

NBER WORKING PAPER SERIES

THE MARGINAL DISUTILITY FROM CORRUPTION IN SOCIAL PROGRAMS:
EVIDENCE FROM PROGRAM ADMINISTRATORS AND BENEFICIARIES

Arya Gaduh
Rema Hanna
Benjamin A. Olken

Working Paper 30905
<http://www.nber.org/papers/w30905>

NATIONAL BUREAU OF ECONOMIC RESEARCH
1050 Massachusetts Avenue
Cambridge, MA 02138
January 2023

We thank our partners at Bappenas and the Ministry of Social Affairs for their support and assistance. We thank our outstanding team for their work on this study, in particular, Colley Windya and Alexa Weiss for excellent research management and Robert Dulin, Michelle Han, Nikhil Kumar, Daniela Paz, Indriani Pratiwi, and Ignatius Sebastian for wonderful research assistance. We thank Pia Raffler and Stefanie Stantcheva for helpful comments. Funding from the Australian Department of Foreign Affairs and Trade and the J-PAL Innovation in Government Initiative is gratefully acknowledged. This randomized control trial was registered in the American Economic Association Registry for randomized control trials under trial AEARCTR-0007010. The views expressed here are those of the authors, and do not necessarily reflect those of the many individuals or organizations acknowledged here or the National Bureau of Economic Research.

At least one co-author has disclosed additional relationships of potential relevance for this research. Further information is available online at <http://www.nber.org/papers/w30905>

NBER working papers are circulated for discussion and comment purposes. They have not been peer-reviewed or been subject to the review by the NBER Board of Directors that accompanies official NBER publications.

© 2023 by Arya Gaduh, Rema Hanna, and Benjamin A. Olken. All rights reserved. Short sections of text, not to exceed two paragraphs, may be quoted without explicit permission provided that full credit, including © notice, is given to the source.

The Marginal Disutility from Corruption in Social Programs: Evidence from Program Administrators and Beneficiaries

Arya Gaduh, Rema Hanna, and Benjamin A. Olken

NBER Working Paper No. 30905

January 2023

JEL No. D73,I38,O15

ABSTRACT

Concerns about fraud in welfare programs common arguments worldwide against such programs. We conducted a survey experiment with over 28,000 welfare program administrators and over 19,000 beneficiaries in Indonesia to elicit the ‘marginal disutility from corruption,’ i.e., the trade-between more generous social assistance and losses due to corruption and fraud. Merely mentioning corruption reduced perceived program success, equivalent to distributing more than 20 percent less aid. However, respondents were not sensitive to the amount of corruption—respondents were willing to trade off\$2 of additional losses for an additional \$1 distributed to beneficiaries. Program administrators and beneficiaries had similar assessments.

Arya Gaduh
Department of Economics
University of Arkansas
402 Business Building
Fayetteville, AR 72701
and NBER
AGaduh@walton.uark.edu

Benjamin A. Olken
Department of Economics, E52-542
MIT
77 Massachusetts Avenue
Cambridge, MA 02139
and NBER
bolken@mit.edu

Rema Hanna
Kennedy School of Government
Harvard University
79 JFK Street
Cambridge, MA 02138
and NBER
Rema_Hanna@hks.harvard.edu

I. INTRODUCTION

A commonly voiced concern about redistributive programs is that funds go missing, due to some combination of leakages and fraud. In the United States, Ronald Reagan, for example, repeatedly talked about “open-ended welfare programs riddled with fraud and inefficiency,” and in his campaign speeches, he frequently came back to a story about a “welfare queen” who “used 80 names, 30 addresses, 15 telephone numbers to collect food stamps, Social Security, veterans’ benefits for four nonexistent deceased veteran husbands, as well as welfare.” (Black and Sprague, 2016). Perhaps due to this type of rhetoric, more than half of Americans in polls believe it is common for people to lie about their eligibility to receive food assistance and other benefits, far exceeding estimates of actual fraud rates (Sanders, 2019).

These concerns are not limited to the United States. For example, in Mexico, President López Obrador cancelled *Prospera*—the latest iteration of the well-regarded *Progres*a program—having “blasted the program for corruption, without providing evidence, and even though a government policy watchdog concluded that the program improved overall student performance” (Ali-Habib and Lopez, 2022). In Pakistan, the then-Chairperson of the country’s leading social assistance program, Dr. Sania Nishtar, established a “Cyber Crime Wing... for the purpose of preventing and controlling incidents where criminal elements commit fraud using digital technology” in the program, vowing that “those involved in the fraud will be dealt with an iron hand.”¹

Policies put in place to address fraud and corruption, however, often face a tradeoff: increasingly strict rules to reduce leakage and fraud may hamper the distribution of assistance to legitimate beneficiaries.² In the Indian state of Jharkhand, for example, requiring biometric authentication for beneficiaries meant that poor, eligible individuals who did not yet have a linked biometric card were 50 percent less likely to be able to receive the benefits they were entitled to (Muralidharan, Niehaus, and Suktankar, forthcoming). Concerns about fraud and the types of people who access social programs can also lead to increasing rules (e.g., criminal checks, drug testing) that may increase stigma around the programs, which may discourage those who need the program from applying. And, more broadly, concerns about fraud can undermine political support, and lead to programs that are also designed to give out less assistance to fewer people.³

In the presence of this tradeoff, decisions about how much weight to put on fraud reduction relative to making programs more generous requires understanding preferences on this tradeoff. That is, suppose that by implementing a new screening technology, or adding a new audit system or administrative check, one can reduce fraud and corruption by \$1, but in doing so, would reduce actual benefits given to beneficiaries by \$2. The decision about whether to implement such a technology or policy depends on social preferences on exactly this tradeoff, which we term the *marginal disutility from corruption*.

¹ <http://www.heartfile.org/dr-sania-nishtar-directs-the-formation-of-cyber-crime-wing-for-ehsaas/?print=print>

² Kleven and Kopczuk (2011) theoretically explore these issues, describing complexity of enrollment processes as a policy tool that the government may care about, and depending on priorities, may choose to reduce enrollment of those who are eligible (i.e., generate exclusion error) in order to reduce inclusion errors.

³ For example, Moro-Egido and Solano-García (2020) show in a correlational analysis that perceived levels of fraud in welfare programs are negatively associated with support for higher taxation and the welfare state.

This paper seeks to measure this tradeoff. To do so, we designed a survey experiment which we implemented with over 45,000 respondents located across Indonesia. The survey experiment was embedded into a large cross-sectional online survey we conducted in cooperation with the Indonesian government during the COVID-19 pandemic to measure changes in economic conditions during the crisis. Respondents consisted of over 28,000 local program administrators of a targeted government assistance program and over 19,000 randomly selected beneficiaries of the program.

In the experiment, respondents were asked to consider a *hypothetical* new government assistance program that provided cash transfers to poor households. Respondents were told what fraction of the assistance reached beneficiaries, what fraction was “missing”, and what fraction remained unspent (these fractions summed to 100 percent). Respondents were also told how happy beneficiaries were with the program. Respondents were then asked to assess, using a 10-point Likert scale, the overall ‘success’ of the program. By randomly varying the percentage of the program that reached beneficiaries and was lost to corruption and fraud—i.e., what fraction was “missing”—as well as beneficiary happiness, we can estimate preferences over these tradeoffs. Specifically, we can estimate indifference curves; that is, what is the tradeoff between delivering more assistance to beneficiaries and additional corruption that holds respondents’ overall assessment of the program constant. This ratio is an estimate of the marginal disutility from corruption.

Using this approach, we have three main findings.

First, we find that the marginal disutility from corruption is surprisingly small. Respondents did, indeed, view programs with more missing funds more negatively. But quantitatively, the degree to which this affected their assessment of program performance was small when compared to more aid reaching beneficiaries. Specifically, we estimate that the marginal disutility from corruption—the ratio between disutility from an additional \$X lost to corruption and additional \$1 distributed to beneficiaries that holds utility constant—is 0.45 on average. This means that, on the margin, marginally changing a program to distribute \$1 more to beneficiaries, even if it led to \$2 more missing funds, would still be considered an improvement. The results are similar if we benchmark corruption to the fraction of beneficiaries that are satisfied with the program results rather than the fraction of the population that receives benefits.

Second, we find that while respondents are relatively insensitive to the *amount* of corruption in a program, they are quite sensitive to the mere *presence* of corruption. One experimental treatment did not mention missing funds and was silent on what happened to the rest of the funds that did not reach beneficiaries. Comparing this treatment to the treatments that did mention missing funds shows that the mere mention of missing funds reduced program assessment substantially. Indeed, to benchmark the magnitude of this decline, it is worth noting that the negative effect of mentioning corruption on the perceived program success is equivalent to decreasing the amount distributed to beneficiaries by 20 percentage points on average in the sample.

Third, we investigate the differences in preferences between program administrators and program beneficiaries. Stricter rules to fight fraud could translate into particularly large reductions in assistance if politicians and bureaucrats prefer inaction lest they are caught in oversight (Leaver,

2009; Shi, 2009). We therefore compare the disutility from corruption between these two different samples—whom one might imagine may be on the extreme about caring about corruption (program administrators) or caring about delivering assistance (program beneficiaries). We cannot reject that the disutilities from corruption—both on the intensive and extensive margin—are the same for these two samples.

We also investigate the degree to which these preferences respond to local conditions, such as unusually high local food prices (which may lead to a greater weight for assistance) or higher rates of estimated leakage in government programs (which could either lead to a greater preference for reducing corruption or, conversely, reflect more local tolerance for it). We find no detectable heterogeneity in the tradeoff between corruption and amount distributed depending on these conditions, though it appears beneficiaries are somewhat more responsive to both amount distributed and amount lost to corruption in areas with high food prices.

In sum, the key finding is that respondents are sensitive to the presence of corruption, but not particularly sensitive to the amount of corruption. This suggests that quantitatively small, but widely reported, incidents of misuse of social assistance funds—such as Reagan’s ‘welfare queen’, or even small anecdotal evidence of corruption—can have disproportionate impacts on program perceptions and popularity.

This paper is related to several literatures. First, this paper relates to the existing literature that establishes the tradeoff between cracking down on fraud and more exclusion error. While some initiatives to reduce fraud can be potentially win-win, saving taxpayers funds without imposing costs on legitimate beneficiaries (for example, see Amna-Rana, et al. (2022) work on the Paycheck Protection Program (PPP) in the United States), many studies find that these kinds of initiatives to reduce fraud (or inclusion error) often indeed increase exclusion error. For example, in addition to the Muralidharan et al. (forthcoming) study discussed above, Alatas et al. (2019) shows that one measure of capture—assistance being diverted to relatives of local officials in Indonesia, which is always mentioned as a key concern of the program—is quantitatively small from a welfare perspective relative to the benefits of just improving the data for targeting aid. In the United States, Meyer and Wu (2021) show that computerization of welfare benefits in Indiana also led to increased exclusion of needy beneficiaries, particularly on recertification for welfare programs. The results here suggest that corruption reductions would need to be extremely large to justify excluding substantial numbers of beneficiaries from programs. Similarly, Gray (2019) and Homonoff and Somerville (2021) show how re-verification procedures in SNAP also result in the loss of eligible households.

Second, this paper relates to the emerging literature on using surveys and survey experiments to estimate preferences. For example, Hvidberg, Kreiner, and Stantcheva (2021) study survey data on Danish people’s views on inequality, linking surveys on preferences to administrative data, and Stantcheva (2021) uses surveys to examine preferences and reasoning about taxation in the United States. While several of these papers experimentally vary information or primes people receive, our paper by contrast uses experiments built into the survey in a cross-subject design to estimate an indifference curve by experimentally varying features of the program that people are asked to evaluate.

The remainder of the paper is organized as follows. Section II describes the empirical design and data collection. Section III presents the estimation approach and results. Section IV concludes.

II. EMPIRICAL DESIGN AND DATA COLLECTION

We begin in Section A by discussing the sample construction. We then discuss the survey experiment that we ran in Section B.

A. Sample Construction

We interviewed local program administrators and beneficiaries of Indonesia’s targeted, conditional cash transfer (CCT) program (*Program Keluarga Harapan*) through an online survey during October-December 2020 to collect information on both the current socio-economic situation and the general public health response to the COVID-19 pandemic.⁴ The survey included questions about the current public health, economic and educational conditions in the village, as well as questions on social program receipt. We embedded the survey experiment (discussed below) into the survey.

A staff roster listed 36,578 local CCT program administrators. The local program administrators work with the beneficiaries within their community to share health and educational information, to validate the conditions of the cash transfer program to determine beneficiary payments, and to help troubleshoot any problems that beneficiaries may have in regard to the program or payments. The administrators were very active in working with beneficiaries; in our survey, beneficiaries reported meeting with their administrators (pre-COVID) about 10 times per year.

To survey the administrators, we first obtained their names and cell phone numbers from the staff roster. To avoid having these administrators consider our SMS messages as spam, the CCT regional coordinators sent the online survey link to the program administrators under their supervision through WhatsApp.⁵ As a result, response was very high, particularly for an online survey: as shown in Appendix Table 1, 28,396 survey responses were recorded. Out of these, 27,034 survey responses provided a name and phone number that matched the staff roster. Note that nationally, coverage was high: at least one administrator was recorded in 467, or 92 percent, of Indonesia’s districts (see Appendix Figure 1).

We used information from the local administrators to both create a random sample of beneficiaries and to obtain their contact information. It is a challenge to random-sample beneficiaries because the central government did not have direct contact information for program beneficiaries. However, each local administrator has a pre-printed list of beneficiaries under their supervision; thus, we first asked the administrator to give us the name, cell phone, and whether the person had a smart phone of the X-th person on the pre-printed list (where X came from a random number generator). We continued in this fashion until either we obtained five beneficiaries with smart phones (which were needed to do the survey) or listed about 20 beneficiaries in total. We then asked the administrators to send the survey link to the beneficiaries in our sample. Out of 80,686

⁴ As a pilot, the surveys were first sent to two provinces (Gorontalo and Sumatera Barat) on October 8th. They were then sent to the rest of the provinces on November 19th. The survey was closed on December 12th, 2020.

⁵ To encourage response, we also sent several follow-up reminders through the WhatsApp group. In addition, we paid both the program administrators and the beneficiaries Rph 25,000 (\$1.71) for their time.

beneficiaries listed on the survey, we confirmed that 48,951 surveys were sent, and we received 19,732 responses (Appendix Table 1). We then validated the names and survey cell phone numbers off the list from program administrators, leaving us with 16,945 responses. Again, there was high geographic spread: with 405 districts, or 79 percent, covered (see Appendix Figure 1 for coverage).

B. Survey Experiment

Our survey experiment was embedded into the online survey described above. As shown in Figure 1, each respondent was given the same hypothetical situation: “Imagine the government started a new program last year to provide cash transfers to poor households. At the end of the year, the government is assessing the success of the program.”⁶

Next, they were shown information on the outcomes of the program. They received one of seven sets of information about the program, as detailed in Table 1 (and shown pictorially in Figure 1). The “base case” explained that 70 percent of the program budget reached beneficiaries, 15 percent was “missing,” 15 percent of the funds were unspent, and that 8 out of 10 beneficiaries were happy with the program. Note that, in Bahasa Indonesia, the word used for missing (“*hilang*”) could also be translated as disappeared and has a clear connotation of corruption in this context. The phenomenon of funds being unspent is common in this context: due to administrative capacity issues, some programs are not able to disburse all allocated funds. Hence, a reasonable way to think about this is that the utilization rate measures the expansiveness of the social assistance program in general.

The other six cases varied components of the base case one at a time:

- **No mention of leakage:** omit information on money unspent and on the money unaccounted for
- **More leakage, holding the amount distributed fixed:** increase the share of missing funds to 25% (and hence 5% unspent)
- **Less leakage, holding the amount distributed fixed:** decrease the share of missing funds to 5% (and hence 25% unspent)
- **Best:** increase the amount distributed to 90% (adjusting unspent and percent missing to both 5%)
- **Worst:** decrease the amount distributed to 50% (adjusting unspent and percent missing to both 25%)
- **Less happy:** decrease the share of happy beneficiaries to 60%

After seeing the information, beneficiaries were asked how they would rate the success of this program with 1 being unsuccessful and 10 being successful.

Appendix Tables 2 and 3 provide a balance check for both program administrators and beneficiaries, respectively. We include the basic demographic variables that we had, along with the questions that came before the survey experiment in the survey. Both tables show balance across the treatment groups.

⁶ Appendix Figure 2 provides the Bahasa Indonesian version shown to participants.

III. FINDINGS

A. Empirical Specification

To understand how the different features of the program drive program satisfaction, we estimate the following equation:

$$\text{Eq 1: } Y_{id} = \beta_0 + \beta_1 \text{Unhappy}_{id} + \beta_2 \text{Distributed}_{id} + \beta_3 \text{Missing}_{id} + \beta_4 \text{CorrSal}_{id} + \alpha_d + \varepsilon_{id}$$

where Y_{id} is the program rating that respondent i in district d gave, ranging from 1 to 10 with higher values indicating a higher rating, Unhappy_{id} is the normalized values of unhappiness observed, Distributed_{id} is the normalized share that reaches beneficiaries, Missing_{id} is the normalized values of “missing funds,” CorrSal_{id} is an indicator variable for whether one saw the corruption measure, and α_d are district fixed effects. Since $\text{Distributed} + \text{Missing} + \text{Unspent} = 100\%$ in all cases, the share unspent is omitted from the regression. Note that we normalize the continuous variables against the base case (that is, we subtract the value in the base case) in order to have the right comparison group for the indicator variable. As we randomized at the individual level, we have robust standard errors.

This specification allows us to compute respondents’ willingness to trade off reaching more beneficiaries with the risk of more corruption. Specifically, this regression allows us to compute the ratio $-\frac{\beta_3}{\beta_2}$, which is the average indifference curve between distributing more to beneficiaries and having more corruption in a program.⁷ That is, the ratio $-\frac{\beta_3}{\beta_2}$ allows us to capture the “marginal disutility from corruption in social programs,” i.e., respondents’ willingness to tolerate additional corruption in return for more generous social assistance.

We then investigate to what degree these preferences—i.e., the relative weight respondents place on more expansive benefits or the control of corruption—are related to local conditions. We specifically investigate three hypotheses. First, we hypothesized that in districts where food prices are high, the amount that beneficiaries receive may be particularly important and salient, and thus preferences may more heavily weight getting assistance to beneficiaries relative to the risk of corruption. To test this, using data from our survey, we calculated the change in prices of rice, beef, cooking oil and eggs since March 2020, took the average change in price, calculated which districts had above median higher changes in price levels, and examined how preferences change in high vs. low price increase areas.

Second, we hypothesized that districts with high levels of leakage in their programs may differ from those with low levels: perhaps the high leakage levels reflect the fact that there is more tolerance for corruption in some areas than others; conversely, perhaps people are more concerned about corruption in areas where it is more prevalent. To examine this, we measure leakage following Olken (2006). Specifically, we use data on receipt of a key in-kind transfer program (*Rastra*) from the national sample survey (SUSENAS) prior to our survey (March 2018).⁸ We

⁷ As we show below, Appendix Table 4 provides a regression of each of the individual six treatments against the base case and finds the same qualitative conclusions.

⁸ *Rastra* is Indonesia’s largest food assistance program of the central government that provided subsidized rice to Indonesia’s poorest households.

combine this with the official administrative allocations of the program to calculate which districts had above median leakages of the transfer program and examined whether high leakage areas had a higher relative tolerance for corruption.

B. *Results – The Marginal Disutility from Corruption*

Table 2 shows the results. In Column 1, we estimate equation 1 for the program administrators, while in Column 2, we do so for the program beneficiaries. In Column 3, we additionally estimate the difference between the administrators’ and beneficiaries’ ratings along the various dimensions of the program. We include the mean rating of the “base” group at the bottom to help provide a sense of magnitude. And, finally, we also include the p-values of the difference in coefficients for: (1) missing versus money unspent (2) missing versus unhappiness and (3) money unspent versus unhappiness.⁹

Importantly, it is worth noting that as best we can tell participants understood the question that we posed to them. First, we piloted the question extensively through a combination of sending the question and calls to help design and choose the best way to present the information. Second, as shown in Appendix Table 4, both types of participants rated the treatment that was strictly the best as higher than the base case, and they rated the treatment that was strictly the worst more negatively than the base case.

We document four key facts. First, we compute respondents’ willingness to trade off reaching more beneficiaries with the risk of more corruption, i.e., the ratio $-\frac{\beta_3}{\beta_2}$. Table 2 suggests that for program administrators, this ratio is 0.45; for program beneficiaries, the ratio is 0.8.

To be precise, consider a program that starts at a base of $x\%$ distributed to the poor, $y\%$ lost to corruption, and the rest unspent. Consider an alternate program that distributed $x+\alpha*\epsilon\%$ to the poor but lost $y+\epsilon\%$ to corruption. These estimates imply that program administrators (beneficiaries) are indifferent between the two programs for $\alpha = 0.45$ (0.8). Put another way, increasing aid by \$1 while increasing the ‘missing’ amount by \$1 would be judged an improvement by both administrators and beneficiaries.¹⁰ In fact, for our overall sample, our point estimates suggest a program that increased aid by \$1 while increasing the ‘missing’ amount by \$2 would in fact be preferred.

These numbers illustrate, while respondents dislike corruption—the sign on corruption is negative, as expected—they are much more tolerant of marginal increases in missing funds than one might have expected. They imply, for example, that a program with 50% lost to corruption would be preferred by beneficiaries to not running the program at all, holding corruption salience fixed. But this relative lack of responsiveness on the intensive margin is not the full story.

⁹ We explore two forms of robustness. First, Table 2 is robust to the exclusion of district fixed effects (Appendix Table 5). Second, we only include surveys that are on our sampling frame (which may be imperfect if, for example it is not updated to include new program administrators); however, our results are robust to the inclusion of unverified survey respondents (Appendix Table 6).

¹⁰ We can statistically rule out a ratio higher than 1.02 for program administrators and 1.17 for the entire sample, suggesting a remarkable tolerance for corruption on the margin relative to additional benefits. For program beneficiaries, where our estimates are noisier, while the point estimate is 0.8, the largest estimate that we can rule out at the 95 percent level is 3.13.

Second, we find that just making corruption salient (i.e., mentioning it at all), dramatically reduces program satisfaction, holding constant the actual reach of the program to beneficiaries and the happiness of the beneficiaries. The effect size is remarkably similar for both the administrators (-0.227) and beneficiaries (-0.236). This is a very large effect relative to the size of some of the other effects we estimate—as a benchmark, the negative effect of mentioning corruption on program satisfaction is equivalent to decreased program satisfaction from decreasing the amount distributed to beneficiaries by 20 percentage points on average in the sample.

Combined with the first result, this suggests that respondents are sensitive to the presence of corruption, but not particularly sensitive to the *amount* of corruption. This suggests that quantitatively small, but widely reported, incidents of misuse of social assistance funds—such as the infamous ‘welfare queen’ anecdote popularized by Ronald Reagan in the United States in his campaigns against more extensive social assistance, or an anecdote about corruption—can have disproportionate impacts on program perceptions and popularity.

Third, both program administrators and beneficiaries cared about overall beneficiary happiness, even holding the amount received and other characteristics of the program constant. For each additional percentage point of unhappy beneficiaries in a program, program administrators rated the programs 0.011 points lower, and beneficiaries rated the programs 0.008 points lower. Again, respondents are relatively more sensitive to program happiness in their assessments of program success than to missing funds: the increase in respondents’ perception of program success from a program that increases beneficiary happiness by 10 percentage points is the same as that from a program with a 20 percentage decrease in missing funds.

Finally, a key reason that we conducted the experiment with both administrators and beneficiaries is that we hypothesized that they may have different views on the attributes that determine a “successful” social program. For example, one could imagine that program administrators may be very concerned about corruption and leakages levels, as they are often judged by this metric—a concern made salient by the arrest of the head of the ministry that administered the CCT for corruption near the end of the survey.¹¹ Alternatively, one could imagine that beneficiaries may care a lot about what they actually receive, and may care relatively less about the amount lost to corruption or other disbursement challenges.

In practice, Table 2 shows that while there are slight differences between views of program administrators and beneficiaries, their outcomes are remarkably similar, and we do not observe large, statistically significant differences between the two (Column 3). Indeed, if anything, the marginal disutility from corruption ($-\frac{\beta_3}{\beta_2}$) is *higher* for beneficiaries than for program administrators. The differences we do observe seem driven if anything by the denominator β_2 , which captures the amount distributed, with administrators being more sensitive to this than beneficiaries (though this difference is not statistically significant)—while the point estimates on β_3 are essentially identical in the two samples (0.005 and 0.004, respectively). We also do not

¹¹ In Appendix Table 7, we investigated whether or not we could look for a discontinuous change in outcomes after this news story broke, but so few of our observations—less than 5 percent—are after this news story broke that we do not have the power to do so, especially after controlling for other secular time trends.

observe differences in their ratings on happiness, whether corruption is measured and so forth. On net, these local program administrators seem to have remarkably similar preferences to those of program beneficiaries.

C. Results—Relationships between Preferences and Economic Conditions

We examine two key forms of heterogeneity: (i) whether respondents live in a district with increasing food prices—where one may expect a higher relative weight on giving out assistance; and (ii) whether the respondents’ district had a high level of leakage in the transfer programs prior to the survey experiment—where one may expect a higher concern about corruption (or conversely, the higher leakage could reflect higher local tolerance for corruption).¹²

Table 3 reports these results, showing that the findings are fairly consistent across these different types of locations.¹³ First, in Columns 1 and 2, we examine whether living in an area where food prices are currently high changes the tradeoffs one makes between corruption and the amount that is distributed. We do find that the beneficiaries in areas with higher food prices valued programs with higher distribution rates more than those in lower food-price areas, as we hypothesized (p-value of 0.043, Column 2). Otherwise, we do not observe large statistically detectable differences in how the beneficiaries or program administrators view the corruption levels or happiness of beneficiaries.

Second, we compared the results in areas with high versus low levels of leakages in transfer programs prior to our experiment. As shown in Columns 3 and 4, on net, we do not find statistically detectable differences between these two areas.

IV. DISCUSSION AND CONCLUSION

Reducing fraud and corruption in social programs often also reduces the ease in which beneficiaries may access the program, and hence the amount that they receive under the program. In this paper, we have examined the tradeoff in preferences over this tradeoff, which we term the *marginal disutility from corruption*.

Running a large survey experiment with 45,000 respondents across Indonesia, we have 3 main findings. First, we find that the marginal disutility from corruption is surprisingly small. Respondents did, indeed, view programs with more missing funds more negatively, but this was small when compared to the value they attached to more aid reaching beneficiaries. Specifically, on the margin, changing a program to distribute \$1 more to beneficiaries raises a respondent’s satisfaction with the program more than reducing missing funds by \$2. Second, we find that while respondents are relatively insensitive to the *amount* of corruption in a program, they are quite sensitive to the mere *presence* of corruption: in fact, just mentioning corruption has the same effect on respondent satisfaction as decreasing the amount distributed to beneficiaries by 20 percentage

¹² In addition, we also examined the heterogeneity of responses by the administrators’ gender in Appendix Table 8, as 51 percent of administrators were women. Women, on average, had lower ratings than men, and rate programs with higher levels of corruption as worse than men. We did not examine the heterogeneity of the beneficiaries since (by construction) 98 percent were women.

¹³ In Appendix Table 9, we also examined heterogeneity by whether the respondent stated that they thought the CCT program should be more lenient in enforcing the conditions, as we pre-specified that we would also explore this heterogeneity. We find no differences for either program administrators or beneficiaries.

points on average in the sample. Third, when we compare program administrators and program beneficiaries, they have remarkably similar preferences.

These results suggest that people are sensitive to the presence of corruption, but not particularly sensitive to the *amount* of corruption. This suggests that quantitatively small, but widely reported, incidents of misuse of social assistance funds—such as Reagan’s ‘welfare queen’ anecdote—can have disproportionate impacts on program perceptions and popularity. Preventing these small occurrences may have disproportionate effects on program perceptions. The results also suggest that people value improvements in services substantially, so that technologies that limit fraud at the expense of increased exclusion error may not improve overall welfare.

Works Cited

- Abi-Habib, Maria, and Oscar Lopez. 2022. “Mexico’s Leader Says Poverty Is His Priority. But His Policies Hurt the Poor.” *The New York Times*, 9–9.
- Alatas, Vivi, Abhijit Banerjee, Rema Hanna, Benjamin A. Olken, Ririn Purnamasari, and Matthew Wai-Poi. 2019. “Does Elite Capture Matter? Local Elites and Targeted Welfare Programs in Indonesia.” *AEA Papers and Proceedings*, 109: 334–39.
- Aman-Rana, Shan, Daniel Gingerich, and Sandip Sukhtankar. 2022. “Screen Now, Save Later? The Trade-Off between Administrative Ordeals and Fraud.” Mimeo.
- Black, Rachel, and Aleta Sprague. 2016. “The Rise and Reign of the Welfare Queen.” *New America Weekly*.
- Gray, Colin. 2019. “Leaving benefits on the table: Evidence from SNAP.” *Journal of Public Economics*, 179: 104054.
- Homonoff, Tatiana, and Jason Somerville. 2021. “Program Recertification Costs: Evidence from SNAP.” *American Economic Journal: Economic Policy*, 13(4): 271–98.
- Hvidberg, Kristoffer B., Claus Kreiner, and Stefanie Stantcheva. 2020. “Social Positions and Fairness Views on Inequality.” NBER Working Paper #28099.
- Jakarta Post. 2021. “Juliari’s ridiculous sentence.” (<https://www.thejakartapost.com/academia/2021/08/24/juliaris-ridiculous-verdict.html>.)
- Kompas. 2022. Corruption of Social Assistance Money, PKH Companion in Banten Demanded 5 Years in Prison (<https://regional.kompas.com/read/2022/09/02/091233478/korupsi-uang-bansos-pendamping-pkh-di-banten-dituntut-5-tahun-penjara>)
- Kleven, Henrik Jacobsen, and Wojciech Kopczuk. 2011. “Transfer Program Complexity and the Take-Up of Social Benefits.” *American Economic Journal: Economic Policy*, 3(1): 54–90.

- Leaver, Clare. 2009. “Bureaucratic Minimal Squawk Behavior: Theory and Evidence from Regulatory Agencies.” *American Economic Review*, 99(3): 572–607.
- Meyer, Bruce D., Derek Wu, Victoria Mooers, and Carla Medalia. 2021. “The Use and Misuse of Income Data and Extreme Poverty in the United States.” *Journal of Labor Economics*, 39(S1): 5-58.
- Moro-Egido, Ana I., and Angel Solano-Garcia. 2020. “Does the perception of benefit fraud shape tax attitudes in Europe?” *Journal of Policy Modeling*, 42(5): 1085–1105.
- Muralidharan, Karthik, Paul Niehaus, and Sandip Sukhtankar. Forthcoming. “Identity Verification Standards in Welfare Programs: Experimental Evidence from India.” *Review of Economics and Statistics*.
- Olken, Benjamin A. 2006. “Corruption and the costs of redistribution: Micro evidence from Indonesia.” *Journal of Public Economics*, 90(4): 853–870.
- Sanders, Linley. 2019. “Americans Believe Benefits Fraud is Common for SNAP.” *YouGovAmerica*.
- Shi, Lan. 2009. “The limit of oversight in policing: Evidence from the 2001 Cincinnati riot.” *Journal of Public Economics*, 93(1): 99–113.
- Stantcheva, Stefanie. 2021. “Understanding Tax Policy: How do People Reason?” *The Quarterly Journal of Economics*, 136(4): 2309–2369.

1 Tables

Table 1: Survey Experiment

Program	Money distributed to the poor	Money unspent	Money that could not be accounted for	Beneficiaries happy with the program
Base Case	70%	15%	15%	80%
Corruption Not Mentioned	70%	-	-	80%
More Corruption	70%	5%	25%	80%
More Unspent	70%	25%	5%	80%
Best	90%	5%	5%	80%
Worst	50%	25%	25%	80%
Less Happy	70%	15%	15%	60%

Note: This table reports the characteristics of the seven different programs that were presented to program administrators and beneficiaries.

Table 2: Survey Experiment Results

	Program Administrators	Program Beneficiaries	All
Outcome: Program Score	(1)	(2)	(3)
Unhappiness with Program (Normalized)	-0.011*** (0.002)	-0.008** (0.004)	-0.011*** (0.002)
Amount Distributed (Normalized)	0.011*** (0.002)	0.005* (0.003)	0.011*** (0.002)
Corruption (Normalized)	-0.005** (0.003)	-0.004 (0.004)	-0.005** (0.003)
Corruption Salient	-0.227*** (0.038)	-0.236*** (0.064)	-0.227*** (0.038)
Beneficiary			1.844*** (0.534)
Unhappiness with Program $\times$ Beneficiary			0.003 (0.004)
Amount Distributed $\times$ Beneficiary			-0.006 (0.004)
Corruption $\times$ Beneficiary			0.001 (0.005)
Corruption Salient $\times$ Beneficiary			-0.009 (0.075)
<i>P-value</i>			
Unhappiness vs. Amount Distributed $\times$ -1	0.916	0.538	0.631
Unhappiness vs. Corruption	0.055	0.483	0.729
Amount Distributed $\times$ -1 vs. Corruption	0.140	0.883	0.543
Observations	26882	14847	41729
Control Mean	8.086	7.887	8.015

Note: This table reports the regression results of three variables indicative of the program's success on the rating score of the program or the program's score. "Corruption Salient" is a dummy that indicates that the percent of money that could not be accounted for was mentioned in the program scenario (i.e., a scenario other than the "Corruption Not Mentioned" scenario was presented). Column 1 reports results for program administrators, Column 2 for beneficiaries, and Columns 3 presents the stacked results for both program administrators and beneficiaries. All regressions include district (*kabupaten*) fixed effects and an indicator for missing values in this variable. Robust standard errors in parenthesis.

* $p < 0.1$ ** $p < 0.05$ *** $p < 0.01$.

Table 3: Heterogeneity

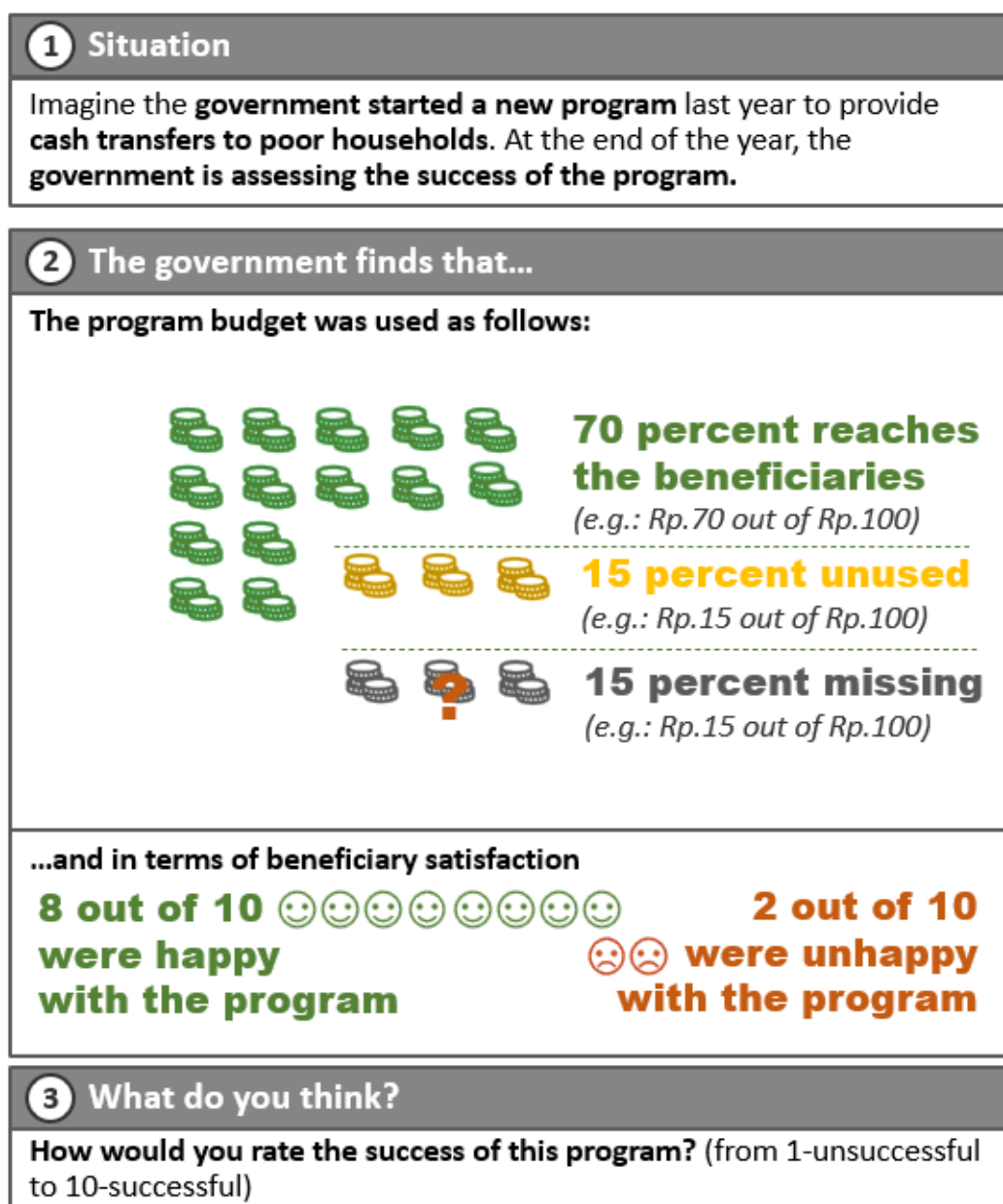
	High Food Prices		High Leakage	
	Administrators (1)	Beneficiaries (2)	Administrators (3)	Beneficiaries (4)
Outcome: Program Score				
Unhappiness with Program (Normalized)	-0.013*** (0.003)	-0.010* (0.005)	-0.015*** (0.003)	-0.011** (0.005)
Amount Distributed (Normalized)	0.010*** (0.003)	-0.002 (0.005)	0.010*** (0.003)	0.006 (0.004)
Corruption (Normalized)	-0.007* (0.004)	-0.014** (0.007)	-0.006 (0.004)	-0.006 (0.006)
Corruption Salient	-0.262*** (0.053)	-0.215** (0.096)	-0.242*** (0.055)	-0.256*** (0.092)
Unhappiness with Program $\times$ Variable	0.003 (0.004)	0.000 (0.007)	0.007 (0.004)	0.006 (0.007)
Amount Distributed $\times$ Variable	0.002 (0.004)	0.014** (0.007)	0.001 (0.004)	-0.001 (0.006)
Corruption $\times$ Variable	0.004 (0.005)	0.013 (0.009)	0.001 (0.005)	0.003 (0.009)
Corruption Salient $\times$ Variable	0.071 (0.076)	-0.031 (0.136)	0.030 (0.076)	0.037 (0.129)
<i>P-value</i>				
Unhappiness vs. Amount Distributed $\times$ -1	0.378	0.175	0.150	0.600
Unhappiness vs. Corruption	0.927	0.299	0.365	0.795
Amount Distributed $\times$ -1 vs. Corruption	0.495	0.077	0.794	0.884
Observations	26868	13276	26863	14791
Control Mean	8.086	7.911	8.085	7.888

Note: High Food Prices indicates whether the average increase in food prices (rice, eggs, beef, and oil) between March 2019 and the time of the survey in the respondent's district is more than the median of the sample. High Leakage indicates whether the gap between the total amount of Rastra subsidy received in the district and the official subsidy allocation for the district in 2018 is above the sample median. Data on Rastra subsidy received come from the March 2018 SUSENAS and data on official subsidy allocations come from Indonesia's Ministry of Social Affairs. All regressions include district (*kabupaten*) fixed effects and an indicator for missing values in this variable. Robust standard errors in parenthesis.

* $p < 0.1$ ** $p < 0.05$ *** $p < 0.01$.

2 Figures

Figure 1: Survey Experiment Base Case



Note: This figure displays the English translation of the base case program scenario that was presented to respondents. All other randomly displayed scenarios were displayed with an identical format. See Appendix Figure 2 for the original Bahasa Indonesian version shown to respondents.

3 Appendix Tables

Appendix Table 1: Survey Response

Variables	Program Administrators (1)	Program Beneficiaries (2)
Had smartphone	–	80,686
Surveys sent	36,578	48,951
Completed and consented surveys	28,396	19,732
Verified surveys	27,034	16,945

Note: This table reports our survey response by type of survey. The survey links were sent via WhatsApp to administrators. To reach beneficiaries, we requested administrators to enter the phone number and name of 5 randomly selected beneficiaries from their list of beneficiaries. The variable “Had smartphone” corresponds to the number of phone numbers entered by them (we gave the option to add "0800000" if the administrators did not have the beneficiary number and we are including these answers here). After entering the information of the selected beneficiaries, we asked administrators to confirm they were able to send the text message with the survey link to beneficiaries. The variable “Surveys sent” in Column 2 corresponds to the number of text messages sent to beneficiaries by administrators. For administrators, it corresponds to the total number in the roster. Finally, the variable “Verified surveys” correspond to answers we were able to match to the corresponding roster. Throughout this analysis, we are only using the matched sample.

Appendix Table 2: Baseline Randomization Check for Facilitators

Variable	Base Case Mean (1)	Facilitators						F-Statistic (8)
		Corruption Not Mentioned (2)	More Corruption (3)	More Unspent (4)	Best (5)	Worst (6)	Less Happy (7)	
Female	0.501	0.013 (0.011)	0.018 (0.011)	0.002 (0.011)	0.021* (0.011)	0.015 (0.011)	0.003 (0.011)	1.132 [0.340]
Lives in Urban Area	0.804	-0.004 (0.007)	0.005 (0.007)	0.001 (0.007)	0.005 (0.007)	0.003 (0.007)	0.011 (0.007)	0.896 [0.497]
Increase in Unemployment	0.889	0.001 (0.007)	0.007 (0.007)	0.010 (0.007)	0.003 (0.007)	-0.003 (0.007)	0.002 (0.007)	0.808 [0.563]
Last PKH Assistance Amount	580.340	2.585 (5.175)	12.754** (5.429)	6.038 (5.168)	3.208 (5.142)	1.560 (5.182)	-1.459 (4.902)	1.521 [0.167]
Difficulty Accessing Health Facilities	0.311	0.011 (0.011)	-0.006 (0.010)	-0.002 (0.010)	-0.009 (0.010)	-0.012 (0.010)	-0.009 (0.010)	1.056 [0.386]

Note: This table provides a baseline balance check for facilitators only. In Columns 2 - 7, we compute the difference in means for each of the treatment scenarios conditional on district fixed effects. Robust standard errors are in parentheses. In Column 8, we compute the F-statistic of the joint orthogonality test across all treatment scenarios, with the p-value in square brackets. "Last PKH Assistance Amount" is in thousands of rupiah.

Appendix Table 3: Baseline Randomization Check for Beneficiaries

Variable	Beneficiaries							
	Base Case Mean (1)	Corruption Not Mentioned (2)	More Corruption (3)	More Unspent (4)	Best (5)	Worst (6)	Less Happy (7)	F-Statistic (8)
Age	38.010	-0.078 (0.194)	-0.081 (0.197)	-0.046 (0.200)	0.096 (0.197)	-0.201 (0.192)	-0.038 (0.196)	0.429 [0.860]
Female	0.978	-0.001 (0.004)	-0.006 (0.005)	0.001 (0.004)	-0.002 (0.004)	0.004 (0.004)	0.000 (0.004)	1.009 [0.417]
Lives in Urban Area	0.793	0.012 (0.008)	0.011 (0.008)	0.015* (0.008)	0.011 (0.008)	0.007 (0.008)	0.004 (0.008)	0.744 [0.614]
Last PKH Assistance Amount	764.737	9.485 (14.352)	-6.349 (14.511)	-2.513 (14.627)	-9.091 (14.430)	-12.577 (14.560)	7.328 (14.788)	0.656 [0.686]
Difficulty Accessing Health Facilities	0.288	0.010 (0.014)	0.014 (0.014)	-0.010 (0.014)	-0.012 (0.014)	-0.015 (0.014)	-0.009 (0.014)	1.296 [0.255]
Difficulty Meeting Basic Needs	0.413	-0.016 (0.014)	-0.003 (0.014)	-0.017 (0.014)	-0.005 (0.014)	-0.011 (0.014)	0.004 (0.014)	0.625 [0.710]
Ate Less In Last Week	0.536	-0.007 (0.014)	-0.013 (0.014)	-0.004 (0.014)	-0.007 (0.014)	-0.017 (0.014)	0.001 (0.014)	0.432 [0.858]
Worked in Last Week	0.486	-0.006 (0.015)	0.001 (0.015)	-0.006 (0.015)	0.010 (0.015)	-0.003 (0.015)	0.016 (0.015)	0.602 [0.729]
Monthly Electricity Bill	90.637	-1.161 (4.363)	-7.207* (3.869)	-7.831** (3.968)	-1.155 (4.365)	-3.082 (4.165)	-3.933 (3.954)	1.282 [0.262]

Note: This table provides a baseline balance check for beneficiaries only. In Columns 2 - 7, we compute the difference in means for each of the treatment scenarios conditional on district fixed effects. Robust standard errors are in parentheses. In Column 8, we compute the F-statistic of the joint orthogonality test across all treatment scenarios, with the p-value in square brackets. "Last PKH Assistance Amount" and "Monthly Electricity Bill" are in thousands of rupiah.

Appendix Table 4: Alternative Specification

Outcome: Program Score	Program Administrators (1)	Program Beneficiaries (2)	All (3)
Corruption Not Mentioned	0.197*** (0.050)	0.240*** (0.085)	0.197*** (0.050)
More Corruption	-0.042 (0.051)	-0.081 (0.089)	-0.042 (0.051)
More Unspent	0.060 (0.050)	0.007 (0.089)	0.060 (0.050)
Best	0.185*** (0.050)	0.200** (0.089)	0.185*** (0.050)
Worst	-0.356*** (0.054)	-0.106 (0.091)	-0.356*** (0.054)
Less Happy	-0.256*** (0.051)	-0.164* (0.090)	-0.256*** (0.052)
Beneficiary			1.775*** (0.536)
Corruption Not Mentioned × Beneficiary			0.043 (0.099)
More Corruption × Beneficiary			-0.039 (0.102)
More Unspent × Beneficiary			-0.053 (0.101)
Best × Beneficiary			0.015 (0.102)
Worst × Beneficiary			0.250** (0.105)
Less Happy × Beneficiary			0.092 (0.104)
<i>P-value</i>			
More Corruption vs. More Unspent	0.042	0.319	0.889
Corruption Not Mentioned vs. Best	0.806	0.635	0.774
Best vs. Worst	0.000	0.001	0.025
Observations	26882	14847	41729
Control Mean	8.086	7.887	8.015

Note: This table reports the regression results of the treatment assigned (one of seven programs presented, see details in Table 1) on the respondent's rating of the program or program score. All coefficients are interpretable relative to Program 1 (Base Case), which is the omitted category. Column 1 reports results for program administrators, Column 2 for beneficiaries, and Column 3 presents the stacked results for both program administrators and beneficiaries. All regressions include district (*kabupaten*) fixed effects and an indicator for missing values in this variable. Robust standard errors in parenthesis.

* $p < 0.1$ ** $p < 0.05$ *** $p < 0.01$.

Appendix Table 5: Replication of Table 2, Omitting District Fixed Effects

Outcome: Program Score	Program Administrators (1)	Program Beneficiaries (2)	All (3)
Unhappiness with Program (Normalized)	-0.011*** (0.002)	-0.009*** (0.003)	-0.011*** (0.002)
Amount Distributed (Normalized)	0.011*** (0.002)	0.006** (0.003)	0.011*** (0.002)
Corruption (Normalized)	-0.005** (0.002)	-0.003 (0.004)	-0.005** (0.002)
Corruption Salient	-0.230*** (0.038)	-0.234*** (0.064)	-0.230*** (0.038)
Beneficiary			-0.166** (0.067)
Unhappiness with Program $\times$ Beneficiary			0.002 (0.004)
Amount Distributed $\times$ Beneficiary			-0.005 (0.004)
Corruption $\times$ Beneficiary			0.002 (0.005)
Corruption Salient $\times$ Beneficiary			-0.004 (0.074)
<i>P-value</i>			
Unhappiness vs. Amount Distributed $\times$ -1	0.898	0.554	0.566
Unhappiness vs. Corruption	0.078	0.279	0.956
Amount Distributed $\times$ -1 vs. Corruption	0.129	0.636	0.729
Observations	26882	14847	41729
Control Mean	8.086	7.887	8.015

Note: See Table 2 for details on the specifications.

* $p < 0.1$ ** $p < 0.05$ *** $p < 0.01$.

Appendix Table 6: Replication of Table 2 with all Respondents, Regardless of Matching

Outcome: Program Score	Program Administrators (1)	Program Beneficiaries (2)	All (3)
Unhappiness with Program (Normalized)	-0.011*** (0.002)	-0.007** (0.003)	-0.011*** (0.002)
Amount Distributed (Normalized)	0.011*** (0.002)	0.005* (0.003)	0.011*** (0.002)
Corruption (Normalized)	-0.005* (0.002)	-0.003 (0.004)	-0.005* (0.002)
Corruption Salient	-0.215*** (0.038)	-0.188*** (0.061)	-0.215*** (0.038)
Beneficiary			1.737*** (0.494)
Unhappiness with Program $\times$ Beneficiary			0.004 (0.004)
Money Unspent $\times$ Beneficiary			-0.006* (0.003)
Corruption $\times$ Beneficiary			0.002 (0.005)
Corruption Salient $\times$ Beneficiary			0.027 (0.072)
<i>P-value</i>			
Unhappiness vs. Amount Distributed $\times$ -1	0.971	0.653	0.716
Unhappiness vs. Corruption	0.047	0.399	0.752
Amount Distributed $\times$ -1 vs. Corruption	0.114	0.708	0.615
Observations	28109	17194	45303
Control Mean	8.082	7.851	7.994

Note: See Table 2 for details on the specifications.

* $p < 0.1$ ** $p < 0.05$ *** $p < 0.01$.

Appendix Table 7: Heterogeneity by Scandal

Outcome: Program Score	After Arrest	
	Administrators (1)	Beneficiaries (2)
Unhappiness with Program (Normalized)	-0.011*** (0.002)	-0.009** (0.004)
Amount Distributed (Normalized)	0.011*** (0.002)	0.006* (0.003)
Corruption (Normalized)	-0.006** (0.003)	-0.003 (0.005)
Corruption Salient	-0.216*** (0.039)	-0.244*** (0.066)
After Arrest	-0.086 (0.196)	-0.211 (0.263)
Unhappiness with Program $\times$ After Arrest	-0.006 (0.012)	0.013 (0.016)
Amount Distributed $\times$ After Arrest	0.009 (0.011)	-0.016 (0.015)
Corruption $\times$ After Arrest	0.020 (0.015)	-0.018 (0.021)
Corruption Salient $\times$ After Arrest	-0.311 (0.214)	0.159 (0.290)
<i>P-value</i>		
Unhappiness vs. Amount Distributed $\times$ -1	0.855	0.919
Unhappiness vs. Corruption	0.174	0.242
Amount Distributed $\times$ -1 vs. Corruption	0.228	0.312
Observations	26882	14847
Control Mean	8.086	7.887

Note: This table uses the same heterogeneity specification as Table 3. The interaction variable "After Arrest" is an indicator for being surveyed after December 6, 2020. In the final rows, we report the following: P-values from a F-test involving the difference between the triple-interaction coefficients, and the Control Mean. All regressions include district (*kabupaten*) fixed effects and an indicator for missing values in this variable. Robust standard errors in parenthesis.

* $p < 0.1$ ** $p < 0.05$ *** $p < 0.01$.

Appendix Table 8: Heterogeneity by Gender

Outcome: Program Score	Administrators (1)
Unhappiness with Program (Normalized)	-0.011*** (0.003)
Amount Distributed (Normalized)	0.014*** (0.003)
Corruption (Normalized)	-0.000 (0.004)
Corruption Salient	-0.208*** (0.056)
Female	-0.283*** (0.069)
Unhappiness with Program $\times$ Female	-0.001 (0.004)
Amount Distributed $\times$ Female	-0.005 (0.004)
Corruption $\times$ Female	-0.009* (0.005)
Corruption Salient $\times$ Female	-0.037 (0.076)
<i>P-value</i>	
Unhappiness vs. Amount Distributed $\times$ -1	0.276
Unhappiness vs. Corruption	0.260
Amount Distributed $\times$ -1 vs. Corruption	0.096
Observations	26882
Control Mean	8.086

Note: This table uses the same heterogeneity specification as Table 3. Beneficiaries are not included as they are all female. All regressions include district (*kabupaten*) fixed effects and an indicator for missing values in this variable. Robust standard errors in parenthesis.

* $p < 0.1$ ** $p < 0.05$ *** $p < 0.01$.

Appendix Table 9: Heterogeneity by Leniency

Outcome: Program Score	High Leniency	
	Administrators (1)	Beneficiaries (2)
Unhappiness with Program (Normalized)	-0.011*** (0.003)	-0.011* (0.006)
Amount Distributed (Normalized)	0.010*** (0.003)	0.005 (0.006)
Corruption (Normalized)	-0.004 (0.004)	-0.008 (0.008)
Corruption Salient	-0.252*** (0.063)	-0.284** (0.113)
High Leniency	0.247*** (0.071)	0.443*** (0.123)
Unhappiness with Program $\times$ High Leniency	0.000 (0.004)	0.004 (0.008)
Amount Distributed $\times$ High Leniency	0.001 (0.004)	0.001 (0.007)
Corruption $\times$ High Leniency	-0.003 (0.005)	0.009 (0.009)
Corruption Salient $\times$ High Leniency	0.042 (0.078)	0.045 (0.137)
<i>P-value</i>		
Unhappiness vs. Amount Distributed $\times$ -1	0.859	0.603
Unhappiness vs. Corruption	0.624	0.703
Amount Distributed $\times$ -1 vs. Corruption	0.783	0.509
Observations	26791	13837
Control Mean	8.084	7.888

Note: This table uses the same heterogeneity specification as Table 3. The interaction variable "High Leniency" is an indicator for a respondent answering above the median on a scale of 1 - 10, on a question regarding how lenient enforcement of the PKH program should be, with 1 being extremely rigid and 10 being extremely lenient. All regressions include district (*kabupaten*) fixed effects and an indicator for missing values in this variable. Robust standard errors in parenthesis.

* $p < 0.1$ ** $p < 0.05$ *** $p < 0.01$.

4 Appendix Figures

Appendix Figure 1: Distribution of Survey Responses by District

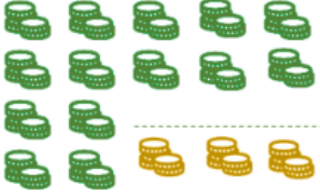
Panel A: Responses from Program Administrators



Panel B: Responses from Program Beneficiaries



Appendix Figure 2: Survey Experiment Base Case

1 Situasi Bayangkan pemerintah memulai program baru tahun lalu yang bertujuan untuk memberikan transfer tunai ke rumah tangga miskin. Di akhir tahun, pemerintah menilai kesuksesan program.
2 Pemerintah menemukan bahwa... Anggaran program digunakan sebagai berikut:  70 persen dana sampai ke penerima manfaat (contoh.: Rp.70 dari Rp.100)  15 persen dana tidak terpakai (contoh: Rp.15 dari Rp.100)  15 persen dana hilang (contoh: Rp.15 dari Rp.100)
...dan dari segi kepuasan penerima manfaat 8 dari 10 orang puas dengan program ini  2 dari 10 orang tidak puas dengan program ini 
3 Bagaimana pendapat Anda? Bagaimana Anda menilai kesuksesan program ini? (dari 1-program tidak berhasil ke 10-program berhasil)

Note: This figure displays the Bahasa Indonesian base case program scenario that was presented to respondents. All other randomly displayed scenarios were displayed with an identical format.