

Worth your Weight: Experimental Evidence on the Benefits of Obesity in Low-Income Countries

Online Appendix

Elisa Macchi

October 20, 2022

A Weight-Manipulated Portraits

To implement the photo-morphing, I work with two photographers who manually create a thinner and fatter version of each portrait using computer software. The originals are 30 Kampala resident portraits (Ugandan nationality) and 4 portraits of white-race individuals. Kampala residents are recruited via focus groups; participants provide written consent and receive a digital copy of their portrait. White-race portraits are computer generated and obtained from an algorithm similar to <https://thispersondoesnotexist.com/>.

Half of the portrayed individuals are women, and the minimum age is 20 years. Portraits are heterogeneous according to initial body size, age, ethnicity, religion, and socio-economic status. After discarding the originals, the final set is composed of 34 weight-manipulated portraits' pairs, each made of the thinner and fatter version of the same portrait (Appendix Figure G.1).

On average, thinner portraits are perceived as normal weight, while fatter portraits are portrayed as obese. To quantify the body mass variation across thinner and fatter portraits, I elicit the portraits' perceived BMI among 10 independent raters (Kampala residents). To rate portraits' perceived BMIs, raters compare each portrait to the figurative Body Size Scale for African Populations developed and validated in Cohen et al. (2015). The portraits' perceived BMIs range from 20 to 44 points. The perceived body mass distribution is plotted in Appendix Figure 1. Appendix Figure G.2 displays the body size scale and the rating procedure. Importantly, none of the thinner portraits are perceived to be underweight ($BMI < 18.5$), and all the fatter portraits are perceived to be obese ($BMI \geq 30$). Thus, the experimental average treatment effect captures the effect of obesity relatively to normal weight, which is estimated in the data using a dummy equal to

1 if the portrait is shown in the obese version.

B Beliefs Experiment

B.1 Respondents' Wards of Residence

The wards are selected at random from the list of all wards in the districts of Kampala, Mukono, and Wakiso (Greater Kampala). The selection is stratified by quintiles of a poverty index at the ward level, which I use to proxy for socioeconomic status for the respondents. I build this ward-level poverty index from Ugandan census data. From the universe of wards in Greater Kampala, I then drop one industrial area, the two richest neighborhoods (Kololo and Muyenga), and the wards counting less than 2% of the population. The final list has 99 wards.

Using ward-level aggregate data from the 2014 Ugandan census (Uganda Bureau of Statistics, 2016), I create a poverty index averaging four variables: share of households with no decent dwelling, share of households living on less than two meals per day, share of households that do not have a bank account, and share of illiterate adults. The poverty index ranges from 5, richest, to 42, poorest, (sd: 5.75). I define poverty index quintiles and randomly select 10 wards from each of the first, third, and fifth quintiles. Appendix Table G.1 provides a list of selected wards and their characteristics.

C Credit Experiment

C.1 Outcome Wording

The outcome wording is as follows: *Approval likelihood*: “Based on your first impression, how likely would you be to approve this loan application? (1–5, not at all likely to extremely likely); *Interest rate*: “If you had to approve this loan application, which interest rate would you charge? (standard, higher, lower, not applicable)”; *Creditworthiness*: “Creditworthiness describes how likely a person is to repay a financial obligation according to the terms of the agreement. Based on your first impression, how would you rate the person’s creditworthiness? (1–5, not at all likely to extremely likely)”; *Financial ability*: “Based on your first impression, how likely do you think this person would be to put the loan money to productive use? (1–5, not at all likely to extremely likely)”; *Information reliability*: “How reliable do you think the information provided by the applicant is? (1–5, not at all reliable to extremely reliable, not applicable if no additional info)”; and *Referral request*: “Based on your first impression, would you like us to refer you to a similar applicant to meet and discuss his/her loan application? (yes/no).”

C.2 Hypothetical Borrower Profiles

Using information from loan officer focus groups and data from 187 real prospective borrowers in Kampala, I build 30 hypothetical profiles. To cross-randomize the information in the applications, I use Python *numpy.random* and the *itertools.cycle* functions. Each profile includes a set of borrower characteristics and the borrower's portrait, selected from the weight-manipulated portrait set (black race only). I stratify the information randomization by body mass and, as the signaling power of body mass might differ for men and women, by gender.

The procedure is as follows. First, the hypothetical borrower's body mass and gender are randomly assigned (male/female, thin/fat). Then, the following happens:

- **Portrait:** Each portrait is randomly selected from the set of 30 black-race original portraits, conditional on gender.
- **Loan profile and reason for loan:** There are three different loan profiles: Ush 1 million, Ush 5 million, and Ush 7 million. The reason for the loan was either business or personal. All loan profiles have a six-month term to maturity, and loans could be personal or for business. Business was left generic, while the reasons for personal loans included home improvements, purchase of land, purchase of an animal, and purchase of an asset (e.g., a fridge or car). Loan profile and reason for loan randomization is stratified by the borrower's gender and body mass.
- **Name, passport ID, nationality, and place of residence:** Name and passport ID are included to increase realism but are blurred. Nationality is always Ugandan as most credit institutions would not issue loans to non-Ugandan citizens. Place of residence is always Kampala as most loan officers would be skeptical about issuing a loan to people living in another city. All applications include a date of birth, where the year of birth is the actual year of birth of the portrayed individual, while month and day are randomly selected. This information was not randomized.
- **Occupation:** The information was randomized conditional on the applicant's gender. Typical female occupations include being an owner of the following: a retail and mobile money shop, a boutique, a jewelry shop, an agricultural produce and drug shop, or a hardware store. Typical male occupations include owning a retail shop and mobile money business, owning a phone accessory and movie shop, selling clothes (owning a boutique), running a poultry and egg business, and running a dairy project. The set of

occupations were vetted in focus groups with loan officers. All the hypothetical loan applicants were self-employed because employees normally have a line of credit with their employer.

- **Monthly income:** Income information is provided in the form of last month’s self-reported revenue and profits, which are randomly assigned conditional on loan profile and the borrower’s gender and body mass and type. I first randomly assign each profile to a type: good (low DTI ratio) or bad (high DTI ratio). I then compute the monthly repayments based on the average interest rate in Kampala and determine monthly profits according to the formula $MonthlyRepayment = X \cdot MonthlyProfits$. If the borrower type is good, X is randomly selected from $[0.3; 0.35; 0.37; 0.4]$; if the borrower type is bad, X is randomly selected from $[0.9; 0.95; 0.97; 1.05]$. Notably, “bad” borrowers are relatively defined and could still be considered for a loan. It is not uncommon to approve loans such that $X = 0.95$ or $X = 1$. This made the profiles realistic: borrowers with no chance of being approved would normally not apply or would lie. Moreover, it raised loan officers’ stakes by showing they could access a good pool of borrowers by participating in the experiment.
- **Collateral:** Collateral is randomly assigned conditional on the borrower’s body mass, gender, and loan profile. For loan profiles of Ush 1 million, the choice is between motorcycle and land title. For loans of Ush 5 million and above, the choice is either car or land title.

The financial information is displayed at the bottom of the loan profile, using the sentence “This applicant is self employed and runs a [occupation type] in Kampala. The applicant claims that the business is going well. Last month, the business revenues amounted to [revenue amount]. The profits were [profit amount]. The applicant could provide a [collateral type] as collateral. Please notice that the information on revenues, profits and collateral are self reported by the applicant, and have not yet been verified.”

C.3 Implementation of Borrower Referrals

To refer loan officers to real borrower referrals that match their preferences, I use their choices in the credit experiment. The matching is borne out of a machine learning algorithm that accounts for all observable characteristics except gender and body mass. I exclude these characteristic to avoid implementing biased referrals, following Kessler et al. (2019). This choice may be seen as deceptive since loan officers may expect body mass or gender to matter. I believe the ethical concerns to be minimal since I do not specify the characteristics based on which

I match borrowers and lenders and since a perfect match would never be feasible and would be justified by the need of avoiding biased credit outcomes.

To implement the procedure, I use R and the code mostly relies on *Tidymodels*.¹

Introduction to the Machine Learning Problem The problem of matching new borrowers with loan officers based on loan officer preferences is a supervised machine learning algorithm problem. Supervised machine learning revolves around the problem of predicting out-of-sample y from in-sample x . One needs to predict loan officers' preferences for new borrowers (out of sample) based on the preferences they expressed on hypothetical borrowers in the credit experiment (in sample). Since my measure of loan officers' preferences is the binary choice of requesting, or not, to meet with the hypothetical borrower, I train a supervised classification algorithm.

To implement this matching, in short, I train a set of competing classification models on the experimental data and select the optimal model to identify loan officers' preferences. Then, I apply it to the new database of real prospective borrowers to predict which borrowers would be more likely to get a meeting with a given loan officer. The real prospective borrowers are 187 Kampala residents who need a loan. For each new borrower, I select the loan officer who has the highest probability of requesting a meeting with that borrower. Finally, the details of the loan officers are communicated to that borrower with a phone call in spring 2020. Depending on the loan officers' stated choice, I refer the borrower to either the institution or a specific loan officer.

Data Description The loan officer preferences data are based on 238 loan officers, evaluating between 4 to 30 applications each. To improve on referral quality, I exclude profiles for which the loan officer has no information on the applicants' financial information. The total number of observations is 4,419.

Machine learning algorithms search automatically for the variables, and interactions among them, that best predict the outcome of interest. One must decide how to select, encode, and transform the underlying variables before they are fed to the machine learning algorithm. I include all loan officers and firm characteristics recorded in the credit experiment. For the borrower characteristics, I include all the characteristics in the profile except 1) gender and body mass because of ethical reasons and 2) occupation, which was elicited as an open question to the new borrowers. Including the occupation information requires making some assumptions on how to code the self-reported occupations of the prospective borrowers, which does not seem worthwhile considering that algorithm performances are quite good even without occupation information.

¹The code is available upon request.

The preferences data include the following:

- Loan officers: age, body mass, gender, education, self-reported financial knowledge, financial knowledge score, experience, role (dummies for manager or owner), employed/self-employed status, monthly income, family members, activities performed, perceived stress of the verification procedure, dummies for factors influencing loan officer choices (age, gender, income, nationality, appearance, education, guarantor, collateral, occupation), number of applicants met daily, number of applicants approved daily, dummies for actions implemented to verify the applicants, performance pay, and relevance of the performance pay.
- Financial institutions: institution name, tier, district, organization size, interest rate for 1 million, 5 million, and 7 million loan types offered.
- Borrowers: age, monthly profits, collateral, loan reason (business, personal), loan amount, place of residence, and nationality.

Moreover, the data include outcome information: loan officers' choice to meet, or not meet, a borrower with similar characteristics (meeting request).

The data on real prospective borrowers come from a subsample of the beliefs experiment respondents. These are 187 individuals from the 511 respondents in the beliefs experiment who said they need a loan and agreed to be contacted with information on where to apply for a loan. The data include age, monthly income, collateral, requested loan amount, requested loan type, requested loan reason, place of residence, and nationality.

Setup and Pre-Processing I split the preferences database into a training set and a test data set, stratifying over the outcome variable. This is because "meeting request" classes in the preferences database are unbalanced: 76% are in class 1 (wants to meet), and 24% are in class 0 (does not want to meet). The test sample contains 20% of the observations. After selecting the relevant variables, I convert the education, financial knowledge, loan amount, and the stress variable to ordered factors as well as convert all string variables and numerical dummies to factor variables.

After the initial pre-processing, each model has its unique pre-processing steps. In *Tidymodels*, these steps are defined in the respective recipe. In most models, I include polynomials of degree 3 for continuous variables (loan officers' and applicants' age, loan officers' body mass, borrower profits). I standardize all predictors and remove those with no variation. When necessary (e.g., in Lasso), I create dummies for all non-continuous predictors and impute all missing values with a nearest neighbor procedure.

Training Process and Model Selection I use the training set to tune the hyper parameters of each model. I first select the models and parameter combinations that result in the highest AUC on the training data set. I then use the test data set to compare the different models and select the preferred model. The performance of the preferred model on unseen data is be assessed on the test data. but before that, I tune the algorithm parameters on the trained data. I use fivefold cross validation and a two-step procedure to find the optimal parameter: first, I use a semi-random set of parameter values for the first grid. In a second step, based on the results from this first grid, I used Bayes optimization to estimate additional models around the parameter combinations that resulted in the highest AUC in the first tuning step.

The models with the highest test AUCs are the gradient boosting classifiers (extreme gradient boosting) followed very closely by a random forest classifier. Gradient boosting models are more complex, require more careful tuning, and are prone to overfitting. Given the limited test data available, I chose to rely on the simpler random forest model. The preferred random forest model is run with the ranger engine and includes polynomial variables for age and BMI of the loan officer as well as age and profits of the applicants. It also imputes missing data using nearest neighbors (three neighbors), uses numeric scores for all ordered categorical variables, and reduces the number of levels of variables by grouping infrequent categories into a new “other” category. I fit the random forest model with optimal parameters one last time to the entire available data.

Matching and Referrals To match borrowers and lenders, I merge the borrower data with the preferences data. Then, I apply the trained model to the merged database to predict a meeting request probability for each borrower-loan officer pair. The result of the classification exercise, the probability score, is a variable between 0 and 1, indicating the probability that a given loan officer would want to meet that applicant. Finally, I select those matches that are classified as positive by the algorithm, and among these I select the best match (the highest probability score). The process is successful, and I obtain a recommendation for each prospective borrower.

C.4 Robustness Checks

No Evidence of Order Effects In the credit experiment, the order of the information treatment is not randomized: loan officers first evaluate profiles without information and later evaluate profiles with self-reported financial information. Randomizing the order may have induced loan officers to think that the amount of information displayed was a strategic choice of the borrower rather than a design

choice. For example, they may have assumed that borrowers who did not present collateral information had no collateral.

At the same time, one may worry that lack of treatment randomization could bias the results, if evaluating an application has spillovers on future evaluations (e.g., if people get tired). To investigate whether this is a relevant concern, I test whether applications presented later to loan officers (within a given arm) are rated systematically differently. I generate a dummy variable that indicates whether a given application was displayed in the first half (1–5) or in the second half (6–10) and test for the heterogeneity by order at baseline, and in the effect of body mass in a regression including both loan officer and information treatment fixed effects. Appendix Table G.4 summarizes the results: there is no evidence of order effects, and, most notably, there is no significant interaction of order with body mass.

Randomization Inference The credit experiment results are consistent, large, and therefore unlikely to have occurred by chance. In this section, I demonstrate this with a simulation exercise following Athey and Imbens (2017), who recommend randomization-based statistical inference for significance tests. This approach calculates the likelihood of obtaining the observed treatment effects by random chance, where the randomness comes from an assignment of a fixed number of units (in our case, high schools) to treatment rather than from the random sampling of a population.

I focus on the main results: the benefits in access to credit in the pooled analysis. Using the experimental data, I re-assign the applications' obesity status using the same procedure used in the original randomization, and I estimate treatment effects based on this reassignment. I repeat this procedure 10,000 times to generate a distribution of potential treatment effects that could be due to baseline differences of applications and loan officers when they are combined together. For each outcome, I calculate the share of the 10,000 simulated treatment-control differences that is larger in absolute value than the difference observed in the actual random assignment discussed throughout the paper. This proportion represents the randomization-based p -value.

The results are summarized in Figure G.5, where I plot the distribution of treatment effects from the 10,000 iterations for a selection of outcomes. The dashed vertical line in each graph plots the actual treatment effect. The analysis confirms that the findings cannot be explained by random differences between the loan officers and applications including a portrait in its obese version.

D Beliefs Accuracy Analysis

D.1 Second Laypeople Sample

In Spring 2020, I ran an additional survey to elicit laypeople’s beliefs on the income distribution by body size in Kampala. This survey was not pre-registered. The initial idea was to interview a random sub-sample of the respondents of the beliefs experiment, via an in-person follow-up survey. Because of COVID-19, this was not feasible. Therefore, we initially switched to an online phone survey. We interviewed 49 respondents of the 511, but quickly realized that this approach made it complicated to refer to visual aids such as the Body Size Scale for African Populations. Because anyway we had to rely on sending these images via WhatsApp, we decided to switch to an online survey. We enroll respondents via WhatsApp, from a sample of Kampala residents who provided their phone numbers to IGREC and agreed to take part in phone and online surveys in the future. Respondents had to provide consent and received a small compensation for completing the survey. We enrolled additionally 79 respondents.

In the analysis, I pool the online and phone samples; Appendix Table G.11 provides the summary statistics for the 124 respondents. In the phone survey, I also elicited willingness to pay for nutritional advice and respondents’ beliefs on reasons for weight gain and weight loss in Kampala. To elicit these beliefs I use an open ended survey questions. Table G.12 tabulates the answers.

E External Validity

E.1 Replication in Malawi

This paper tests a theory—that obesity is perceived as a signal of wealth—whose processes are defined in general terms and is therefore is likely to find application in contexts characterized by a similar stage in the nutritional transition, that is, with a similar positive BMI and wealth correlation. To investigate the external validity of these findings, I conduct a similar, smaller-scale survey experiment with 241 women in rural Malawi. Different from the Ugandan survey experiment, the Malawi survey exploits only two portraits (1 man and 1 woman), for a total of four photo-morphed pictures. I elicit only second-order beliefs (not incentivized). For each picture, the respondents are asked to guess how many out of 10 people would rate the individual as wealthy, would rate the individual as beautiful, would give credit to the individual, would go on a date with the person, or would respect the individuals’ admonitions.

Obese individuals are around 30 percentage points more likely to be perceived as wealthy and slightly more likely to be perceived as creditworthy. Similarly,

the effects on other outcomes are not statistically significant (Appendix Figure G.6). Comparative with the Ugandan sample, the Malawi sample is substantially poorer and less educated. These results, combined with the extensive qualitative literature showing evidence of positive perception of fat bodies across developing countries and, in the past, in Europe or the US, suggests that obesity is perceived as a signal of wealth in poor countries in general.

E.2 Replication on Amazon MTurk (US)

To further investigate the external validity of the results, I investigate whether obesity is exploited as a wealth signal in a high-income country setting. First, since obesity and wealth are negatively correlated in rich countries today, obesity would be a signal of being poor. Most notably, however, if the results on the asymmetric information mechanism are correct, one should not expect people to rely much on appearance because of the existence of better verification technologies.

To test for these predictions, I replicate the beliefs experiment on Amazon MTurk in Spring 2020. I select respondents to be US residents. I recruit 37 respondents, each rating 3 portraits for a total of 111 observations. This is a small sample, but a similar-sized pilot in Uganda was able to detect statistically significant effects of obesity on wealth beliefs. Each respondent rates each portrait both in terms of first- and second-order beliefs, and their answers are not incentivized.

Respondents rate portraits in terms of nine characteristics; seven traits (wealth, beauty, health, life expectancy, self-control, ability, and trustworthiness) are the very same as in the original beliefs experiment. The remaining two allow me to measure obesity premium or penalty in credit markets: creditworthiness and willingness to lend money. All responses are on a scale from 1 to 4, as in the original experiment. Appendix Figure G.7 shows the results. Obese portraits are associated with worse ratings along all outcomes. The difference in ratings, however, is not statistically different from zero except for beauty. The effects are also in smaller in magnitude as compared to the Ugandan experiment. I interpret these results as suggestive that obesity is stigmatized in the US context, but it is not exploited as a wealth signal as in poor countries, likely because of lower asymmetric information problems.

F Sugar Beverage Tax and Weight Gain Benefits

Building on Allcott et al. (2019), henceforth ALT, I now describe how accounting for the obesity benefits can affect the calibration of obesity prevention policies by focusing on the optimal sugar beverage tax example. ALT develops a theoretical framework for optimal sin taxes and exploits it to estimate the optimal soda tax

in the US. The strength of this framework is that it delivers empirically implementable sufficient statistics formulas for the optimal commodity tax, which can be estimated in a wide variety of empirical applications. To estimate how accounting for obesity benefits would affect the optimal sugar tax (beverages) in the Ugandan context, I proceed in two steps: (1) I exploit equation (1) to estimate a benchmark for the Ugandan sugar tax in the absence of monetary obesity benefits, and I (2) introduce obesity benefits and investigate how the tax is affected.

The equation for the optimal sin tax in the ALT framework (given a fixed income tax) is

$$t \approx \frac{\bar{\gamma}(1 + \sigma) + e - \frac{p}{\bar{s}\bar{\zeta}^c}((Cov[g(z); s(z)] + A)}{1 + \frac{1}{\bar{s}\bar{\zeta}^c}((Cov[g(z); s(z)] + A)}, \quad (1)$$

where $A = E(\frac{T'(z(\theta))}{1-T'(z(\theta))}\zeta_z(\theta)\bar{s}(\theta)\epsilon(\theta))$. $\bar{\gamma}$ is the bias, σ is the redistributive effect of the corrective motive, e measures the externality from the sin good consumption, $g(z)$ are welfare weights, $T(z)$ is the income tax, $\bar{\zeta}^c$ is the compensated price elasticity, and ζ_z is the compensated elasticity of income relative to the marginal tax.

The Ugandan context differs from the US one for three main reasons. First, own survey data show that in Uganda, contrary to the US, soda consumption correlates positively with income. It follows that a sugar beverage tax is not regressive. Thus, $\sigma \leq 0$ and the correlation between welfare weights and sugary beverage consumption is negative. Second, health care cost externalities are likely lower because of the absence of a large health care system. Finally, there is low-state capacity to collect taxes. Because of these three differences, I make the following parametric assumptions: 1) $\sigma = 0$, 2) $e = 0$, and 3) $A = 0$. Thus, the equation for the optimal tax for Uganda simplifies to

$$t_{uga} \approx \frac{\bar{\gamma} - \frac{p}{\bar{s}\bar{\zeta}^c}(Cov[g(z); s(z)])}{1 + \frac{1}{\bar{s}\bar{\zeta}^c}(Cov[g(z); s(z)])}. \quad (2)$$

How do obesity benefits enter the optimal sugar beverage tax? My results show there are two types of benefits, social and financial. The social benefits are that sugary beverage consumption increases people's BMI and higher BMI individuals are perceived as wealthier. The financial benefits are that obese people have easier access to credit or other monetary returns.

Social benefits enter the utility function and are captured in the elasticity of sugar beverage consumption in Equation (2). As far as monetary benefits are concerned, this is equivalent to a subsidy in sugar beverage consumption equal to the expected returns per unit consumed ($p' = p - E(b)$). The optimal sugar

beverage tax accounting for financial benefits is

$$t_{uga}^b \approx \frac{\bar{\gamma} - \frac{(p-E(b))}{\bar{s}\zeta^c} (Cov[g(z); s(z)])}{1 + \frac{1}{\bar{s}\zeta^c} (Cov[g(z); s(z)])}. \quad (3)$$

The effect of financial benefits on the tax depends on $(Cov[g(z); s(z)])$, that is, the correlation between welfare weights and sugar beverage consumption. When $(Cov[g(z); s(z)]) > 0$, like in the US where poor people (higher welfare weights) consume more soda on average, the larger the financial benefits, the higher the optimal tax. When $(Cov[g(z); s(z)]) < 0$, like in Uganda where rich people (lower welfare weights) consume more soda, the larger the financial benefits, the lower the optimal tax.

References

- Allcott, Hunt, Benjamin B Lockwood, and Dmitry Taubinsky**, “Regressive sin taxes, with an application to the optimal soda tax,” *The Quarterly Journal of Economics*, 2019, *134* (3), 1557–1626.
- Athey, Susan and Guido W Imbens**, “The econometrics of randomized experiments,” in “Handbook of economic field experiments,” Vol. 1, Elsevier, 2017, pp. 73–140.
- Cohen, Emmanuel, Jonathan Y. Bernard, Amandine Ponty, Amadou Ndao, Norbert Amougou, Rihlat Saïd-Mohamed, and Patrick Pasquet**, “Development and validation of the body size scale for assessing body weight perception in African populations,” *PLoS ONE*, 2015, *10* (11), 1–14.
- Kessler, Judd B, Corinne Low, and Colin D Sullivan**, “Incentivized resume rating: Eliciting employer preferences without deception,” *American Economic Review*, 2019, *109* (11), 3713–44.
- Uganda Bureau of Statistics**, “The National Population and Housing Census 2014 – Main Report, Kampala, Uganda,” Technical Report 2016.

G Online Appendix Figures and Tables

Table G.1: Randomly Selected Wards in Greater Kampala for Recruiting for Beliefs Experiment Sample

District	Subcounty	Ward	Pop. share (%)	Poverty index	Quintile
Kampala	Kawempe Division	Makerere University	0.25	5	1
Kampala	Nakawa Division	Kiwatule	0.75	12	1
Kampala	Kawempe Division	Makerere II	0.66	13	1
Kampala	Nakawa Division	Bukoto II	1.01	13	1
Kampala	Rubaga Division	Lubaga	0.99	13	1
Kampala	Nakawa Division	Mutungo	2.87	14	1
Kampala	Central Division	Bukesa	0.40	15	1
Kampala	Makindye Division	Luwafu	0.87	15	1
Kampala	Makindye Division	Salaama	1.47	15	1
Kampala	Central Division	Kamwokya II	0.83	18	3
Kampala	Kawempe Division	Kanyanya	1.19	18	3
Kampala	Kawempe Division	Kawempe II	1.03	18	3
Kampala	Kawempe Division	Mpererwe	0.27	18	3
Kampala	Nakawa Division	Butabika	0.87	18	3
Kampala	Nakawa Division	Mbuya I	1.13	18	3
Kampala	Rubaga Division	Kabowa	1.76	18	3
Kampala	Kawempe Division	Wandegeya	0.32	23	5
Kampala	Central Division	Kisenyi II	0.37	25	5
Kampala	Makindye Division	Katwe II	0.60	26	5
Mukono	Central Division	Namumira Anthony	0.93	18	3
Wakiso	Nansana Division	Nansana West	1.08	15	1
Wakiso	Nansana Division	Kazo	1.48	18	3
Wakiso	Ndejje Division	Ndejje	2.28	18	3
Wakiso	Kasangati Town Council	Kiteezi	0.741	22	5
Wakiso	Kasangati Town Council	Wattuba	0.61	22	5
Wakiso	Kasangati Town Council	Kabubbu	0.61	25	5
Wakiso	Kasangati Town Council	Nangabo	0.39	26	5
Wakiso	Kasangati Town Council	Katadde	0.36	33	5
Wakiso	Mende	Bakka	0.28	41	5
Wakiso	Mende	Mende	0.25	42	5

Notes: The table shows the wards visited to recruit respondents for the beliefs experiment. Wards characteristics are from the main report of the National Population and Housing Census 2014. The selection procedure is described in Appendix B.1.

Table G.2: Heterogeneity in Obesity Wealth-Signaling Value

	(1)	(2)	(3)
	Wealth	Wealth	Wealth
Obese	0.600 (0.074)	0.548 (0.193)	0.732 (0.078)
Male	0.070 (0.076)		
Obese $\times$ Male	0.042 (0.099)		
Age		0.011 (0.004)	
Obese $\times$ Age		0.002 (0.005)	
Additional wealth signal			0.652 (0.194)
Obese $\times$ Additional wealth signal			-0.184 (0.108)
Observations	1,699	1,699	1,699

Note: Data are from the beliefs experiment. The table summarizes the wealth-signaling value of obesity by portrait's gender (column 1), portrait's age (column 2), and presence of an additional wealth signal in the portrait's description (column 3). In column 3, the additional wealth signal can be either "living in a slum" or "owning a car" or "owning a land title". *Wealth* measures the first-order beliefs on the portrayed individual's wealth (rating on 1 to 5 scale, standardized). All regressions include respondent fixed effects. Standard errors clustered at the respondent level in parentheses.

Table G.3: Hypothetical Borrower Profiles Content

Information	Randomization	Conditionality	Options
Body mass	Randomized		<i>High</i> <i>Low</i>
Gender	Stratified by BM		<i>Male</i> <i>Female</i>
Picture	Stratified by BM	Women Men	<i>Pic n1 to n15</i> <i>Pic n16 to n30</i>
Loan profile	Stratified by BM and gender		<i>Ush 1 million</i> <i>Ush 5 million</i> <i>Ush 7 million</i>
Reason for loan	Stratified by BM and gender		<i>Business</i> <i>Home improvement</i> <i>Purchase of animal</i> <i>Purchase of land</i> <i>Purchase of asset</i>
Date of birth	Non-randomized	Based on picture's age	
Residence	Non-randomized		<i>Kampala</i>
Nationality	Non-randomized		<i>Ugandan</i>
Occupation	Stratified by BM	Women	<i>Retail shop and mobile money</i> <i>Boutique (sells clothes)</i> <i>Jewelry shop</i> <i>Produce and drug shop</i> <i>Hardware store</i>
		Men	<i>Retail and mobile money shop</i> <i>Phone and movies shop</i> <i>Poultry and eggs business</i> <i>Boutique (sells clothes)</i> <i>Diary project</i>
Income	Stratified by BM and gender		<i>High</i> <i>Low</i>
Monthly profits		Low debt-to-income ratio	$DTI = [30, 35, 37, 40]$
Revenues = 3.5 profits	Not randomized	High debt-to-income ratio	$DTI = [90, 95, 97, 1.05]$
Collateral	Strat. by BM and gender	Ush 7 or 5 million Ush 1 million	<i>Car</i> <i>Land title</i> <i>Motorcycle</i> <i>Land title</i>

The table summarizes the procedure for building the hypothetical profiles. The content information comes from real prospective borrowers and typical loan profiles from focus groups with loan officers.

Table G.4: Obesity Premium by Profiles' Rating Order

	(1) Approval likelihood	(2) Financial ability	(3) Credit- worthiness	(4) Referral request
Obese	0.111 (0.037)	0.099 (0.032)	0.080 (0.033)	0.017 (0.012)
Second half	-0.006 (0.036)	-0.020 (0.035)	-0.024 (0.032)	-0.005 (0.013)
Obese $\times$ Second half	0.037 (0.061)	0.043 (0.050)	0.040 (0.047)	0.006 (0.021)
Observations	6,645	6,645	6,645	6,645

Note: Data are from the credit experiment. *Obese* is a dummy equal to one if the borrower profiles included the obese version of the original picture. *Second half* is a dummy equal to one if the profile was the 5th to the 10th profile rated, within each arm. Regressions include loan officer and information arm fixed effects. Standard errors clustered at the loan officer level in parentheses.

Table G.5: Earnings Premium in Credit Experiment

	(1) Approval likelihood	(2) Financial ability	(3) Credit- worthiness	(4) Referral request
Profits Ush mil(.)	0.125 (0.010)	0.097 (0.009)	0.076 (0.009)	0.055 (0.009)
Observations	4,566	4,566	4,566	4,566

Note: Data are from the credit experiment. *Profits* is a continuous variable indicating the self-reported profits (Ush million) reported on the profile and applies only to profiles randomly selected to display additional information. Outcomes are standardized. Regressions include loan officer fixed effects. Standard errors clustered at the loan officer level in parentheses.

Table G.6: Obesity Premium by Timing of Financial Information Provision

	(1) Approval likelihood	(2) Financial ability	(3) Credit- worthiness	(4) Referral request
Obese	0.233 (0.041)	0.174 (0.036)	0.160 (0.041)	0.030 (0.015)
Sequential information	0.191 (0.048)	0.124 (0.041)	0.130 (0.047)	0.008 (0.023)
All information at once	0.203 (0.057)	0.103 (0.046)	0.091 (0.051)	0.035 (0.024)
Obese $\times$ Sequential information	-0.135 (0.049)	-0.077 (0.044)	-0.089 (0.051)	-0.002 (0.021)
Obese $\times$ All information at once	-0.167 (0.056)	-0.082 (0.047)	-0.089 (0.053)	-0.027 (0.018)
Observations	6,645	6,645	6,645	6,645
p -value: Obese $\times$ Sequential information = Obese $\times$ All information at once	0.541	0.911	0.994	0.166

Note: Data are from the credit experiment. The estimation focuses on profiles that displayed additional financial information. *Obese* is a dummy for the profile being associated with a fatter weight-manipulated portrait. *Sequential information* indicates that the baseline information is shown first and then the financial information is provided. *All information at once* indicates that both baseline and financial information is shown immediately. The excluded category are profiles where picture and demographic information are not shown. Regressions include borrower profile and loan officer fixed effects. Standard errors clustered at the loan officer level in parentheses.

Table G.7: Credit Experiment Likelihood Ratios

Outcome	Rate obese	Rate non-obese	Ratio
<i>No financial information</i> [2,079]			
Approval likelihood ≥ 4	20.78 %	14.86 %	1.4
Creditworthiness ≥ 4	11.89 %	8.86 %	1.34
Financial ability ≥ 4	24.34 %	20.26 %	1.2
Referral request = 1	73.49 %	70.5 %	1.04
<i>Financial information</i> [4,566]			
Approval likelihood ≥ 4	23.44 %	21.59 %	1.09
Creditworthiness ≥ 4	12.8 %	10.38 %	1.23
Financial ability ≥ 4	22.07 %	19.76 %	1.12
Referral request = 1	74.04 %	72.48 %	1.02

Note: Data are from the credit experiment. The panel above reports data for profiles that were randomized not to display borrower financial information, while the panel below focuses on the profiles that displayed financial information. For the categorical variables, a rating equal to four meant “very high or likely,” while a rating equal to five meant “extremely high or likely.”

Table G.8: Obesity Premium for Male Loan Officers Rating Male Borrowers

	(1) Approval likelihood	(2) Financial ability	(3) Credit- worthiness	(4) Referral request
Profile BMI (Obese)	0.196 (0.042)	0.143 (0.045)	0.145 (0.044)	0.089 (0.042)
Observations	1,977	1,977	1,977	1,977

Note: Data are from the credit experiment. The sample is restricted to male loan officers rating male borrower profiles. Outcomes are standardized. Standard errors clustered at the loan officer level in parentheses.

Table G.9: Obesity Premium Heterogeneity by Loan Officer Characteristics

Approval likelihood	(1)	(2)	(3)	(4)	(5)	(6)	(7)	(8)	(9)
	Age	BMI	Education	Experience	Days verify	Gender	Owner	Performance pay: Any	Performance pay: Sales volume
Obese	-0.034 (0.095)	0.154 (0.103)	0.106 (0.222)	0.090 (0.026)	0.041 (0.041)	0.071 (0.024)	0.104 (0.020)	0.047 (0.065)	0.123 (0.023)
Obese $\times$ Age	0.005 (0.003)								
Obese $\times$ BMI (loan officer)		-0.002 (0.004)							
Obese $\times$ Education (years)			0.000 (0.014)						
Obese $\times$ Experience (years)				0.007 (0.006)					
Obese $\times$ Days/week to verify information					0.033 (0.016)				
Obese $\times$ Male						0.064 (0.037)			
Obese $\times$ Owner							0.036 (0.061)		
Obese $\times$ Performance pay								0.068 (0.068)	
Obese $\times$ Performance pay: Sales volume									-0.045 (0.041)
Observations	5,363	6,645	6,645	6,645	5,469	6,645	6,645	6,645	6,645

Note: Data are from the credit experiment. The table summarizes the heterogeneity analysis in the obesity premium by loan officers characteristics and reports the interaction effects of each corresponding saturated model. The outcome is *Approval likelihood* (1–5 scale, standardized), the perceived likelihood of approving the loan application. *Obese* is a dummy equal to one if the profile displays the borrower portrait in the obese version. All regressions include borrower profile fixed effects. Standard errors clustered at the loan officer level in parentheses.

Table G.10: Obesity Premium Heterogeneity: R-Squared Analysis

Dep var: Obesity premium	(1) Approval likelihood	(2) Financial ability	(3) Credit- worthiness	(4) Referral request
Residual premium (T)	0.197 (0.090)	0.251 (0.090)	0.098 (0.092)	0.173 (0.087)
Earnings, self-reported (E)	0.156 (0.109)	0.332 (0.147)	0.179 (0.163)	0.189 (0.118)
Car collateral (E)	-0.055 (0.059)	0.101 (0.071)	-0.067 (0.076)	-0.048 (0.057)
Land collateral (E)	0.088 (0.055)	-0.113 (0.067)	0.045 (0.072)	0.043 (0.053)
Constant	0.127 (0.035)	0.116 (0.038)	0.118 (0.039)	0.035 (0.032)
Observations	238	238	238	238
R2	0.041	0.060	0.020	0.038

Note: Data are from the credit experiment. The table summarizes the results of a multivariate regression to investigate the extent to which the variance of the obesity premium can be explained by variation in observable borrower financial characteristics and variation in the residual premium, conditional on learning about a borrower self-reported characteristics. The data is from the credit experiment. The regressions are estimated at the loan officer level. The dependent variable is the estimated obesity premium for each access to credit outcome. The residual premium is the estimated obesity premium for the given outcome, conditional on providing additional financial information. *Earnings*, *Land collateral*, and *Car collateral* are the estimated effects on the given access to credit outcome of self-reported earnings, car collateral, and land collateral. Regressions include fixed effects for the set of portraits evaluated, and control for borrower age and gender.

Table G.11: Summary Statistics: Belief Accuracy Sample

VARIABLES	(1) mean	(2) sd	(3) p50
Gender: Female	0.61	0.49	1.00
Age: 18 to 24	0.25	0.43	0.00
25 to 35	0.49	0.50	0.00
35 to 44	0.18	0.38	0.00
55 to 64	0.04	0.19	0.00
Education: Primary school	0.02	0.13	0.00
Secondary school	0.11	0.31	0.00
Professional degree	0.65	0.48	1.00
Some college	0.02	0.13	0.00
Two year degree	0.21	0.41	0.00
Personal income: Far below average	0.11	0.31	0.00
Moderately below average	0.07	0.26	0.00
Slightly below average	0.23	0.42	0.00
Average	0.28	0.45	0.00
Slightly above average	0.12	0.33	0.00
Moderately above average	0.14	0.35	0.00
Far above average	0.05	0.23	0.00
Personal income (month, Ush million)	0.66	0.71	0.40
BMI	26.62	6.72	25.84

Note: The table displays summary statistics for the 124 Kampala residents who were part of the beliefs accuracy survey. Because of COVID-19, the survey was run partly on the phone (49) and partly online (79). Body-mass index values are self-reported using the Body-Size Scale for African Populations of Cohen et al. (2015).

Table G.12: Reasons Why People Think Other People Want to Gain or Lose Weight in Kampala

Want to gain	Want to lose
To be more respected and look presentable in the society. They want to appear wealthy and command that respect of economic bulls So that they appear attractive and respected. Its common for unmarried people. Most ladies don't want to introduce slimy men (...) To look wealthy To be respected in public Most of them say fat people are respected on account that they are loaded (they have money) Just like myself, they feel you can look cash but after gaining the weight you start battling to reduce it In Kampala its commonly known that people with money have the weight (...) Respect Prestige. Fat people are respected even in terms of finances Financial-such other people should look at them as wealthy To look rich and show that they doing well financially To look more representable and wealthy Fat people are assumed to have money and are respected Peer pressure fit in community To be more respected They are ignorant It just happens as they Eat fatty foods and do not do exercise To gain respect Earn more respect, self confidence They want to be seen as different and attractive Get respect in community To look rich To gain more respect from people around them So that they can look good with some weight To fit in community So that they can respect them Gain more respect Fit in group Get more respect To earn more respect To gain more respect Due to Inferiority complex So that they don't under rate them To earn more respect To earn more respect So that they can be more attractive So that they can be respected Earn more respect, to gain some big status	To avoid diseases like pressure To maintain healthy living. Overweight make ones body vulnerable to diseases like pressure Sexual pleasure. Slender people enjoy sex very well as compared to overweight people To avoid diseases To easily do work without getting tired To be healthy. You know very fat people are easily attacked by diseases like the heart disease To live healthier To look smarter though most times weight people don't want to lose weight. (...) To avoid diseases like pressure and other heart related diseases To be more healthy To be more fit Feeling to appear healthy To be healthy and lighter Overweight is associated with diseases so most people do it to prevent easy attacks Be fit for some jobs To be healthy and fit People may mistake n you to be wealth Avoid sickness related to over weight Avoid sickness associated with over weight Fighting the attack of diseases and be more flexible To be more flexible and attractive Get rid of sickness associated with obesity Healthier To be more flexible, and to be in good shape To fight disease attack Fit in community To look more attractive Avoid diseases like pressure and diabetes Fit in society pear pressure Fear to sicknesses Fighting not to get diseases To be in shape and flexible Portability To fight disease and look attractive They don't want to be attacked by diseases and be fit Fear of getting diseases Not to get diseases To be in good shape They look more flexible

Note: Data from the laypeople phone sample ($N = 49$), with 10 missing responses. Each respondent is asked an open-ended question on reasons for why people in Kampala may want to gain weight and may want to lose weight.

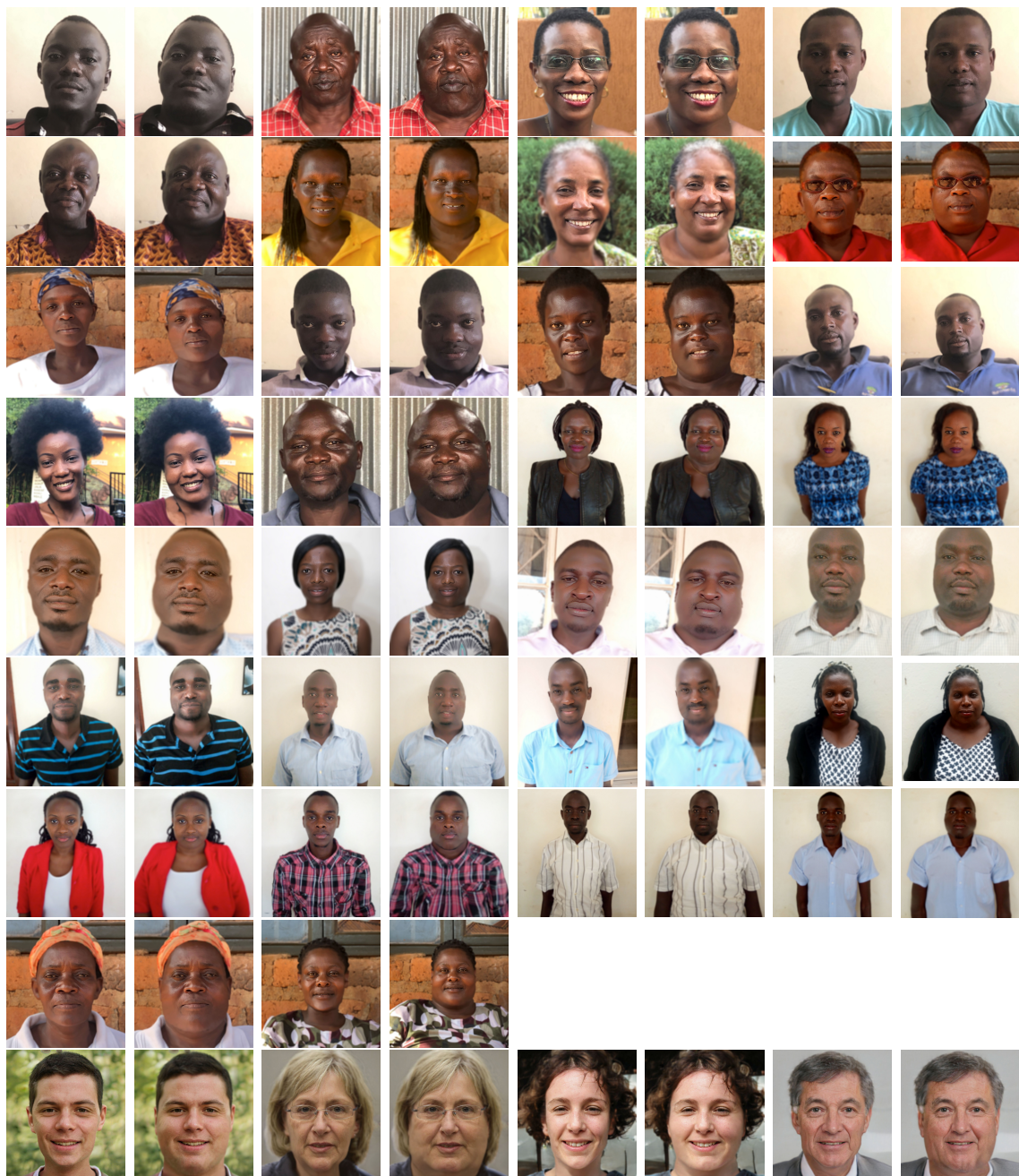


Figure G.1: Weight-Manipulated Portraits

Note: The figure displays the 34 manipulated portraits used in the analysis. The original portraits (not displayed) have been manually manipulated by two experts using a photo-morphing software to create thinner and fatter versions. The black-race original portraits are of Kampala residents, and the white-race original portraits are computer generated.

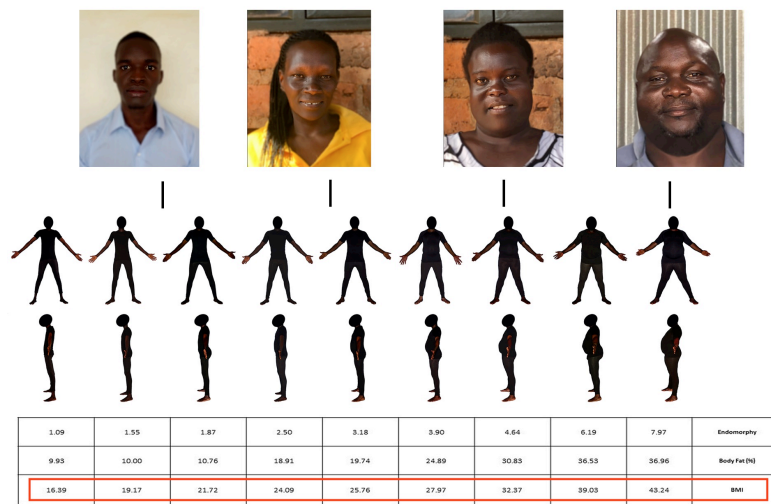


Figure G.2: Linking Weight-Manipulated Portraits to a Perceived BMI Value

Note: Ten independent Ugandan raters match each weight-manipulated portrait using the Body Size Scale for African Populations, developed and validated by Cohen et al. (2015). I take the ratings average at the portrait level and compute the corresponding BMI using the conversion model.

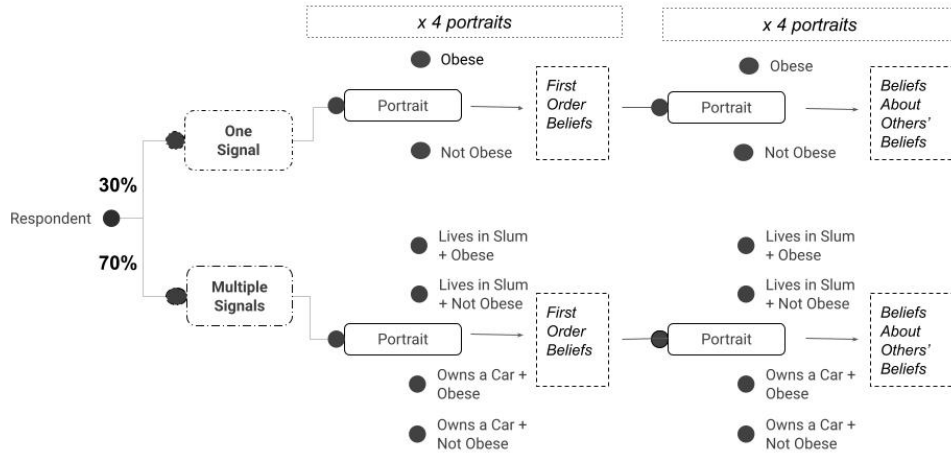


Figure G.3: Beliefs Experiment Design

Note: The graph summarizes the beliefs experiment design. Respondents rate four portraits each along with seven characteristics in random order. Portraits are selected from the weight-manipulated portrait set and are randomly displayed in the obese or non-obese version. Body mass randomization is at the respondent portrait level. Respondents can be assigned either to the “one-signal” arm to see the portrait and learn only the individual’s age. Respondents assigned to the “multiple signals” arm learn about asset ownership (car or land title— rich type) or place of residence (whether the person lives in a slum—poor type). The four portraits are first rated in terms of first-order beliefs (non-incentivized) and then in terms of beliefs about others’ beliefs (incentivized).

		Borrower's Body Mass (Portrait)			
Degree of Asymmetric Information	Demographics + loan profile information [10 profiles]	Obese		Not-obese	
	+ self-reported financial information [20 profiles]	Obese / Low DTI	Obese / High DTI	Not-obese / Low DTI	Not-obese / High DTI

Figure G.4: Credit Experiment Design

Note: The figure outlines the credit experiment design. Loan officers each evaluate 30 hypothetical borrower profiles, which include a portrait. For each borrower profile, a loan officer is randomly assigned to see the portrait either in the non-obese or obese version. The borrower BMI is cross-randomized with the amount of information provided. The first 10 applications evaluated display the borrower's picture plus demographics and loan profile details: reason for loan, type of loan, and loan amount. The last 20 applications evaluated display in addition self-reported financial information: revenue, profits, collateral, and occupation. Profits were randomized to induce a high bad or low debt-to-income ratio (DTI).

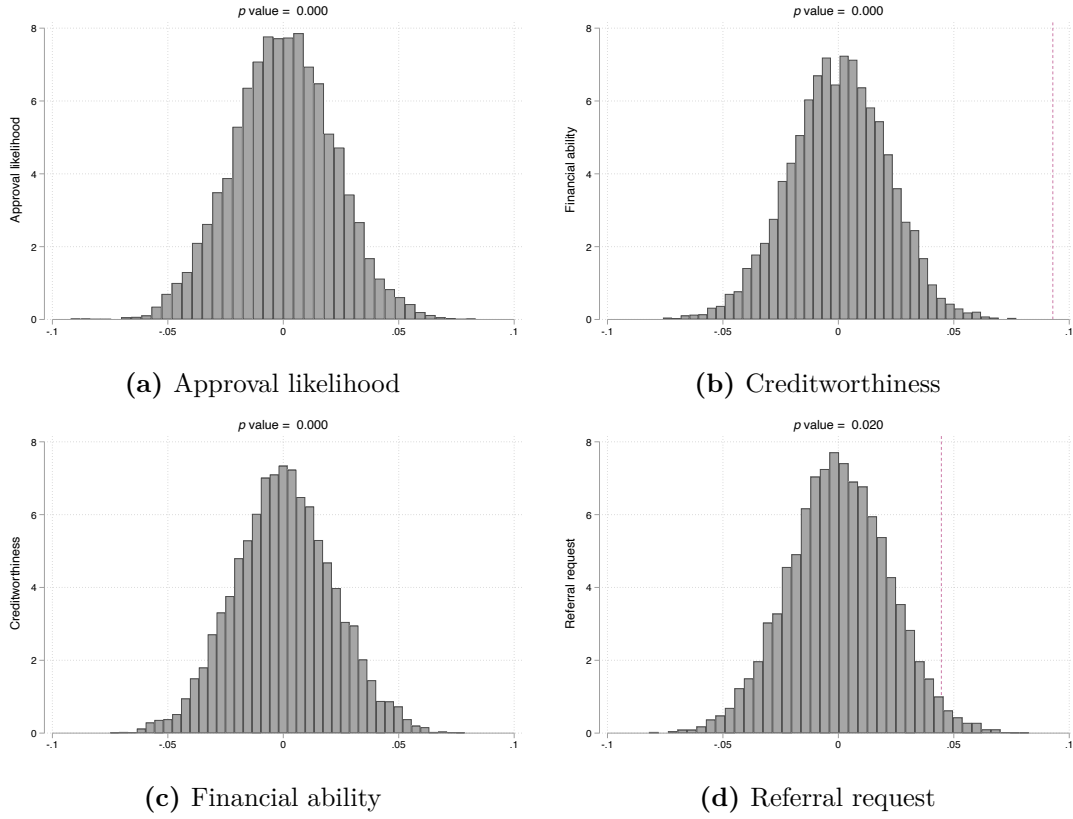


Figure G.5: Randomization Inference Exercise for Obesity Premium

Note: The figure shows a simulation exercise following Athey and Imbens (2017). Outcome variables are standardized. Each simulated treatment effect comes from randomly assigning profiles to the "obese" treatment using the same randomization algorithm used for the true assignment and then running a regression of the outcome on the "obese" status, including borrower profile and loan officer fixed effects. The dashed line is the estimated effect. The reported p -value is calculated as the number of simulated effects greater in absolute value than the estimated effect.

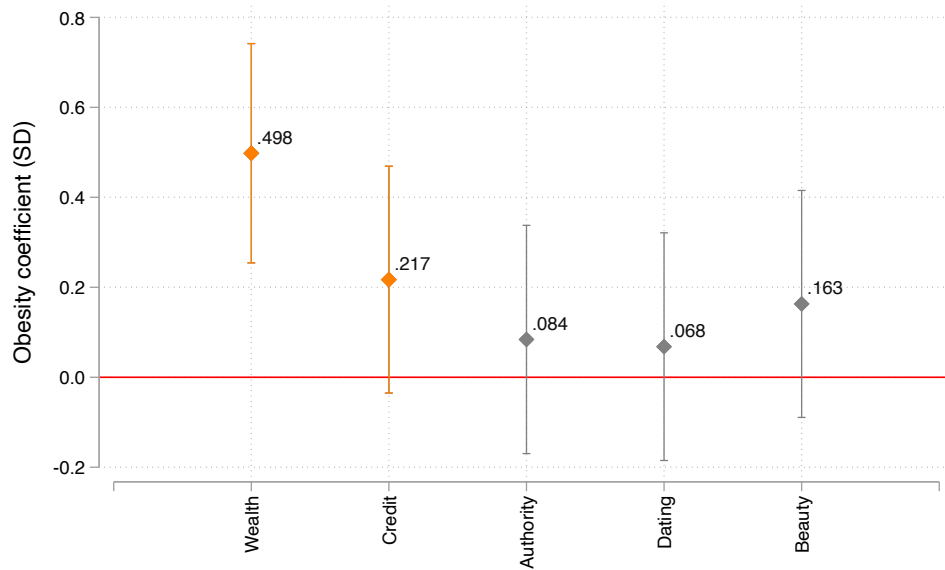


Figure G.6: Beliefs Experiment Replication in Malawi

Note: The figure shows the results from a small-scale experiment in rural Malawi to investigate external validity on a rural, poorer sample. The respondents are 241 women. The paradigm is conceptually equivalent to the beliefs experiment. The main difference is that a) women rate one picture each and b) the portraits are portrait drawings from Project Implicit instead of portraits. I use two pairs of fat/thin drawing portraits, one male and one female. The outcomes measured are second-order beliefs elicited using the wording: "How many out of 10 individuals would...: 1) rate the individual as wealthy, 2) lend money, 3) listen to a monition, 4) go on a date, or 5) rate the individual as attractive."

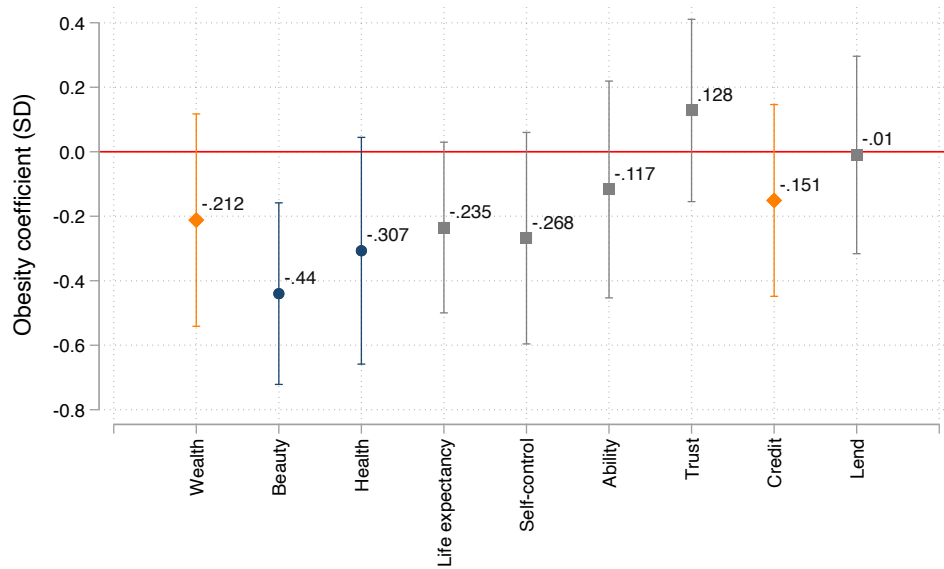


Figure G.7: Beliefs Experiment Replication on Amazon MTurk

Note: The figure plots first-order beliefs from a beliefs experiment on Amazon MTurk. The survey involves 37 respondents, for a total 111 portrait evaluations. This is a small sample, but a similar-sized pilot in Uganda had produced statistically significant results. The ratings are elicited on a 1–4 scale, using the same wording as in the original experiment. Portraits are randomly shown either in the obese or non-obese version, stratified by race (black, white). The results show that people appear to engage in (negative) obesity discrimination and second-order beliefs are aligned with first-order beliefs.

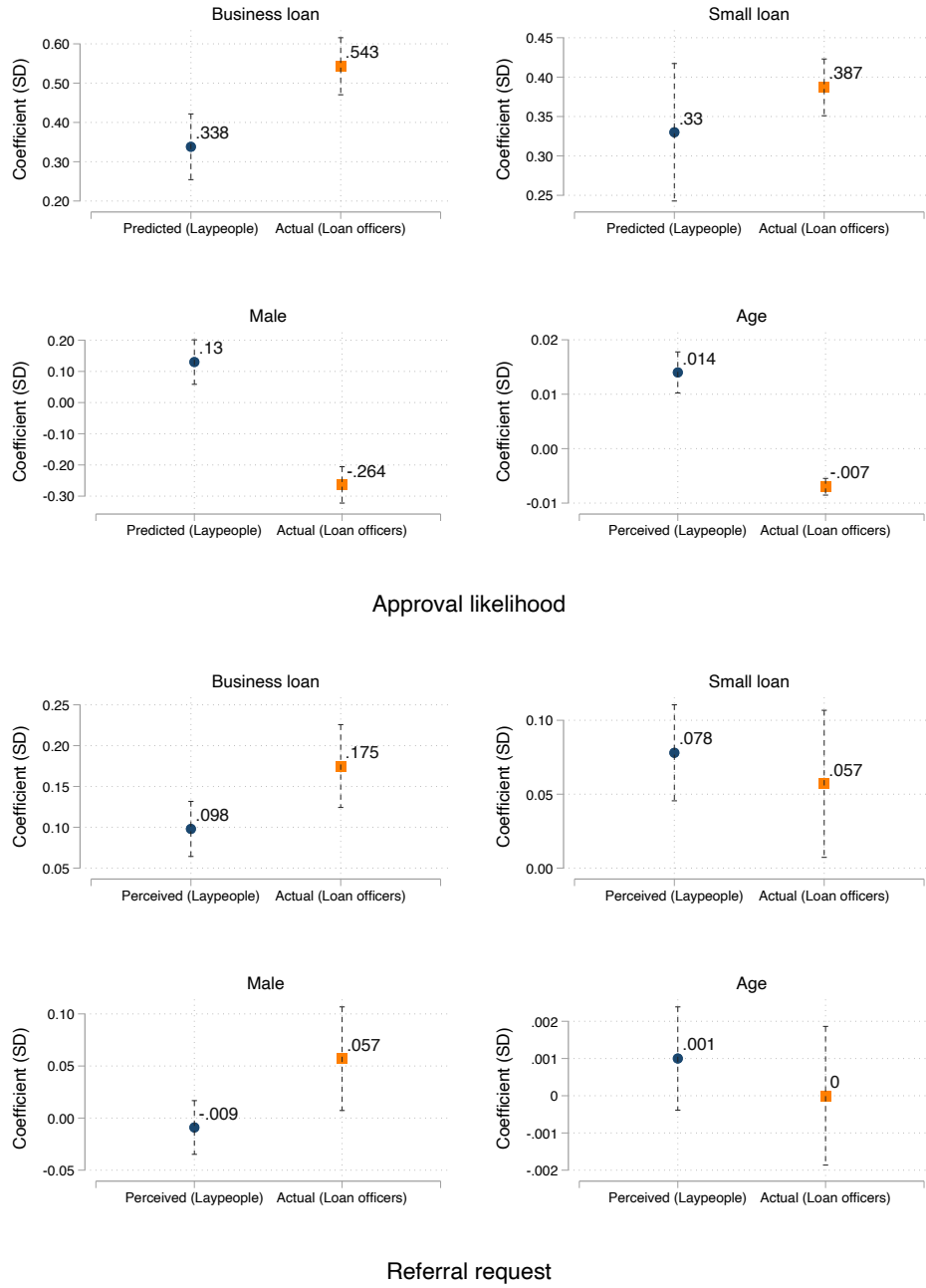


Figure G.8: Predicted vs. Actual Effects of Non-Financial Profile Characteristics

Note: The figure plots laypeople's guesses of the effect of each baseline characteristic on credit outcomes in the borrower profiles, and the actual coefficient in the credit experiment. The laypeople respondents are the same respondents of the beliefs experiment.