

Deciding in the Dark: Social Norms Guiding Risky Decisions for Others*

J Dustin Tracy[†] Prithvijit Mukherjee[‡]

September, 2026

Abstract

How does one make a decision on behalf of another when all options have risk, and one is not sure of the other's preferences? We experimentally investigate whether social norms exist and whether they guide Selectors' choices among sets with uncertain outcomes, and whether norm adherence is rewarded. Independent Evaluators first rate the social appropriateness of choices between risky prospects, and whether those choices deserve a reward or a punishment. Selectors then choose between these prospects for Recipients, with treatments varying whether Recipients can provide monetary rewards or punishments. We find that Evaluators coordinate on social norms, and Selectors follow them, selecting more appropriate options even in the absence of monetary incentives. In contrast, Recipients base rewards and punishments primarily on realized outcomes rather than norm compliance, rewarding lucky norm violations while punishing unlucky norm adherence.

Key Words: Social norms, Decision-making for others, Laboratory experiments, Decision-making under risk

JEL Codes: C9, D63, D81, D90, G41

1 Introduction

Individuals often make decisions for others without fully knowing the recipient's preferences. A friend plans a surprise; a professor setting a syllabus chooses readings and assignments without knowing what each student expects; a doctor lays out a course of treatment; a policymaker proposes a program without knowing how those it affects would weigh its risks; a fund manager allocates someone's retirement account. Many of these choices are also inherently stochastic. The surprise might be a meal at a restaurant. The experienced utility of patrons with heterogeneous preferences might be thought of as corresponding to the number of stars in their reviews. The distribution of star levels therefore represents the probability of a new diner's (the friend's) experienced utility payout. Similarly, one can imagine students from previous semesters rating a particular reading. Many medical treatments are only successful for a known proportion of patients. These situations share two features: the choice is made without full knowledge of the recipient's preferences, and the choice itself carries risk. Risk also creates a wedge between when a choice is made and when it is judged. The Selector chooses between two options without knowing the outcome, while the Recipient evaluates that choice only after the outcome is realized.

Contracts are often put in place to align incentives. But what happens when enforcing them is costly, or

*We thank Erik Kimbrough, Erin Krupka, and Tim Salmon for their input and comments on the design; and Kevin James for his help programming the software. Any mistakes are our own.

[†]Augusta University

[‡]Bryn Mawr College

when it is too cumbersome to spell out every scenario in advance?¹ What then guides the decision-making process? Despite the ubiquity and importance of deciding for others, we lack a clear understanding of what guides these choices. The existing literature offers competing explanations. Traditional principal-agent models assume decision-makers maximize their own utility, leading to moral hazard when interests diverge (Jensen and Meckling, 1976). A meta-analysis of the experiments studying decisions for others finds people may be more risk-taking when deciding for others than for themselves (Polman and Wu, 2020), though findings remain mixed (Eriksen et al., 2020). The social values hypothesis proposes that decisions for others reflect cultural norms about appropriate risk-taking in specific contexts (Stone and Allgaier, 2008, Dore et al., 2014). Yet these frameworks struggle to explain a puzzling regularity: decision-makers often appear conscientious when choosing for others, yet face evaluation based primarily on outcomes beyond their control.

This paper investigates three questions: 1) whether Selectors exert the time and effort to adhere to these norms; 2) whether social norms predict behavior in settings in which one makes a decision for another and; 3) whether adherence to these norms is rewarded and deviance is punished. We design a controlled experiment where Selectors choose between two risky prospects for Recipients, systematically varying whether Recipients can provide financial rewards or punishments.

To the best of our knowledge, ours is the first paper to elicit norms from independent Evaluators using the coordination method developed by Krupka and Weber (2013) in a stochastic setting. For each choice, we measure: 1) its social appropriateness, and 2) whether choosing it deserves reward or requires punishment. We elicit these norms in two settings: in the “Cold” setting, Evaluators assess the choices without knowing the outcome of the lottery, while in the “Hot” setting they assess the choices knowing the outcome. This allows us to test whether Selectors follow these norms and whether Recipients actually reward norm compliance versus favorable outcomes.

Our experiment comprises three treatments: *Self*, *Other*, and *Consequences*. In the *Self* treatment, the Selectors make choices for themselves, so Recipient is not a separate role. In the *Other* treatment, the Selectors also make choices for the Recipients; however, the Recipients cannot affect the Selectors’ bonuses. In the *Consequences* treatment, Selectors make choices for the Recipients, who after learning the choice and the lottery realization, can make costly adjustments to the Selectors’ bonus.

The *Self* treatment serves two purposes. First, it provides a point of reference for how people choose when only their own payoff is at stake. Second, it sets up a check on whether norms matter at all: if there is no variation between choices made for oneself and those made for others, then there is no behavior that we need norms to explain. That norms do predict choices in the *Other* and *Consequences* treatments tells us that deciding for someone else is governed by something beyond one’s own preferences.

In the *Other* treatment, there is no possibility of enforcing the norms, so any compliance is due to a non-pecuniary motivation to follow the norms. In the *Consequences* treatment, stricter compliance to norms, could be attributed to anticipation of enforcement of the norm.

We find that the norms elicited from independent Evaluators predict Selectors’ choices. In the *Other* treatment, Selectors are more likely to choose the lottery rated as more socially appropriate. In the *Consequences* treatment, they are more likely to choose the lottery rated as more deserving of reward and less deserving of punishment. These relationships remain when the two options’ ratings are combined into a single measure of their normative difference. The results indicate that shared norms provide guidance when Selectors make risky choices for Recipients whose preferences they do not observe.

Because Selectors follow social norms consistently, there are few cases of strong norm violations, which limits our power to test norm enforcement by the Recipients.

The Recipients’ adjustments of bonuses are driven more by realized outcomes than by whether the Selector chose the option Evaluators deem deserving of a reward. Recipients are more likely to reward Selectors whose chosen lottery happens to yield a high payoff, and to withhold reward or punish when the payoff

¹Unlike a principal-agent game, there is no contract in which the Recipient has instructed the Selector on how the decisions should be made. We thank Tim Salmon and other participants at the 2025 ESA North American Meeting for making this point.

is low—even though the Selector has no control over the realization. Norm compliance at the moment of choice matters far less to Recipients than the outcome they ultimately receive.

Our study contributes to three strands of the literature. First, we add to the literature measuring social norms. The Krupka and Weber (2013) coordination game has become the standard tool for identifying shared normative standards (Charness et al., 2025), but it has been applied almost exclusively to deterministic choices, where the mapping from action to outcome is immediate and certain. Whether norms govern choices under risk, where the decision-maker selects a *lottery* rather than an *allocation*, remains unexplored. We extend the method to stochastic environments by eliciting appropriateness ratings for choices between risky prospects, and we elicit these ratings on a seven-point scale that affords finer discrimination while retaining a neutral midpoint. Crucially, we distinguish two normative dimensions that prior work has left together: *social appropriateness* and whether a choice *requires reward or deserves punishment*. Though correlated, the two carry distinct predictive power for behavior and for enforcement responses.

Second, we contribute at the intersection of the principal-agent and risk-taking for others literatures. The canonical principal-agent model (Jensen and Meckling, 1976) attributes moral hazard to divergent incentives, largely assuming the agent knows the principal’s objectives. In many delegation settings, neither condition holds: explicit contracts are absent, and agents cannot observe the principal’s risk preferences (Stone and Allgaier, 2008, Eriksen et al., 2020). The risk-taking literature compounds the puzzle, with some studies finding greater risk-taking on behalf of others (Chakravarty et al., 2011, Polman, 2012, Agranov et al., 2014, Pollmann et al., 2014), others finding greater risk aversion (Charness and Jackson, 2009, Bolton and Ockenfels, 2010, Pahlke et al., 2012), and still others finding no systematic difference (Harrison et al., 2013, Luzuriaga et al., 2017, Barrafreem and Hausfeld, 2020); meta-analysis suggests the net effect is small (Polman and Wu, 2020). We show that norms elicited from independent Evaluators predict Selectors’ choices, serving as a coordination device in delegation relationships where contracts and preference information are both absent.

Third, we contribute to work on enforcement and outcome-based evaluation. Rubin and Sheremeta (2016) show that when stochastic shocks enter a delegation setting, principals reward realized output rather than effort, even with full information about the shocks. We find the same pattern governs reward and punishment here: Recipients respond primarily to lottery realizations rather than to whether Selectors chose the normatively endorsed option. This is consistent with outcome bias in moral judgment (Baron and Hershey, 1988) and helps explain why shared standards may be insufficient to sustain compliance in repeated delegation relationships.

The remainder of the paper proceeds as follows. Section 2 develops the theoretical framework and derives the hypotheses. Section 3 describes the experimental design and norm-elicitation procedure. Section 4 presents the results. Section 5 discusses their implications.

2 Background and Hypotheses

We formalize the decisions made by the Selector for a Recipient as a sequential game. The Selector chooses between two risky prospects (lotteries) that determine the Recipient’s payoff. The game begins with a default option (A). The Selector may exert effort to investigate an alternative option (B). If the Selector does not investigate, the default option (A) is implemented. Importantly, one option can be selected through inaction, while the other requires action.

We begin by outlining theoretical predictions under the standard assumptions for subgame-perfect Nash equilibrium, where the Selectors and Recipients are purely self-interested for the three treatments.

In the *Self* treatment, the Selector chooses between lotteries solely to maximize their own expected utility. Since no strategic interaction is present, behavior depends entirely on the Selector’s risk preferences and beliefs about the payoff distributions of options (A) and (B). The decision to exert effort to investigate option (B) is optimal if and only if the expected utility gain from improved information exceeds the time cost of effort. Standard theory therefore yields no sharp prediction beyond heterogeneity driven by individual risk attitudes and beliefs.

In the *Other* treatment, the Selector chooses a lottery on behalf of a Recipient while receiving a fixed payoff that is independent of the lottery’s realization. Under the assumption of purely self-interested preferences, the Selector derives no utility from the Recipient’s payoff. As a result, exerting effort to investigate option (B) is strictly dominated, and the Selector optimally allows the default option (A) to be implemented. The unique subgame-perfect Nash equilibrium predicts complete inertia.

From a standard economic perspective with purely self-interested Selectors, the predictions in the *Consequences* treatment are also straightforward. Applying backward induction, a Recipient will never choose to reward or punish the Selector, as doing so is costly and occurs after both the Selector’s decision and the realization of the lottery in a one-shot game. Anticipating the absence of any credible sanctions or rewards, a rational Selector has no incentive to incur effort costs to investigate option (B) and instead allows the default option (A) to be implemented. The subgame-perfect Nash equilibrium therefore predicts complete inertia: no exploration by Selectors and no rewarding or punishment by Recipients.

There is substantial evidence that shows actual behavior systematically deviates from these predictions. To understand these deviations, we turn to Adam Smith’s insights about moral behavior and social standards. We follow Vernon Smith’s interpretation of how these principles manifest in economic interactions. A key insight from *The Theory of Moral Sentiments* (Smith, 1759) is that individuals are guided not only by self-interest but also by a desire to act in ways that an ‘impartial spectator’ would approve. This spectator represents the internalized judgment of society about appropriate conduct.

To elicit the views of the ‘impartial spectator’, we adapt the coordination game approach developed by Krupka and Weber (2013), extending it to a setting with stochastic outcomes. We elicit ratings from independent Evaluators about the social appropriateness of Selectors’ decisions and whether these decisions deserve reward or punishment.² This coordination game method has been successfully applied across various economic contexts, including dictator games (Krupka and Weber, 2013), ultimatum games (Smith and Wilson, 2018), trust games (Smith, 2020), and public goods games (Kimbrough and Vostroknutov, 2016).³

We expand on the framework for eliciting norms in three ways. First, (Krupka and Weber, 2013) used a four-point scale to elicit norms; we expand this to a seven-point scale, allowing Evaluators to express neutrality. The seven-point rating scale also parallels the seven options the Recipients have for adjusting the Selectors’ bonuses $\{-60, -40, -20, 0, 20, 40, 60\}$ ¢. This richer scale enables more nuanced judgments, particularly important in settings where the appropriateness of risky decisions may be ambiguous. These ratings serve as normative benchmarks for evaluating actual behavior.

Second, we elicit normative judgments both *before* and *after* the realization of lottery outcomes. “Cold” ratings are elicited from Evaluators, who assess decisions without knowing the realized payoffs, while “Hot” ratings are elicited from Evaluators, who observe outcomes before evaluating the same decisions.

Third, we distinguish between two dimensions of normative judgment: social appropriateness and deserved punishment or reward. Social appropriateness captures whether a decision aligns with shared standards of appropriate conduct. In contrast, deserved punishment or reward reflects whether the decision warrants consequences, which may depend not only on the decision but also on its realized outcome.

This distinction allows us to explore a central tension in stochastic environments: individuals might judge a risky choice as normatively appropriate *ex ante*, yet become more willing to endorse punishment when an unfavorable outcome materializes. Conversely, successful outcomes may attenuate demands for sanction even when the underlying decision violates shared standards. By eliciting both measures in Cold and Hot conditions, we can separate stable normative approval from outcome-contingent enforcement responses.

In the presence of norms, we hypothesize two mechanisms through which appropriate behavior might emerge in the absence of enforceable contracts. First, social norms can guide decision-making in settings where formal incentives are incomplete or absent, by shaping both what individuals believe others typically do and what they believe others think one ought to do (Vostroknutov, 2020, Gorges and Nosenzo, 2022).

²For each lottery pair, Evaluators rate on a seven-point scale whether: whether: (i) the Selector’s lottery choice is {“Extremely”, “Moderately” or “Slightly”} socially appropriate or inappropriate or neither; and (ii) whether the Selector’s choice requires reward or deserves punishment. For analysis, we treat both ratings as continuous variables ranging from -3 to 3.

³Please see Charness et al. (2025) for a review of the literature.

Second, the possibility of reward or punishment, even when personally costly, might reinforce adherence to these norms by providing extrinsic support for norm-consistent behavior, as shown in experiments on norm enforcement and sanctioning in public-goods and related games (Fest   and Garrouste, 2014, Andreoni and Bernheim, 2009).

The literature on making risky choices for others shows mixed results, some studies find more risk-taking (Chakravarty et al., 2011, Polman, 2012), others find more risk aversion (Charness and Jackson, 2009, Bolton et al., 2015), and some find no difference (Harrison et al., 2013, Barrafre   and Hausfeld, 2020). A meta-analysis by Polman and Wu (2020) finds only a very small indication of more risk-taking behavior when making decisions for others. These inconsistent findings suggest that understanding the role of social norms may be crucial for predicting behavior in our setting.

The three treatments are designed to test how social norms influence behavior: Self, Other, and Consequences. The Self treatment provides a point of reference for how Selectors choose when only their own payoff is at stake, and social norms play no role in shaping behavior. The Other treatment isolates the role of social norms by having the Selector make decisions that affect another individual's payoff without facing material consequences. In contrast, the Consequences treatment introduces the possibility of monetary rewards or punishments, enabling us to assess how behavior changes when norm violations can be rewarded or sanctioned financially.

Selectors receive a fixed payment for making decisions. Since deliberation carries an implicit opportunity cost of time, we expect systematic differences in decision-making effort across treatments.

In the Self treatment, Selectors are familiar with their own preferences and can decide quickly. In contrast, in the Other and Consequences treatments, Selectors must consider what is socially appropriate or anticipate how another individual might evaluate their decision, which increases cognitive demands.

Effort is also reflected in the choice whether to reveal Option B. In the Self treatment, if the default is "good enough," a Selector might be less inclined to search. However, when choosing for another, the moral weight of the "impartial spectator" or the risk of sanctioning may drive a more thorough investigation of the available alternatives. This yields the following two hypotheses.

Hypothesis 1.a *Selectors will spend the least amount of time making decisions in the Self treatment, followed by the Other treatment, and the most time in the Consequences treatment.*

Hypothesis 1.b *Selectors are more likely to exert effort to investigate the alternative (Option B) in the Other and Consequences treatments compared to the Self treatment.*

Once the Selector has deliberated and, if applicable, investigated the alternative, they must choose a lottery. The remaining decision reflects how Selectors translate norms and incentives into final choices. In the Other treatment, Recipients cannot reward or punish the Selector. As a result, any deviation from the default option must be motivated by internalized social norms or beliefs about what constitutes an appropriate choice. In contrast, the Consequences treatment introduces the possibility of monetary rewards and punishments, which might reinforce norm-consistent behavior by attaching material consequences to perceived appropriateness. This yields the following two hypotheses:

Hypothesis 2.a *When making decisions for a Recipient in the Other treatment, Selectors are more likely to choose the option that is ranked as more socially appropriate.*

Hypothesis 2.b *When making decisions for a Recipient in the Consequences treatment, Selectors are more likely to choose the option that is judged to deserve a higher reward.*

Finally, we consider the Recipient's behavior. While standard theory predicts zero sanctioning, the 'impartial spectator' framework suggests that Recipients will use their payoff to signal approval or disapproval. We expect this to be a nuanced adjustment based on the Selector's intent (the choice made) rather than just the luck of the draw (the outcome).

Hypothesis 3.a *In the Consequences treatment, Recipients' bonus decisions to reward or punish will be guided by whether the Selector's choice aligns with prevailing norms.*

Hypothesis 3.b *In the Consequences treatment, Recipients will reward or punish only active choices.*

3 Design

The experiment centers on a sequential game in which a Selector chooses between two risky lotteries (Options A and B) on behalf of a Recipient. In the sessions, we used neutral terms, e.g., “another participant”. However, throughout the paper, we use the terms Recipient, Selector, and Evaluator. We use a between-subjects design where we cross the three treatments—Self, Other, and Consequences—with Aligned and Unaligned defaults (where alignment refers to whether the default option favors risk neutrality). Table 1 describes the flow of the experiment, the actions of the Selector, and the choices the Recipient has in the Consequences and the Other treatment.⁴

Table 1: Experimental Design

Role		Treatment	
	Self	Other	Consequences
Actions			
Selector		Makes 6 choices between pairs of prospects	
Nature		Selects one choice and realizes value of chosen prospect	
Recipient	NA	Sees outcome	Adjusts bonus
		Rates Social Appropriateness	See outcome
		Rates Deserve Punish/Reward	Re-adjusts bonus
Payments (in addition to endowment)			
Selector	Realized prospect	Random prospect	Adjusted bonus
Recipient	NA	Realized prospect	Realized prospect

Option A is the default, requiring no effort to select, while Option B is revealed only after a deliberate investigation by the Selector. Figure A.5 shows an example of the Selector decision screen. Selectors had to click on the image with the question mark. Once the image was clicked, it changed to a plot akin to the one for Option A, and the text alongside the image appeared. Selectors could only select Option A before the terms of Option B were revealed. After revealing Option B, either could be selected. An example like this was presented to all participants as part of the instructions. They could not proceed until they had clicked the question mark to ensure they understood. For half the Selectors, the default was **Aligned** with the selection that maximized the Recipients’ expected payouts, i.e., Option A is α , and Option B is β . For the other half of Selectors, the default was **Unaligned** with the Recipients’ interest, i.e., Option A is β , and Option B is α .

The experiment consists of six pairs of lotteries, as detailed in Table 2. Option α was designed to be the more rational (risk-neutral) choice, though to varying degrees. The lottery pairs capture different trade-offs between risk and return:

Pair 1 presents a clear dominance relationship: Option α first-order stochastically dominates Option β , offering a higher expected value (\$0.90 vs. \$0.70) and superior outcomes in all states. This pair serves to identify the strength of social norms in favor of a stochastically dominant lottery.

Pair 2 presents a trade-off where Option α offers a higher expected value (\$1.01 vs. \$0.90) but with greater variance, including a possible zero payout. This pair allows us to test the social appropriateness of prioritizing higher expected value over increased variance.

⁴Subject instructions are provided in Appendix A.1.

Table 2: Lottery Pair Details

Pair	Option α					Option β					Note
	Pr. 1	Pay 1	Pr. 2	Pay 2	EV	Pr. 1	Pay 1	Pr. 2	Pay 2	EV	
1	0.75	\$1.15	0.25	\$0.15	\$0.90	0.75	\$0.90	0.25	\$0.10	\$ 0.70	FOSD
2	0.75	\$1.35	0.25	\$0.00	\$1.01	0.75	\$1.10	0.25	\$0.30	\$ 0.90	↓ EV & Variance
3	0.75	\$0.96	0.25	\$0.32	\$0.80	0.75	\$1.00	0.25	\$0.20	\$ 0.80	EV Equal, ↑ Var.
4	0.75	\$1.05	0.25	-\$0.30	\$0.71	0.75	\$0.80	0.25	\$0.00	\$ 0.60	↓ EV, $pay_2 < 0$
5	0.50	\$1.80	0.50	\$0.40	\$1.10	0.75	\$1.10	0.25	\$0.30	\$ 0.90	↓ EV & pr(Min)
6	0.25	\$1.45	0.75	\$0.45	\$0.70	0.75	\$0.90	0.25	\$0.10	\$ 0.70	↓ Max & pr(Min)

Pair 3 holds the expected value constant at \$0.80 while varying risk: Option α offers lower variance than Option β . This tests whether risk reduction itself generates normative pressure when expected returns are equal.

Pair 4 replicates the risk-return structure of Pair 2 but shifts all payouts down by \$0.30, introducing the possibility of losses relative to the \$0.50 endowment. This allows for an investigation into whether the presence of potential losses strengthens or weakens social norms for higher expected value.

Pair 5 modifies the risk-return trade-off by equalizing the probabilities of high and low outcomes ($p = 0.5$) for Option α , while maintaining a higher expected value (\$1.10 vs. \$0.90) than Option β . This pair tests the robustness of the normative endorsement of expected value maximization in a balanced probability environment.

Pair 6 maintains equal expected values (\$0.70) but reverses the typical risk pattern: Option α has a lower maximum payout and a higher probability of the low outcome compared to Option β . This serves as an exploratory case for social norms when risk and returns do not provide a clear anchor.

After the experiment, all participants completed a ten-question Big Five Personality Test.⁵ The experiment is preregistered on aspredicted.org as #84401. Hypotheses 1.a and 2.a were suggested by audience members at conference presentations. It was programmed in oTree (Chen et al., 2016) and run on prolific.com with a sample of adults residing in the United States.

3.1 Evaluators

We ran the Evaluator role first to check that there was convergence on social norms before we ran the other roles. We elicited both before and after realization of the lottery, under the belief that “Cold” ratings from Evaluators, who did not know the realization, might be better predictors of the action of Selectors, who also did not know the outcomes; while “Hot” ratings from Evaluators who knew the realization might better predict Recipients’ adjustments to the Selector bonuses, as the Recipients would know the realization when they made the final decision.

Evaluators rated 24 decisions,⁶ both Options for all 6 Aligned and 6 Unaligned pairs. For the Evaluators, the pairs were ordered; however, Evaluators systematically started in different places in that order. One twelfth saw the unaligned pairs in order and then the aligned pairs in order. Another twelfth saw pairs 2 through 6 of the unaligned, then the aligned pairs in order, and then unaligned pair 1. We continued this pattern. This was done so that, if we needed to cut ratings, which seemed fatigued; we would still have the same number of ratings for each. Hot Evaluators only rated one outcome per decision, so we doubled the number of Evaluators.⁷ All Evaluators were asked to rate each decision on two scales: one, how socially appropriate it was; two, whether the decision deserved punishment or required reward. Both were on

⁵<https://www.ocf.berkeley.edu/~johnlab/bfi.htm> Big Five scores had no predictive power for decision in our study, so are not presented.

⁶We recorded time spent on each page and also checked if the ratings of participants who finished particularly quickly showed a lack of variation across decisions. The authors thank Erin Krupka for these design suggestions. We detected no evidence of decision fatigue, so we used all responses.

⁷We had targeted 240.

seven-point scales between “Extremely Socially Inappropriate” to “Extremely Socially Appropriate” with zero as neither socially appropriate nor inappropriate; and “Deserves Extreme Punishment” to “Require Extreme Reward”. Figure A.6 shows an example of a decision screen. Evaluators could earn up to \$6.00 in bonus payments.⁸ In our final sample, we have 120 Cold Evaluators and 243 Hot Evaluators.

3.2 Selectors

Selectors choose one lottery out of each pair for a total of six decisions. They saw the pairs in a random order. Figure A.7 shows an example decision screen after the Selector had clicked on Option B’s image and revealed its terms. The two default treatments crossed with the three payment treatments yielded six cells. There were 120 Selectors per cell and 720 in total. Selectors were paid \$0.60 for completing the six decisions. Bonus payments varied by treatment. For the Self treatment, in which they made choices for themselves, one pair was randomly selected, and the Selectors were paid a realization of their chosen lottery plus an endowment of \$0.50. For the Other treatment, in which they knew one of their decisions would determine a Recipient’s pay, they were paid the same as a randomly drawn Recipient. For the Consequences treatment, the default bonus pay was \$0.60, but it could be adjusted by the Recipient and varied from \$0.00 to \$1.20.

3.3 Recipients

There were 480 Recipients, one for each Selector (aside from those in the Self treatment). They earned \$0.30 for participating and were paid a bonus of a realization of the lottery selected for them plus an endowment of \$0.50. In the Other treatment, Recipients learned the Selector’s choice and their own pay and then made the same ratings of their Selectors’ decisions as the Evaluators.

In the Consequences treatment, Recipients learned the Selector’s decision and then were asked to select a bonus for the Selector for each of the two possible payment outcomes. The default bonus to the Selector was \$0.60. However, in \$0.20 increments, Recipients could decrease the bonus to \$0.00 or increase it to \$1.20, at a cost to Recipients of one-fifth of the (absolute value of the) adjustment. Figure A.8 shows an example of this decision screen. On the next screen, Recipients saw the realization of the lottery and resulting pay and were allowed to change the adjustment to the bonus. Figure A.9 shows an example of this decision screen. The decision from the previous screen for the lottery realization was pre-checked, but any adjustment could be chosen.

4 Results

We report results for each of the three roles, Evaluators, Selectors, and Recipients, in turn, in the following subsections.

4.1 Evaluators

We analyze the responses from the Evaluators in detail in a related paper (Tracy and Mukherjee, 2026). In the present paper, we simply report mean ratings and differences in ratings between options, which are independent variables in our subsequent analysis of Selector and Recipient decisions (Table 3).

Pair 1 as expected, the stochastically dominant option is viewed as more socially appropriate than the dominated lottery, both when the default is stochastically dominant and when it is not. The stochastically dominant option is very strongly rated as requiring reward, and the stochastically dominated option is rated as deserving punishment. The ratings were stronger when the dominated option was not the default.

Pair 2 we find only weak evidence that the option with the higher expected value is seen as more socially appropriate, and the effect is sensitive to the default. When the higher-EV option is the default, choosing it is more socially appropriate; when it is not the default, sticking with the lower-EV lottery is more socially

⁸One decision was randomly chosen for each Evaluator; for each scale, if their rating matched the modal response, they earned \$3.00.

Table 3: Mean Cold Ratings from Evaluators

Pair	Social Appropriateness						Punishment / Reward					
	Aligned			Unaligned			Aligned			Unaligned		
	α	β	Δ	α	β	Δ	α	β	Δ	α	β	Δ
1	1.76	-0.99	2.75	1.71	-0.41	2.12	1.30	-0.82	2.12	1.46	-0.38	1.84
2	0.74	0.57	0.18	0.40	0.85	-0.45	0.39	0.33	0.06	0.35	0.57	-0.22
3	0.66	0.68	-0.02	0.51	0.72	-0.21	0.32	0.48	-0.16	0.34	0.49	-0.15
4	0.25	0.82	-0.57	-0.24	0.88	-1.12	0.13	0.53	-0.40	-0.25	0.68	-0.92
5	1.41	-0.07	1.48	1.12	-0.19	1.32	0.96	-0.13	1.09	0.95	-0.20	1.15
6	0.34	0.32	0.02	0.36	0.47	-0.12	0.17	0.23	-0.07	0.35	0.14	0.21
In Aligned elicitations, α was the default. In Unaligned elicitations, β was the default. $\Delta = \alpha - \beta$												

appropriate. Both options are very weakly rated as requiring reward. The option which was the default had the higher rating. Differences were negligible.

Pair 3 we find no evidence that reducing variance is more socially appropriate when expected value is held constant. When the higher-variance option is the default, however, keeping it is more socially appropriate than switching to the lower-variance option. Both options are weakly rated as requiring reward. Option α had the higher rating, regardless of default status. Differences were small.

Pair 4 the option with the negative payoff is judged to be less socially appropriate, and this effect is stronger when that option is not the default. In that case, switching to the negative-payoff lottery is especially socially inappropriate. The option with the negative payoff is very weakly rated as requiring reward when it is the default and weakly rated as deserving punishment when an active choice. The other option is rated as requiring reward.

Pair 5 similar to Pair 1, choosing the option with the higher expected value is perceived as more socially appropriate in both aligned and unaligned default conditions. The option with the higher expected value is strongly rated as requiring reward, while the option with the lower expected value is very weakly rated as deserving punishment. Default status had minimal effect on ratings.

Pair 6 we find no clear norm favoring either option. When the default is unaligned, however, choosing the default option appears slightly more socially appropriate. Both options are very weakly rated as requiring reward. The option which was the default had the higher rating. Differences were negligible.

4.2 Selectors

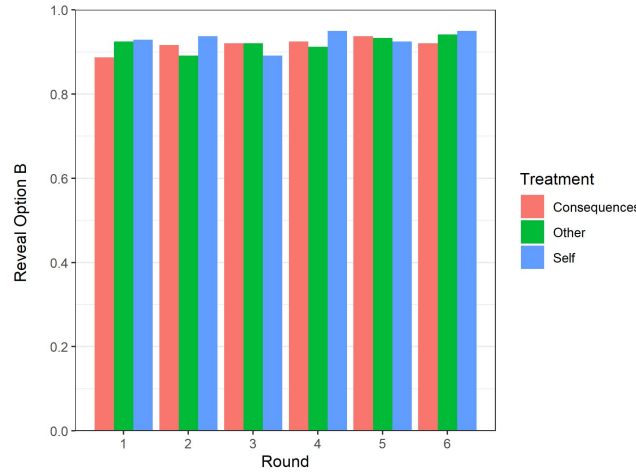
In the two treatments where Selectors makes a choice for the Recipient, standard theory predicts no effort on the part of Selector. Additionally, there is potential for decision fatigue, which may lead Selectors to choose not to reveal Option B.

Figure 1 graphs, by treatment and across rounds, the proportion of Selectors who clicked on the stand-in image of Option B to reveal the true image and lottery terms. There is no decline across rounds, nor is there any difference by treatment. Selectors made just as much effort across all treatments. Prolific provides elapsed time between when a participant accepts an invitation until they submit a completion code.

The mean time for the Self treatment is 187 seconds, whereas mean time for the Consequences and Other treatments is 231 and 226 seconds, indicating Selectors spent more time deliberating decisions that impacted Recipients' pay than ones that impacted their own pay. Wilcoxon rank-sum (Mann-Whitney) tests indicate that all three differences are statistically significant, p -values < 0.001 . This provides evidence against the *homo economicus* prediction of inaction, and in favor of Hypotheses 1a & b

Result 1.a *Consistent to Hypothesis 1.a, Selectors spend more time when making selections for others as for themselves. Time does not substantial increase with the Recipients' ability to impose financial consequences.*

Figure 1: Proportion of Selectors Revealing Option B by Treatment



Result 1.b *Contrary to Hypothesis 1.b, Selectors spend the same amount of effort in revealing the second option when making selections for others as for themselves. Effort does not substantial increase with the Recipients' ability to impose financial consequences.*

Figure 2 illustrates that there is a default effect— α is more likely to be selected if it is the default—for the first four pairs. The opposite is true for the last two pairs. In some cases, the effect is attributable to Selectors not revealing the details of Option B. For example, in Pair 1, the bars that include the “Not revealed” cases are about the same height as the bars to their left, meaning basically all the effect is through the Selectors not revealing Option α . However, there are pairs in which the difference between selection rates remains even after taking into account whether Option B was revealed. For example, in Pair 2, for the Consequences treatment, the not revealed cases could only explain about half the gap. A Wilcoxon rank-sum (p -value = < 0.001) indicates the default is more likely to be chosen. We find that Selectors chose the risk-neutral option in 55.7% of lotteries in the Other treatment, compared with 50.4% in the Consequences treatment. The difference is statistically significant (p -value = 0.004).

Figure 2: Proportion of Selectors choosing Option α by Treatment

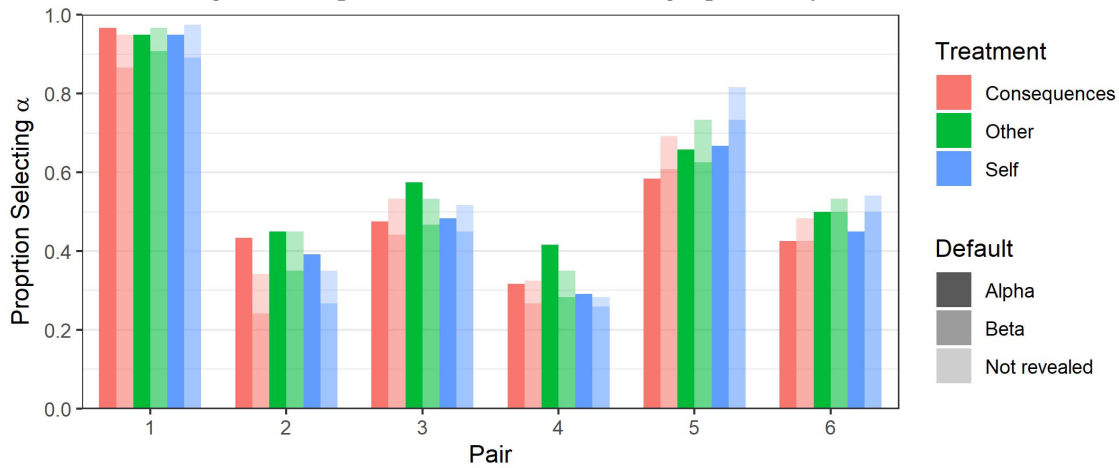
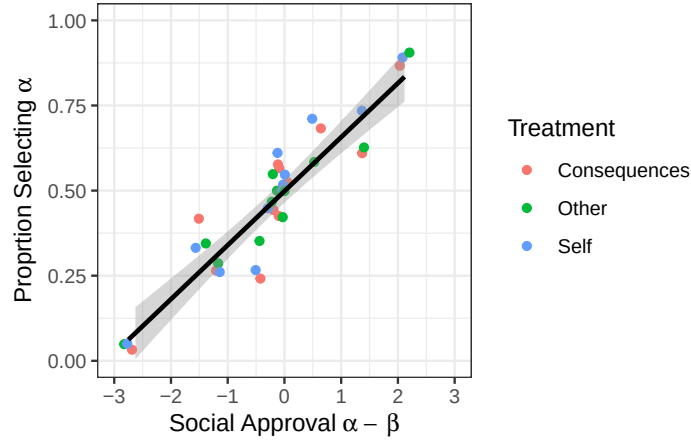


Figure 3 shows that in the aggregate there is a strong relationship between Cold Evaluators' ratings of social approval and the likelihood that an option is chosen by Selectors. It plots the $\alpha - \beta$ differences (Δ s from Table 3) against the proportion of Selectors choosing option α .

Figure 3: Proportion of Selectors' Selection Option α by Social Approval Rating



We next test if this relationship holds with panel logit regressions with standard errors clustered at the individual level. In the Other treatment, the Recipients can neither punish nor reward the Selectors, so the norms regarding punishment and reward are not applicable. Whereas if the Selectors wanted to do right by the Recipients, they would follow norms regarding social appropriateness.

In contrast, in the Consequences treatment, where the Recipients can adjust the Selectors' bonus, norms regarding what required reward or deserved punishment might guide Recipients' bonus decisions, and anticipating this, Selectors' option choice. Whereas concerns about social appropriateness would be secondary. Given this and our preregistered hypotheses, we present separate regressions for the two treatments with distinct independent variables. The overlap in norms permits pooling of the treatments. Estimates from those regressions are similar. h (Please see appendix Table A.3.)

Table 4 reports marginal effects from logit regressions testing how norms for social appropriateness impact the likelihood that a Selector in the Other treatment chooses Option α . Cold ratings (from Evaluators who did not know the lottery realizations) of the social appropriateness of the selection are a good predictor of choices. Higher ratings of Option α make selecting α more likely. Conversely, higher ratings of Option β make selecting α less likely. In specification 2, the ratings of α and β to create a single variable (Δ Columns of Table 3 and the X-Axis of Figure 3); the results are consistent. Of note, the *Aligned* variable is not significant in the two specifications, indicating that the asymmetry in how an option was rated was presented as Option A (the default) versus Option B, the default effect is captured by the norms elicited from the Evaluators.

Table 4: Marginal Effects from Logit Regression on Selectors' Likelihood to Pick Option α in Other Treatment

	(1)		(2)	
Aligned	0.0114	(0.0231)	0.00901	(0.0233)
Social Appropriate, Opt. α	0.136**	(0.0445)		
Social Appropriate, Opt. β	-0.200***	(0.0501)		
Diff in Soc. App. ($\alpha - \beta$)			0.167***	(0.00910)
Observations	1440		1440	
ChiSquared	177.0		178.5	
LogLikelihood	-876.9		-877.2	
NumberIDs	240		240	

Clustered robust Std. Err. in parentheses

* $p < 0.05$, ** $p < 0.01$, *** $p < 0.001$

Table 5 reports marginal effects from logit regressions testing how norms for deserving punishment or requiring reward impact the likelihood that a Selector in the Consequences treatment chooses Option α . The structure is akin to Table 4. Norms are a good predictor of the Selectors' choice in this treatment as well. Higher ratings that Option α requires reward make selecting α more likely. Conversely, higher ratings that Option β requires reward make selecting α less likely. In specification (2), we combine the ratings of α and β into a single difference measure (Δ Columns of Table 3 and X-axis of Figure 3); the results are consistent. As in Table 4, the *Aligned* variable is not statistically significant in either specification, indicating that the asymmetry in whether an option was presented as Option A (the default) or Option B is captured by the elicited norms rather than reflecting an independent default effect.

Table 5: Marginal Effects from Logit Regression on Selectors' Likelihood to Pick Option α in Consequences Treatment

	(1)		(2)	
Aligned	0.0296	(0.0258)	0.0369	(0.0254)
Req. Reward/Dev. Punish, Opt. α	0.128*	(0.0551)		
Req. Reward/Dev. Punish, Opt. β	-0.319***	(0.0603)		
Diff in Req. Rew./Dev. Pun. ($\alpha - \beta$)			0.217***	(0.00968)
Observations	1440		1440	
ChiSquared	256.3		247.9	
LogLikelihood	-867.6		-868.7	
NumberIDs	240		240	
Clustered robust Std. Err. in parentheses				
* $p < 0.05$, ** $p < 0.01$, *** $p < 0.001$				

Result 2.a *Consistent with Hypothesis 2.a, Selectors in the Other treatment are more likely to choose the option that is ranked as more socially appropriate.*

Result 2.b *Consistent with Hypothesis 2.b, Selectors in the Consequences treatment are more likely to choose the option that is judged to deserve a higher reward.*

4.3 Recipients

Figure 4 shows the paucity of instances of options with low social approval being selected, particularly when the norm is strong (the difference in rating of the two options is large). Of the 240 decisions, there were only 39 where the Selector chose an option that rated at least 0.25 worse than the other option. Only 34 were at least 0.5 worse. The lack of many strong norm violations limits our power to test norm enforcement by the Recipients.

Table 6 reports mean adjustments by pair, alignment, option $\{\alpha, \beta\}$, and outcome $\{\text{low, high}\}$. Note that despite that bonus adjustments are costly to the Recipients, there are very few cases where the mean adjustment was zero. More than two-thirds of Recipients (164 of 240) made a costly adjustment. The vast majority (130 of 164) were positive adjustments.

Our hypothesis had been that bonus adjustments would reward norm compliance and punish non-compliance. Figure 5 plots final bonus adjustments versus norm compliance. We use ratings from Hot Evaluators who knew the lottery realization to give norms the best shot. The horizontal axis is the difference between the rating of the option chosen by the Selector and the rating of the option not chosen. Dot size represents the number of Recipients making a particular choice. For example, the left-most dot in the top row represents the two Recipients whose Selector chose the option that was rated 1.46 worse than the alternative, and increased the Selector's bonus by 60¢, whereas the dot immediately to the right represents only a single Recipient, who also increased the bonus by 60¢ to a Selector who chose a worse-rated option. The fit line shows a slight relationship between ratings and bonus adjustment. Finding only weak support for our hypothesis, we explore another explanation.

Figure 4: Frequencies of Selection by Options Social Appropriateness

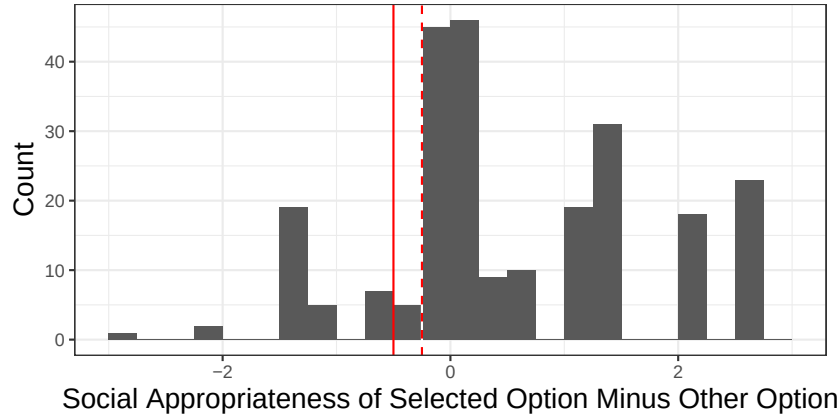


Table 6: Mean Bonus Adjustment
Unaligned Aligned

Pair	α		β		α		β	
	Low	High	Low	High	Low	High	Low	High
1	-2.86	29.09	0.00	20.00	20.00	28.18	-20.00	
2	-20.00	20.00	0.00	30.00	-30.00	16.67	20.00	16.00
3	-40.00	18.18	-8.57	10.00	10.00	30.00	13.33	22.86
4	-20.00	20.00	-4.00	8.57	-30.00	40.00	0.00	10.00
5	11.67	20.00	-10.00	22.22	6.67	40.00	0.00	24.00
6	4.00	40.00	-6.67	25.00	9.09	40.00	-10.00	8.57

Figure 5: Relationship between Bonus Adjustment and Hot Ratings

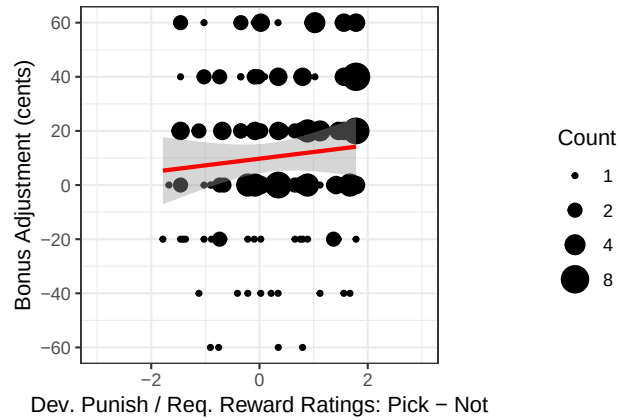


Figure 6 explores the relationship between the realization of the lottery and the bonus the Recipient allocates to the Selector. It has the same vertical axis as Figure 5; however, the horizontal axis is the realization of the lottery the Selector chose for the Recipient (the Recipient's additional earnings from the lottery). The fit line reveals a much stronger relationship. Recipient bonus decisions to Selectors depend more on realizations of the lotteries (over which the Selector has no influence) than on which lottery the Selector chose. For example, in Pair 1 Unaligned, two Selectors selected β , the dominated default option, in the case in which the realization was the higher prize the principal sent a bonus of 20¢; in the case the realization was the low

prize, the principal did not adjust the bonus.⁹

Figure 6: Relationship between Bonus Adjustment and Lottery Realization

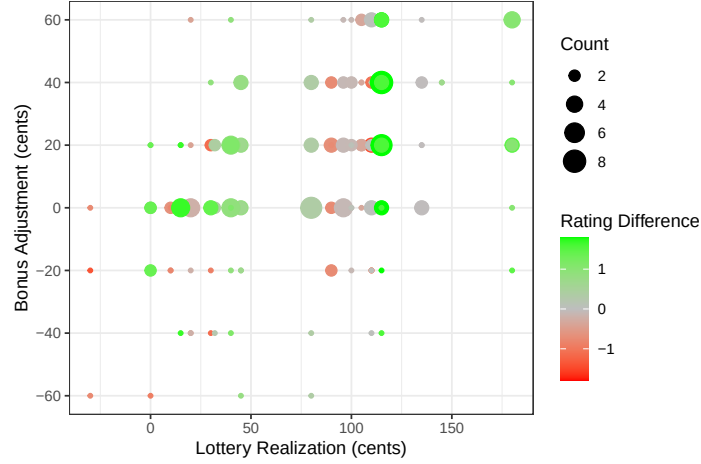


Table 7 reports regressions of Recipients' implemented bonus adjustments. Positive values correspond to bonuses, while negative values correspond to punishments. The first two variables are based on Evaluators' Hot ratings of whether the selected option deserves punishment or requires reward. *Lottery Prize* is the realized value of the selected lottery, and *Selected B* indicates that the Selector chose the non-default option, i.e., made an active choice. Column (1) shows that both normative assessments are related to Recipients' bonus decisions. Actions judged to deserve punishment are associated with lower bonus adjustments, while actions judged to require reward are associated with higher adjustments. The punishment estimate is larger, suggesting that actions viewed as deserving punishment have a stronger influence on Recipients' decisions, while the reward estimate is more precisely estimated. Column (2) adds the realized value of the lottery. Once the realized outcome is included, neither norm variable remains statistically significant, while the lottery prize strongly predicts bonus adjustments. This suggests that although normative assessments have some explanatory power, realized outcomes dominate Recipients' decisions. Column (3) further controls for whether the Selector chose the non-default option. The coefficient on *Selected B* is not statistically significant, and the remaining estimates are largely unchanged. Column (4) confirms the coefficient on *Selected B* is not statistically significant even when it is the sole regressor. Overall, Recipients' bonus decisions are driven more by realized outcomes than by Evaluators' assessments of whether the selected option deserved reward or punishment.

Result 3.a *Contrary to Hypothesis 3.a, lottery realizations have a more profound impact on bonuses to Selectors than social norms regarding Selectors' actions*

Result 3.b *Contrary to Hypothesis 3.b, lottery realizations have a more profound impact on bonuses to Selectors than whether the Selector made an active choice.*

⁹There were more (16) times Selectors made this choice; however, only two were randomly selected for payment.

Table 7: Regression on Recipients Hot Bonuses to Selectors

	(1)	(2)	(3)	(4)
Deserves Punishment, Pick Hot	-2.439* (0.958)	-1.281 (0.833)	-1.281 (0.851)	
Requires Reward, Pick Hot	0.670*** (0.172)	0.191 (0.186)	0.190 (0.185)	
Lottery Prize		1.071*** (0.182)	1.070*** (0.182)	
Selected B			-0.0894 (0.158)	-0.109 (0.174)
Constant	0.149 (0.180)	-0.295 (0.181)	-0.251 (0.194)	0.764*** (0.118)
Observations	240	240	240	240
Rsquared	0.103	0.205	0.206	0.00165

Standard errors in parentheses

* $p < 0.05$, ** $p < 0.01$, *** $p < 0.001$

5 Discussion

In this paper, we study how social norms shape the decisions people make on behalf of others when the Recipient’s preferences are unknown. We study two environments: one in which the Recipient cannot punish or reward the Selector for the choice made for them, and one in which the Recipient can punish or reward the Selector at a cost. We find that when Recipients have no way to reward or punish the Selector’s choice, Selectors choose the option that Evaluators rated as more socially appropriate. When Recipients can punish, Selectors choose the option that Evaluators rated as deserving a reward. Overall, social norms predict the Selector’s choices. Choices in the Other treatment are on average more risk-neutral than in the Consequences treatment, and the treatment with enforcement leads on average to lower total social welfare.

Because Selectors chose the options rated more socially appropriate, Recipients had few opportunities to enforce the norm. Recipients’ decisions to reward or punish were driven by the outcome of the lottery. We cannot conclude with certainty that outcome realization would remain the primary driver of bonus decision when there was greater need for norm enforcement. But this leaves us with an interesting dichotomy: Selectors make the more risk-averse choice for the Recipient in the Consequences treatment, and the Recipients tend to reward the Selector for landing the higher lottery prize.

This dichotomy points to something deeper about delegation. Norms govern the choice before the outcome is known, but consequences are handed out after. The Selector is judged on a decision made under uncertainty, while the Recipient reacts to a result that luck largely determined. Future research could explore independent evaluator approval of Recipients’ decisions.

We think this helps explain familiar behavior outside the lab. A friend picks the safe, familiar restaurant and passes up a new experience. A professor who is concerned about being penalized during student evaluations might assign readings that are a bit dated but known to work. In contrast, a professor who is less concerned with student evaluations might assign a “riskier” reading. Similarly, a career official might select a bolder, more effective policy than an elected official who might play it safe, fearing the consequences of any unfavorable outcomes in the next election cycle. A physician who fears being blamed for a bad outcome may lean toward the conservative treatment, even when the bolder option would, on average, serve the patient better. A financial advisor may hold a client in safe assets for the same reason.

In each case, the Selector picks the safer option, which can exclude a chance at the highest possible payoff for the Recipient. The goal is to keep the Recipient from ending up with a bad result and punishing the Selector for that bad result. The Recipient rewards the outcome, not the intention: the biggest rewards go to the choices that happen to land a high payoff. So the Selector’s aim to play it safe and the Recipient’s

reward, based on the realized outcome, can end up working against each other. The central result is clear. When there is no contract and no way to learn what the other person wants, shared norms give Selectors a standard to choose by, and they use it. That is where norms do their work: in guiding the choice, but not in rewarding or punishing norm compliance and deviation.

Our study contributes to the literature on social norms by eliciting norms in a setting where actions are not deterministic, which creates a wedge between *intention* and *outcome*. Norm compliance predicts behavior well when it comes to intention, but we find little evidence that norms are enforced once outcomes are realized.

It also extends the literature on decision-making for others by offering a channel beyond risk preferences alone. We find that the context in which a person chooses for others matters, and that norms help explain the Selector's choices. Asking for bonus adjustment conditional on lottery outcomes before revealing the outcome allowed us to gather additional data points, but could have either dampened responses to the outcome or created a perception of experimenter demand to have bonuses based on lottery outcome. Future designs could vary by treatment the timing of when bonus decisions are made. Decisions made before the outcome was revealed might reward norm compliance. If our results were driven by a perception of experimenter demand, the post-realization treatment would correct for that. Alternatively, this design could increase the relationship between prize and bonus.

Finally, our experimental environment is intentionally stylized. Participants are assigned to roles without contextual framing, history, or relational cues. In many real-world settings, however, choices for others occur within richer institutional and social contexts. Selectors may be agents operating under incomplete contracts, intermediaries subject to reputational concerns, or individuals performing a favor rather than exercising formal authority. Introducing such context may strengthen, weaken, or fundamentally reshape how norms guide the selection process. Context might also impact the enforcement process. A conscious rather than random role assignment and even partial guidance might mean greater penalties for deviation. Exploring how contextual framing interacts with norm adherence is a promising direction for future research. Another potential avenue of investigation is how partial instructions impact norms for unanticipated situations. For example, what if there are instructions not to take a 0.75p of \$1.50 (vs. \$1, but none for 0.75p of \$2; would the instruction impact the norm?

Another implication of making these choices in an intentionally stylized environment is that we specifically selected participants living in the United States. Norms are shaped by culture, so the judgments we elicit here need not hold elsewhere. Participants in other countries, working from their own norms, might coordinate on very different standards. Testing whether a model of general norms holds across cultural settings is a promising direction for future research.

Given the paucity of observation for the Recipient to enforce the norms. Future research could increase the number of selections deviating from the norm (perhaps through increasing the Selectors' effort cost) or increase the power to detect norm enforcement (perhaps by decreasing the variation in lottery outcomes).

Taken together, our findings suggest that norms can meaningfully influence behavior even when they are not enforced and when violations carry no material consequences. Understanding when norms bind, when they unravel, and how they are interpreted by different parties remains crucial for explaining behavior in organizations, contracts, and informal economic interactions.

References

- Agranov, Marina, Alberto Bisin, and Andrew Schotter**, "An experimental study of the impact of competition for other people's money: the portfolio manager market," *Experimental Economics*, 2014, 17 (4), 564–585.
- Andreoni, James and B Douglas Bernheim**, "Social image and the 50–50 norm: A theoretical and experimental analysis of audience effects," *Econometrica*, 2009, 77 (5), 1607–1636.
- Baron, Jonathan and John C Hershey**, "Outcome bias in decision evaluation," *Journal of Personality and Social Psychology*, 1988, 54 (4), 569–579.
- Barrafrem, Kinga and Jan Hausfeld**, "Tracing risky decisions for oneself and others: The role of intuition and deliberation," *Journal of Economic Psychology*, 2020, 77, 102188.
- Bolton, Gary E and Axel Ockenfels**, "Betrayal aversion: Evidence from brazil, china, oman, switzerland, turkey, and the united states: Comment," *American Economic Review*, 2010, 100 (1), 628–33.
- , —, and **Julia Stauf**, "Social responsibility promotes conservative risk behavior," *European Economic Review*, 2015, 74, 109–127.
- Chakravarty, Sujoy, Glenn W Harrison, Ernan E Haruvy, and E Elisabet Rutström**, "Are you risk averse over other people's money?," *Southern Economic Journal*, 2011, 77 (4), 901–913.
- Charness, Gary and Matthew O Jackson**, "The role of responsibility in strategic risk-taking," *Journal of Economic Behavior & Organization*, 2009, 69 (3), 241–247.
- , **Eugen Dimant, Uri Gneezy, and Erin Krupka**, "Experimental methods: Eliciting and measuring social norms," *Journal of Economic Behavior & Organization*, 2025, 237, 107187.
- Chen, Daniel L., Martin Schonger, and Chris Wickens**, "oTree—An open-source platform for laboratory, online, and field experiments," *Journal of Behavioral and Experimental Finance*, March 2016, 9, 88–97.
- Dore, Rebecca, Eric R Stone, and Christy M Buchanan**, "Decisions for others involving risk," [Journal], 2014.
- Eriksen, Kristoffer W, Ola Kvaløy, and Miguel Luzuriaga**, "Risk-taking on behalf of others," *Journal of Behavioral and Experimental Finance*, 2020, 26, 100283.
- Festré, Agnès and Pierre Garrouste**, "Somebody may scold you! A dictator experiment," *Journal of Economic Psychology*, 2014, 45, 141–153.
- Görges, Lisa and Daniele Nosenzo**, "Measuring Social Norms in Economics: Why It Is Important and How It Is Done," *Analyse & Kritik*, 2022, 44 (2), 285–314.
- Harrison, Glenn W, Morten I Lau, E Elisabet Rutström, and Marcela Tarazona-Gómez**, "Preferences over social risk," *Oxford Economic Papers*, 2013, 65 (1), 25–46.
- Jensen, Michael C and William H Meckling**, "Theory of the firm: Managerial behavior, agency costs and ownership structure," *Journal of Financial Economics*, 1976, 3 (4), 305–360.
- Kimbrough, Erik O and Alexander Vostroknutov**, "Norms make preferences social," *Journal of the European Economic Association*, 2016, 14 (3), 608–638.
- Krupka, Erin L and Roberto A Weber**, "Identifying social norms using coordination games: Why does dictator game sharing vary?," *Journal of the European Economic Association*, 2013, 11 (3), 495–524.
- Luzuriaga, Miguel et al.**, "Taking Risk with Other People's Money: Does Information about the Others Matter?," *Review of Behavioral Economics*, 2017, 4 (2), 107–133.
- Pahlke, Julius, Sebastian Strasser, and Ferdinand M Vieider**, "Risk-taking for others under accountability," *Economics Letters*, 2012, 114 (1), 102–105.

- Pollmann, Monique MH, Jan Potters, and Stefan T Trautmann**, “Risk taking by agents: The role of ex-ante and ex-post accountability,” *Economics Letters*, 2014, 123 (3), 387–390.
- Polman, Evan**, “Self–other decision making and loss aversion,” *Organizational Behavior and Human Decision Processes*, 2012, 119 (2), 141–150.
- **and Kaiyang Wu**, “Decision making for others involving risk: A review and meta-analysis,” *Journal of Economic Psychology*, 2020, 77, 102184.
- Rubin, Jared and Roman Sheremeta**, “Principal–Agent Settings with Random Shocks,” *Management Science*, April 2016, 62 (4), 985–999.
- Smith, Adam**, *The Theory of Moral Sentiments*, H. G. Bohn, 1759. Google-Books-ID: 96Wpjt4H4wC.
- Smith, Vernon L.**, “Trust, reciprocity, and social history: New pathways of learning when max U (own reward) fails decisively,” 2020.
- Smith, Vernon L and Bart J Wilson**, “Equilibrium play in voluntary ultimatum games: Beneficence cannot be extorted,” *Games and Economic Behavior*, 2018, 109, 452–464.
- Stone, Eric R and Lauren Allgaier**, “The social values hypothesis,” *[Journal]*, 2008.
- Tracy, J. Dustin and Prithvijit Mukherjee**, “Norms of Norms,” 2026. Working paper.
- Vostroknutov, Alexander**, “Social norms in experimental economics: Towards a unified theory of normative decision making,” *Analyse & Kritik*, 2020, 42 (1), 3–40.

Appendices

A.1 Software Screen Shots

Figure A.1: Instructions for Evaluators

Instructions

Only in 'Hot' Treatment

In this survey you are being asked to rate 24 decisions **with outcomes**. All the decisions are being made by other participants in the survey, and will be used to determine bonus payments for participants in another role. For each decision, please rate that decision for how socially acceptable it is and whether it deserves punishment or requires reward. In each decision, Participant #2 starts with Option A. The terms of A are known to Participant #1. Participant #1 can do nothing and opt for A. Discovering the terms of Option B requires clicking on its image. Participant #2 will be paid \$0.50 plus the realization by a random draw of whichever option Participant #1 picks.

You will be paid the completion fee if you make all 24 decisions. In addition, after everyone has completed the survey, we will randomly pick one decision for each participant rated (including you) for payment. We will compare your rating of that decision to the most popular (modal) rating of that decision for everyone completing the survey. If your rating of appropriateness and how deserving punishment or requiring reward is the same as the average rating, you will be paid \$3.00. If you match both, you will be paid \$6.00.

This means you will earn more if your ratings are closer to what you think other people would respond.

Next

Figure A.2: Instructions for Selector

Instructions

In this survey, you are being asked to make 6 decisions for another participant. In each decision, the other participant starts with \$0.50 and Option A. You know the terms of Option A and can pick Option A by clicking on its button. You can also switch the other participant to Option B. However, you need to click on its image to see its terms before you can choose it. After seeing the terms of Option B you can choose either option. The payout of both options depends on a random draw.

You will be paid the completion fee if you make all 6 decisions. After everyone has completed the survey, we will randomly pick one decision for each participant (including you) for payment. We will pay the other participant based on whatever option you choose and the corresponding random draw.

You start with a bonus of \$0.60. After the other participant sees your decision (and resulting pay after the random draw) for him or her, he or she can penalize or reward you. It cost the other participant \$0.04 for each \$0.20 of reward or penalty. After rewards or penalty your bonus can be between \$0.00 and \$1.20.

Not in No Cosequences Treatment

Next

Figure A.3: Instructions for Recipients in Consequences Treatment

Instructions

In this survey, another participant made 6 decisions, knowing one of them would be randomly selected to determine your earnings. In each decision, you start with \$0.50 and Option A. The other participant knows the terms of Option A and can pick Option A by clicking on its button. He or she can also switch you to Option B. However, he or she needs to click on its image to see its terms before he or she can choose it for you. After seeing the terms of Option B you he or she choose either option. The payout of both options depends on a random draw.

One decision was randomly picked for each participant (including you) for payment. We will pay you based on whatever option (A or B) the other participant chose and the corresponding random draw. You will see both options and the other participant's choice. There are two possible outcomes for each choice.

The participant who made a choice for you starts with a bonus of \$0.60. You can penalize or reward that participant. It cost you \$0.04 for each \$0.20 of reward or penalty. After rewards or penalty the other participant's bonus can be between \$0.00 and \$1.20. After you see the other participant's decision for you (but before you learn the outcome), we will ask you to choose a reward or penalty for both outcomes. After you see the outcome and your resulting pay, you can change the penalty or reward.

Next

Figure A.4: Instructions for Recipients in Other treatment

Instructions

In this survey, another participant made 6 decisions, knowing one of them would be randomly selected to determine your earnings. In each decision, you start with \$0.50 and Option A. The other participant knows the terms of Option A and can pick Option A by clicking on its button. He or she can also switch you to Option B. However, he or she needs to click on its image to see its terms before he or she can choose it for you. After seeing the terms of Option B you he or she can choose either option. The payout of both options depends on a random draw.

One decision was randomly picked for each participant (including you) for payment. We will pay you based on whatever option (A or B) the other participant chose and the corresponding random draw. You will see both options and the other participant's choice. There are two possible outcomes for each choice.

After you see the other participant's decision for you (but before you learn the outcome), we will ask you to rate that decision for how socially acceptable it is and whether it deserves punishment or requires reward. After you see the outcome and your resulting pay, you can change how you rate that decision.

Next

Figure A.5: Example of Strong Default

Example decision screen

This an example of Participant 1's decision screen.

Participant 1 can select Option A immediately, but cannot select Option B without clicking on the question mark. After revealing Option B, either option can be selected.

Please click on the question mark.

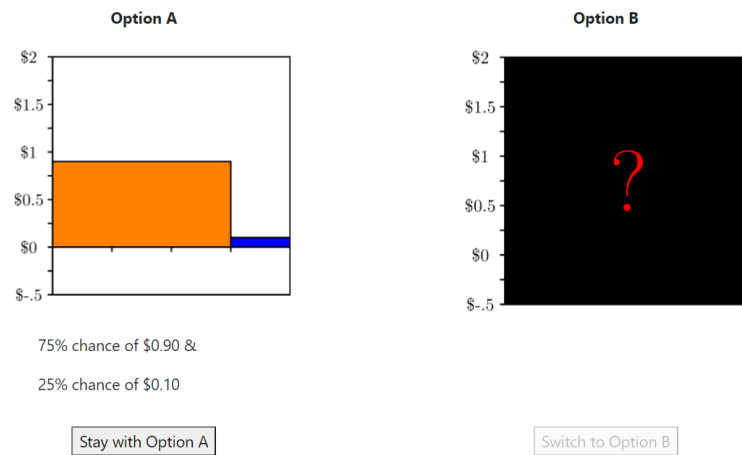


Figure A.6: Screen for Evaluators to Rate Selector Decision

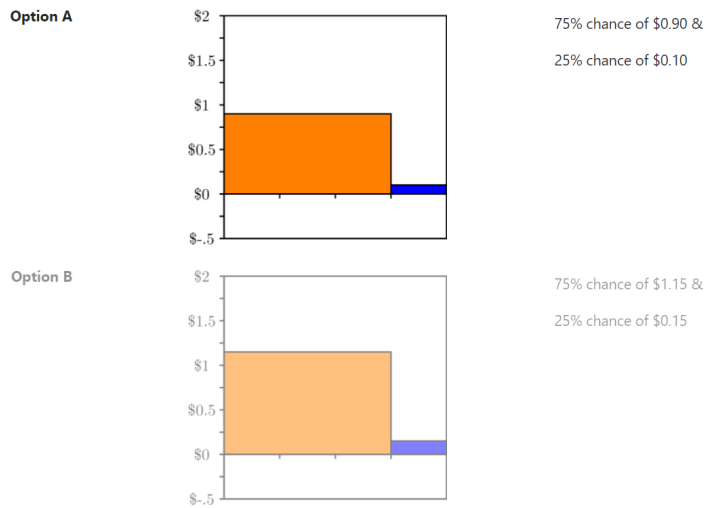
Please rate Participant #1's decision.

Participant #2 started with Option A and an endowment of \$0.50.

Participant #1 did NOT switch Participant #2 to Option B.

Only for 'hot' not 'cold' Judges

As a result of this decision and a random draw, Participant #2 will earn an extra \$0.90.



Please rate social appropriateness of Participant #1's decision.

Socially Inappropriate			Neither	Socially Appropriate		
Extremely	Moderately	Slightly		Slightly	Moderately	Extremely
☐	☐	☐	☐	☐	☐	☐

Please rate the extent to which Participant #1's decision deserves punishment or requires reward.

Deserves Punishment			Neither	Requires Reward		
Extreme	Moderate	Slight		Slight	Moderate	Extreme
☐	☐	☐	☐	☐	☐	☐

Next

Figure A.7: Selector Decision Screen, Pair 5

The other participant starts with \$0.50 and Option A but you can switch him or her to Option B

Option B cannot be chosen until you click on the question mark.

If this decision is chosen for payment, the other participant will be paid based on whatever option you choose and the corresponding random draw.

Only displayed in Consequences Treatment

You start with a bonus of \$0.60. After the other participant sees your decision (and resulting pay after the random draw) for him or her, he or she can penalize or reward you. It cost the other participant \$0.04 for each \$0.20 of reward or penalty. After rewards or penalty your bonus can be between \$0.00 and \$1.20.

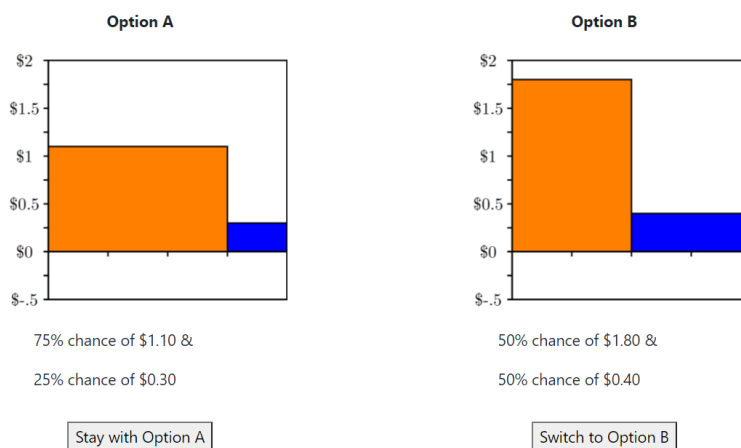
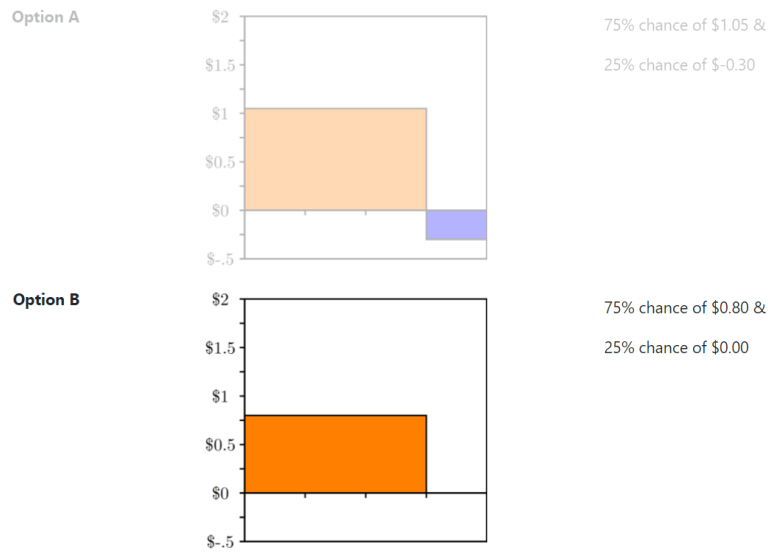


Figure A.8: Recipient Cold Bonus Amount Decision Screen

Results

You started with Option A and \$0.50

The other person choose Option B, therefore you will be paid based on Option B.



If your pay from the Option is \$0.80,
how much would you reward or penalize the person who choose which Option (A or B) you would get?

	Penalty				Reward		
to other Person	-60¢	-40¢	-20¢	0¢	+20¢	+40¢	+60¢
	☐	☐	☐	☐	☐	☐	☐
Cost to you	-12¢	-8¢	-4¢	0¢	-4¢	-8¢	-12¢

If your pay from the Option is \$0.00,
how much would you reward or penalize the person who choose which Option (A or B) you would get?

	Penalty				Reward		
to other Person	-60¢	-40¢	-20¢	0¢	+20¢	+40¢	+60¢
	☐	☐	☐	☐	☐	☐	☐
Cost to you	-12¢	-8¢	-4¢	0¢	-4¢	-8¢	-12¢

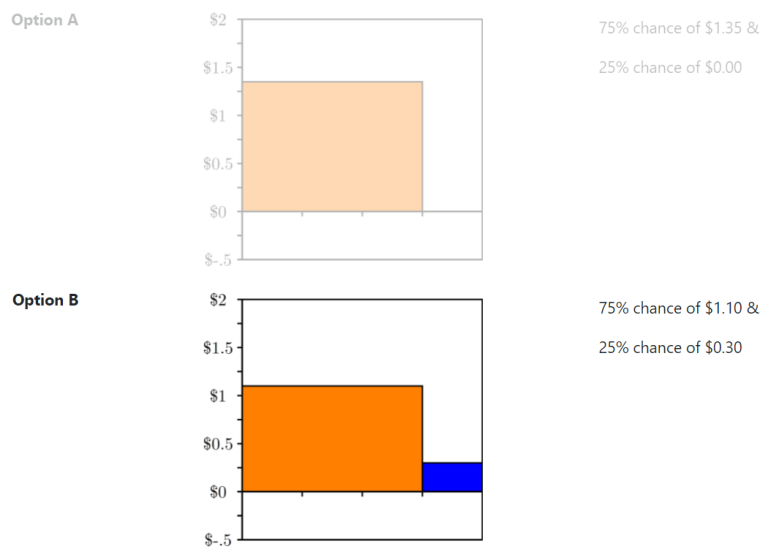
Next

Figure A.9: Recipients Hot Bonus Amount Decision Screen

Results

You started with Option A and \$0.50

The other person choose Option B, therefore you will be paid based on Option B.



The random number drawn for you was 30 (out of 100), so your pay from the Option is \$1.10.

Your choice for this outcome is below. You can adjust your choice to reward or penalize the person who choose which Option (A or B) you would get.

	Penalty				Reward		
to other Person	-60¢	-40¢	-20¢	0¢	+20¢	+40¢	+60¢
	☐	☐	☒	☐	☐	☐	☐
Cost to you	-12¢	-8¢	-4¢	0¢	-4¢	-8¢	-12¢

Next

A.2 Participant Demographics

Table A.1: Participant Characteristics

Age	Mean	SD	Obs
	31.52	12.60	1534
Sex	Female	Male	Prefer not to say
	1033	505	7

Table A.1 reports the age and sex according to data available from subjects' Prolific profiles. Not all information was available from all profiles.

A.3 Additional Results

A.3.1 Evaluators

Table A.2: Mean Reward Punishment Rating (Hot)
Unaligned Aligned

Pair	α		β		α		β	
	Low	High	Low	High	Low	High	Low	High
1	1.24	1.88	-0.43	0.32	0.99	1.70	-0.59	-0.08
2	-0.47	1.10	0.94	1.13	-0.23	0.96	0.68	0.98
3	0.65	0.75	0.43	0.84	0.56	0.92	0.15	0.82
4	-0.84	0.58	0.53	0.92	-0.31	0.44	0.45	0.78
5	1.02	1.71	0.13	0.25	1.18	1.39	0.06	0.37
6	0.86	1.08	0.20	0.40	0.78	1.03	-0.02	0.30

A.3.2 Selectors

Table A.3: Marginal Effects from Logit Regression on Selectors' Likelihood to Pick Option α

	(1)	(2)	(3)	(4)
Aligned	0.00579 (0.0173)	0.00147 (0.0175)	0.0367 (0.295)	0.0412* (0.0170)
Consequences	-0.0518** (0.0169)	-0.0519 (0.420)	-0.0518 (0.424)	-0.0518** (0.0170)
Social Appropriate, Opt. α	0.118*** (0.0304)			
Social Appropriate, Opt. β	-0.232*** (0.0323)			
Diff in Soc. App. ($\alpha - \beta$)		0.173 (1.419)		
Req. Reward/Dev. Punish, Opt. α			0.158 (1.311)	
Req. Reward/Dev. Punish, Opt. β			-0.279 (2.314)	
Diff in Req. Rew./Dev. Pun. ($\alpha - \beta$)				0.214*** (0.00737)
Observations	2880	2880	2880	2880
Chi Squared	496.7	504.3	491.2	494.2
Log Likelihood	-1741.1	-1742.5	-1744.0	-1744.8
Number of IDs	480	480	480	480

Clustered robust Std. Err. in parentheses

* $p < 0.05$, ** $p < 0.01$, *** $p < 0.001$

A.3.3 Recipients

Table A.4: Mean Reward Punishment Rating (Hot) from Recipients
Unaligned Aligned

Pair	α		β		α		β	
	Low	High	Low	High	Low	High	Low	High
1.00	1.00	2.00			0.00	1.25		-1.00
2.00	-0.50	1.73	-0.25	1.33	-0.25	2.00	0.00	1.11
3.00	0.00	0.17	0.00	0.67	0.67	1.00	1.33	1.40
4.00		1.67	0.40	1.80	-2.00	1.60	-1.25	1.73
5.00	1.00	1.33	-1.00	0.00	-0.14	1.00	0.00	1.40
6.00	0.12	1.62	0.50	1.00	0.33	0.25	1.00	2.00