We conduct a counterfactual analysis for a scenario that differs from the implemented experiment, where we deployed farmers to

the target and uniform groups. Instead, here we conduct counterfactual analysis for the scenarios where the entire sample of farmers (600,000 farmers) would be called according to the targeted policy, and there is no uniform call time group.

Figure 6: Relaxing the Budget Constraint



Notes: The off-policy evaluation is conducted using data collected in the uniform randomization arm during weeks 1,2 and 3 of the experiment. $\hat{\pi}$ is estimated using Algorithm 1, with modified constraints. Panel (a) shows the gains in the pickup for the sample as the bandwidth is expanded from calling 10,000 farmers in an hour to 90,000 farmers. Panel (b) and (c) show the gains by subgroups.

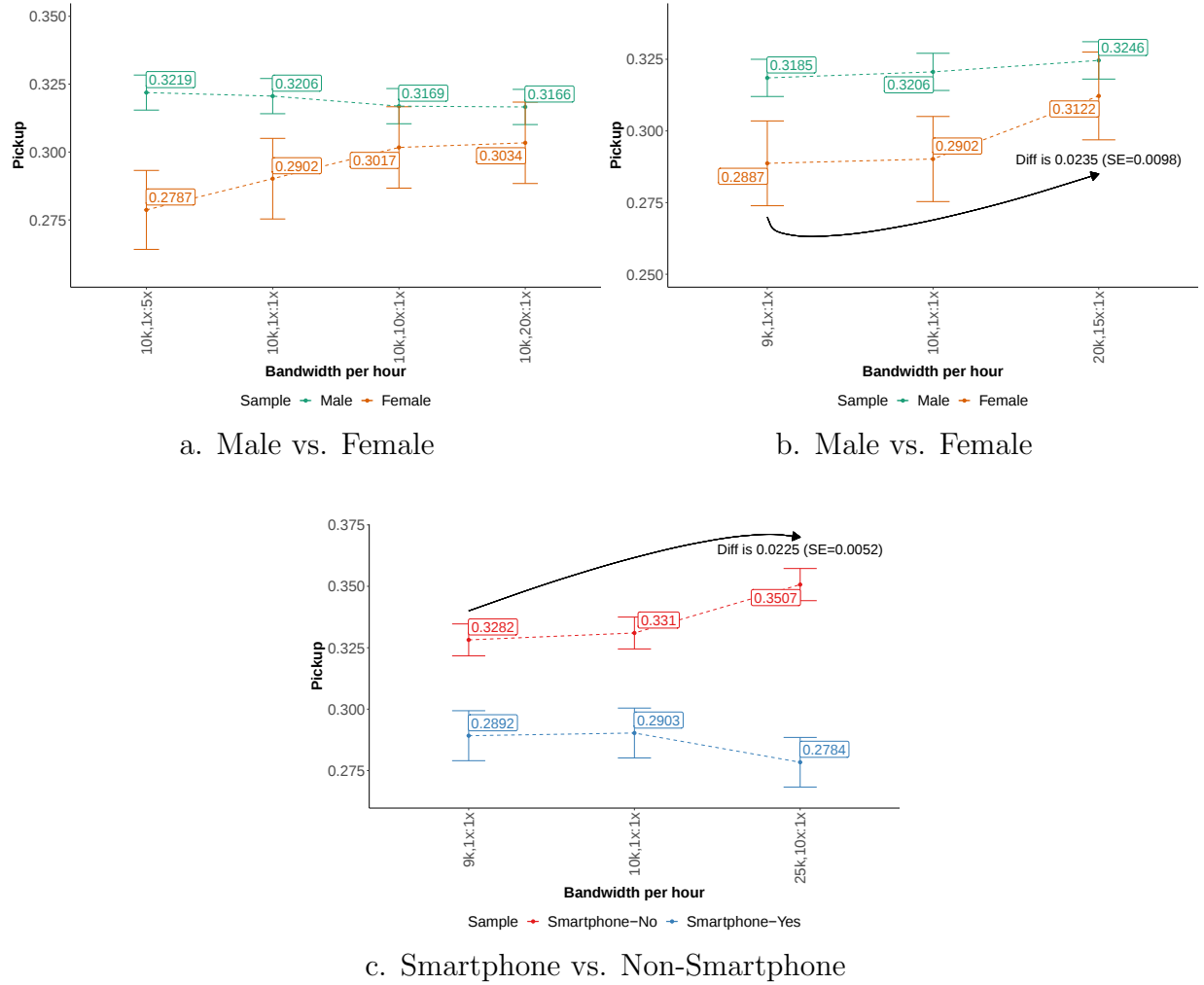
We find substantial counterfactual gains in pickup for the targeted policy by relaxing the budget constraint. As the budget constraint is expanded from 10,000 to 90,000, the estimate for $\hat{V}^{\text{Off}}(\hat{\pi}, \mathcal{S}_U)$ increases from 0.3151 (SE=0.0037) to 0.3349 (SE=0.0037). The

increase for all farmers is shown in Panel (a) in Figure 6. In Panel (b) in Figure 6, we show the changes for female and male farmers. Finally, in Panel (c), we illustrate the increase separately for smartphone and non-smartphone users. Expanding bandwidth improves call pickup rates for all subgroups. However, since all types of farmers are weighed equally in the optimization, we need large increases in bandwidth to substantially improve the pickup for vulnerable groups such as female farmers and non-smartphone users.

An alternative way to improve the engagement of vulnerable groups with the agricultural advisory service is to attach higher weights to the vulnerable groups in the constrained optimization step. In fact, there is a substantial overlap between the top pickup call times for male and female farmers, as seen in Panel (b) and (c) of Figure 2. This suggests the possibility that an algorithm that is designed to maximize overall pickup rates could assign high-value slots to men; there may be a trade-off in employing a welfare-sensitive algorithm that puts a greater weight on reaching female farmers. This motivates the next analysis, where we vary the weights by farmers' characteristics. We first draw a random sample of 600,000 farmers from the datasets collected using uniform randomization in weeks 1, 2, and 3 of the experiment. We estimate $\hat{\mu}$ as before, modifying the constrained optimization step of Algorithm 1 to add weights as follows. Let G be the group indicator random variable with support in $\mathcal{G}$ and g be the realization of this random variable. X follows a mixture distribution $X|G = g \sim F_g$. The population objective function in this case is $V(\pi, w, \mathcal{S}^{eval}) = \sum_g w(g) \mathbb{E}_{X \sim F_g} [\mathbb{E}_Y Y_t[(\pi(X))]]$, where w is the weight function $w : \mathcal{G} \rightarrow [0, 1]$ s.t. $\sum_{g \in \mathcal{G}} w(g) = 1$.

For any given set of weights, we estimate an optimal policy for the corresponding weighted objective function. Panel (a) in Figure 7 illustrates three scenarios for varying weighing schemes for male and female farmers. Here, we keep the bandwidth constant at 10,000 calls per hour and vary weights only. In the first case, males and females are weighed equally. In the second case, females are given 10 times the weight as males. Finally, in the last case, females are given 20 times the weight as males. The constraints are tight and

Figure 7: Weighing Subgroups Differently in Constrained Optimization



Notes: The uniform randomization data for weeks 1,2 and 3 of the experiment are used for these off-policy evaluations. Panel (a) shows the scenarios where the bandwidth is kept constant at 10,000 farmers, and we vary the weights across the three scenarios between male and female farmers. The value of targeted policy is denoted where we attach weights for subgroups is $V(\pi, w, \cdot) = \sum_g w(g) E_{X \sim F_g} E(Y(\pi(X)))$ as defined in Section 6.4. Panel (b) shows an expansion in bandwidth from 9,000 to 20,000 but attaches a higher weight to females than males. Finally, Panel (c) shows gains as we expand the bandwidth but put higher weight on non-smartphone users.

binding in the optimization, which means increasing the relative weight on females moves more female farmers into near-optimal slots. We can also calculate the “real world” sense of trade-offs. Since 18% of our sample is female (sample size is 600,000 for this exercise), moving from scenario 1 to scenario 3 enables us to reach 1,426 more women at the expense of 1,968 fewer men.

Moreover, we explore the scenarios where the constraints are relaxed, but while doing that, we put higher weights on the vulnerable groups. For instance, in Panel (b) we expand the bandwidth from 9,000 to 20,000 farmers per hour. However, we weigh the females 15 times higher than the males in the constrained optimization step. This counterfactual analysis shows that female pickup increases by 2.35 pp (SE=0.98). Similarly, in Panel (c), we expand the bandwidth from 9,000 to 25,000 but weigh the non-smartphone users more than the smartphone users. We find that this improves the pickup of non-smartphone users by 2.25 pp (SE=0.52). This translates into reaching an additional 8,505 non-smartphone users at the cost of reaching 2,398 fewer smartphone users. Attaching a higher weight to non-smartphone users as more slots are added could be beneficial as they are likely to benefit more from this advisory service since they might have limited access to other sources of information.

7 Discussion

This project develops, implements, and evaluates a recommendation system designed to increase engagement in an agricultural advisory service in rural areas in Odisha, India. Our research design, which tests estimated targeted policies against a control group policy of uniform random assignment, allows us to perform both on- and off-policy evaluation. We find that personalizing call timings can increase engagement meaningfully, with estimated off-policy gains as high as 8 percent. These gains come at virtually no programmatic cost: While estimating the targeted policy requires LASSO regression and a complex integer

programming problem, the computational cost is negligible, even in a setting with almost one million weekly users.

Our paper serves as an important proof of concept, showing that advanced recommendation systems typically employed in apps or web settings can work with technology as simple as automated voice telephone calls, even in data-poor environments.

Our paper also highlights some of the key trade-offs between on- and off-policy evaluation. Off-policy evaluation has the significant virtue of flexibility. Because our control group was called at randomly assigned times, we were able to tackle a range of important questions. We were able to precisely quantify the equity-efficiency trade-off that the organization would face if it placed higher welfare weights on reaching women farmers, and we show that recent data is more predictive of future farmer behavior than data just a few weeks older.

In contrast, on-policy evaluation has significantly more statistical power. It also helps us to identify the possible existence of technology or preference shocks that degrade the performance of targeted policies out of sample. We show several ways of modifying the policy to account for the possibility of such shocks. First, attaching a higher weight to training samples closest to the test data week provides robust policies for future weeks. This is because agriculture is a seasonal activity, and the weeks closest to the test data capture the likelihood of such shocks better than older weeks. We also find that the nonprofit's default policy of scheduling follow-up calls 24 hours later and, if necessary, 48 hours later can reduce the apparent effect of shocks. Gains could be even larger if the timing of the follow-up calls were personalized.

Engagement and learning new information are the initial key steps in improving farmers' productivity and welfare. As shown in [Fabregas et al. \(2019\)](#), sharing agricultural information using digital technologies leads to a 22% increase in farmers' odds of adopting recommended agricultural inputs and a 4% increase in farmers' yield. While our experiment focuses on measuring the impact on farmers' engagement with the digital service, future

work will evaluate the impact of customizing content on farmers' practices and welfare.

References

- Agrawal, K., Athey, S., Kanodia, A., and Palikot, E. (2022). Personalized recommendations in edtech: Evidence from a randomized controlled trial. *arXiv preprint arXiv:2208.13940*.
- Aher, S. B. and Lobo, L. (2013). Combination of machine learning algorithms for recommendation of courses in E-Learning System based on historical data. *Knowledge-Based Systems*, 51:1–14.
- Aker, J. C., Ghosh, I., and Burrell, J. (2016). The promise (and pitfalls) of ICT for agriculture initiatives. *Agricultural Economics*, 47(S1):35–48.
- Anderson, J. R. and Feder, G. (2004). Agricultural extension: Good intentions and hard realities. *The World Bank Research Observer*, 19(1):41–60.
- Araya, R., Menezes, P. R., Claro, H. G., Brandt, L. R., Daley, K. L., Quayle, J., Diez-Canseco, F., Peters, T. J., Cruz, D. V., Toyama, M., et al. (2021). Effect of a digital intervention on depressive symptoms in patients with comorbid hypertension or diabetes in Brazil and Peru: Two randomized clinical trials. *JAMA*, 325(18):1852–1862.
- Ascarza, E. (2018). Retention futility: Targeting high-risk customers might be ineffective. *Journal of Marketing Research*, 55(1):80–98.
- Athey, S., Karlan, D., Palikot, E., and Yuan, Y. (2022). Smiles in profiles: Improving fairness and efficiency using estimates of user preferences in online marketplaces. *arXiv preprint arXiv:2209.01235*.
- Athey, S. and Wager, S. (2021). Policy learning with observational data. *Econometrica*, 89(1):133–161.

- Beretta, E., Santangelo, A., Lepri, B., Vetrò, A., and Martin, J. C. D. (2019). The invisible power of fairness. how machine learning shapes democracy. In *Canadian Conference on Artificial Intelligence*, pages 238–250. Springer.
- Berman, M. and Fenaughty, A. (2005). Technology and managed care: patient benefits of telemedicine in a rural health care network. *Health Economics*, 14(6):559–573.
- Cole, S. and Fernando, A. N. (2021). ‘Mobile’izing agricultural advice: Technology adoption, diffusion, and sustainability. *Economic Journal*, 131(633):192–219.
- Dammert, A. C., Galdo, J., and Galdo, V. (2013). Digital labor-market intermediation and job expectations: Evidence from a field experiment. *Economics Letters*, 120(1):112–116.
- Dammert, A. C., Galdo, J., and Galdo, V. (2015). Integrating mobile phone technologies into labor-market intermediation: A multi-treatment experimental design. *IZA Journal of Labor & Development*, 4(1):1–27.
- Duflo, E. (2017). The economist as plumber. *American Economic Review*, 107(5):1–26.
- Ekeland, A. G., Bowes, A., and Flottorp, S. (2010). Effectiveness of telemedicine: A systematic review of reviews. *International Journal of Medical Informatics*, 79(11):736–771.
- Fabregas, R., Kremer, M., and Schilbach, F. (2019). Realizing the potential of digital development: The case of agricultural advice. *Science*, 366(6471):eaay3038.
- Gabel, S. and Timoshenko, A. (2022). Product choice with large assortments: A scalable deep-learning model. *Management Science*, 68(3):1808–1827.
- Gill, G. J. et al. (1991). *Seasonality and Agriculture in the Developing World: A problem of the poor and the powerless*. Cambridge University Press.

- Hitsch, G. J. and Misra, S. (2018). Heterogeneous treatment effects and optimal targeting policy evaluation. Working Paper.
- Hoxby, C. M. (2014). The economics of online postsecondary education: MOOCs, nonselective education, and highly selective education. *American Economic Review*, 104(5):528–33.
- Jambor, T. and Wang, J. (2010). Optimizing multiple objectives in collaborative filtering. In *Proceedings of the fourth ACM conference on Recommender systems*, pages 55–62.
- Knight, H., Jia, R., Ayling, K., Bradbury, K., Baker, K., Chalder, T., Morling, J. R., Durrant, L., Avery, T., Ball, J. K., et al. (2021). Understanding and addressing vaccine hesitancy in the context of COVID-19: Development of a digital intervention. *Public Health*, 201:98–107.
- Kuang, K., Xiong, R., Cui, P., Athey, S., and Li, B. (2020). Stable prediction with model misspecification and agnostic distribution shift. In *Proceedings of the AAAI Conference on Artificial Intelligence*, volume 34, pages 4485–4492.
- Lee, J., Foschini, L., Kumar, S., Juusola, J., Liska, J., Mercer, M., Tai, C., Buzzetti, R., Clement, M., Cos, X., et al. (2021). Digital intervention increases influenza vaccination rates for people with diabetes in a decentralized randomized trial. *NPJ digital medicine*, 4(1):1–8.
- Mazumder, R., Hastie, T., and Tibshirani, R. (2010). Spectral regularization algorithms for learning large incomplete matrices. *The Journal of Machine Learning Research*, 11:2287–2322.
- Owusu, V., Donkor, E., and Owusu-Sekyere, E. (2018). Accounting for the gender technology gap amongst smallholder rice farmers in northern Ghana. *Journal of Agricultural Economics*, 69(2):439–457.

- Portugal, I., Alencar, P., and Cowan, D. (2018). The use of machine learning algorithms in recommender systems: A systematic review. *Expert Systems with Applications*, 97:205–227.
- Quisumbing, A. R. and Pandolfelli, L. (2010). Promising approaches to address the needs of poor female farmers: Resources, constraints, and interventions. *World Development*, 38(4):581–592.
- Rabanser, S., Günnemann, S., and Lipton, Z. (2019). Failing loudly: An empirical study of methods for detecting dataset shift. *Advances in Neural Information Processing Systems*, 32.
- Rambachan, A., Kleinberg, J., Ludwig, J., and Mullainathan, S. (2020). An economic perspective on algorithmic fairness. In *AEA Papers and Proceedings*, volume 110, pages 91–95.
- Rodriguez-Segura, D. (2022). Edtech in developing countries: A review of the evidence. *The World Bank Research Observer*, 37(2):171–203.
- Samin, H. and Azim, T. (2019). Knowledge based recommender system for academia using machine learning: A case study on higher education landscape of Pakistan. *IEEE Access*, 7:67081–67093.
- Seymen, S., Abdollahpouri, H., and Malthouse, E. C. (2021). A constrained optimization approach for calibrated recommendations. In *Fifteenth ACM Conference on Recommender Systems*, pages 607–612, Amsterdam Netherlands. ACM.
- Simester, D., Timoshenko, A., and Zoumpoulis, S. I. (2020). Efficiently evaluating targeting policies: Improving on champion vs. challenger experiments. *Management Science*, 66(8):3412–3424.

- Steck, H. (2018). Calibrated recommendations. In *Proceedings of the 12th ACM Conference on Recommender Systems*, pages 154–162, Vancouver British Columbia Canada. ACM.
- Swanson, B. E. and Davis, K. (2014). Status of agricultural extension and rural advisory services worldwide. *GFRAS Summary Report*.
- Vemireddy, V. and Pingali, P. L. (2021). Seasonal time trade-offs and nutrition outcomes for women in agriculture: Evidence from rural India. *Food policy*, 101:102074.
- Voss, C., Schwartz, J., Daniels, J., Kline, A., Haber, N., Washington, P., Tariq, Q., Robinson, T. N., Desai, M., Phillips, J. M., et al. (2019). Effect of wearable digital intervention for improving socialization in children with autism spectrum disorder: A randomized clinical trial. *JAMA pediatrics*, 173(5):446–454.
- Xiao, L., Min, Z., Yongfeng, Z., Zhaoquan, G., Yiqun, L., and Shaoping, M. (2017). Fairness-aware group recommendation with Pareto-efficiency. In *Proceedings of the Eleventh ACM Conference on Recommender Systems*, pages 107–115.
- Yang, J., Eckles, D., Dhillon, P., and Aral, S. (2020). Targeting for long-term outcomes. *arXiv preprint arXiv:2010.15835*.
- Yang, K. and Stoyanovich, J. (2016). Measuring fairness in ranked outputs. *arXiv:1610.08559 [cs]*.
- Zehlike, M., Bonchi, F., Castillo, C., Hajian, S., Megahed, M., and Baeza-Yates, R. (2017). FA*IR: A fair top-k ranking algorithm. In *Proceedings of the 2017 ACM on Conference on Information and Knowledge Management*, pages 1569–1578. *arXiv:1706.06368 [cs]*.
- Zhou, Z., Athey, S., and Wager, S. (2022). Offline multi-action policy learning: Generalization and optimization. *Operations Research*.

Appendix

A Example Script of Agricultural Advisory Messages

This section provides two examples of agricultural advisory messages that were sent through the extension service to the farmers in our multistage experiment sample.

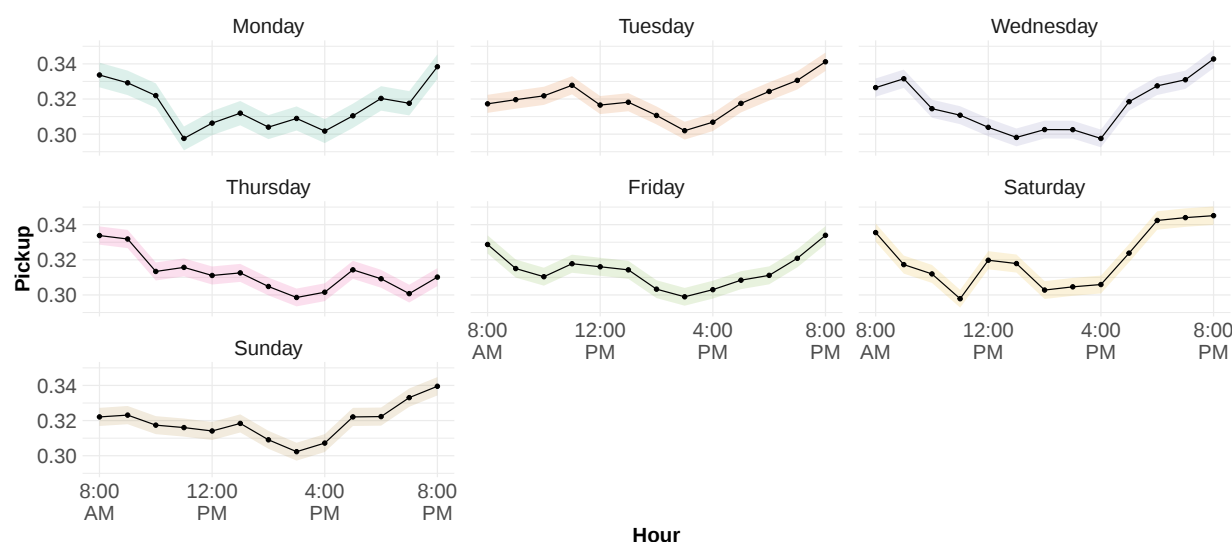
Example 1. Advisory on pest management: *Namaskar. Today we will discuss about neck blast disease and its management in paddy crop. Due to high relative humidity and differential day and night temperature Neck Blast disease incidence can be seen in paddy crops. To manage these diseases, first drain out excess water from the paddy field. Spray Hexaconazole (Contaf Plus/Hexadhan Plus/Trigger Pro) @ 400-ml/acre or Azoxystrobi+ Difenconazole (Amistar Top/ Chemistar /Karishma) @ 200-ml/acre or Tebuconazole + Trifloxystrobin (Nativo) @ 80-gram/acre. Thank you and remember that if you have questions about this advisory or need more information, you can call the hotline number on 155333.*

Example 2. Advisory on basal fertilizer application for transplanting: *Namaskar. Today we will share a few tips for applying basal fertilizers correctly for farmers who are growing HYV paddy. If you have not yet done so, we advise you to complete your transplanting by August 15th. Before applying fertilizers at the time of sowing, you should determine what kind of soil you have. This is because the fertilizers recommended are different for different soil. You should apply 35 kg DAP, 30 kg Potash, and 8 kg Urea per acre during the last puddling. Remember again that you should apply 35 kg DAP, 30 kg Potash, and 8 kg Urea per acre at the time of last puddling. Please remember 1acre=25 guntha. However, you should apply Potash in two equal splits at the basal and PI stages if you have sandy soil. Also, do not forget that zinc deficiency is the most widespread micronutrient disorder in*

paddy, affecting plant growth. For this reason, we suggest that you can apply 10 kg of Zinc Sulphate micronutrient per acre based on a soil test report or if your soil is deficient in zinc. Thank you and remember that if you have questions about this advisory or need more information, you can call the hotline number on 155333, and we will provide this message to your mobile via SMS.

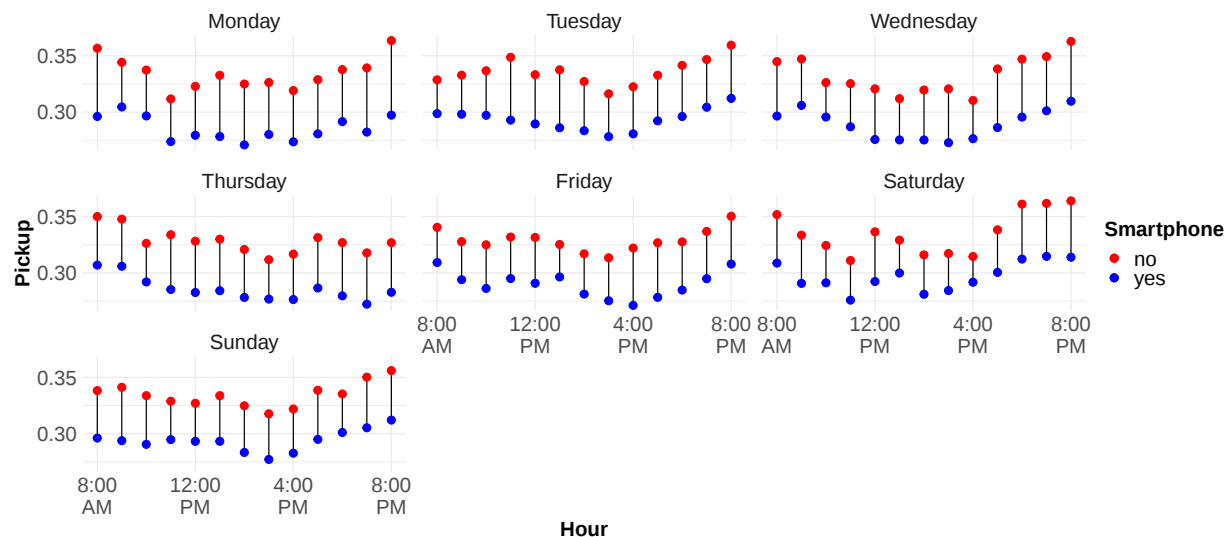
B Additional Tables and Figures

Figure A1: Pickup by 91 Treatment Arms

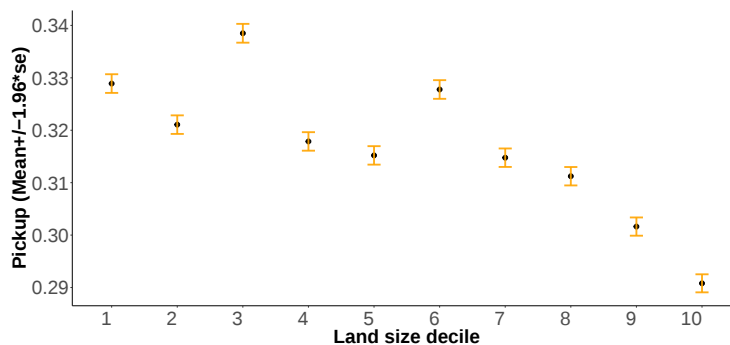


Notes: We pool the data collected using uniform randomization in weeks 1, 2, and 3 of the experiment (see Table 1 for sample size). It illustrates the variation in engagement between the 91 call times, which is a combination of the day of the week and the hour of the day (Mean $\pm$ 1.96SE).

Figure A2: (a) Pickup by time and smartphone, (b) Pickup by the distribution of land size.



a. Variation in Pickup by Smartphone Dummy for 91 Treatment Arms



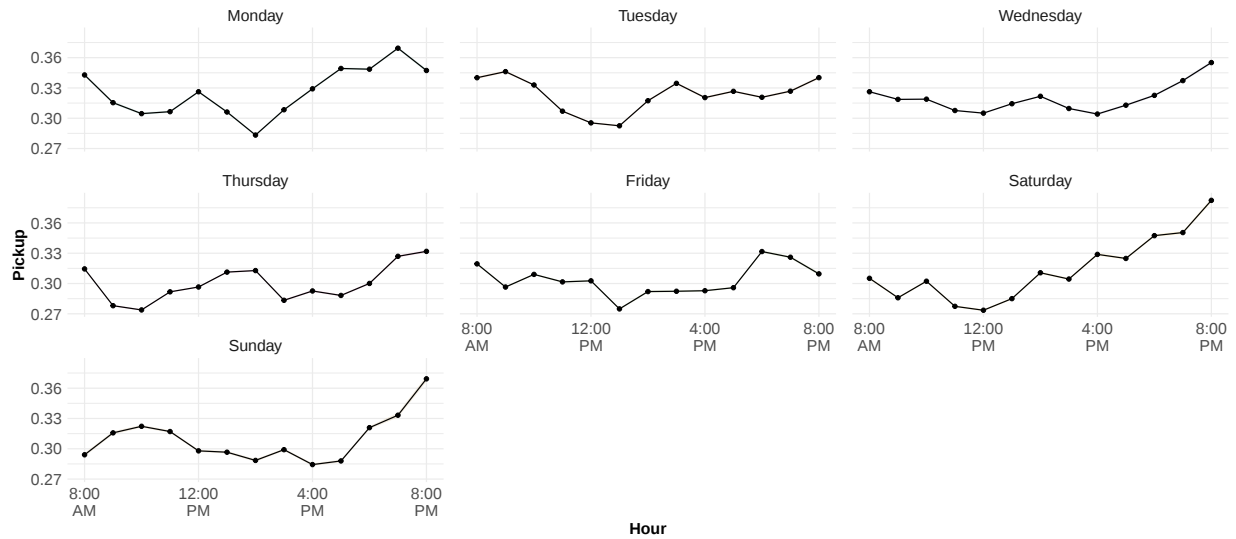
b. Pickup by Land size

Notes: For this figure, we pool the data collected using uniform randomization in weeks 1, 2, and 3 of the experiment. Panel (a) illustrates the smartphone users' and nonusers' gap in engagement with the service over the 91 call times. Panel (b) shows the variation in pickup by land size (Mean +/- 1.96se). We divide the sample of farmers into deciles based on their land size. For each decile, we compute the mean and standard error for pickup.

Table A1: Summary Statistics

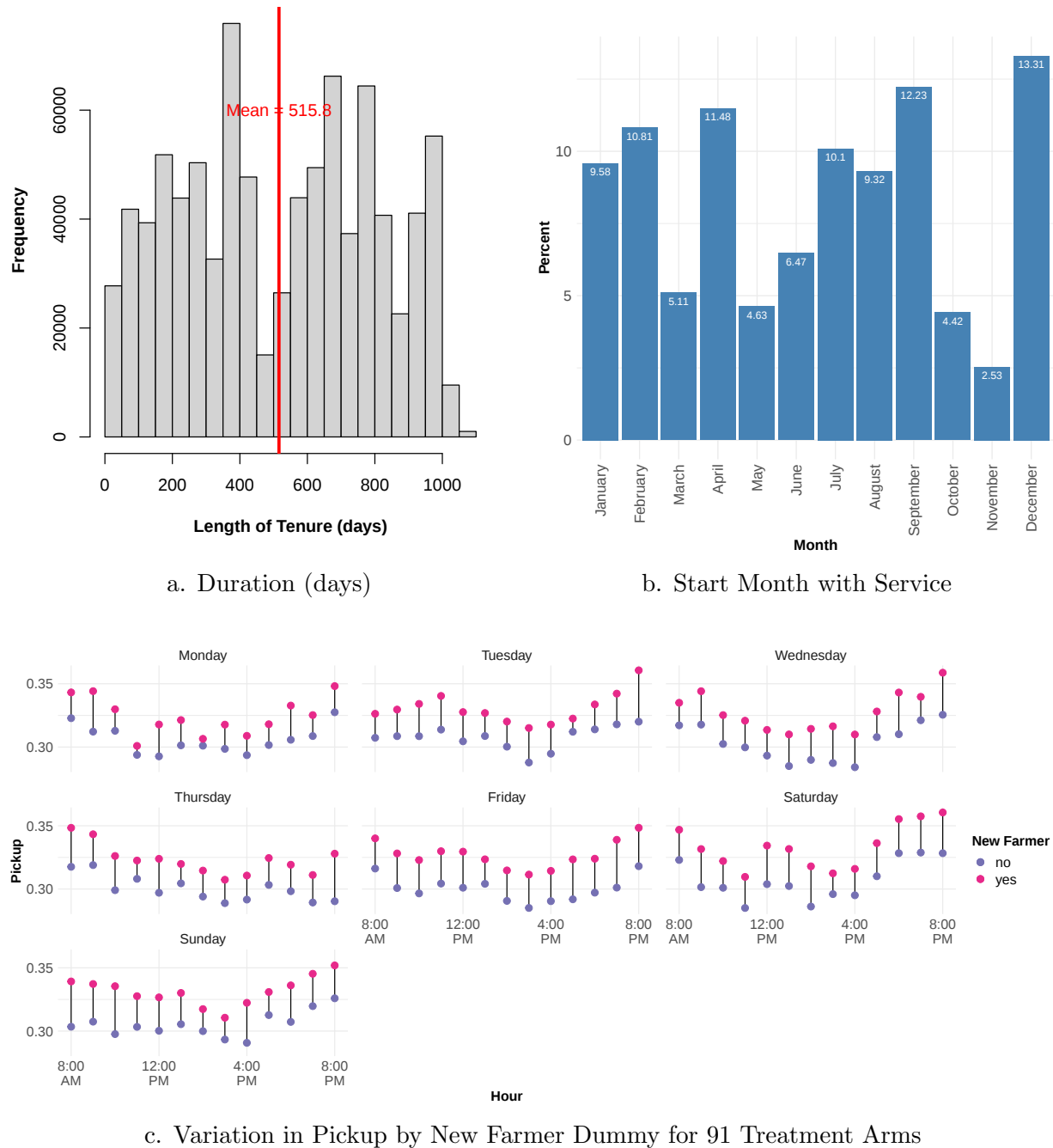
Variable	Mean	Std.	Median
A. Outcome Variable			
Pick-up	0.318	0.283	0.250
B. Covariates			
Female	0.181	0.385	0.000
Smartphone	0.374	0.484	0.000
Irrigation	0.444	0.497	0.000
Landsize	1.286	1.881	0.809
N	884,194		

Notes: We also include the district of residence and historical engagement as covariates for estimating $\hat{\pi}(x_i, \cdot)$.

Figure A3: Pickup by 91 Treatment Arms

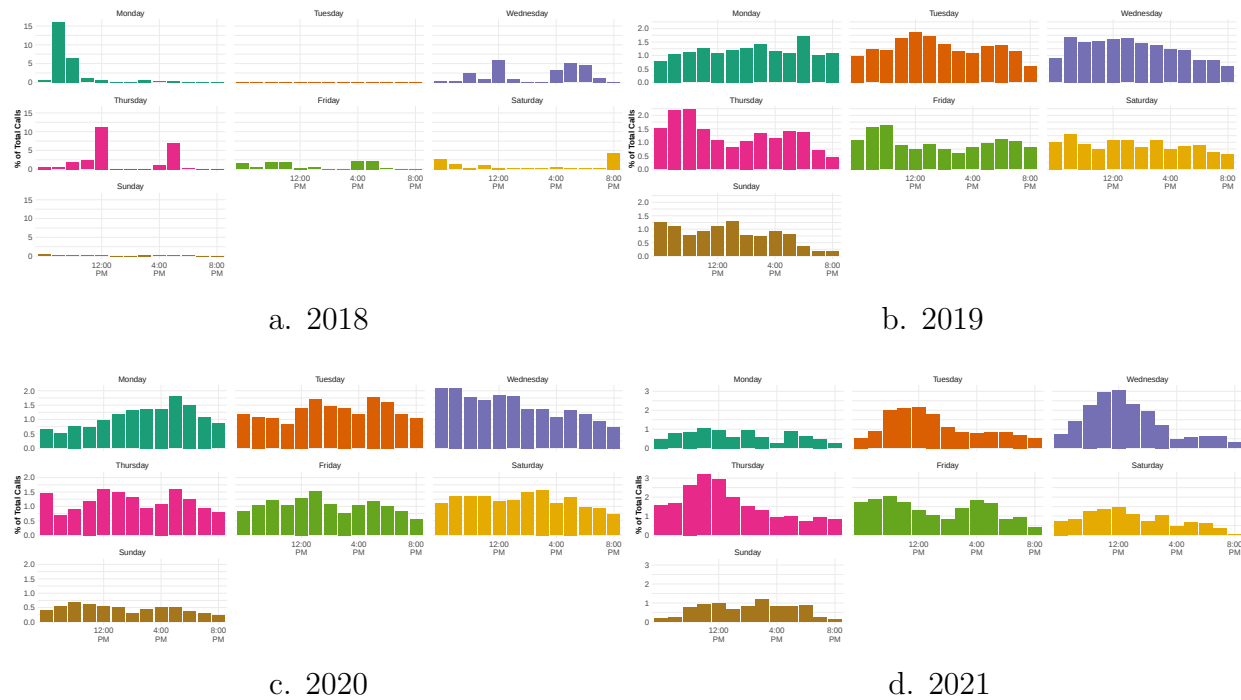
Notes: We pool the data from July 2018 to September 2021 for this figure.

Figure A4: (a) Duration with service, (b) Start month of service, (c) Pickup by time and tenure.



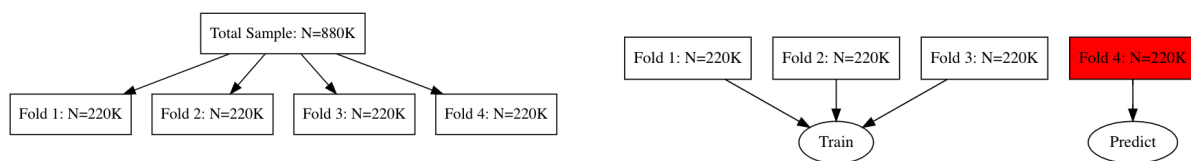
Notes: Panel (a) illustrates the distribution of the length of the farmers' tenure with the service. Panel (b) shows the starting month with the service for the sample. For Panel (c), we split the sample into two groups based on their length of tenure. The new farmer dummy takes a value of 1 if the length is lower or equal to the mean length of 516 days and 0 otherwise.

Figure A5: Number of push calls sent during each call hour in the past



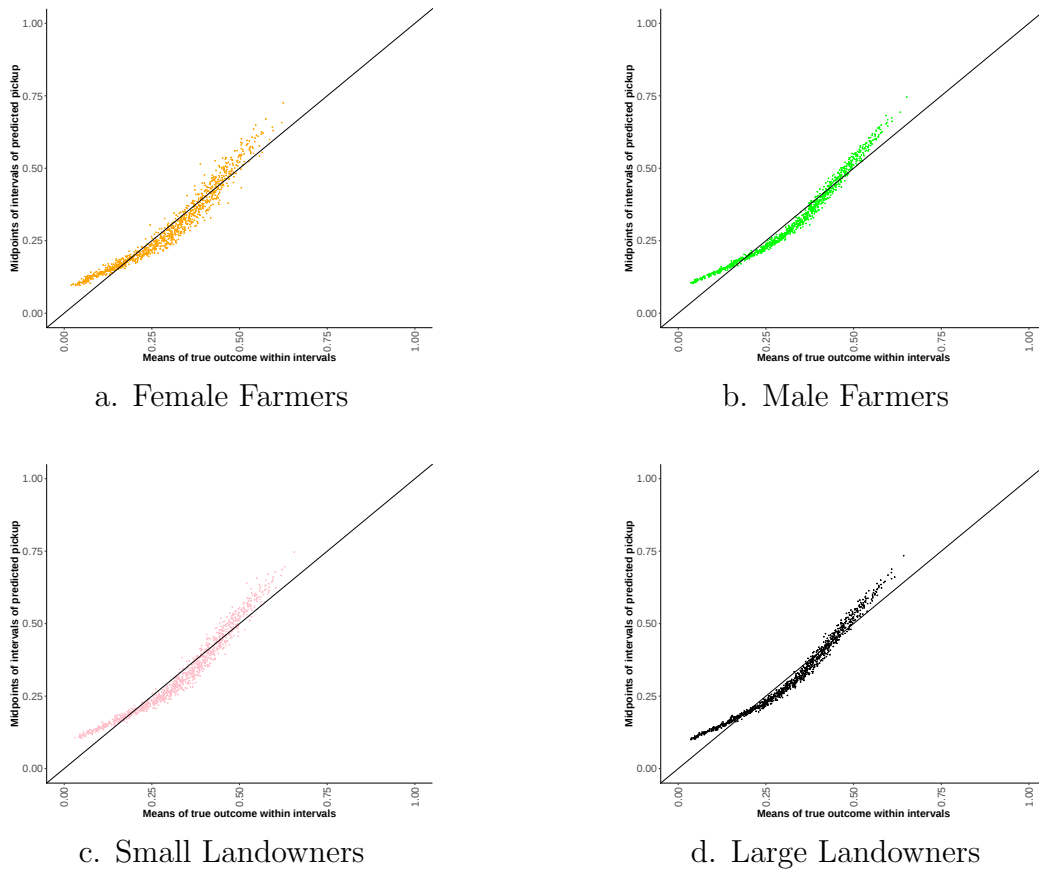
Notes: These figures show the number of push calls sent between July 2018 and September 30, 2021. It illustrates that push calls used to be sent in an ad hoc manner in the past.

Figure A6: Process of Cross-Fitting on Data (No. of Folds=4).



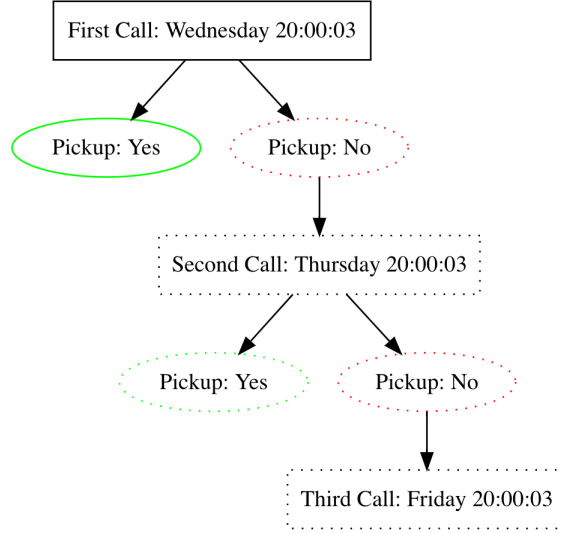
Notes: The figure on the left shows the process of splitting the sample into four folds. The figure on the right shows that the same data is not used for estimation and evaluation. To predict the pickup for the farmers in fold four, the engagement data for farmers in the remaining three folds are used. The same step is repeated for the other three folds.

Figure A7: Calibration Plots for LASSO by Farmer Covariates



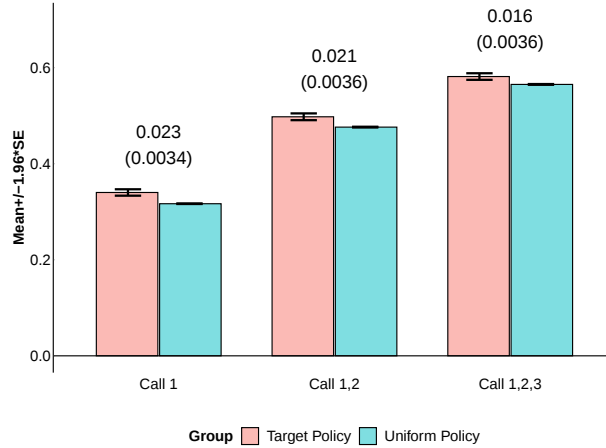
Notes: Panel (a) shows the calibration plots between the true and predicted outcome for the female farmers. Panel (b) shows the calibration plots for the male farmers. Panel (c) shows the calibration plot for the small landowners whose land size is below or the same as the 25th percentile of the land-size distribution. Panel (d) shows the calibration plot for large landowners. Large landowners are defined as those farmers whose land holding is above 25th percentile of the land-size distribution.

Figure A8: Repeat Calls from the Call Center



Notes: This figure illustrates the process of follow-up calls from the call center. If the farmer does not pick up the call on the first attempt, the call center makes a call exactly 24 hours after the first call. Moreover, if the farmer does not pick up the call on the second attempt, the third call is made 24 hours after the second call.

Figure A9: Impact of Targeting First Call on Overall Farmer Engagement: Off-Policy Evaluation

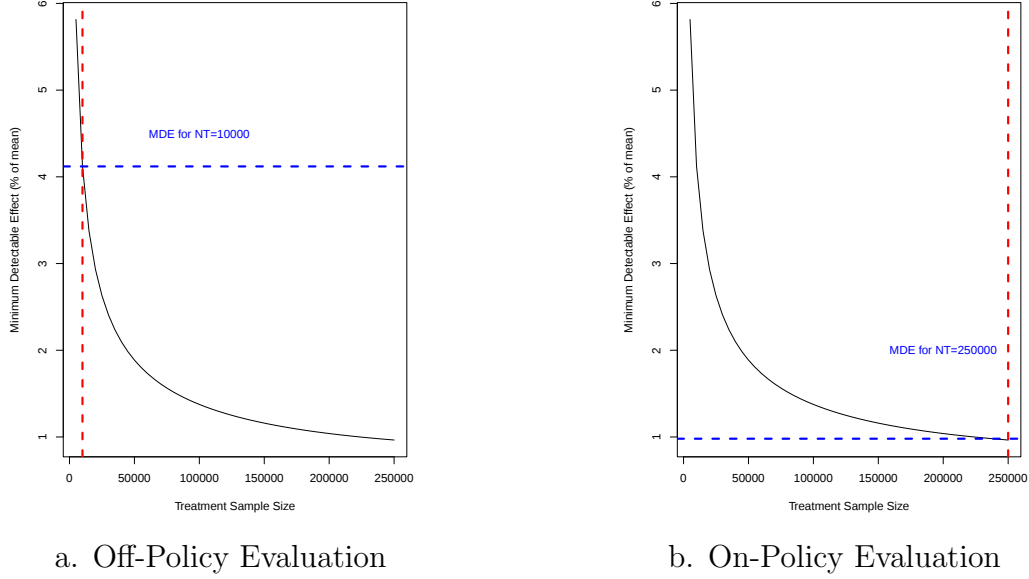


Notes: The training and evaluation data for estimating the value of $\hat{\pi}$ counterfactually and the uniform policy uses the uniform randomization data for week 1 and week 2. We estimate $\hat{\pi}$ on the first calls for the two weeks of data using cross-fitting (following Algorithm 1) and counterfactually evaluate $\hat{\pi}$ using the estimator defined in Section 5.2. We estimate $\hat{V}^{\text{Off}}(\hat{\pi}, \mathcal{S}_{\mathcal{U},1,2})$, where the outcome changes to the total pickup for call 1, call 2 for second pair of bars. The evaluation dataset for the $\hat{\pi}$ group corresponds to the subset of people whose actual assignment matched the first call $\hat{\pi}$. We estimate $\hat{V}^{\text{Off}}(\hat{\pi}, \mathcal{S}_{\mathcal{U},1,2,3})$, where the outcome changes to the overall pickup for call 1, 2, 3 for right most pair of bars.

C Power Calculations

This section provides an outline of the power calculations that we did before starting the multistage experiment. The goal of this exercise is to assess the minimum detectable effect for the off-policy and on-policy evaluations that we intended to do for each week. For each N , we estimate the minimum detectable effect for the differences in the estimate of the value of targeted policy and uniform policy under off-policy and on-policy evaluations. For the off-policy evaluation, let us assume that 900,000 farmers are allocated to uniform randomization in a week. This implies about $1/91 \times 900,000$ would be receiving calls as per the targeted policy under the uniform randomization data (refer to simplified setting in Section 5.2 for details on the fraction). Varying the sample size for those getting called according to $\hat{\pi}$ and those according to the uniform policy, we illustrate the relationship between the minimum detectable effect and the sample size. Figure A10 shows that under the sample sizes close to the ones we see in our project, the MDE for off-policy evaluation is much higher than the MDE for on-policy evaluation. This is happening as we can allocate a higher sample size (250,000) to receive calls according to the targeted policy under on-policy evaluation.

Figure A10: Power Calculations



Notes: Panel (a) shows the relationship between the MDE and sample size for the off-policy evaluation. Panel (b) shows the relationship between MDE and sample size for the on-policy evaluations.

D Simplified Setting

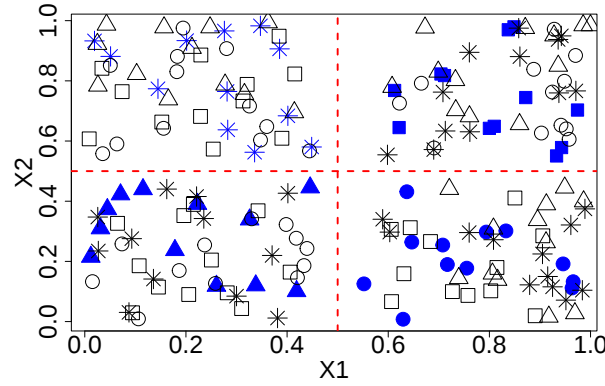
D.1 Off-Policy Evaluation

The concept of off-policy evaluation can be explained using a simplified setting with four treatment arms and two covariates. Assume two uniformly distributed covariates $X = [X_1, X_2]$ and the outcome variable is Y . Let there be only four treatments (W_1, W_2, W_3, W_4). The two policies $\mathcal{U}$ and $\hat{\pi}$ are defined below.

$$\mathcal{U} = \begin{cases} W_1, & p = 0.25 \\ W_2, & p = 0.25 \\ W_3, & p = 0.25 \\ W_4, & p = 0.25 \end{cases} \quad \hat{\pi} = \begin{cases} W_1, & x_1 > 0.5, x_2 > 0.5 \\ W_2, & x_1 > 0.5, x_2 < 0.5 \\ W_3, & x_1 < 0.5, x_2 < 0.5 \\ W_4, & \text{otherwise} \end{cases}$$

Under the uniform randomization, every individual could be allocated to any of those four treatments with a 0.25 probability. Under the targeted policy, individuals in the top left quadrant, top right quadrant, bottom left quadrant, and bottom right quadrant would be assigned to treatment W_4 , W_1 , W_3 , and W_2 , respectively (Figure A11). Because of uniform randomness, $\frac{N}{\text{No. of arms}}$ individuals under the uniform randomization are actually assigned to the treatment arm they would have received if the targeted policy $\hat{\pi}$ was deployed. Those individuals are highlighted using blue points in the figure, and they would be the treated individuals in the off-policy evaluation.

Figure A11: Off-Policy Evaluation: Simplified Setting



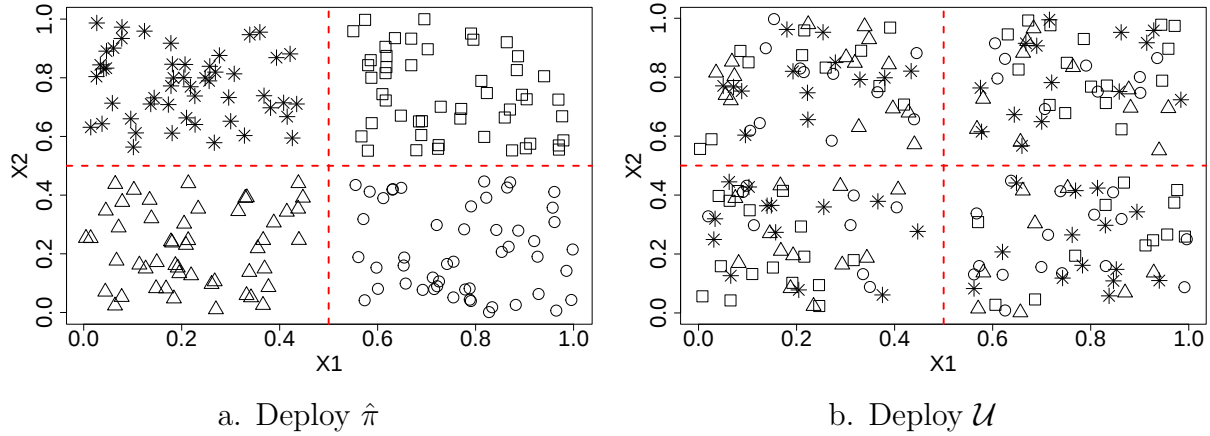
Notes: Data is collected using the uniform randomization, where probability of allocation to each arm is 0.25. 25% of farmers get the same assignment under $\mathcal{U}$ as they would have received under $\hat{\pi}$ (blue points).

D.2 On-Policy Evaluation

Using the same simplified setting with two covariates, we illustrate graphically the concepts of on-policy evaluation used in this study. Figure A12 illustrates the policy for the above setting. The difference in expectation for the subset of individuals who should get W_1 based on $\hat{\pi}$ is obtained by comparing the outcomes for the upper right quadrant in the uniform policy group and the targeted policy group. In other words, the covariate space is kept the same and only the policy is varied between the two groups. Similarly, the differences can be computed for the other arms. Additionally, the overall benefit of deploying $\hat{\pi}$ over

uniform policy can be computed by comparing outcomes in Panel (a) with the outcomes in Panel (b).

Figure A12: On-Policy Evaluation: Simplified Setting



Notes: Here, individuals are randomly allocated between the group that receives the treatment according to $\hat{\pi}$ (left) and uniform randomization (right).

E Estimated Policies: Implementation Phase

In this section, we provide details on the implementation phase of the targeted policies. While collecting the data using the uniform randomization, we simultaneously developed the technology for deployment of targeted policies. The goal of this exercise was to test the technological constraints with the advisory system and update the targeted policies according to our learning. For the second week, we estimated and evaluated a crude policy ($\hat{\pi}_A$), where the treatment arms were morning, afternoon, and evening trained on uniform randomization data in week 1 (refer to Figure 1).

Next, as we collected more data using the randomization between 91 call times for additional weeks, we updated the policy ($\hat{\pi}_B$) to take into account the hour of the call along with the day of the week. We could only test the updated policy for a limited number of days for the week of November 1st to 7th (week 4) because, during this week, the call center

only operated for five days instead of seven (this week had two major festivals, Diwali and Bhai Dooj). Next, we tested another intermediate policy ($\hat{\pi}_C$) in week 5. This was the first week that we could implement the targeted policy with 91 treatment arms for all seven days. However, during this time, we were also trying to learn about the capacity constraints and bandwidth limits for the call center. We deployed and tested the intermediate policy for week 5 but with tighter constraints. For the last week, week 6, we deployed the policy for all seven days and relaxed the constraint further to accommodate additional farmers in their best-predicted call times ($\hat{\pi}_D$). Ideally, we would want to implement a few more targeted policies taking into account the possibility of shocks, but we used off-policy evaluation to estimate and evaluate those additional policies. We only had time to deploy a few limited, targeted policies during the six weeks.

Table A2: Implementation Phase (Week 2)

Data Collection Method	$\hat{\pi}_A$	Uniform Policy	Difference
All	0.3030 [0.0009]	0.3178 [0.0006]	-0.0148 [0.0011]
N	265,188	616,656	

Notes: This table shows the on-policy evaluation for $\hat{\pi}_A$ in week 2. The farmers are randomized between two groups. Group A gets called according to $\hat{\pi}_A$ and group B gets called according to uniform randomization. Sample means are used to estimate the value of $V(\pi_A, \mathcal{S}^{\text{eval}})$ and $\bar{V}(\mathcal{U}, \mathcal{S}^{\mathcal{U}})$.

Table A3: Implementation Phase (Week 4)

Data Collection Method	$\hat{\pi}_B$	Uniform Policy	Difference
All	0.3125 [0.0011]	0.3172 [0.0006]	-0.0047 [0.0013]
N	167,995	707,644	

Notes: This table shows the on-policy evaluation for $\hat{\pi}_B$ in week 4. The farmers are randomized between two groups. Group A gets called according to $\hat{\pi}_B$ and Group B gets called according to uniform randomization. Sample means are used to estimate the value of $V(\pi_B, \mathcal{S}^{\text{eval}})$ and $\bar{V}(\mathcal{U}, \mathcal{S}^{\mathcal{U}})$.

Table A4: Implementation Phase (Week 5)

Data Collection Method	$\hat{\pi}_C$	Uniform Policy	Difference
All	0.3133 [0.0011]	0.3195 [0.0006]	-0.0063 [0.0011]
N	234,276	624,801	

Notes: This table shows the on-policy evaluation for $\hat{\pi}_3$ in week 5. The farmers are randomized between two groups. Group A gets called according to $\hat{\pi}_C$ and Group B gets called according to uniform randomization. Sample means are used to estimate the value of $V(\pi_C, \mathcal{S}^{\text{eval}})$ and $\bar{V}(\mathcal{U}, \mathcal{S}^{\mathcal{U}})$.

F Alternative Specifications

F.1 Functional Form of the Treatment Variables

For the main analysis, the treatment effects are incorporated as 91 treatment dummies in our specifications. We estimate an alternative specification here. We define the hour variable as a continuous variable and use a polynomial function. There are also seven dummies, one for each day of the week. The model incorporates interactions of the hour

variable and its higher-order terms with the days of the week. The modified model is provided below.

$$\begin{aligned} \text{logit}(p_{ij}) = & \beta_1 X_i + \beta_2 \text{hour}_j + \beta_3 \text{hour}_j^2 + \beta_4 \text{hour}_j^3 + \beta_5 \text{day}_j + \eta_1 \text{hour}_j \text{day}_j + \eta_2 \text{hour}_j^2 \text{day}_j + \\ & \eta_3 \text{hour}_j^3 \text{day}_j + \gamma_1 X_i \text{hour}_j + \gamma_2 X_i \text{hour}_j^2 + \gamma_3 X_i \text{hour}_j^3 + \gamma_4 X_i \text{day}_j + \delta_1 X_i \text{hour}_j \text{day}_j + \\ & \delta_2 X_i \text{hour}_j^2 \text{day}_j + \delta_3 X_i \text{hour}_j^3 \text{day}_j \end{aligned}$$

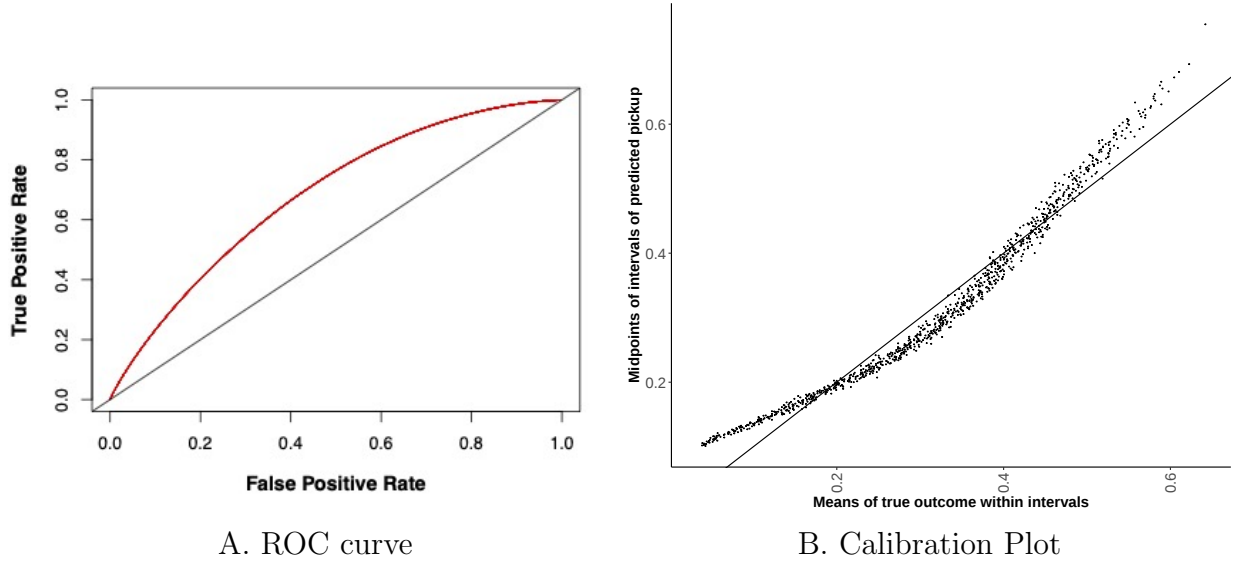
The parameters of the model are estimated using LASSO. We do not penalize the coefficients on the treatment variables, which include $\{\beta_1, \beta_2, \beta_3, \beta_4, \beta_5, \eta_1, \eta_2, \eta_3\}$. For the other coefficients, we choose a regularization parameter using cross-validation.

Here, we assess the out-of-sample prediction accuracy of this model using cross-fitting. The sample is divided in K folds, and the prediction accuracy for farmers in fold k is assessed by estimating the model parameters on all farmers except k . This is repeated K times to estimate the predictions for farmers in every fold. The prediction accuracy and model fit of this alternative specification is compared with Equation 1. The AUC is comparable for the two specifications, 0.683. The ROC and calibration plots are provided below.

F.2 Hierarchical LASSO

In this section, we estimate a hierarchical LASSO model. First, Equation 4 is estimated. In this specification, the regularization parameters vary for the main effects and interaction of the covariates with the treatment dummies. Additionally, a constraint is introduced such that the penalty on the interactions is larger than the penalty on the main effects.

Figure A13: Out-of-Sample Predictive Performance



Notes: (a) shows the ROC curve for the out-of-sample prediction. The AUC is 0.683. (b) shows the calibration plot of true and predicted pickup on the test dataset.

$$\begin{aligned} \text{logit}(p_{ij}) = & X_i\beta + \sum_j \delta_{1j}\text{Hour}_j + \sum_j \delta_{2j}\text{day}_j + \sum_j \gamma_{1j}X_i\text{hour}_j + \\ & \sum_j \gamma_{2j}X_i\text{day}_j + \sum_j \gamma_{3j}X_i\text{hour}_j\text{day}_j \end{aligned} \quad (4)$$

The objective of the hierarchical LASSO is to minimize the following-

$$\begin{aligned} & -l(\theta, X, \text{hour}, \text{day}) + \lambda_1\|\beta\| + \lambda_2\|\gamma_1\| + \lambda_2\|\gamma_2\| + \lambda_3\|\gamma_3\| \\ & \text{s.t.} \\ & \lambda_1 \leq \lambda_2 \leq \lambda_3 \end{aligned} \quad (5)$$

We set up the optimization problem as stated in Equation 5 and solve for the regularization parameters that minimize the objective and satisfy the constraint. The *nloptr* package in R is used for optimal penalty parameters. The out-of-sample accuracy for this

specification is compared with the baseline model. The AUC is 0.68, which is comparable to the AUC of the baseline specification, where we use the same penalty for all the variables except the ones on the treatment dummies.

Next, the method is modified to allow for the regularization parameters to vary by the degree of the polynomial terms and the interaction coefficients. The hour variable in the above setup has polynomial terms of degree one, two, and three, and there are interactions of the hour variable with day and other covariates. Consequently, there are five regularization parameters $\lambda_1, \dots, \lambda_5$, and the objective of LASSO is to minimize the following.

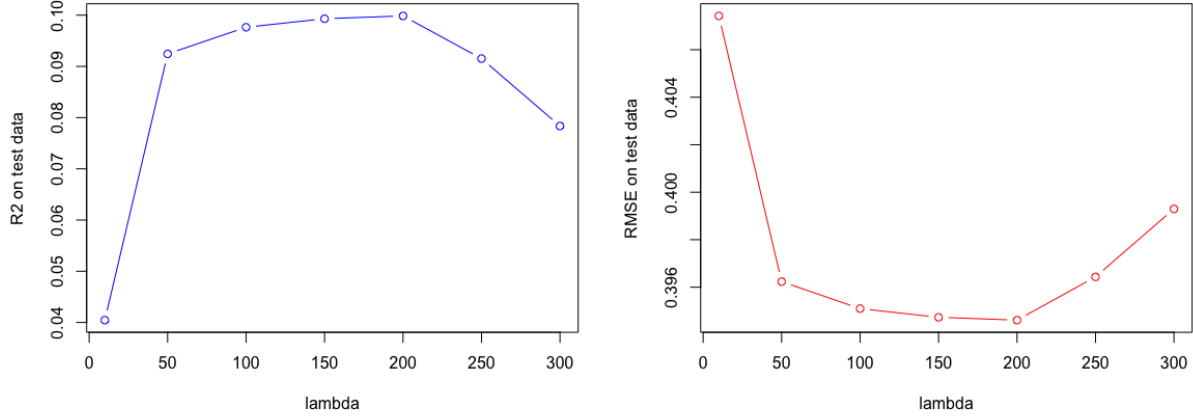
$$\begin{aligned}
 & -l(\theta, X, T) + \lambda_1 \|\beta_1\| + \lambda_2 \|\gamma_1\| + \lambda_2 \|\gamma_4\| + \lambda_3 \|\gamma_2\| + \lambda_3 \|\delta_1\| + \lambda_4 \|\gamma_3\| + \lambda_4 \|\delta_2\| + \lambda_5 \|\delta_3\| \\
 & \text{s.t.} \\
 & \lambda_1 \leq \lambda_2 \leq \lambda_3 \leq \lambda_4 \leq \lambda_5
 \end{aligned} \tag{6}$$

The penalty on the coefficients increases with the increasing degree of the interactions of the coefficients with treatment variables. The out-of-sample accuracy of this model is comparable to the one we estimated in Equation 1.

F.3 Matrix Completion for Historical Data

We describe the matrix completion exercise of the historical data in this section. Note that the historical data for the experimental sample is available from July 2018 to September 2021. However, we do not have the historical engagement for every farmer for each of the 91 call times, as calls were scheduled in an ad hoc manner in the past. The non-missing entries constitute about 49% of the total number of entries. Hence, a matrix completion algorithm is needed to complete the past engagement matrix. The Soft-Impute algorithm described

Figure A14: Nuclear norm penalty for the Matrix Completion on Past Engagement



a. Out of Sample R^2

b. RMSE

Notes: (a) shows the R^2 on the test data. (b) shows the RMSE corresponding to different values of the nuclear norm penalty.

in Mazumder et al. (2010) is used for the matrix completion exercise. This algorithm uses nuclear norm regularization for matrix completion.

As described in Mazumder et al. (2010), the optimization problem for completing the historical engagement matrix Y_{history} is provided below. Note that the matrix Y_{history} has dimension $N \times J$ where each (i, j) element corresponds to the historical engagement of farmer i in call time j .

$$\min_M \frac{1}{2} \|P_{\Omega}(Y_{\text{history}} - M)\|_F^2 + \lambda \|M\|_*$$

, where $\|M\|_*$ is the sum of singular values of M and $P_{\Omega}(Y_{\text{history}})$ is the projection matrix with observed elements in Y_{history} .

The regularization parameter is chosen using the following steps. The matrix elements are split 80:20 into training and test entries. The non-missing entries in the training matrix that are part of the test IDs are made NAs for this exercise. We choose several values of λ and do the matrix completion using Soft-Impute. Both the R-MSE and out-of-sample R^2

were computed for the test entries. The λ corresponding to the lowest R^2 was chosen as the optimal λ . Panels (a) and (b) of Figure A14 display these measures for the different values of λ . The matrix completion exercise provides a 10% improvement over imputing the missing entries with the mean of the matrix. The predicted past engagement is used as a covariate for estimating and evaluating targeted policies on the uniform randomized data.