

NBER WORKING PAPER SERIES

IS MONEY OVERRATED? MISPERCEIVED SATISFACTION FROM INCOME

Ricardo Perez-Truglia
Rafael Macedo Rubião

Working Paper 35423
<http://www.nber.org/papers/w35423>

NATIONAL BUREAU OF ECONOMIC RESEARCH
1050 Massachusetts Avenue
Cambridge, MA 02138
July 2026

We are thankful for comments by colleagues and seminar participants. The experiment was pre-registered at the AEA RCT Registry (AEARCTR-0018251) and was reviewed and approved in advance by the Institutional Review Board at UCLA. We are grateful for financial support from the Behavioral Lab at UCLA's Anderson School of Management. The views expressed herein are those of the authors and do not necessarily reflect the views of the National Bureau of Economic Research.

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 They have not been peer-reviewed or been subject to the review by the NBER Board of Directors that accompanies official NBER publications.

© 2026 by Ricardo Perez-Truglia and Rafael Macedo Rubião. 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.

Is Money Overrated? Misperceived Satisfaction from Income
Ricardo Perez-Truglia and Rafael Macedo Rubião
NBER Working Paper No. 35423
July 2026
JEL No. C93, D91, I31

ABSTRACT

People often make important choices in pursuit of higher income, but do they place too much weight on income when deciding what path to take? We propose a simple model of misspecified learning in which individuals overestimate the marginal satisfaction from income—the expected effect of an income increase on overall life satisfaction—and therefore give income too much importance in their decisions. Guided by this model, we designed a pre-registered experiment that elicits beliefs about the marginal satisfaction from income and randomly assigns scientific evidence about it. Respondents substantially overestimate the marginal satisfaction from income both for themselves and for others, but especially for themselves. These biases shrink when respondents are exposed to scientific evidence, and the effects persist one month later. To study whether these beliefs matter for behavior, we measure income-versus-non-income trade-offs using job-choice scenarios tailored to each respondent and a real-world decision the respondent is facing. To elicit the latter, we make a methodological contribution: an AI-led interview method that moves beyond static, predefined survey instruments by combining the flexibility of qualitative interviewing with the discipline of closed-ended survey measurement. We find that biased beliefs are consequential: after learning that income matters less for life satisfaction than they initially thought, respondents place less weight on income in their decisions.

Ricardo Perez-Truglia
University of California, Los Angeles
and NBER
ricardotruglia@gmail.com

Rafael Macedo Rubião
University of California, Los Angeles
rafael.rubiao.phd@anderson.ucla.edu

An online appendix is available at
<http://www.nber.org/data-appendix/w35423>
A randomized controlled trials registry entry is available at
<https://www.socialscisceregistry.org/trials/18251>

1 Introduction

People routinely make important decisions that require trading off income against other things: whether to take a new job, move to a new city, work longer hours, delay retirement, choose a major, or go to graduate school. These decisions implicitly require answering questions such as: how would I feel if I had more money but had to work longer hours? Or more money but less time with my family? The difficulty is that people observe only the life they choose, not the counterfactual. Someone who chooses the extra hours does not know how they would have felt had they chosen more leisure instead. Still, individuals can draw on experience to form beliefs about the effects of income on their well-being, and use those beliefs to guide future decisions. Biases in these beliefs may therefore lead to suboptimal choices.

In economic models, the marginal rate of substitution represents how individuals trade off income against other considerations. For example, to decide whether to work an extra hour, an individual would compare the marginal utility from the additional hourly income with the marginal disutility from working the extra hour, and choose based on which is larger. Ideally, we would measure people’s perceptions of these marginal utilities directly. However, we cannot ask such questions because utility is a unit-free abstraction created by economists. By contrast, we draw on the long tradition in the social sciences of measuring subjective well-being. It is important to note, however, that our framework does not require utility and life satisfaction to be equivalent. Instead, life satisfaction could simply be one of many inputs into the utility function that individuals seek to maximize (Becker and Rayo, 2008; Benjamin et al., 2014).

A key object of our analysis is the perceived *marginal satisfaction from income*, which we define as the expected effect of a 20% increase in income on overall life satisfaction, where the latter is measured on a scale from 0 (extremely dissatisfied) to 100 (extremely satisfied).¹ We elicit beliefs about this perceived effect both for respondents themselves and for the average person. For example, a perceived marginal satisfaction from income of 5 implies that an individual expects a 20% income increase to raise life satisfaction by 5 points. Our central question is whether individuals hold accurate beliefs about their own marginal satisfaction from income. If they overestimate this effect for themselves, for example, they may place excessive weight on income when making choices that involve trade-offs between income and other things, such as career choices, work effort, and family time.

Despite a large literature estimating the relationship between income and subjective well-being, there is almost no research on what people believe about this relationship, and no

¹We could have chosen any number other than 20%, but we chose this one for practical reasons, as explained in Section 3.3.

research on how those misperceptions may lead to suboptimal economic behavior. This gap matters because perceived, rather than realized, marginal satisfaction from income is what individuals use *ex ante* when deciding how to trade off income against other valuable outcomes. We therefore study whether individuals systematically misperceive the marginal satisfaction from income, whether they update these beliefs after receiving scientific evidence, and whether this updating changes their choices.

To motivate the empirical analysis and guide the experimental design, we develop a simple model of beliefs about the marginal satisfaction from income. Individuals are uncertain about both their own long-run marginal satisfaction from income and the population-wide average, and they learn from both their own experiences and information about peers. We formalize this as a hierarchical Bayesian learning model with misspecification. Changes in income generate short-run changes in satisfaction, but part of these gains fades as reference income adjusts. The key learning friction is that individuals fail to account for this adaptation.

The model makes three predictions. First, individuals hold upward-biased beliefs about the marginal satisfaction from income, with stronger bias in beliefs about themselves than in beliefs about the average person. Second, accurate information about the average person’s marginal satisfaction from income lowers beliefs about both the average person and the individual’s own marginal satisfaction from income, with a larger reduction in beliefs about the average person. Third, by lowering the perceived satisfaction gains from income, the information treatment makes respondents less likely to choose higher-income options.

Next, we conduct a pre-registered online information-provision experiment to test these predictions. The baseline survey begins by eliciting respondents’ prior beliefs about the marginal satisfaction from income, both for themselves and for the average person. Respondents are then randomly assigned to the information treatment or to a control condition. The information treatment provides a population-level signal: it presents correlational and causal evidence indicating that a 20% increase in income raises life satisfaction by about one point on a 0–100 scale, and briefly discusses why the effect may be modest, including adaptation and rising aspirations. Afterward, we re-elicite these beliefs, allowing us to measure how respondents update them. Finally, to see whether belief changes influence choices, we include modules where respondents face trade-offs between income and non-income considerations.

For this portion of the experiment, before the information-treatment stage, we first gathered information about the respondent’s current or most recent job. Using this information, we created three job-choice scenarios tailored to each respondent. In one scenario, the respondent faced a trade-off between earnings and work hours: one job paid more but required longer work hours. In another scenario, they encountered a trade-off between income and commute time. In a third, they faced a trade-off between income and hours of sleep. The out-

come of interest is whether the respondent wants to choose the higher-income option, coded on a 1–4 scale so that higher values indicate greater willingness to choose the higher-income option.

In addition to eliciting what the respondent would choose in each scenario (the “choice framing”), which is the primary object of interest, we also elicited the option that each respondent thought would make them more satisfied with their lives (the “satisfaction framing”). Every respondent saw both sets of questions, one set with the choice framing and the other with the satisfaction framing, in randomized order. This variation allows us to separately identify the effects on utility-maximizing and satisfaction-maximizing behavior (Becker and Rayo, 2008; Benjamin et al., 2014).

While these job-choice scenarios were tailored to respondents, they did not involve decisions respondents were actually facing. Thus, we incorporated a real-world decision outcome. We developed a mixed-method procedure that uses an AI-led interview to construct a quantitative question tailored to the respondent. We asked respondents about a decision they were considering that involved an income-related trade-off. Respondents then chatted with the AI interviewer, who requested the important details of the situation each respondent was facing. The AI interviewer then created a question tailored to the respondent involving a binary choice between Option A and Option B, where the options traded off income against a non-income consideration. The outcome of interest is whether the respondent wants to choose the higher-income option, again coded on a 1–4 scale so that higher values indicate greater willingness to choose the higher-income option. Because the higher-income option is not fixed in advance in these respondent-generated decisions, we identify it with the help of a large language model (LLM). Our main analysis focuses on the subset of decisions for which we can construct a reliable real-world decision outcome (1,711 observations).²

The trade-offs behind the real-world decisions are concentrated around a few intuitive margins. On the income side, respondents most often consider higher-paying jobs or working more hours, followed by decisions such as relocating, freelancing, or returning to work. On the non-income side, these choices are most often weighed against family time, free time, schedule flexibility, and proximity to family and friends. To provide a concrete example, one respondent was deciding between moving to a different city (Option A) and staying in their current city (Option B), where the trade-off was that the job in the new city would offer higher pay but would come at the cost of being far away from friends. For another respondent, the relevant decision was whether to work extra hours to help pay for an expensive wedding, at the cost of having less time with family and risking burnout.

²Among the decisions left out, for example, some involve a trade-off between immediate financial benefits and future financial benefits.

To measure the persistence of the treatment effects and track what respondents did with their real-world decisions, we invited respondents to a follow-up survey about one month after the baseline. The follow-up re-elicited beliefs about the marginal satisfaction from income for respondents themselves and for the average person, repeated the choice-framing job-choice questions, and asked respondents whether they had made the real-world decision discussed in the baseline survey and, if so, what they decided. By the follow-up, 71.5% had made the decision, and among them the option they preferred at baseline matched their realized choice in 74.3% of cases.

We recruited 3,002 respondents via Prolific in April 2026. The survey included multiple quality checks, and the vast majority of subjects passed them. After applying sample restrictions, such as excluding respondents with extreme prior beliefs, our baseline sample contains 2,775 respondents. About 79.4% of these subjects completed the follow-up survey. Subjects in the treatment and control arms are well balanced on pre-treatment characteristics, and response rates to the follow-up survey are balanced as well. As is common in online samples, respondents are somewhat younger, more educated, and more left-leaning than the general U.S. population.

We start by describing prior beliefs about the marginal satisfaction from income. On average, respondents believe that a 20% increase in their own income will increase their own life satisfaction by 6.80 points. Given average baseline life satisfaction of 62.1 points, this belief implies an income-satisfaction elasticity of 0.55 ($= \frac{6.80}{\frac{62.1}{0.2}}$)—that is, each 1% increase in income increases satisfaction by 0.55%. Consistent with the view that respondents do not know the magnitude of this effect precisely, most respondents report uncertainty about their responses. There is substantial dispersion in beliefs about the average person’s marginal satisfaction from income: while the average respondent expects a 20% income increase to raise the average person’s life satisfaction by 5.20 points, the 25th percentile is 3 and the 75th percentile is 7. Moreover, we provide suggestive evidence that these differences in prior beliefs are meaningful. First, these beliefs are highly persistent over time: using the control arm, whose subjects did not receive any information, we find test-retest correlations of 0.727 for beliefs about the average person’s marginal satisfaction from income and 0.871 for beliefs about the respondent’s own marginal satisfaction from income when measured twice within the same survey, and 0.342 and 0.460, respectively, when measured one month apart. Respondents’ self-reported confidence in these beliefs also suggests that the measures contain meaningful signal. Second, these perceptions correlate with behavior: respondents who expect lower marginal satisfaction from income for themselves are less willing to choose the higher-income option, placing less emphasis on income in trade-offs.

Our empirical findings align closely with the model’s predictions. In line with the model’s

first prediction, we find that—relative to the scientific benchmark—the average respondent holds upward-biased beliefs about the marginal satisfaction from income for both the average person and themselves, and that the bias is more pronounced for beliefs about themselves.³ Take prior beliefs about the average person, for example. The scientific evidence we provided to respondents suggests that a 20% income increase would increase satisfaction by just 1 point. In comparison, respondents expect the same income increase to raise the average person’s life satisfaction by 5.20 points, or more than five times the benchmark.

Consistent with the model’s second prediction, when presented with new information, respondents revise beliefs about the marginal satisfaction from income for both themselves and the average person in the direction of the information, and they update more strongly about the average person than about themselves. For beliefs about themselves, the average posterior belief is 4.23 in the treatment arm versus 6.95 in the control arm (difference, $p < 0.001$). For beliefs about the average person, the average posterior belief is 2.50 in the treatment arm versus 5.81 in the control arm (difference, $p < 0.001$). Moreover, this belief-updating is significantly persistent, with about a third of these treatment effects persisting a month later, suggesting that subjects were genuinely incorporating what they had learned.

Next, we examine the effect of beliefs about one’s own marginal satisfaction from income on decision-making. As a first step, we estimate a simple linear regression of willingness to choose the higher-income option on the respondent’s posterior belief about their own marginal satisfaction from income. This specification uses both experimental and non-experimental variation in beliefs, so the estimates may suffer from omitted-variable bias. At the same time, using all available variation maximizes statistical power. We find a positive, statistically significant relationship: when facing a choice with a trade-off between income and a non-income consideration (e.g., leisure), respondents with higher posterior beliefs about their own marginal satisfaction from income are more willing to choose the higher-income option. This is true for each of the three job-choice scenarios as well as for the real-world decisions. These associations are economically meaningful as well. For example, for the choice trading off higher income versus lower work hours, a 1-point reduction in the respondent’s posterior belief about their own marginal satisfaction from income is associated with a reduction in willingness to choose the higher-income option by 0.023 standard deviations ($p < 0.001$). The results are in the same order of magnitude for each of the other two job-choice scenarios and for the real-world decision, with effects ranging from 0.020 to 0.034 standard deviations and always statistically significant. The results are also similar when respondents were asked to

³For brevity, we use “upward-biased” and related language as shorthand for beliefs that exceed the scientific-evidence benchmark used in the experiment. For beliefs about the respondent’s own marginal satisfaction from income, interpreting this gap as an average bias requires an additional assumption—see Section 5.2 for more details.

make satisfaction-maximizing choices.

Next, we estimate the effects of the information treatment on choices. We find significant effects and in the expected direction. Relative to the control arm, treated subjects (who receive evidence that the average person’s marginal satisfaction from income is lower than they thought) become less willing to choose the higher-income option. For example, consider the scenario involving 20% higher earnings at the expense of 20% more work hours. When asked what they would choose, respondents in the treatment arm are less inclined to choose the higher-income option: the information treatment reduces willingness to choose the higher-income option by 0.089 standard deviations, an effect that is statistically significant ($p=0.020$). The effects are in the same order of magnitude for the other two job-choice scenarios and for the real-world decision, and are also similar across the choice framing and the satisfaction framing. However, due to lower statistical precision, these coefficients are not always statistically significant. The consistent pattern across outcomes and framings supports the interpretation that lower beliefs about one’s own marginal satisfaction from income shift respondents away from higher-income options.

To put the magnitude of the effects on behavior in perspective, we combine the estimated effects on beliefs and on choices. More precisely, we use a simple two-stage least squares (2SLS) model that estimates the effect of beliefs about one’s own marginal satisfaction from income on choices by using only the variation in beliefs driven by random assignment to the information treatment. As a result, and unlike the OLS estimates, the 2SLS estimates can be interpreted as causal effects. The results from the 2SLS model are similar in magnitude to the results from the OLS model discussed above, but less precisely estimated—as expected, given that it leverages a subset of the variation in beliefs. For example, for the choice trading off higher income versus lower work hours, a 1-point reduction in the respondent’s posterior belief about their own marginal satisfaction from income is associated with a statistically significant ($p<0.05$) reduction in willingness to choose the higher-income option by 0.033 standard deviations. The 2SLS point estimate of 0.033 standard deviations is somewhat larger in magnitude than the corresponding estimate of 0.023 standard deviations from the OLS model, but we cannot reject the null hypothesis that they are equal. While the 2SLS estimates for the other choices and the alternative framing are not always statistically significant, they always have the same sign and expected order of magnitude, which is reassuring.⁴

Our paper relates to and contributes to several strands of literature. First, we relate to the extensive literature on the effects of income on subjective well-being. These studies focus on estimating the actual causal effect of income on life satisfaction (Apouey and Clark, 2015;

⁴Additionally, we estimated the 2SLS model with the responses to the follow-up survey, but the standard errors are three times as large, meaning that we are underpowered to detect effects.

Lindqvist et al., 2020) and the underlying mechanisms, such as hedonic adaptation (Di Tella et al., 2010; Clark and Georgellis, 2013), income aspirations (Stutzer, 2004), and income comparisons (Luttmer, 2005; Clark et al., 2008). The contribution of this paper is to shift the focus from the *actual* marginal satisfaction from income, which has been the focus of the existing literature, to the *perceived* marginal satisfaction from income. We are the first to measure these perceptions and to estimate their causal effects on economic behavior.⁵

This paper also relates to the longstanding literature in psychology and economics on imperfect self-knowledge and affective forecasting errors. Seminal contributions include Wilson and Gilbert (2003) on affective forecasting, Loewenstein et al. (2003) on projection bias, and Kahneman and Thaler (2006) on experienced versus decision utility. Existing empirical evidence in this literature often asks respondents to predict their future happiness after a life event and later compares those predictions with realized happiness. This type of evidence has two limitations that may help explain why it has received less attention in economics. First, it is not causal. We address this limitation by developing an experimental framework that identifies causal effects. Second, from an economic perspective, forecasting errors are relevant only insofar as they lead to suboptimal choices, yet the existing evidence does not measure behavior. We contribute to this literature by developing an empirical framework that links misperceptions directly to economic decision-making.

Our paper also contributes to an emerging literature that uses LLMs for survey measurement. Our methodological contribution is to use AI-led interviews at the intersection of qualitative and quantitative methods. Recent studies propose using LLMs to conduct interviews and collect qualitative data at scale (Chopra and Haaland, 2023; Geiecke and Jaravel, 2026).⁶ We use AI-led interviews for a different purpose: the AI interviewer first interviews respondents about a real decision they are facing and then uses that information, in real time, to construct a closed-ended survey question tailored to the respondent. Our approach moves beyond static, predefined survey instruments by combining the flexibility of qualitative interviewing with the discipline of closed-ended survey measurement.

The rest of the paper proceeds as follows. Section 2 presents the conceptual framework. Section 3 describes the research design. Section 4 discusses implementation details and measurement. Section 5 presents the main results. The last section concludes.

⁵For the first part—measuring the perceived effect of income on happiness—there are two exceptions, both from the psychology literature: Aknin et al. (2009) asked respondents to predict the happiness of others, and their own happiness, at different income levels; and Cone and Gilovich (2010) asked respondents to rank the income–happiness correlation among a set of other empirical correlations.

⁶LLMs have also been used for other purposes, such as experimental interventions. For example, Costello et al. (2024) use an LLM to persuade respondents through dialogue.

2 Conceptual Framework

This section presents a simple model that motivates and informs the experimental design. We consider a model of *hierarchical Bayesian learning with misspecified learning*. The model is *hierarchical* because the individual’s own marginal satisfaction from income is modeled as a draw from a population distribution, so beliefs about the individual’s own marginal satisfaction from income and about the population-wide average marginal satisfaction from income are jointly updated. It is *Bayesian* because the individual combines prior beliefs with noisy estimates using Gaussian updating. It is *misspecified* because the individual’s estimate neglects adaptation.

2.1 Income and Satisfaction

Let subscript i denote individuals and subscript t denote time periods. Let $Y_{i,t}$ denote life satisfaction; for brevity, we sometimes refer to it simply as satisfaction. Let $X_{i,t}$ be the individual’s income. This model revolves around the Marginal Satisfaction from Income (MSI). More precisely, there are two objects of interest:

- Individual i ’s own marginal satisfaction from income, which is determined by parameter β_i . For short, we refer to this as the self-MSI.
- The population-wide average marginal satisfaction from income, which is determined by parameter $\bar{\beta} > 0$. For short, we refer to this as the average-MSI.

The individual begins with a prior over average-MSI,

$$\bar{\beta} \sim \mathcal{N}(\bar{\beta}_0, s_0^2), \quad (1)$$

where true average-MSI satisfies $\bar{\beta} > 0$. Conditional on $\bar{\beta}$, the individual’s true self-MSI is drawn from the population:

$$\beta_i \mid \bar{\beta} \sim \mathcal{N}(\bar{\beta}, \sigma_{\bar{\beta}}^2). \quad (2)$$

2.2 Adaptation

The key feature of the true satisfaction process is that satisfaction depends not only on current income, but also on the gap between current income and a gradually adjusting reference income, $R_{i,t}$. As discussed below, reference income is an umbrella term meant to capture several phenomena documented in the literature on psychology and economics, such as hedonic adaptation and income aspirations.

The true satisfaction process is

$$Y_{i,t} = \alpha_i + (1 + \theta)\beta_i X_{i,t} - \theta\beta_i R_{i,t} + \varepsilon_{i,t}, \quad \theta > 0, \quad (3)$$

where $X_{i,t}$ is current income, $R_{i,t}$ is reference income, and $\varepsilon_{i,t}$ is a mean-zero shock. Reference income adjusts gradually to past income as follows:

$$R_{i,t} = (1 - \lambda)R_{i,t-1} + \lambda X_{i,t-1}, \quad \lambda \in (0, 1]. \quad (4)$$

We assume that individuals' income follows a random walk:

$$X_{i,t} = X_{i,t-1} + \nu_{i,t}, \quad \nu_{i,t} \sim \text{i.i.d. } (0, \sigma_\nu^2). \quad (5)$$

Throughout, income innovations are exogenous: $\nu_{i,t}$ is i.i.d. with finite fourth moments and is independent of the full sequence of satisfaction shocks $\{\varepsilon_{i,s}\}_s$ and of the lagged income and reference-income history. Satisfaction shocks $\varepsilon_{i,t}$ are i.i.d. and independent of the income process. For peers, the same data-generating process applies, with peer self-MSI parameters $\beta_j \mid \bar{\beta}$ drawn independently from $\mathcal{N}(\bar{\beta}, \sigma_\beta^2)$.

Under this normalization, a permanent increase in income has an immediate effect of $(1 + \theta)\beta_i$ on satisfaction and a long-run effect of β_i once reference income catches up. Thus β_i is directly the individual's long-run self-MSI, while the short-run effect is inflated by the factor $(1 + \theta)$.

This adaptive reference income is intended to capture several phenomena that have been widely discussed in the psychology and economics literature, some of which are mentioned explicitly on page 3 of the information treatment. For example, the reference point could partly reflect income aspirations: as individuals become richer and adjust to their new circumstances, what they consider a sufficient income rises along with their actual income. A closely related phenomenon is income comparison: as individuals become richer, they tend to interact with richer people, and these new comparisons may raise their income aspirations (Clark and Senik, 2010). The reference point could also reflect what is known as hedonic adaptation. Individuals adapt to repeated stimuli, so while a new experience, such as driving a new car, may feel especially rewarding at first, the feeling fades with repeated exposure until, after some time, driving the new car feels much like driving the old one. This view is supported by evolutionary explanations and neuroscience findings (Rayo and Becker, 2007), and has also been documented using life-satisfaction data (e.g., Di Tella et al., 2010; Galiani et al., 2018).

2.3 Misspecified Bayesian Learning

Let $\Delta X_{i,t} \equiv X_{i,t} - X_{i,t-1}$ and $\Delta Y_{i,t} \equiv Y_{i,t} - Y_{i,t-1}$ denote changes in income and satisfaction. To form beliefs about self-MSI and average-MSI, the individual uses two data sources. First, the individual observes own data: a time series that yields T observations on $(\Delta X_{i,t}, \Delta Y_{i,t})$. Second, the individual observes peer data: analogous first-difference observations for each of N peers, with each peer panel containing T_P observations. Because β_i is itself drawn from the population distribution, both sources are informative about both objects: own data are directly informative about self-MSI but also informative about average-MSI, while peer data are directly informative about average-MSI but also informative about self-MSI.

To formalize how individuals extract information from these data, we assume they summarize each data source by a first-difference slope. The motivation is that levels of satisfaction are hard to interpret: in any given year, income may change alongside many other events that affect health, relationships, family, leisure, and other dimensions of life. Short-run changes are often more salient and easier to connect to particular events. For example, someone may remember feeling especially satisfied after receiving a large bonus, or during an expensive vacation. From such episodes, they may infer that income raises satisfaction. For the individual's own history, let $\hat{b}_i^\Delta$ denote the first-difference slope relating $\Delta Y_{i,t}$ to $\Delta X_{i,t}$. For peer j , define the analogous slope using $\Delta Y_{j,t}$ and $\Delta X_{j,t}$, and let $\bar{b}_N^\Delta$ denote the average of these slopes across peers. The first-difference slopes are meant to capture this type of reasoning parsimoniously: individuals need not literally run regressions, but may use a heuristic to infer causal relationships from memories about changes in income and satisfaction.

Under the true data-generating process, first-difference regressions identify *short-run* responses to income shocks rather than long-run effects. In large samples of own and peer data, the key implication is

$$\text{plim } \hat{b}_i^\Delta = (1 + \theta)\beta_i, \quad \text{plim } \bar{b}_N^\Delta = (1 + \theta)\bar{\beta}. \quad (6)$$

The definitions of the estimators, the derivation of (6), and the corresponding Gaussian approximations are given in Appendix A.1.

The individual neglects adaptation and therefore interprets these regression coefficients as noisy signals of the long-run MSI objects β_i and $\bar{\beta}$. We assume that, when updating, the individual uses fixed subjective signal variances $S_i > 0$ and $S_p > 0$; the formal subjective signal structure is given in Appendix A.2. Here, S_p is interpreted as the variance of a peer-level slope signal after accounting for the finite length T_P of each peer panel. The objective Gaussian approximations in Appendix A.1 motivate this signal formulation, but the subjective variances S_i and S_p need not coincide with the objective variances Ω_i and Ω_p .

Thus, the bias in beliefs arises not from non-Bayesian updating, but from Bayesian up-

dating under a misspecified model of how income affects satisfaction over time.

Given the prior (1)–(2) and the subjective Gaussian signal structure, beliefs are updated in the standard Gaussian way. Define precisions

$$\kappa_0 \equiv \frac{1}{s_0^2}, \quad h \equiv \frac{1}{\sigma_\beta^2}, \quad \tau_i \equiv \frac{T}{S_i}, \quad \tau_p \equiv \frac{N}{S_p}. \quad (7)$$

The posterior mean self-MSI belief is a weighted average of the own-data regression estimate, the peer-data regression estimate, and the prior mean:⁷

$$\tilde{\beta}_i = w_i \hat{b}_i^\Delta + w_p \bar{b}_N^\Delta + w_0 \bar{\beta}_0. \quad (8)$$

Likewise, the posterior mean average-MSI belief is

$$\tilde{\bar{\beta}} = c_i \hat{b}_i^\Delta + c_p \bar{b}_N^\Delta + c_0 \bar{\beta}_0. \quad (9)$$

The exact formulas for the weights (w_i, w_p, w_0) and (c_i, c_p, c_0) , as well as the derivation of (8)–(9), are given in Appendix A.2.

Equations (8)–(9) make clear that both own and peer data affect both beliefs. Because β_i is drawn from the population distribution, own data update average-MSI beliefs, while peer data update self-MSI beliefs.

2.4 Predictions

The model delivers three predictions that guide the empirical analysis. The first concerns average pre-information bias, the second concerns belief updating after accurate information, and the third links belief updating to choices. Let $\bar{B}^{self}$ denote the average bias in self-MSI beliefs, and let $\bar{B}^{avg}$ denote the average bias in average-MSI beliefs:

$$\bar{B}^{self} \equiv \mathbb{E}_{\beta_i | \bar{\beta}} [\mathbb{E}[\tilde{\beta}_i - \beta_i | \bar{\beta}, \beta_i]], \quad \bar{B}^{avg} \equiv \mathbb{E}_{\beta_i | \bar{\beta}} [\mathbb{E}[\tilde{\bar{\beta}} - \bar{\beta} | \bar{\beta}, \beta_i]]. \quad (10)$$

Appendix A.2 shows that these average biases can be decomposed as

$$\bar{B}^{self} = \underbrace{w_0(\bar{\beta}_0 - \bar{\beta})}_{\text{prior misperception}} + \underbrace{(w_i + w_p)\theta\bar{\beta}}_{\text{average adaptation neglect}}, \quad (11)$$

$$\bar{B}^{avg} = \underbrace{c_0(\bar{\beta}_0 - \bar{\beta})}_{\text{prior misperception}} + \underbrace{(c_i + c_p)\theta\bar{\beta}}_{\text{average adaptation neglect}}. \quad (12)$$

The first component is prior misperception, captured by $\bar{\beta}_0 - \bar{\beta}$. The second is adaptation neglect, captured by the terms proportional to $\theta\bar{\beta}$. In the predictions below, we focus on the benchmark case with unbiased priors, $\bar{\beta}_0 = \bar{\beta}$, which isolates the role of adaptation

⁷We assume $\tau_i + \tau_p > 0$, otherwise the agent has no informative data at all and the model collapses to pure prior beliefs.

neglect in generating average overestimation. This assumption is not meant to rule out prior misperceptions in practice. When $\bar{\beta}_0 \neq \bar{\beta}$, the prior-bias terms enter additively in (11) and (12), shifting mean beliefs in the direction of the prior bias, with magnitudes governed by the prior weights w_0 and c_0 . The adaptation-neglect component, however, remains positive and has the same interpretation.

Prediction 1. On average, individuals overestimate both self-MSI and average-MSI, with weakly stronger overestimation of self-MSI. Under unbiased priors,

$$\bar{B}^{self} = (1 - w_0)\theta\bar{\beta} > 0, \quad \bar{B}^{avg} = (1 - c_0)\theta\bar{\beta} > 0, \quad (13)$$

$$\bar{B}^{self} - \bar{B}^{avg} = (c_0 - w_0)\theta\bar{\beta} \geq 0, \quad (14)$$

with strict inequality in the second line whenever $\tau_i > 0$ and $\kappa_0 > 0$.

Proof: See Appendix A.3.

The intuition is that people may form beliefs from vivid memories of material gains: how good it felt to receive a large raise or promotion, or to drive a new car home for the first time. First-difference regressions capture similar short-run satisfaction gains from income, part of which fades as individuals adapt. If people extrapolate from these short-run gains to long-run life satisfaction, both own data and peer data can raise average beliefs above the corresponding long-run effects.

To see why the self-MSI bias can be larger, consider the benchmark case with no own data, so that $\tau_i = 0$. In that case, the individual learns only from peer information, and self-MSI beliefs are formed entirely by projecting the inferred average-MSI onto oneself. As a result, the average bias in self-MSI beliefs coincides exactly with the average bias in average-MSI beliefs. Starting from this benchmark, introducing own data increases both biases, because first-difference regressions on own data are also upward biased by adaptation neglect. However, the increase is larger for self-MSI beliefs than for average-MSI beliefs. The reason is that own data provide direct evidence about β_i , but only indirect evidence about $\bar{\beta}$, since an individual's self-MSI is only one draw from the population distribution. Consequently, adaptation neglect in own learning raises both beliefs, but it raises self-MSI beliefs more strongly.

Next, consider the information provided in the experiment. Let m denote a signal about true average-MSI, interpreted by the individual as

$$m \mid \bar{\beta} \sim \mathcal{N}(\bar{\beta}, S_m), \quad \tau_m \equiv \frac{1}{S_m}, \quad S_m > 0. \quad (15)$$

Let $\tilde{\beta}_i(m)$ and $\tilde{\tilde{\beta}}(m)$ denote posterior mean beliefs after observing this signal, holding fixed the individual's own and peer data. Define the average belief changes induced by the signal as

$$\Delta^{self}(m) \equiv \mathbb{E}[\tilde{\beta}_i(m) - \tilde{\beta}_i], \quad \Delta^{avg}(m) \equiv \mathbb{E}[\tilde{\tilde{\beta}}(m) - \tilde{\tilde{\beta}}], \quad (16)$$

where the expectation averages over pre-information data and over $\beta_i \mid \bar{\beta}$. The benchmark accurate signal corresponds to the realization $m = \bar{\beta}$.

Prediction 2. On average, accurate information lowers both average-MSI and self-MSI beliefs, with a larger reduction in average-MSI beliefs. Under unbiased priors and the benchmark accurate signal $m = \bar{\beta}$,

$$\Delta^{avg}(\bar{\beta}) < 0, \quad \Delta^{self}(\bar{\beta}) < 0, \quad |\Delta^{avg}(\bar{\beta})| \geq |\Delta^{self}(\bar{\beta})|. \quad (17)$$

The last inequality is strict whenever $\tau_i > 0$.

Proof: See Appendix A.4.

The intuition is that, under unbiased priors, adaptation neglect makes average-MSI beliefs too high. Accurate information therefore moves average-MSI beliefs downward. Self-MSI beliefs also fall because the individual's self-MSI is drawn from the population distribution, but own data make self-MSI beliefs less responsive to population-level information, making this average reduction smaller.

We close the model by linking beliefs to choices. Let a_i denote an income-generating action, such as hours worked, effort devoted to obtaining a higher-paying job, or the choice of a higher-paying career path. To keep notation consistent with the setup above, income is $X_i(a_i)$, with $X'_i(a_i) > 0$. Let $\mathcal{S}_i(a_i; m)$ denote the individual's perceived life satisfaction from income after observing signal m :

$$\mathcal{S}_i(a_i; m) = \tilde{\beta}_i(m) X_i(a_i). \quad (18)$$

The individual values life satisfaction, but life satisfaction is not the only input into utility. In general, utility can be written as $U_i(\mathcal{S}_i(a_i; m), X_i(a_i), a_i)$, where the first two arguments enter positively and the income-generating action enters negatively through its direct cost. A simple separable specification is

$$V_i(a_i; m) = \tilde{\beta}_i(m) X_i(a_i) + g_i(X_i(a_i)) - \phi_i(a_i), \quad (19)$$

where $g'_i(X) > 0$ captures other benefits of income and $\phi'_i(a) > 0$ captures the direct cost of the income-generating action. Assume that $V_i(a_i; m)$ is strictly concave in a_i , so that the optimal choice $a_i^*(m)$ is unique and interior. Let Δa_i denote the post-information minus

pre-information change in the income-generating action.

Prediction 3. On average, accurate information lowers the income-generating action. Under unbiased priors, the benchmark accurate signal $m = \bar{\beta}$, and a common local responsiveness of the income-generating action to perceived self-MSI, for small belief changes,

$$\mathbb{E}[\Delta a_i] < 0. \quad (20)$$

Proof: See Appendix A.5.

The intuition is that accurate information lowers self-MSI beliefs on average. This lowers the perceived benefit of income-generating actions, while leaving their direct cost unchanged, so individuals choose less of the income-generating action on average.

3 Research Design

3.1 Overview

A sample survey instrument for the baseline survey is provided in Appendix H, and the corresponding instrument for the follow-up survey is provided in Appendix I.⁸ Both surveys are described in more detail below. The baseline survey began with questions used to construct job-choice scenarios tailored to each respondent and with an AI-led interview used to elicit a real-world income-related decision. We then elicited prior self-MSI and average-MSI beliefs. Respondents were randomly assigned either to the information treatment or to a control condition. After this information-treatment stage, we re-elicited the same beliefs and measured choices involving trade-offs between income and non-income considerations. The follow-up survey, conducted about one month later, re-elicited the main belief measures, repeated the choice-framing job-choice measures, and asked respondents about the real-world decision discussed at baseline.

3.2 Prior Average-MSI and Self-MSI Beliefs

Before any information was provided, we elicited respondents' prior self-MSI and average-MSI beliefs. More precisely, we elicited the perceived effect of a 20% income increase on life satisfaction on a 0–100 scale. Respondents may be largely unfamiliar with reporting life

⁸The actual survey varied across respondents because some questions depended on earlier answers and because key elements were randomized, including treatment assignment and framing order.

satisfaction on such a scale, so they may not have an intuitive sense of how large a one-point satisfaction increase is. This is analogous to asking respondents to evaluate prices in a currency different from the one they are most accustomed to. While we cannot eliminate this type of measurement error entirely, we can at least take steps to mitigate it. The first step was to ask them to report their overall life satisfaction using the same 0–100 scale, so they had the opportunity to familiarize themselves with the concept and the scale. The second step takes inspiration from the analogy with a foreign currency: even if the exchange rate is known, having some benchmark prices may help respondents get an intuitive sense of the scale of the foreign currency. To that end, immediately before the prior-belief elicitation, we showed respondents some benchmarks for the effects of non-income changes on life satisfaction, based on scientific evidence: the death of a spouse is associated with a 7.9-point reduction in life satisfaction in the year following the event; becoming unemployed is associated with a 5.2-point reduction; getting married is associated with a 4.6-point increase; and deactivating Facebook for four weeks increases life satisfaction by about 3.1 points during that period.⁹

We elicited two types of prior beliefs: average-MSI beliefs and self-MSI beliefs. For average-MSI beliefs, respondents considered an average person who earns \$50,000 per year and reports life satisfaction of 75 out of 100. The question then asked how much an unexpected 20% income increase, corresponding to an additional \$10,000 in income, would change that person’s life satisfaction over the next 12 months. For self-MSI beliefs, respondents were asked the analogous question about themselves: how much an unexpected 20% increase in their own income would change their own life satisfaction over the next 12 months.

In both cases, respondents answered using the following options: a decrease of 1 point or more, no change, an increase of 1 point, an increase of 2 points, continuing in one-point increments up to an increase of 14 points, and an increase of 15 points or more. In a pilot study, we randomized respondents between this multiple-choice format and an open-ended response. The distribution of answers was similar across formats, suggesting that the measure is not sensitive to this design choice. We used the multiple-choice format in the experiment because we thought it would be easier for respondents to answer. To keep the baseline survey short, we elicited confidence in these beliefs only in the shorter follow-up survey.

3.3 Information Treatment

After prior beliefs were elicited, subjects were randomly assigned to the treatment arm or the control arm, each with a 50% probability. Following best practices in information-provision experiments, we made the randomization explicit: respondents were told that some subjects

⁹The source for the first three effects is [Clark and Georgellis \(2013\)](#), and the source for the last one is [Allcott et al. \(2020\)](#).

would be randomly selected to receive the information treatment and that they would learn on the next screen whether they had been selected.¹⁰ Subjects assigned to the treatment arm received a short information module summarizing the scientific evidence on the effect of income on life satisfaction. The central message of this module was that the effect of a 20% increase in income on life satisfaction is positive but modest, approximately one point on a 0–100 life-satisfaction scale. Respondents in the control arm saw an otherwise identical survey, except without the information module.

The information module was split across three consecutive screens, reproduced in Figure 1. This structure allowed respondents to digest the evidence gradually. The first screen focused on the empirical correlation between income and life satisfaction. It presented evidence from the 2018 General Social Survey, a large U.S. survey that asked respondents about both household income and life satisfaction. The module explained that, after converting life satisfaction to a 0–100 scale and estimating the relationship between life satisfaction and log income, a 20% increase in income is associated with an increase in life satisfaction of approximately one point. This first screen was meant to make the magnitude transparent, but it was framed explicitly as correlational evidence. This magnitude is consistent with estimates from other large surveys; we chose the 2018 GSS only for reasons of reproducibility and convenience.¹¹

The goal of the second screen was to show that, while the correlations reported in the first screen do not imply a causal effect, there is causal evidence that supports the same finding. It discussed [Lindqvist et al. \(2020\)](#), who use the random allocation of lottery prizes as a natural experiment to study the long-run effects of wealth on psychological well-being. The module explained that the study followed lottery winners for many years after the lottery event and found persistent effects on overall life satisfaction. When the lottery windfalls are converted into an equivalent annual income increase, the implied effect of a 20% income increase is also close to one point on a 0–100 life-satisfaction scale. Thus, the causal evidence presented in the module delivered a similar message to the raw correlation: income matters for life satisfaction, but the marginal effect is modest.

The third screen was focused on causal mechanisms, that is, the scientific evidence that may help explain why MSI is lower than some respondents expect. It discussed the phe-

¹⁰This explicit randomization is meant to prevent subjects from making any unnecessary inferences from seeing the information module: e.g., respondents could think that they were shown the module because their prior beliefs were inaccurate.

¹¹For example, using Gallup data on more than 450,000 U.S. respondents, [Kahneman and Deaton \(2010\)](#) find that life evaluation (the 0–10 Cantril ladder, an evaluative measure akin to life satisfaction) rises steadily with log income: between household incomes of roughly \$75,000 and \$105,000 it increases by about 0.23 ladder points, implying an effect of about one point on a 0–100 scale for a 20% income increase—the same magnitude as our benchmark.

nomenon of hedonic adaptation, with a link to [Di Tella et al. \(2010\)](#): people may initially feel better after an income gain, but part of that effect can fade as they adjust to the new circumstances. This screen also discussed the role of rising aspirations, with a link to [Stutzer \(2004\)](#): as income increases, the amount people view as sufficient may also rise. Finally, the module mentioned that while income can relieve financial constraints, in the end, life satisfaction may depend more strongly on non-income factors such as health, relationships, work stress, loneliness, or family conflict.

The information treatment is designed to provide the empirical counterpart of the accurate signal in the model. It provides information about average-MSI, not a personalized estimate of each respondent’s self-MSI. The information module cited the underlying research studies and used neutral language: it summarized evidence on the income–satisfaction relationship without recommending any particular belief, choice, or lifestyle.

In many information-provision experiments, the information treatment can be reduced to a single familiar statistic. For example, in studies of inflation expectations, respondents are typically accustomed to seeing official inflation statistics. In that context, the information treatment may simply provide the official inflation rate over the previous year ([Cavallo et al., 2017](#)). Our setting is different. Most respondents are unlikely to have encountered scientific estimates of the income-satisfaction relationship before, and a bare statement that a 20% income increase raises life satisfaction by about one point may therefore be difficult to interpret or trust. For this reason, the information module presented not only the headline estimate, but also the evidence behind it. The goal was to make the signal credible and interpretable while keeping the message neutral.

3.4 Posterior Beliefs

Immediately after the information-treatment stage, we re-elicited self-MSI and average-MSI beliefs. Before doing so, the survey clarified that all respondents were given the opportunity to reassess their answers, regardless of their prior responses and regardless of whether they had received the information treatment. This clarification was intended to reduce the concern that respondents would interpret the repeated question as a signal that their initial answer had been incorrect.

Comparing posterior beliefs between the treatment and control arms allows us to test whether information about average-MSI shifts both average-MSI and self-MSI beliefs, as predicted by the model. In the follow-up survey, we elicited these two beliefs again, to measure whether any effects of the information on beliefs persisted one month later.

3.5 Job-Choice Scenarios

After eliciting posterior beliefs, we measured a series of post-treatment outcomes, focusing on job-choice scenarios and real-world decisions involving trade-offs between income and non-income considerations. To tailor the job-choice scenarios to each respondent, we asked four questions near the beginning of the survey about the respondent’s current or most recent job: occupation or job title, annual earnings, typical work hours, and typical commute length.¹² We used this information to construct job-choice scenarios that would be realistic for each respondent.

More precisely, each respondent faced three job-choice scenarios. In each scenario, one option was similar¹³ to the respondent’s current or recent job, while the other option paid more but required giving up a non-income consideration. The three trade-offs were higher income versus a longer commute, higher income versus longer working hours, and higher income versus fewer hours of sleep.

It is easiest to explain these scenarios using one concrete example. Panel A of Figure 2 shows the scenario corresponding to the trade-off between income and hours worked. Respondents saw a table with a side-by-side comparison between two job offers: both had the same position title, but relative to the offer on the left, the offer on the right had 20% higher earnings but 20% higher work hours. Respondents had to choose one of four responses: “Definitely choose Job A”, “Possibly choose Job A”, “Possibly choose Job B”, or “Definitely choose Job B”.

The other two job-choice scenarios use the same format. In the income-versus-commute scenario, one job offered 20% higher earnings but had an 80% longer commute. In the income-versus-sleep scenario, one job offered 30% higher earnings but reduced sleep by 25%, from 8 hours to 6 hours.

For each scenario, respondents answered two sets of questions, in randomized order. The main set—the choice framing—asked which of the two jobs the respondent would choose in each of the three scenarios. The alternative set—the satisfaction framing—asked which option would make the respondent more satisfied with their life overall. The use of both framings allows us to separately identify the effects of the information on utility-maximizing and satisfaction-maximizing behavior (Becker and Rayo, 2008; Benjamin et al., 2014). In the follow-up survey, we re-elicited the three job-choice scenarios, but for the sake of brevity, we did so only for the choice framing.

¹²Respondents who reported not currently working were asked about their most recent job instead.

¹³We say “similar” instead of “identical” because we rounded the default option’s wages to the nearest dollar and commute time to the nearest hour, and the sleep-time numbers are fully hypothetical since we did not ask how many hours individuals sleep.

3.6 Real-World Decisions

Respondents were asked about a decision that they were currently facing or would soon face, in which they were weighing an income consideration against a non-income consideration. Respondents were asked to describe that decision in their own words. They then participated in a short AI-led interview, in which the AI interviewer asked questions to clarify the relevant alternatives and the trade-off involved. For more details about how the AI-led interview was conducted, see Appendix B.3. At the end of the interview, the AI interviewer generated a personalized binary choice and shared it with the respondent. More precisely, the AI interviewer showed a description of Option A, a description of Option B, and a description of the underlying trade-off. After seeing this information, respondents were asked to confirm whether this information accurately reflected the real-world decision they were facing. We also asked them when they expected to make this decision.

Most importantly, we asked respondents which of the two options, Option A or Option B, they were more inclined to choose. For this elicitation, we closely followed the same style used for the job-choice scenarios. More precisely, we created a table that summarized all the information relevant to the decision. An example of this table is shown in Panel B of Figure 2. This respondent was deciding whether to work extra hours (Option A) or not (Option B). The trade-off was described as making some extra money for their wedding versus having less time with their family and risking burnout. In the main choice framing, subjects selected one of four responses: “Definitely choose Option A”, “Possibly choose Option A”, “Possibly choose Option B”, or “Definitely choose Option B”. As in the job-choice module, we also asked the same question but under the satisfaction framing, in randomized order.

The follow-up survey included a few questions related to this real-world decision. Respondents were asked whether they had made the real-world decision discussed in the AI-led interview and, if so, what the outcome was. If they had not made a decision yet, we asked them which of the two options they intended to choose in the future.

3.7 Additional Outcomes

There is an additional behavioral outcome that we pre-registered as one of the main outcomes of interest. Intuitively, since subjects are completing surveys for pay on the Prolific platform, the number of surveys that they complete could be seen as a measure of the amount of work they do. Thus, our original intention was to measure whether the treatment affected the number of surveys that they subsequently completed. At baseline, we asked respondents how many Prolific surveys they aspired to complete over the next 30 days; in the follow-up survey, they reported how many surveys they had completed over the previous 30 days.

In practice, however, we found that these responses are highly unreliable, as subjects have massive misperceptions about how many surveys they actually completed. For this reason, we treat the reported survey completion data as too unreliable for the main analysis—see Appendix G for more details.

The baseline survey also included a few questions about income aspirations: the income respondents thought would allow them to live comfortably, the annual income they aspired to earn, and the income they would need in order to be satisfied with their life. These secondary outcomes were intended to assess the role of the aspirations channel; the results are reported in Appendix F.

4 Implementation Details

4.1 Fielding the Baseline Survey

The experiment was pre-registered at the AEA RCT Registry (AEARCTR-0018251) and was reviewed and approved in advance by the Institutional Review Board at UCLA. We recruited respondents for the baseline survey between April 2 and April 9, 2026. Subjects were recruited via Prolific. We advertised the study as “Scientific study”. We limited eligibility to U.S. residents and required respondents to complete the survey on a laptop or desktop computer. 3,002 respondents completed the baseline survey.¹⁴ The median completion time was 18.1 minutes. About 88.3% of respondents reported that the survey was easy or very easy.

4.2 AI-Use and Data Quality Checks

The baseline survey included several measures to validate the quality of the survey data. The Prolific platform has its own measures in place to detect and prevent AI use in surveys (Gordon, 2026). According to independent accounts, at the time that we fielded the survey the estimated share of Prolific participants using AI was reported to be very small (e.g., Çelebi et al., 2026; Affonso, 2026). However, out of an abundance of caution, we followed best practices by including our own independent checks, as in Çelebi et al. (2026). The main video-code AI-use check is based on showing a short video and asking respondents to enter the numbers shown in the video. The video-code check was passed by 95% of subjects, so the 5% who failed the check provide an upper bound for the potential share of AI-assisted responses, as humans could fail this check if they are not paying close attention. We also included CAPTCHA verification, and the mean reCAPTCHA score was 0.972 on a 0–1 scale, where

¹⁴Some respondents started the baseline survey but did not complete it and therefore are not part of our main analysis sample—for more details, see Appendix C.2.

higher values indicate a greater likelihood that the respondent is human. We also included a couple of standard attention checks at the end of the survey, which were passed by 98.0% of subjects. As another data quality check, the distribution of completion times suggests there is no concern about substantial speed-taking: only 0.9% of respondents completed the survey in under seven minutes.

4.3 Sample Restrictions

Out of an abundance of caution, we exclude the minority (5.0%) of subjects who failed the video-code check. Additionally, in information-provision experiments, it is important to handle outlier beliefs appropriately. Some respondents may report wildly inaccurate guesses not because they genuinely hold such extreme beliefs, but because they misunderstood the question or were not paying sufficient attention.¹⁵ If respondents are confused about the basic concept being measured, it is hard to interpret what the effects of information may be. To reduce sensitivity to outliers, we follow standard practice in information-provision experiments and exclude respondents with the most extreme prior beliefs or misperceptions (e.g., [Fuster et al., 2022](#); [Cullen and Perez-Truglia, 2022](#); [Giaccobasso et al., 2025](#)). Because prior beliefs in our study are measured discretely, we drop respondents who selected one of the two most extreme response options: “Decrease life satisfaction by at least 1 point” and “Increase life satisfaction by at least 15 points”.¹⁶ In the baseline specification, we take a conservative approach and exclude outliers in prior average-MSI beliefs.¹⁷ Importantly, since prior beliefs were elicited before treatment assignment, excluding respondents on the basis of their priors does not compromise the validity of the experimental variation. After applying these two restrictions, the final analytic sample contains 2,775 respondents. We also show robustness checks with alternative sample restrictions and without the prior-belief restriction altogether.

4.4 Descriptive Statistics and Randomization Balance

Column (1) of Table 1 provides summary statistics for the sample: 49% of respondents are female, 64% are White, 76% are employed, average household income is approximately

¹⁵For example, if a respondent reports a recent annual inflation rate of 50% in a country where the true rate is 5%, it is likely that the respondent misunderstood the question or lacks a basic understanding of percentages.

¹⁶Throughout the analysis, we code the two endpoint categories as -1 and +15, respectively. Thus, the belief measures are effectively bottom- and top-coded at -1 and +15 satisfaction points.

¹⁷For respondents who do not report an extreme average-MSI belief but do report an extreme self-MSI belief, our interpretation is that they understood the question but believe they themselves have extreme preferences.

\$91,010, and average baseline life satisfaction is 62.1 on the 0–100 scale. Columns (2) and (3) of Table 1 provide a balance test. These columns show average characteristics in the control and treatment arms, respectively. Column (4) reports the p-value from a two-sided test of equality of means, with standard errors shown in parentheses below each mean. The two arms are very similar across pre-treatment characteristics, including age, gender, race, education, household income, political affiliation, baseline life satisfaction, and employment status. None of the variable-by-variable differences is statistically significant at the 5% level, and the joint test of equality across all balance variables has a p-value of 0.847.

As is common in online samples, our subject pool is somewhat different from the general U.S. population: respondents are somewhat younger, more educated, and more left-leaning—see Appendix C.1 for more details.

4.5 Job-Choice Scenarios

The job-choice module used respondents’ own employment information to construct job-choice scenarios. These scenarios could only be constructed for respondents who were employed or who had held a job in the past and who entered valid information; these respondents make up a strong majority (98.1%). The average subject has annual earnings of \$60,438, works 35.8 hours per week, and has a daily commute time of 31.4 minutes.

The outcomes are coded so that higher values correspond to greater willingness to choose the higher-income option. For example, in the income-vs-work-hours scenario, Offer B was the one that offered higher earnings but longer hours. Thus, the dependent variable is coded as follows: 1 if the respondent chooses “Definitely Offer A”, 2 if “Possibly Offer A”, 3 if “Possibly Offer B”, and 4 if “Definitely Offer B”. We coded the other two job-choice scenarios accordingly. We also constructed a job-choice index that averages across the three individual outcomes.

In each of the three scenarios, the two offers were reasonably balanced in relative attractiveness on average. In the income-versus-work-hours scenario, the average outcome in the control group is 2.45. Since this value is close to 2.5 (the middle of the 1–4 scale), it means that both offers were about equally attractive on average. For the income-versus-commute scenario, the average outcome is 1.92, meaning that the offer with higher income was, on average, a bit less attractive. For the income-versus-sleep scenario, the average outcome is 2.81, meaning that the offer with higher income was, on average, a bit more attractive.¹⁸

¹⁸These averages correspond to the (main) choice framing. The averages are slightly smaller for the satisfaction framing—see the bottom of Table 2 for more details.

4.6 Real-World Decision Module

A manual inspection of conversations between respondents and the AI interviewer suggests that this module went smoothly—for samples of the full interview transcripts, see Appendix B.4. Additionally, in the feedback form at the end of the survey, some subjects explicitly mentioned that they enjoyed the conversation with the AI interviewer. Objective metrics reinforce this view. The AI interviewer was able to produce a description of the respondent’s decision for the vast majority (98.4%) of respondents.¹⁹ When shown this description, 88.4% of respondents rated it as accurate, 10.7% rated it as somewhat accurate, and the remaining 0.9% rated it as inaccurate.

The real-world decisions varied widely from respondent to respondent. We provide a more comprehensive analysis later, but we can start with a few interesting examples. One respondent was weighing whether to work longer hours to pay for an expensive wedding. Another was choosing between a higher-paying career in law and the pursuit of a passion for the circus. One respondent was considering delaying cancer treatment to avoid a loss of income. Another was deciding whether to have a second child, knowing that it would affect her earnings. One respondent was contemplating leaving an unhappy marriage despite the financial hardship that would likely follow.

Constructing the outcome variable for the real-world decisions is more challenging. For the job-choice scenarios, the outcome variable is straightforward: because we constructed the scenarios, we know exactly which of the two offers has higher earnings. For the real-world decisions, however, the choice options were generated in real time by the AI interviewer. As a result, we need to determine ex post whether the two options involve a clear trade-off between income and non-income considerations and, if so, which option offers higher income. For brevity, throughout the rest of the document we use terms such as “income considerations” and “higher-income”, but these terms should be understood broadly. In most cases, they refer to choices that directly involve higher income, such as choosing a higher-paying job or increasing income by working additional hours. In other cases, however, they refer to choices that improve the respondent’s financial situation without necessarily increasing income, such as moving to a state with a lower cost of living or choosing a job with more stable income.

Consider the previous example where Option A is working extra hours to pay for the wedding and Option B is not working extra hours. The trade-off is that working the extra hours would provide more income but leave less time to spend with family and also risk burnout. In this case, we observe a clear trade-off between income and non-income considerations,

¹⁹The AI interviewer was prompted to return an exit code in case the respondent refused to engage in a conversation. Additionally, for a few observations the AI was unable to produce valid output, most likely due to a failed connection to the server.

and it is also clear that Option A is the one that offers higher income. As a result, for this respondent the dependent variable takes the value 1 if the respondent chooses “Definitely Option B”, 2 if “Possibly Option B”, 3 if “Possibly Option A”, and 4 if “Definitely Option A”. However, for other subjects, Option B may be the one that offers higher income. For some decisions, it may be ambiguous which of the two options offers higher income, or there may not be a clear trade-off to begin with.

We used an LLM to evaluate respondents’ decisions observation by observation, supplemented by human coding for the most difficult cases. Appendix E.2 describes the prompts and coding schemes and provides examples; below we provide a brief summary. For each respondent, the LLM was asked to identify which of the two options offered higher income and to indicate whether it was highly certain in that classification. The LLM could also indicate that there was no clear trade-off between income and non-income considerations, or that it was ambiguous which option offered higher income. In addition, the LLM rated the difficulty of the classification. For the 211 most difficult cases, a human coder reassessed the LLM’s classifications.

In the baseline specification, we apply the following exclusion criteria. We start by dropping the 1.6% of observations for which the AI interviewer did not produce valid output. We then drop the observations for which the respondent rated the AI output as somewhat accurate (10.7%) or inaccurate (0.9%). We drop decisions for which there was no clear trade-off between income and non-income considerations, or for which it was ambiguous which of the two options offered higher income (4.1%). Finally, we drop decisions where a trade-off was identified but the LLM reported low confidence about which of the two options offered higher income (25.8%). The combination of these exclusion criteria means that the number of observations for the real-world outcomes is significantly lower than the corresponding number for the job-choice scenarios—1,711 versus 2,721, or 37.1% lower. This was expected: before conducting the survey, we anticipated that the outcome for the real-world decision would be missing for a significant share of respondents, since at any given point in time not every household is facing a real-world decision that involves a clear trade-off between income and non-income considerations. The lower sample size implies that the results for real-world decisions are estimated less precisely. We also report sensitivity analyses using alternative exclusion criteria.

The LLM classification is subject to a common challenge when using LLMs to classify observations: rerunning the same prompts may not always produce identical labels. To address that challenge, we provide a fully reproducible cross-check. We take the 1,711 observations from the baseline specification and identify which of the two options offers higher income using a rule-based keyword classifier. This classifier does not use an LLM or any output

from the LLM classifier; instead, it is applied directly to the raw description of the decision. The keyword classifier agrees with the LLM classifier on which option offers higher income in 97.2% of observations; see Appendix E.2 for more details.

Next, we provide descriptive analysis of the real-world decisions subjects were facing. In terms of timing, the typical decision horizon was about one month. More precisely, when asked one month later in the follow-up survey, 71.5% of subjects reported that they had already made a decision.²⁰ The real-world decisions appear to be fairly balanced: at baseline, the average outcome in the control group is 2.47. Since this value is close to 2.5, the midpoint of the 1–4 scale, this means that the higher-income and lower-income options were, on average, about equally attractive.²¹ Between the baseline and follow-up surveys, many things can happen that may sway subjects in one direction or the other. Thus, while we do not expect choices to be identical between baseline and follow-up, it is still useful to compare them. Among respondents who had reached a decision at follow-up, the option they preferred at baseline coincides with what they ultimately chose one month later in 74.3% of cases.

Next, to provide descriptive statistics about the types of income and non-income considerations subjects were trading off, we use BERTopic, a fully unsupervised topic model (Grootendorst, 2022), in combination with an LLM classifier. Appendix E.1 describes the full procedure and provides additional results. The results are presented in Figure 3. Panel A summarizes the income considerations. The most common category is a higher-paying job, which accounts for about 27% of respondents. The second-largest category, at about 19%, consists of respondents considering working more—through extra hours, additional shifts, or a second job. The remaining categories are smaller and more dispersed, including relocating for higher pay, freelance or gig work, returning to work, starting work, taking time off, working from home, promotions, and retirement-related decisions. Panel B summarizes the non-income considerations. The most common category is family time, which accounts for about 31% of respondents. The second-largest category, at about 18%, is free time, rest, and hobbies. The other categories account for more modest shares and cover a range of considerations, including schedule flexibility, proximity to family and friends, mental health, work-life balance, and shorter commutes. Panel C reports the most common pairs of income and non-income considerations. The most frequent trade-off is working more versus time for rest and hobbies, which accounts for 8.3% of respondents. The next most common pairs

²⁰In comparison, in the baseline survey 33.1% of respondents expected to make a decision within one month. The discrepancy is likely due to the fact that many of these were not one-off decisions but recurring ones: among respondents who had not yet reached a decision at follow-up, 93.1% reported that the decision was still relevant.

²¹This average, 2.47, corresponds to the choice framing; the corresponding value is slightly smaller, 2.31, under the satisfaction framing.

are taking a higher-paying job versus spending time with family, at 7.5%, and working more versus spending time with family, at 7.3%.

4.7 Fielding the Follow-Up Survey

Respondents were invited to participate in the follow-up survey approximately one month after the baseline survey. About 79.4% of subjects in the baseline sample responded to the follow-up survey, and these rates were similar between treatment and control groups.²² The follow-up survey was substantially shorter than the baseline survey, with a median duration of about 6.2 minutes. 94.9% of respondents rated it as very easy or somewhat easy to complete. The follow-up survey also included the video-code check and attention checks, and the vast majority of subjects passed them.²³

5 Empirical Results

5.1 Perceived MSI

Before testing the predictions of the model, we start by providing some descriptive evidence about self-MSI and average-MSI beliefs. Panel A of Figure 4 shows the distribution of prior beliefs. For ease of comparison, Panel A overlays the histograms of prior average-MSI beliefs (green bars) and self-MSI beliefs (brown bars). There is substantial dispersion in beliefs across subjects. For average-MSI beliefs, the median is 5, with a 25th percentile of 3 and a 75th percentile of 7—the dispersion is broadly similar for self-MSI beliefs. Comparing the two histograms also shows that the distribution of self-MSI beliefs is shifted to the right relative to the distribution of average-MSI beliefs: self-MSI beliefs are on average 1.60 points higher than average-MSI beliefs (6.80 vs. 5.20).²⁴ The binned scatterplot from Panel B of Figure 4 shows that prior average-MSI and self-MSI beliefs are highly correlated, but far from perfectly so: their correlation coefficient is 0.636 ($p < 0.001$).²⁵

²²79.5% in the control arm versus 79.1% in the treatment arm, with the difference being statistically indistinguishable from zero ($p = 0.795$). The pre-treatment characteristics are also balanced across treatment arms within the follow-up sample—see Appendix C.4. Follow-up subjects are a bit different from the baseline respondents: they are older, more educated, more active on Prolific, and have somewhat lower baseline average-MSI and self-MSI beliefs; see Appendix C.3 for more details.

²³The follow-up survey included the same video-code check used at baseline, which records typing behavior to identify non-human response patterns. 96.4% of follow-up respondents passed this check.

²⁴A small part of this gap is mechanical because we’re dropping extreme priors in average-MSI beliefs but not in self-MSI beliefs, but the gap is still similar otherwise.

²⁵This high correlation is consistent with the learning framework from Section 2, according to which respondents use their available data to learn jointly about self-MSI and average-MSI.

One common concern with survey data in general, and thus specifically with these measures of prior beliefs, is the role of measurement error. The last four panels of Figure 4 provide a sense of whether these belief measures contain meaningful signal rather than just noise—Panels C and D correspond to average-MSI beliefs while Panels E and F correspond to self-MSI beliefs. These panels are based exclusively on subjects from the control arm. In Panel C, for example, subjects were asked the same question about average-MSI twice during the baseline survey, but did not receive any information between the two elicitations. For that reason, if subjects are certain about average-MSI, we would expect them to provide the same answer consistently the two times. On the other extreme, if the answers to these two questions were uncorrelated, that would indicate that they are pure noise. Under the classical-measurement-error model, this test-retest correlation (or reliability coefficient) can be interpreted as the share of observed variation attributable to signal rather than noise. Within the baseline survey, the test-retest correlations are high: 0.727 for average-MSI beliefs and 0.871 for self-MSI beliefs. Thus, this simple reliability calculation suggests that about 73% of the variation in average-MSI beliefs and 87% of the variation in self-MSI beliefs reflects a meaningful signal. These reliabilities are comparable to, or higher than, standard benchmarks from survey measurement. For example, [Krueger and Schkade \(2008\)](#) report test-retest correlations of about 0.59 for life satisfaction and up to 0.90 for education and income.

At the same time, the within-baseline correlations may overstate reliability if some respondents remembered their earlier answers and repeated them for consistency. To address this concern, Panels D and F of Figure 4 report the correlation between baseline prior beliefs and the same belief measures elicited in the follow-up survey about one month later. This longer-horizon comparison has the opposite limitation: it mechanically understates reliability as a measure of signal, because respondents may genuinely revise their beliefs between the baseline and follow-up surveys. That is, over the course of a month, respondents may have encountered new information about the relationship between income and life satisfaction and update their beliefs accordingly.²⁶ With this caveat in mind, the one-month test-retest correlations remain meaningfully positive: 0.342 for average-MSI beliefs and 0.460 for self-MSI beliefs. Thus, even using this more demanding comparison, the beliefs still contain substantial signal.

For a complementary measure of reliability, we use the question from the follow-up survey about how confident respondents were in their average-MSI and self-MSI beliefs. On the one extreme, if someone answered “not at all confident” one could worry that their answer is a

²⁶In addition, for self-MSI beliefs, the fundamentals may themselves change: for example, a respondent who receives a large hospital bill or experiences another financial shock may have a genuinely different MSI one month later.

wild guess and thus purely noise. On the other extreme, respondents answering “extremely confident” would be expected to provide a meaningful response. The fact that the distribution of responses leans heavily towards high confidence reinforces the evidence from the test-retest correlations that the self-MSI and average-MSI belief measures contain significant signal. For example, in the control group 24.9% reported being extremely confident in their self-MSI belief, 45.0% very confident, 24.5% moderately confident, 5.0% slightly confident, and 0.6% not at all confident. Consistent with the somewhat higher test-retest correlation for self-MSI beliefs than for average-MSI beliefs, self-reported confidence was also higher for self-MSI beliefs than for average-MSI beliefs.²⁷ In light of our model, one interpretation of this finding could be that individuals have more precise information about their own satisfaction and income than about other individuals’ satisfaction and income.

As an additional test of whether the variation in prior beliefs across subjects is meaningful, we examine whether those beliefs can predict respondents’ behavior. The results are presented in Figure 5. Based on subjects from the control group only, each panel shows a binned scatterplot of the relationship between the prior self-MSI belief and one of the choices elicited in the survey. Panel A of Figure 5, corresponding to the income-versus-work-hours scenario, shows that respondents with higher prior self-MSI beliefs are more willing to choose the higher-income option. More precisely, a 1-point increase in the prior self-MSI belief is associated with a statistically significant ($p < 0.001$) increase in this behavioral outcome by 0.027 (on the 1–4 scale). Panels B and C show quantitatively and qualitatively consistent results for the income-versus-commute and income-versus-sleep choices, and Panel D shows that the results are also qualitatively and quantitatively consistent for the real-world decision. We provide a more in-depth discussion of this relationship in Section 5.4 below. In the meantime, the fact that differences in self-MSI beliefs across individuals predict differences in intended behavior is suggestive that this self-MSI belief measure is capturing something genuine about individual preferences and is far from pure noise.

5.2 Biases in Perceived MSI

This section tests the model’s first prediction: that individuals hold upward-biased self-MSI and average-MSI beliefs, with the bias weakly larger for self-MSI beliefs. Strictly speaking, measuring overestimation requires knowing the true average-MSI. Of course, this object is not known with certainty. We therefore use the scientific evidence shown to respondents as our benchmark estimate of average-MSI. For brevity, we describe beliefs above this benchmark as overestimates, but the precise comparison is between respondents’ beliefs and the scientific

²⁷In the control group, 5.2% reported being extremely confident about their average-MSI beliefs, 23.0% very confident, 45.6% moderately confident, 23.8% slightly confident, and 2.5% not at all confident.

evidence, not between respondents’ beliefs and an observed ground truth. On the other hand, this caveat is less important for the other two predictions. There, the key question is not about what happens when respondents are shown the ground truth, but what happens when they are shown an accurate *signal*. Thus, we do not need the scientific evidence to be the ground truth, we just need to assume it is an accurate *signal*.

The key evidence is shown in Panel A of Figure 4. The overlapping histograms show the distribution of prior average-MSI and self-MSI beliefs, respectively. In turn, the vertical line corresponds to the signal about average-MSI provided in the treatment message, which has a value of 1.

We start with average-MSI beliefs first, because they can be compared more directly to respondents’ priors. The vast majority (94%) of subjects’ priors are above the benchmark of +1 from the scientific evidence. If we compare the mean prior average-MSI belief (5.20) to the benchmark (+1), this implies an average overestimation of 4.20 points.²⁸ This average overestimation is not only statistically significant ($p < 0.001$) but also economically meaningful: the average prior is 5.20 times the scientific benchmark.

To assess the bias in self-MSI beliefs, ideally we would observe the *true* MSI of each respondent and compare it to the respondent’s prior self-MSI belief. Of course, it’d be impossible to estimate the true MSI of each respondent with any reasonable degree of certainty. However, while we cannot estimate the individual bias of each respondent, with a simple assumption, we can easily estimate the *average* bias in self-MSI beliefs across respondents. To be precise, borrowing the notation from the model in Section 2, let $\tilde{\beta}_i$ denote respondent i ’s prior self-MSI belief and let β_i denote respondent i ’s true self-MSI. If we average the individual biases, $\tilde{\beta}_i - \beta_i$, across respondents, we obtain the following:

$$\frac{1}{N} \sum_i (\tilde{\beta}_i - \beta_i) = \frac{1}{N} \sum_i \tilde{\beta}_i - \frac{1}{N} \sum_i \beta_i \quad (21)$$

The key assumption we need is that the +1 benchmark from the scientific evidence is an unbiased estimate of the average MSI in the subject pool. In other words, we need to assume that the MSI of the average subject is comparable to the average MSI from the scientific evidence. Indeed, this is an assumption we can test with income and life satisfaction data we collected in our survey, and this assumption seems to be broadly consistent with the data—for more details, see Appendix C.5.

In sum, while we cannot assess biases at the individual level, we can assess the average bias in self-MSI beliefs by just comparing the mean prior self-MSI belief to the +1 benchmark from the scientific evidence. Panel A of Figure 4 shows that the mean prior self-MSI belief

²⁸Misperceptions about average-MSI are consistent with a large body of evidence that individuals hold systematic misperceptions about others (Bursztyn and Yang, 2022).

is 6.80, implying an average bias of 5.80 points. As with average-MSI beliefs, the average bias in self-MSI beliefs is statistically significant and economically meaningful. Moreover, consistent with the first prediction of the model, the average bias in self-MSI beliefs (5.80 points) is (weakly) larger than the average bias in average-MSI beliefs (4.20 points).

5.3 Belief Updating and Persistence

In this section, we test the second prediction of the model: that the provision of accurate information lowers, on average, both average-MSI and self-MSI beliefs, with a larger reduction in average-MSI beliefs. The key evidence is presented in Figure 6, which shows the distribution of posterior beliefs by treatment status, separately for average-MSI and self-MSI beliefs. Panel A shows posterior average-MSI beliefs. The mean posterior average-MSI belief is 5.81 in the control arm and 2.50 in the treatment arm, implying a treatment effect of -3.319 points. Panel B shows the analogous comparison for posterior self-MSI beliefs: the mean is 6.95 in the control arm and 4.23 in the treatment arm, implying a treatment effect of -2.721 points. This pattern is the empirical counterpart of Prediction 2: under unbiased priors, accurate information about average-MSI should lower both average-MSI and self-MSI beliefs, with a larger reduction in average-MSI beliefs. The sizable belief updating is consistent with high engagement with the information treatment. All treated subjects played the video, as required to advance to the next screen; 88% watched it essentially to the end, and the median treated respondent spent 233 seconds on the information-treatment screens. A minority of treated subjects also clicked “Show more” to see additional details.²⁹

Panels C and D of Figure 6 show the corresponding follow-up beliefs elicited about one month later. This follow-up comparison provides a more demanding test of belief change: the immediate posterior elicitation took place just after the information-treatment stage, when the signal was highly salient, while the follow-up beliefs capture the component of updating that respondents retain after leaving the survey. The effects remain sizable. About one-third of the initial treatment effect persists one month later (35% of the baseline effect on average-MSI beliefs and 37% of the effect on self-MSI beliefs), and treated subjects continue to report lower average-MSI and self-MSI beliefs.³⁰ Thus, the follow-up evidence suggests that the information treatment produced durable belief updating, rather than only short-

²⁹Among treated respondents in the main analytic sample, 21.1% clicked “Show more” on treatment page 1 and 10.6% clicked “Show more” on treatment page 2. Overall, 44.7% clicked “Show more” on at least one of the three information-module “Show more” links. A negligible share clicked on the external links: 0.7% clicked the *Review of Economic Studies* link and 0.4% clicked the *New York Times* link.

³⁰The follow-up survey also asked respondents how confident they were in their average-MSI and self-MSI beliefs. We do not find significant evidence that the treatment increased confidence in these beliefs; for details, see Appendix D.1.

run responses to freshly presented information. The magnitude is also in line with prior information-provision experiments. For example, [Cavallo et al. \(2017\)](#) find that roughly 45–48% of belief updating about inflation expectations persisted in follow-up surveys conducted two to four months later.

Appendix [D.2](#) (Figure [D.1](#)) examines heterogeneity in belief updating by respondents’ prior beliefs. The baseline survey suggests some heterogeneity: respondents with higher prior beliefs appear to update more strongly in the direction of the information. However, the follow-up survey suggests that most of this heterogeneity was spurious or transitory, as the persistent treatment effects one month later are more similar across the prior-belief distribution.

5.4 Effects of Information on Choices

In this section, we test the third prediction of the model: that the provision of accurate information lowers, on average, the income-generating action. The key evidence is presented in Figure [7](#), which compares the average belief and choice outcomes between treatment and control groups. Each panel plots the mean outcome for the control and treatment arms, with 90% confidence intervals shown in brackets. Each panel shows the corresponding average treatment effect, with the robust standard error in parentheses and the p-value. While the goal of this figure is to summarize the effects on choices, Panel A begins with posterior beliefs as a benchmark, and Panels B through F show the effects on the different behavioral outcomes.

The results underlying Panel A of Figure [7](#) have been discussed above, but to summarize: the information treatment reduces average-MSI beliefs by 3.32 satisfaction points on the 0–100 scale. We report this effect for transparency, but for respondents’ own decisions, the relevant belief is the posterior self-MSI belief. The information treatment also reduces self-MSI beliefs by 2.72 satisfaction points. Both effects are highly statistically significant, with $p < 0.001$.

Panel B of Figure [7](#) corresponds to the decision involving 20% higher earnings at the expense of 20% more work hours. When asked what they would choose, subjects in the treatment arm are less inclined to choose the higher-income option: the mean falls from 2.45 in the control arm to 2.36 in the treatment arm, on the 1–4 scale coded so that 1 indicates definitely choosing the lower-income option and 4 indicates definitely choosing the higher-income option. This corresponds to an estimated effect of -0.097 ($p = 0.020$). In the satisfaction-framing version, the decline is very similar, from 2.36 to 2.27, with an estimated effect of -0.090 ($p = 0.027$). This suggests that the information treatment reduces willingness to trade away hours for higher income in both framings.

Panel C of Figure 7 considers the analogous trade-off involving commuting time. Here, again, the treatment moves responses away from the higher-income option. In the choice-framing version, the mean falls from 1.92 to 1.85, a statistically significant ($p=0.064$) effect of -0.071. In the satisfaction-framing version, the mean falls from 1.88 to 1.81, also a statistically significant ($p=0.056$) effect of -0.071. Panel D of Figure 7 shows the trade-off involving sleep. The point estimates again go in the same negative direction, but they are smaller and less precisely estimated. In the choice-framing version, the treatment reduces the outcome from 2.81 to 2.75, an effect in the same order of magnitude (-0.066) but borderline statistically insignificant ($p=0.130$). In the satisfaction-framing version, the mean falls from 2.65 to 2.61, an effect that is still negative but smaller (-0.033) and statistically insignificant ($p=0.447$). To maximize statistical power, Panel E of Figure 7 averages the three outcomes from Panels B through D. The treatment reduces the choice-framing index from 2.39 to 2.32, a statistically significant ($p=0.007$) effect of -0.078. For the satisfaction-framing index, the mean falls from 2.30 to 2.23, a statistically significant effect of -0.065 ($p=0.027$).

Panel F of Figure 7 corresponds to the real-world decision. The estimated treatment effects are still negative and in the same order of magnitude as for the other outcomes in Panels B–E, but less precisely estimated. For the choice-framing version of the real-world decision, the mean falls from 2.47 to 2.41, an effect of -0.066 that is not statistically significant ($p=0.171$). For the satisfaction-framing version, the mean falls from 2.31 to 2.21, a statistically significant effect of -0.096 ($p=0.047$). The signs are consistent with the job-choice outcomes, though the estimates are noisier.

Taken together, the pattern across Panels B through E of Figure 7 is reassuring. Even when some outcome-specific estimates are not statistically significant, the effects always point in the same direction and are in the same order of magnitude. Overall, the evidence suggests that the treatment shifted respondents away from higher-income options.

A more straightforward way to interpret the findings is to compare the size of the treatment-induced change in self-MSI with the corresponding change in choices. Table 2 provides this comparison. More precisely, Panel A estimates the relationship between posterior self-MSI ($MSI_{i,post}^{self}$) and choice outcome k (C_i^k) using the following OLS model:

$$C_i^k = \mu_0^k + \mu_{MSI}^k \cdot MSI_{i,post}^{self} + u_i^k \quad (22)$$

For example, in column (1) of Panel A of Table 2, we estimate a regression of the choice-framing work-hours trade-off outcome on posterior self-MSI. Variation in self-MSI is both experimental and non-experimental. By combining the experimental and non-experimental variation, the OLS estimate maximizes statistical power. However, to the extent that the

non-experimental variation in self-MSI is correlated with other unobservable determinants of choice, the OLS coefficient would suffer from omitted-variable bias. We therefore view Panel A as a stepping stone and turn to the experimental estimates later.

We start by discussing the results from the choice framing. In column (1) of Table 2, the coefficient implies that a 1-point decrease in self-MSI is associated with a 0.025-point decrease in willingness to choose higher income over lower work hours, measured on the 1–4 scale. Relative to the control-group standard deviation, this corresponds to about 0.023 standard deviations. Column (2) shows a similar relationship for the commuting trade-off: a 1-point decrease in self-MSI is associated with a 0.020-point decrease in willingness to choose higher income over a shorter commute, or about 0.020 standard deviations. Column (3) shows the largest outcome-specific estimate: a 1-point decrease in self-MSI is associated with a 0.034-point decrease in willingness to choose higher income over more sleep, equal to about 0.030 standard deviations. Consistent with these component outcomes, column (4) shows that the index of the three job-choice outcomes declines by 0.026 points, or about 0.034 standard deviations, for each 1-point decrease in self-MSI. Finally, column (5) shows a similar association for the real-world decision outcome: a 0.027-point decline, or about 0.027 standard deviations.

Overall, Panel A of Table 2 suggests that a lower self-MSI is associated with lower willingness to choose the higher-income option. The satisfaction-framing results are very similar to the choice-framing results, both in magnitude and in statistical significance. The corresponding coefficients are 0.026 for work hours, 0.020 for commute, 0.031 for sleep, 0.026 for the index, and 0.032 for the real-world decision outcome. In standard-deviation units, these correspond to approximately 0.024, 0.020, 0.028, 0.034, and 0.033 standard deviations, respectively.

Panel B of Table 2 estimates the effect of posterior self-MSI on behavior using a 2SLS specification, with the treatment indicator (T_i) as the excluded instrument. This specification isolates the component of posterior self-MSI shifted by the randomized information treatment, rather than by pre-existing differences across respondents. Because treatment was randomly assigned, this induced variation in posterior self-MSI is exogenous. The causal interpretation still requires an exclusion restriction: the information treatment must affect choices only through changes in posterior beliefs ($MSI_{i,post}^{self}$).³¹

Starting with the choice-framing outcomes, column (1) of Table 2 shows that a 1-point decrease in posterior self-MSI causes a 0.035-point decrease in willingness to choose higher

³¹As in standard 2SLS settings, when treatment effects vary across individuals, the estimate captures a local average treatment effect of beliefs, in the sense of Imbens and Angrist (1994). In this context, that means the estimated treatment effect is weighted more heavily toward individuals whose beliefs changed the most in response to the information.

income over lower work hours on the 1–4 scale. Relative to the control-group standard deviation, this corresponds to about 0.033 standard deviations. Column (2) shows a similar effect for commuting: a 1-point decrease in posterior self-MSI causes a 0.026-point decrease in willingness to choose higher income over a shorter commute, or about 0.025 standard deviations. Column (3) shows an effect of 0.024 points for sleep, equal to about 0.021 standard deviations. Column (4), which averages the three job-choice outcomes, implies an effect of 0.028 points, or about 0.037 standard deviations. Finally, column (5) shows a smaller effect for the real-world decision outcome, with a 0.024-point decline, or about 0.024 standard deviations.

Under the satisfaction framing, the 2SLS estimates are again positive, implying that reductions in posterior self-MSI reduce respondents’ expected satisfaction from the higher-income option. The results under the satisfaction framing are also similar in magnitude to the choice-framing estimates. The corresponding 2SLS coefficients are 0.033 for work hours, 0.026 for commute, 0.012 for sleep, 0.024 for the index, and 0.035 for the real-world decision outcome. In standard-deviation units, these correspond to approximately 0.031, 0.026, 0.011, 0.031, and 0.035 standard deviations, respectively. The significance pattern is also similar under the satisfaction and choice framings, with the exception of the real-world decision, which is statistically significant under the satisfaction framing but not under the choice framing.

The comparison between Panels A and B of Table 2 is useful. The key trade-off is that Panel B has a cleaner causal interpretation, because it isolates the component of posterior self-MSI shifted by the randomized information treatment, but it is also less precisely estimated because it uses only a small share of the overall variation in self-MSI: the shocks induced by the experiment. Despite this loss of precision, the estimates in Panels A and B are remarkably similar in sign and magnitude. This is reassuring. While some specific coefficients are statistically insignificant, the fact that the estimates line up across outcomes, framings, and specifications suggests a stable pattern: lower posterior self-MSI shifts respondents away from higher-income options. Appendix D.4 confirms that the job-choice index estimates hold up across a range of alternative samples and specifications. Regarding the real-world decisions, however, the results depend more strongly on some of the criteria used for coding this outcome, for example, excluding observations with ambiguous trade-offs—for more details, see Appendix D.5.

Next, we evaluate whether the experimental effects persisted at follow-up. The relevant results are presented in Panel C of Table 2. This panel is analogous to Panel B, but both the outcome and the endogenous variable are measured at follow-up. Specifically, the dependent variable is willingness to choose the higher-income option at follow-up, the endogenous vari-

able is posterior self-MSI measured at follow-up, and the instrument is again the randomized treatment indicator. Thus, Panel C asks whether the experimentally induced changes in self-MSI continue to translate into changes in choices at follow-up.

Two features of the follow-up data reduce precision. First, attrition lowers sample sizes. For the job-choice scenarios, the sample size declines by 20.6 percent, from 2,721 respondents at baseline to 2,161 at follow-up. For the real-world decision outcome, it falls by 21.4 percent, from 1,711 respondents at baseline to 1,345 at follow-up. Second, the first stage is weaker at follow-up. Although the treatment effect on self-MSI remains statistically significant, it is only partially persistent: the effect falls from 2.72 points at baseline to about 0.99 points at follow-up, roughly one-third of the baseline effect. Together, the smaller sample and weaker first stage make the follow-up 2SLS estimates substantially less precise.

With those caveats in mind, the point estimates in Panel C of Table 2 remain broadly consistent with Panel B. In column (1) of the table, the coefficient is 0.031 in Panel C, compared with 0.035 in Panel B. These magnitudes are quite similar. However, the standard error rises sharply: from 0.015 in Panel B to 0.046 in Panel C, roughly three times as large. This means that, at follow-up, we would not be able to detect effects of the size estimated in Panel B. Taken together across all outcomes, the follow-up estimates are broadly consistent with the baseline causal estimates, but too noisy to draw strong conclusions about persistence. The results are therefore best read as suggestive: they do not contradict the baseline pattern, but they also do not allow us to estimate the follow-up effects with enough precision.

Lastly, a common concern when estimating 2SLS models is weak-instrument bias (Stock et al., 2002). Given the strong effect of the treatment on posterior self-MSI, this is unlikely to be a concern in our application. For a formal assessment, the bottom of Table 2 reports the Kleibergen-Paap rk Wald F-statistics for each 2SLS specification, a standard measure of instrument strength.³² Following the guideline of Staiger and Stock (1997), F-statistics of 10 or higher indicate that weak identification is not a serious concern. Our reported values are well above this threshold, confirming that weak instruments are not a concern. Appendix D.3 shows that this first-stage effect of treatment on posterior self-MSI is large and highly significant across all alternative samples and specifications we consider.

³²Although the conventional rule of thumb is based on the Cragg-Donald statistic, which assumes homoskedastic errors, Baum et al. (2007) recommend using the Kleibergen-Paap statistic as its robust counterpart.

6 Conclusions

We conducted a pre-registered information-provision experiment to study whether individuals misperceive MSI and whether these misperceptions affect choices. We present a simple model of misspecified learning that makes a number of predictions. Guided by that model, the experiment elicits self-MSI and average-MSI beliefs and provides randomly assigned scientific evidence about average-MSI. The results are consistent with the predictions of the model. Respondents hold upward-biased average-MSI and self-MSI beliefs, with larger bias in self-MSI beliefs. When presented with scientific evidence showing that the effect of income on life satisfaction is more modest than they expected, respondents revise their beliefs downward. These belief changes are not merely transitory: a meaningful share persists one month later. Most importantly, the belief changes translate into choices. Treated respondents become less willing to choose the higher-income option when doing so requires giving up other dimensions of well-being such as hours worked, commute time, or sleep.

Taken together, the results suggest that misperceptions about the income–satisfaction relationship can be consequential for economic decision-making. Individuals appear to place too much weight on income partly because they overestimate how much additional income will improve their life satisfaction. Correcting this belief shifts choices away from higher-income options, which suggests that beliefs about subjective well-being are not only measurable and malleable, but also behaviorally relevant.

References

- Affonso, F. M. (2026). Brief commentary: A framework for detecting ai agents in online research. *Journal of Consumer Research*, ucag006.
- Aknin, L. B., M. I. Norton, and E. W. Dunn (2009). From Wealth to Well-Being? Money Matters, but Less Than People Think. *The Journal of Positive Psychology* 4(6), 523–527.
- Allcott, H., L. Braghieri, S. Eichmeyer, and M. Gentzkow (2020). The Welfare Effects of Social Media. *American Economic Review* 110(3), 629–676.
- Apouey, B. and A. E. Clark (2015). Winning Big but Feeling No Better? The Effect of Lottery Prizes on Physical and Mental Health. *Health Economics* 24(5), 516–538.
- Baum, C. F., M. E. Schaffer, and S. Stillman (2007). Enhanced Routines for Instrumental Variables/GMM Estimation and Testing. *The Stata Journal* 7(4), 465–506.
- Becker, G. S. and L. Rayo (2008). Comment on “Economic Growth and Subjective Well-Being: Reassessing the Easterlin Paradox” by Betsey Stevenson and Justin Wolfers. *Brookings Papers on Economic Activity* 39(1), 88–95. Spring.
- Benjamin, D. J., O. Heffetz, M. S. Kimball, and A. Rees-Jones (2014). Can Marginal Rates of Substitution Be Inferred from Happiness Data? Evidence from Residency Choices. *American Economic Review* 104(11), 3498–3528.
- Bursztyn, L. and D. Y. Yang (2022). Misperceptions About Others. *Annual Review of Economics* 14(1), 425–452.
- Cavallo, A., G. Cruces, and R. Perez-Truglia (2017). Inflation Expectations, Learning, and Supermarket Prices: Evidence from Survey Experiments. *American Economic Journal: Macroeconomics* 9(3), 1–35.
- Çelebi, C., C. Exley, S. Harris, H. Kivimäki, M. Serra-Garcia, and J. Yusof (2026). Mission Possible: The Collection of High-Quality Online Data. *CESifo Working Paper No. 12535*.
- Chopra, F. and I. Haaland (2023). Conducting Qualitative Interviews with AI. CESifo Working Paper 10666, CESifo.
- Clark, A. E., P. Frijters, and M. A. Shields (2008). Relative Income, Happiness, and Utility: An Explanation for the Easterlin Paradox and Other Puzzles. *Journal of Economic Literature* 46(1), 95–144.
- Clark, A. E. and Y. Georgellis (2013). Back to Baseline in Britain: Adaptation in the British Household Panel Survey. *Economica* 80(319), 496–512.
- Clark, A. E. and C. Senik (2010). Who Compares to Whom? The Anatomy of Income Comparisons in Europe. *The Economic Journal* 120(544), 573–594.
- Cone, J. and T. Gilovich (2010). Understanding Money’s Limits: People’s Beliefs About the Income-Happiness Correlation. *The Journal of Positive Psychology* 5(4), 294–301.
- Costello, T. H., G. Pennycook, and D. G. Rand (2024). Durably Reducing Conspiracy Beliefs Through Dialogues with AI. *Science* 385(6714), eadq1814.
- Cullen, Z. and R. Perez-Truglia (2022). How Much Does Your Boss Make? The Effects of Salary Comparisons. *Journal of Political Economy* 130(3), 766–822.
- Di Tella, R., J. Haisken-De New, and R. MacCulloch (2010). Happiness Adaptation to Income and to Status in an Individual Panel. *Journal of Economic Behavior & Organization* 76(3), 834–852.
- Fuster, A., R. Perez-Truglia, M. Wiederholt, and B. Zafar (2022). Expectations with Endogenous Information Acquisition: An Experimental Investigation. *Review of Economics and Statistics* 104(5), 1059–1078.
- Galiani, S., P. J. Gertler, and R. Undurraga (2018). The Half-Life of Happiness: Hedonic Adaptation in the Subjective Well-Being of Poor Slum Dwellers to the Satisfaction of Basic Housing Needs. *Journal of the European Economic Association* 16(4), 1189–1233.
- Geiecke, F. and X. Jaravel (2026). Conversations at Scale: Robust AI-Led Interviews with a Simple

- Open-Source Platform. *The Review of Economic Studies*. Forthcoming.
- Giaccobasso, M., B. Nathan, R. Perez-Truglia, and A. Zentner (2025). Where Do My Tax Dollars Go? Tax Morale Effects of Perceived Government Spending. *American Economic Journal: Applied Economics* 17(4), 223–259.
- Gordon, A. (2026). Authenticity Checks. <https://www.prolific.com/resources/authenticity-checks-how-we-tested-the-most-accurate-method-for-identifying-agentic-ai>, accessed May 18, 2026.
- Grootendorst, M. (2022). BERTopic: Neural Topic Modeling with a Class-Based TF-IDF Procedure. *arXiv preprint arXiv:2203.05794*.
- Imbens, G. W. and J. D. Angrist (1994). Identification and Estimation of Local Average Treatment Effects. *Econometrica* 62(2), 467–475.
- Kahneman, D. and A. Deaton (2010). High income improves evaluation of life but not emotional well-being. *Proceedings of the National Academy of Sciences* 107(38), 16489–16493.
- Kahneman, D. and R. H. Thaler (2006). Anomalies: Utility Maximization and Experienced Utility. *Journal of Economic Perspectives* 20(1), 221–234.
- Krueger, A. B. and D. A. Schkade (2008). The Reliability of Subjective Well-Being Measures. *Journal of Public Economics* 92(8–9), 1833–1845.
- Lee, D. S. (2009). Training, wages, and sample selection: Estimating sharp bounds on treatment effects. *The Review of Economic Studies* 76(3), 1071–1102.
- Lindqvist, E., R. Östling, and D. Cesarini (2020). Long-Run Effects of Lottery Wealth on Psychological Well-Being. *The Review of Economic Studies* 87(6), 2703–2726.
- Loewenstein, G., T. O’Donoghue, and M. Rabin (2003). Projection Bias in Predicting Future Utility. *The Quarterly Journal of Economics* 118(4), 1209–1248.
- Luttmer, E. F. P. (2005). Neighbors as Negatives: Relative Earnings and Well-Being. *The Quarterly Journal of Economics* 120(3), 963–1002.
- Rayo, L. and G. S. Becker (2007). Evolutionary Efficiency and Happiness. *Journal of Political Economy* 115(2), 302–337.
- Staiger, D. and J. H. Stock (1997). Instrumental Variables Regression with Weak Instruments. *Econometrica* 65(3), 557–586.
- Stock, J. H., J. H. Wright, and M. Yogo (2002). A Survey of Weak Instruments and Weak Identification in Generalized Method of Moments. *Journal of Business & Economic Statistics* 20(4), 518–529.
- Stutzer, A. (2004). The Role of Income Aspirations in Individual Happiness. *Journal of Economic Behavior & Organization* 54(1), 89–109.
- Wilson, T. D. and D. T. Gilbert (2003). Affective Forecasting. In M. P. Zanna (Ed.), *Advances in Experimental Social Psychology*, Volume 35, pp. 345–411. Academic Press.

Figure 1: Treatment Screens

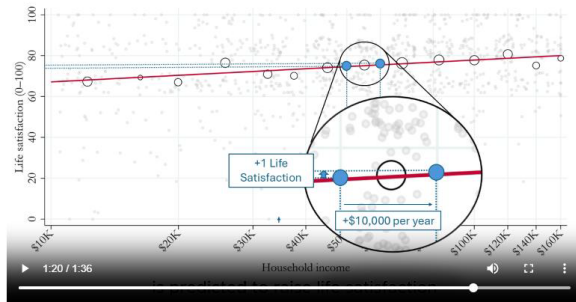
PANEL A: First Treatment Screen

Income and Life Satisfaction

There is a large academic literature in economics and other social sciences examining the effects of income on subjective well-being, including happiness and life satisfaction. Across many studies conducted over several decades, a common conclusion is that income does affect happiness and life satisfaction, but that the effect is modest in size.

To illustrate the kind of relationship these studies typically find, we present evidence from the General Social Survey, a large U.S. survey that interviews thousands of Americans each year. In these data, a 20% increase in annual income—such as from \$50,000 to \$60,000—is associated with a 1-point increase in life satisfaction on a 0–100 scale in the same year.

Please take a moment to watch the video below for a more detailed explanation:



PANEL B: Second Treatment Screen

What Science Says on the Causal Effects of Income on Life Satisfaction

In the previous page, we showed you a simple correlation using real survey data. However, more rigorous analyses that approximate scientific experimental studies find surprisingly similar numbers.

A landmark study by Economics Professors Lindqvist, Östling, and Cesarini study the effects of lottery winnings on life satisfaction. Their [research study](#) was peer-reviewed and published at *The Review of Economic Studies*, which is widely regarded as one of the best journals in the economics profession. It has also been featured in international media outlets such as [The New York Times](#). Exploiting the random allocation of lottery prizes, they track winners for up to 22 years. They find that large windfalls lead to persistent increases in overall life satisfaction 5–22 years after the event, with no evidence of the effects vanishing over time. By converting the lottery prizes effects into an equivalent annual income increase lasting 20 years, a 20% raise in yearly income causes life satisfaction to increase by almost 1 point (on the 0–100 scale), which is similar to our estimates from the previous page.



PANEL C: Third Treatment Screen

Why Higher Income Doesn't Always Bring More Satisfaction

Social scientists have conducted hundreds of studies to understand why the effect of income on life satisfaction is often modest. Many explanations have been proposed; below are some common explanations.

1 Hedonic Adaptation

People often get used to improvements.

As we earn more, our sense of "enough" also increases.



One reason is that people often adapt to improvements in their circumstances. A raise, bonus, or windfall may feel significant at first, but over time a higher level of consumption becomes the new normal. In the well-being literature, this process is usually called hedonic adaptation. A new car, for example, may be exciting initially, but that satisfaction tends to fade, and after a few months one may feel much the same as while driving the old car.

[Read more](#)

2 Rising Aspirations

As we earn more, our sense of "enough" also increases.



A closely related explanation is that income aspirations rise with income itself. Survey-based research shows that as people earn more, they also tend to report a higher level of income as "sufficient." Intuitively, for someone earning \$50,000 a year, \$100,000 may seem more than sufficient; yet if that person eventually reaches \$100,000, it may not be long before \$200,000 comes to seem like what they need to live comfortably.

[Read more](#)

3 Limited Impact on Life Satisfaction

Money can't fix everything.

Money can't fix everything.



A third reason is that higher income mainly relieves certain problems, especially financial strain, rather than transforming every aspect of life. More income can reduce worry about bills, debt, and emergencies, which can in turn improve life satisfaction. But many other important determinants of life satisfaction — such as health, relationships, working conditions, loneliness, purpose, and family conflict — do not disappear simply because income rises. Extra income can therefore help, but only to a limited extent.

[Read more](#)

Notes: The information treatment consisted of three consecutive screens, with one screen shown in each panel. For full-page screenshots of each screen, see Appendix H.

Figure 2: Examples of Choice Elicitation Modules

PANEL A: Job-Choice Module



Please consider the following two job opportunities. One offers higher pay but requires longer working hours. Aside from the differences summarized below, the two options are identical in all other respects.

If you were limited to these two options, which do you think you would choose?

	Offer A	Offer B
Position Title	Lead Operations Processor	Lead Operations Processor
Weekly Pay	\$1,173	\$1,408
Weekly Hours Worked	40 hours	48 hours
	☐ definitely choose A	☐ possibly choose A
	☐ possibly choose B	☐ definitely choose B

PANEL B: Real-World Decision Module



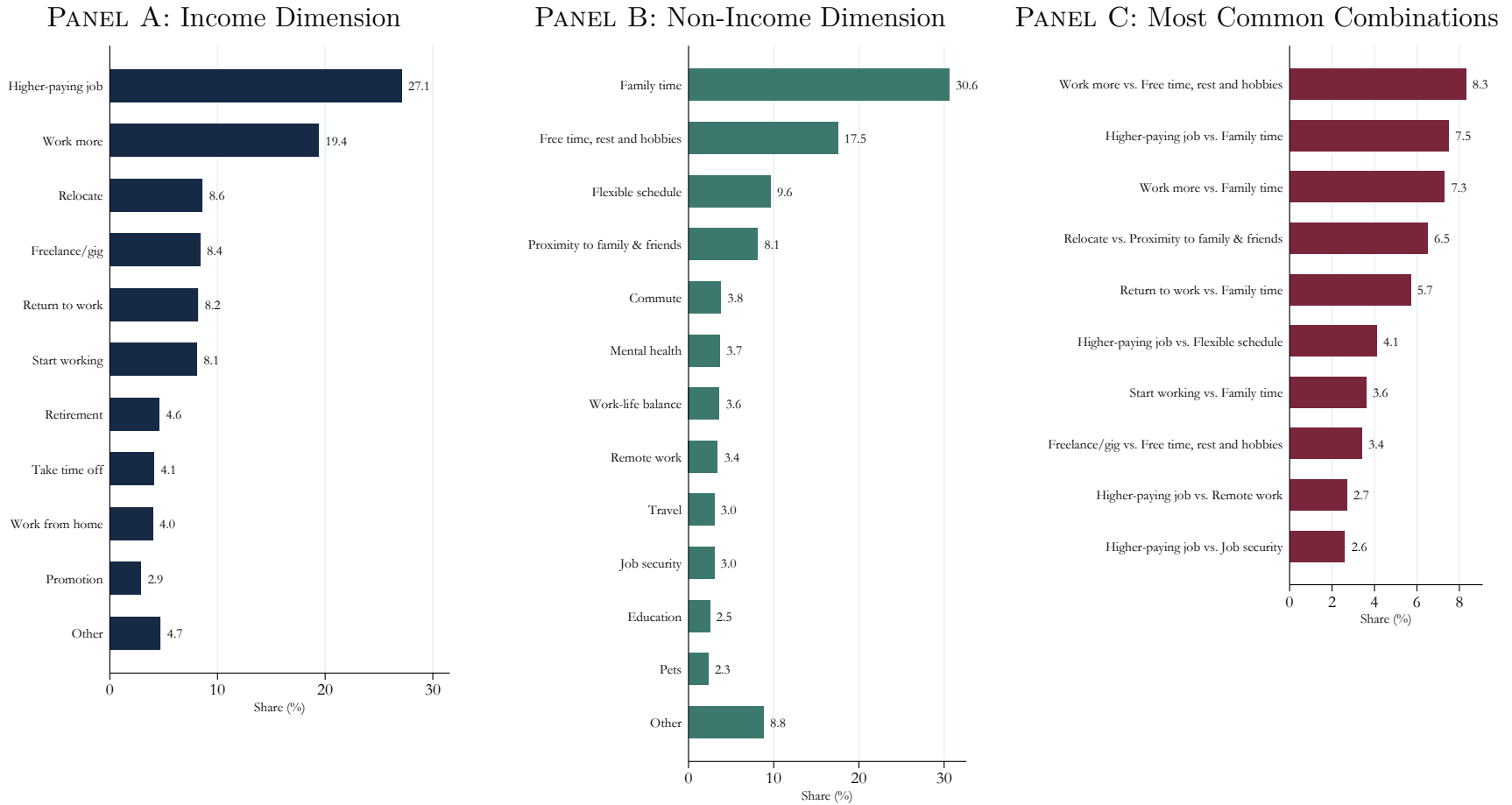
Let's start with the decision you discussed with the AI assistant earlier. If you work extra hours, you would earn more money for the wedding, but you would spend less time with your family and risk burnout.

If you were limited to these two options, which do you think you would choose?

Option A: Work extra hours for more wedding money	Option B: Do not work extra hours
☐ definitely choose A	☐ possibly choose B
☐ possibly choose A	☐ definitely choose B

Notes: This figure shows examples of the two choice elicitation modules from a real survey participant in our main sample. Panel A shows the hours worked job-choice scenario. Panel B shows the real-world decision that was generated from information provided by the respondent, which in this case also involves extra hours worked.

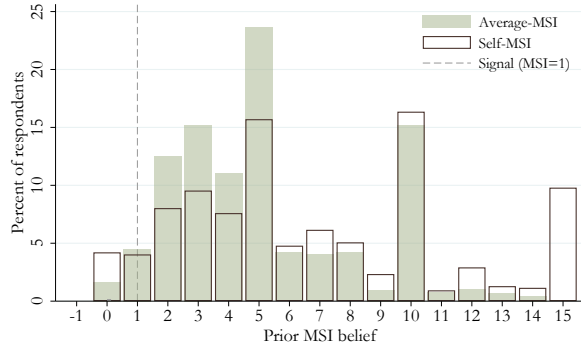
Figure 3: What Real-World Decisions Weigh Against Income



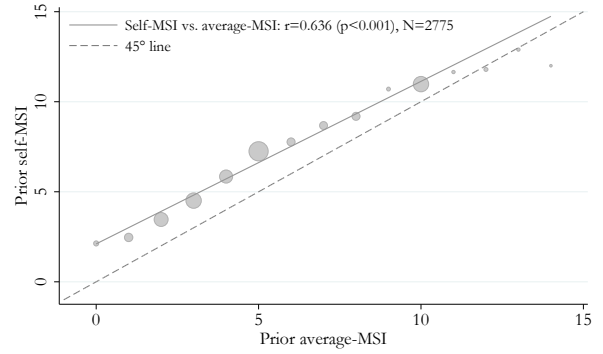
Notes: Distribution of the real-world decisions in the baseline analysis sample by the income consideration at stake (Panel A) and the non-income consideration weighed against income (Panel B). For each decision we ask a language model (gpt-5-mini) to describe each side in the respondent's own words; the descriptions are then clustered with an unsupervised topic model (BERTopic), and closely related topics are grouped into the categories shown. Appendix E.1 gives the full procedure and the complete, unaggregated topics (Tables E.1–E.2). Panel C shows the ten most common combinations of an income and a non-income category, which together account for about 52% of the decisions with an identified theme on both sides. Categories under 2% are grouped into “Other” in Panels A–B. Bars are the percentage of decisions (Panel A: 1,710; Panel B: 1,708; Panel C: 1,498 with a theme on both sides).

Figure 4: Prior Beliefs and Measurement

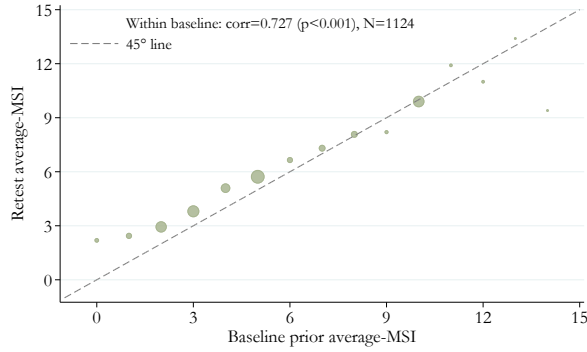
PANEL A: Distribution of Prior Beliefs



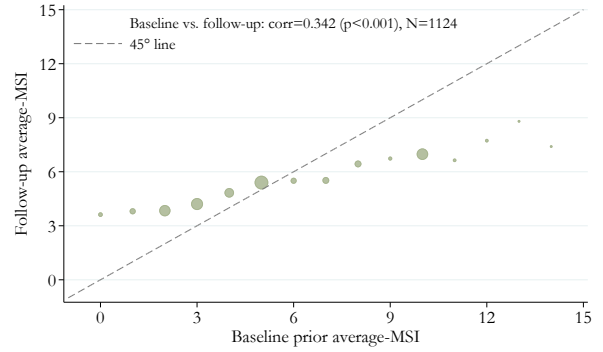
PANEL B: Self-MSI vs. Average-MSI



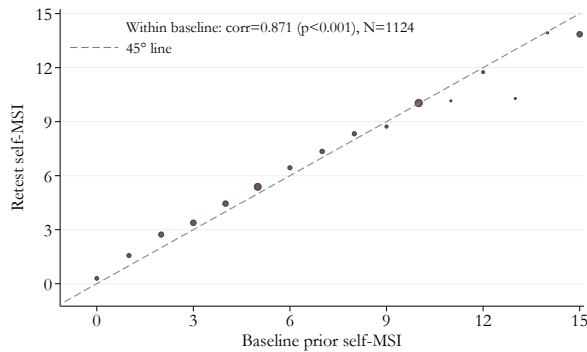
PANEL C: Test-Retest, Average-MSI (Within Baseline)



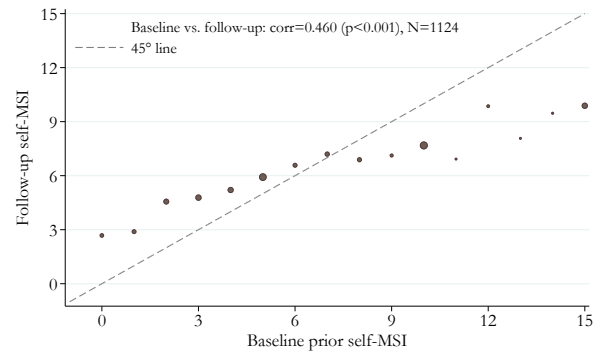
PANEL D: Test-Retest, Average-MSI (Baseline vs. Follow-Up)



PANEL E: Test-Retest, Self-MSI (Within Baseline)



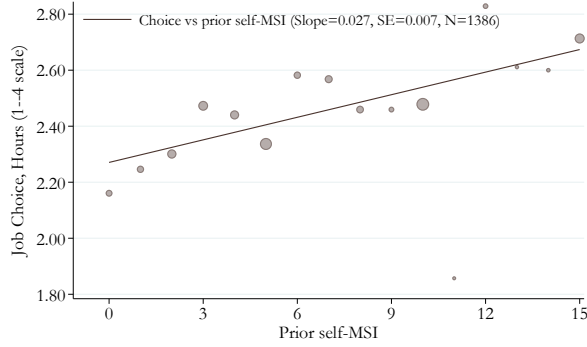
PANEL F: Test-Retest, Self-MSI (Baseline vs. Follow-Up)



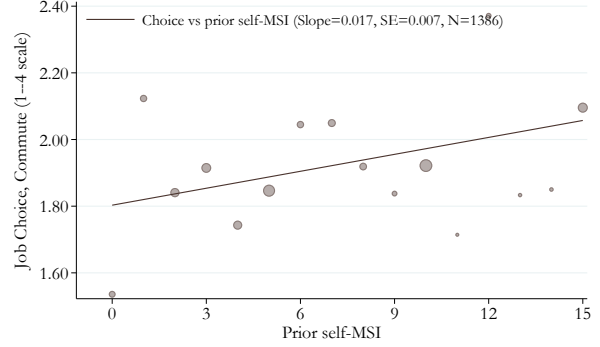
Notes: Panel A reports the distribution of prior average-MSI and self-MSI beliefs. Panel B compares prior self-MSI beliefs to prior average-MSI beliefs. Panels C–F report test-retest checks for the belief measures, based exclusively on control-arm respondents who also completed the follow-up survey. Panels C and D correspond to average-MSI beliefs and Panels E and F to self-MSI beliefs. Panels C and E compare the two baseline elicitations (within-baseline test-retest); Panels D and F compare the baseline elicitation to the same belief measured in the follow-up survey about one month later. Each legend reports the correlation coefficient, its p-value, and the number of observations.

Figure 5: Prior Self-MSI Beliefs and Choices in the Control Arm

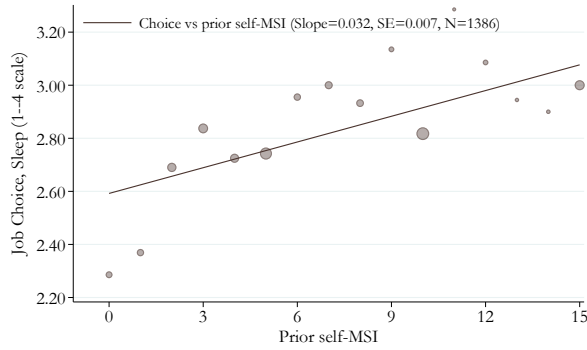
PANEL A: Income-versus-Hours-Worked



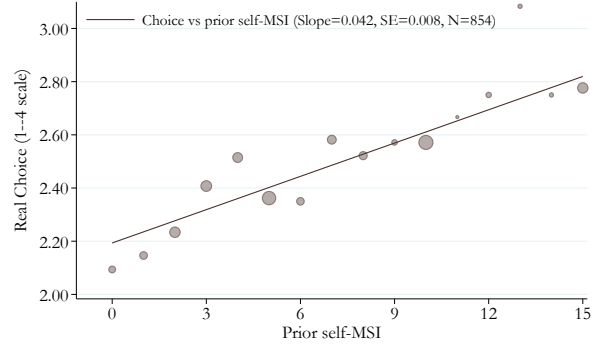
PANEL B: Income-versus-Commute



PANEL C: Income-versus-Sleep



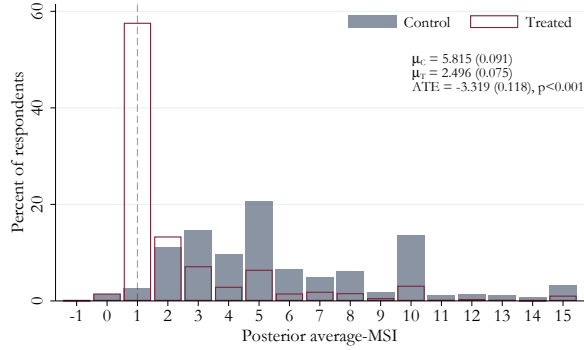
PANEL D: Real-World Decision



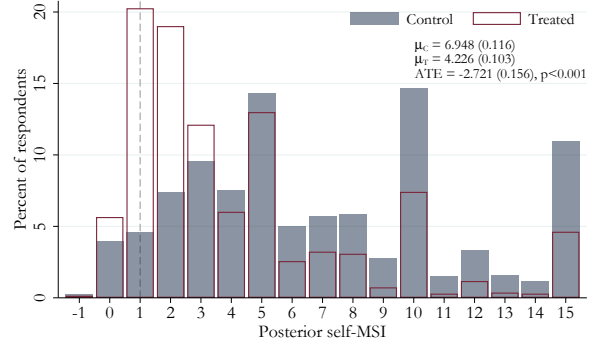
Notes: This figure reports binned scatterplots for subjects in the control arm. Panels A–C relate prior self-MSI beliefs to the three job-choice outcomes. Panel D relates prior self-MSI beliefs to the real-world decision outcome. Higher values of the choice outcomes indicate greater willingness to choose the higher-income option. The choice outcomes are on the raw 1–4 scale.

Figure 6: Posterior Beliefs by Treatment Status

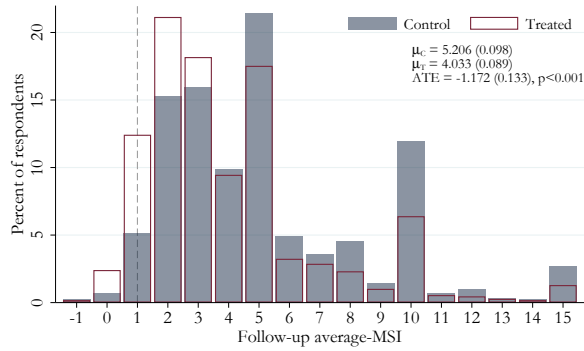
PANEL A: Posterior Average-MSI Beliefs



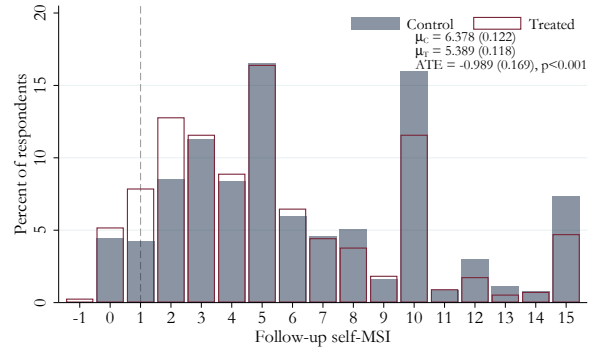
PANEL B: Posterior Self-MSI Beliefs



PANEL C: Follow-Up Average-MSI Beliefs

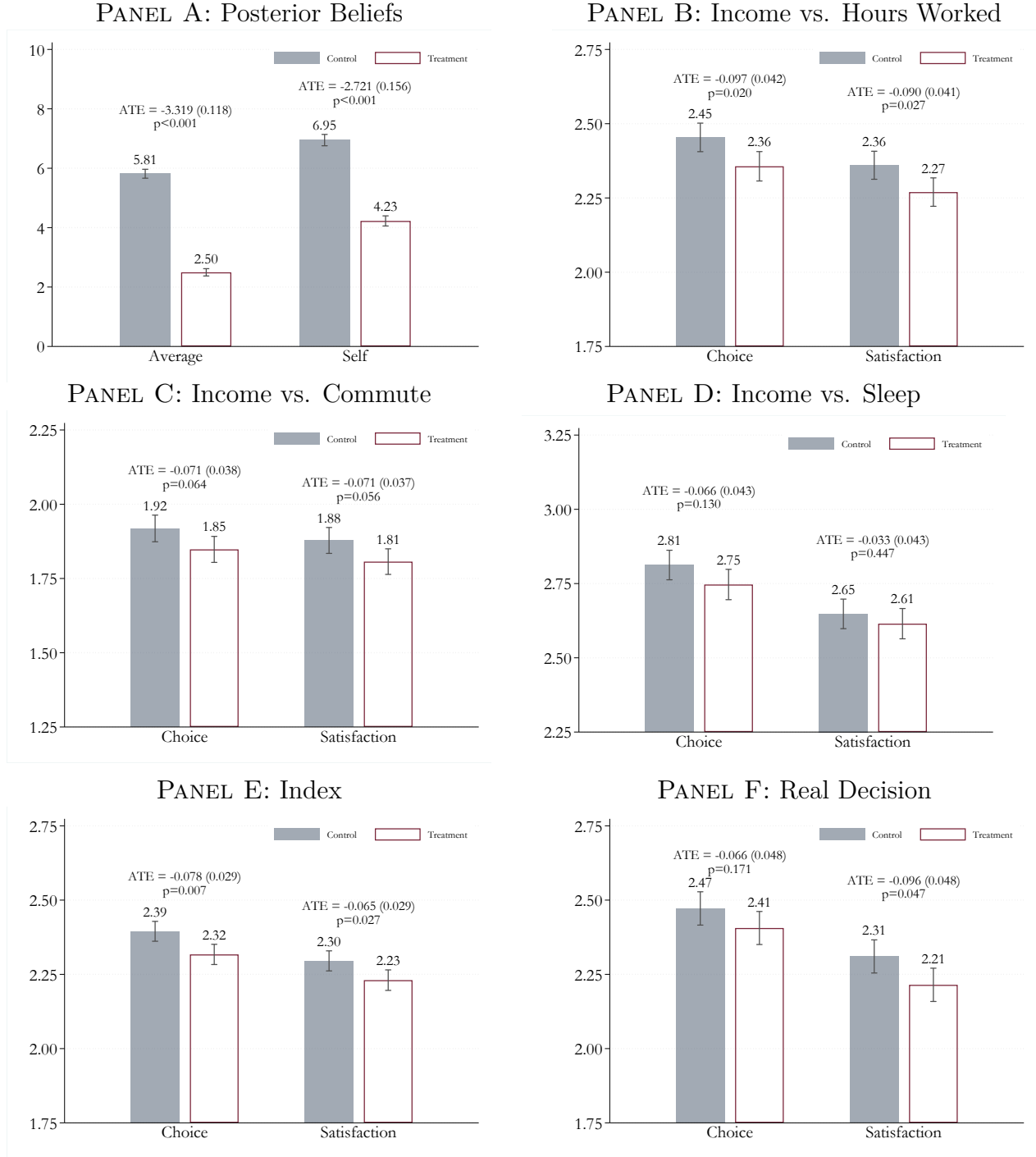


PANEL D: Follow-Up Self-MSI Beliefs



Notes: Panels A and B report the distribution of posterior average-MSI and self-MSI beliefs, elicited immediately after treatment, by treatment status. Panel A shows average-MSI beliefs, and Panel B shows self-MSI beliefs. Panels C and D report the corresponding posterior beliefs at the one-month follow-up. The dashed vertical line marks the treatment signal (MSI = 1). Within each panel we report the control-group mean (μ_C), the treated-group mean (μ_T), and the average treatment effect (ATE = $\mu_T - \mu_C$), with robust standard errors in parentheses and the associated p-value.

Figure 7: Average Treatment Effects of the Information



Notes: Bars are group means with 90% confidence intervals; each annotation gives the average treatment effect, robust standard error in parentheses, and p-value. In Panel A, the outcomes are the posterior beliefs about average-MSI and self-MSI, in life-satisfaction points (0–100 scale). Panels B–E are the job-choice income trade-offs against work hours, commute, sleep, and their index; Panel F is the real-world decision. These behavioral outcomes are codes on a 1–4 scale so that a higher value corresponds to a preference for the higher-income option.

Table 1: Randomization Balance and Sample Means

	Sample means			p-value (4)
	All (1)	Control (2)	Treated (3)	
Panel A: Pre-treatment Characteristics				
Age 18-34	0.430 (0.009)	0.434 (0.013)	0.426 (0.013)	0.683
Age 35-54	0.408 (0.009)	0.401 (0.013)	0.415 (0.013)	0.458
Age 55+	0.162 (0.007)	0.165 (0.010)	0.159 (0.010)	0.659
Female (0/1)	0.488 (0.009)	0.482 (0.013)	0.495 (0.014)	0.497
White (0/1)	0.644 (0.009)	0.655 (0.013)	0.633 (0.013)	0.217
Edu. HS or some college	0.309 (0.009)	0.302 (0.012)	0.315 (0.013)	0.471
Edu. Assoc./BSc degree	0.504 (0.010)	0.505 (0.013)	0.504 (0.014)	0.960
Edu. Postgraduate	0.187 (0.007)	0.193 (0.011)	0.181 (0.010)	0.429
HH income (1,000s)	91.010 (1.126)	91.232 (1.627)	90.779 (1.553)	0.840
Democrat (0/1)	0.441 (0.009)	0.438 (0.013)	0.444 (0.013)	0.745
Republican (0/1)	0.235 (0.008)	0.235 (0.011)	0.236 (0.012)	0.964
Baseline life sat (0-100)	62.126 (0.449)	61.875 (0.637)	62.387 (0.632)	0.569
Employed (0/1)	0.758 (0.008)	0.747 (0.012)	0.770 (0.011)	0.147
Prior average-MSI	5.205 (0.057)	5.178 (0.080)	5.233 (0.081)	0.629
Prior self-MSI	6.803 (0.081)	6.808 (0.114)	6.798 (0.115)	0.950
No. of surveys, past 30 days	137.183 (3.369)	131.713 (4.651)	142.858 (4.880)	0.098
Panel B: Attrition				
Responded to Follow-up	0.794 (0.008)	0.795 (0.011)	0.791 (0.011)	0.795
Observations	2,775	1,413	1,362	

Notes: Main analytic sample: excludes respondents who failed the video-code check and respondents with extreme prior average-MSI. Column (1) is the pooled-sample mean; columns (2) and (3) are the control and treated means; column (4) is the two-sided p-value for the treated-minus-control difference (OLS with robust standard errors). Standard errors are reported in parentheses below each mean. The bottom row reports observation counts in each subsample. A joint test of treatment status on all balance covariates yields $F = 0.62$ ($p=0.847$).

Table 2: The Effect of Self-MSI on Choices: OLS and 2SLS Estimates

	Trade-off Income vs. Hours of...				
	Work (1)	Commute (2)	Sleep (3)	Index (4)	Real Dec. (5)
Panel A: OLS					
Choice-Framing	0.025*** (0.005)	0.020*** (0.005)	0.034*** (0.005)	0.026*** (0.003)	0.027*** (0.006)
Satisfaction-Framing	0.026*** (0.005)	0.020*** (0.004)	0.031*** (0.005)	0.026*** (0.004)	0.032*** (0.006)
Panel B: 2SLS					
Choice-Framing	0.035** (0.015)	0.026* (0.014)	0.024 (0.016)	0.028*** (0.010)	0.024 (0.017)
Satisfaction-Framing	0.033** (0.015)	0.026* (0.014)	0.012 (0.016)	0.024** (0.011)	0.035** (0.017)
Panel C: 2SLS at Follow-Up					
Choice-Framing	0.031 (0.046)	0.003 (0.042)	0.040 (0.046)	0.025 (0.031)	-0.013 (0.073)
Observations					
Baseline	2,721	2,721	2,721	2,721	1,711
Follow-Up	2,161	2,161	2,161	2,161	1,345
Dep. Var. Mean (Control Group)					
Choice-Framing	2.454	1.918	2.812	2.395	2.472
Satisfaction-Framing	2.360	1.878	2.648	2.295	2.310
Dep. Var. SD (Control Group)					
Choice-Framing	1.087	1.015	1.124	0.758	0.996
Satisfaction-Framing	1.068	0.987	1.127	0.763	0.990
Kleibergen-Paap rk Wald F statistic					
Baseline, Choice-Framing	307.0	307.0	307.0	307.0	198.1
Baseline, Satisfaction-Framing	307.0	307.0	307.0	307.0	198.1
Follow-up, Choice-Framing	34.8	34.8	34.8	34.8	22.8

Notes: Each column is an outcome on a 1–4 scale, coded so that 1 indicates definitely choosing the lower-income option and 4 indicates definitely choosing the higher-income option: the job-choice income trade-off against work hours (1), commute (2), sleep (3), their index (4), and the real-world decision (5). Entries are coefficients on posterior self-MSI, with robust standard errors in parentheses. Panel A is OLS; Panel B is 2SLS instrumenting posterior self-MSI with the randomized treatment; Panel C repeats the 2SLS at the one-month follow-up, with both the outcome and self-MSI measured at follow-up. Each panel reports rows for the choice and satisfaction framings (Panel C, choice framing only). Bottom rows give the sample size, the control-group dependent-variable mean and standard deviation, and the Kleibergen–Paap *rk* Wald *F*-statistic for instrument strength. Sample: main analytic sample. * $p < 0.10$, ** $p < 0.05$, *** $p < 0.01$.