

Optimal nudging

Frederick Callaway¹ (fredcallaway@princeton.edu)

Department of Psychology, Princeton University

Mathew D. Hardy¹ (mdhardy@princeton.edu)

Department of Psychology, Princeton University

Thomas L. Griffiths (tomg@princeton.edu)

Departments of Psychology and Computer Science, Princeton University

¹These authors contributed equally

Abstract

People's judgments and decisions often deviate from classical notions of rationality, incurring costs to both themselves and to society. Previous research has proposed that the cost of these biases can be reduced by redesigning decision problems based on theories of human decision making. These modifications—or *nudges*—can have dramatic results and have been successfully applied to variety of domains. However, the formal underpinning of nudge theory is limited, and it is not always clear what the effect of a nudge will be before it is implemented. As a result, designing nudges can be time consuming and error-prone. In this paper, we propose an automatic method for deriving *optimal nudges*. The method is based on a resource-rational model, which assumes that people make decisions in a way that achieves a near-optimal tradeoff between the cost and benefits of deliberation. We then frame nudges as modifications to the costs of different cognitive operations, encouraging the cognitively frugal decision maker to consider some problem features over others. As a proof of concept, we apply the method to the Mouselab process-tracing paradigm, finding that optimal nudges lead participants to make better decisions with less cognitive effort.

Keywords: nudging, decision support, decision making, resource rational analysis

Introduction

How do we choose when to recycle, where to invest our savings, or what to buy at a cafeteria? Investigating the biases that characterize these decisions is a central focus of psychology and behavioral economics. These biases are not only theoretically important in that they violate classical assumptions about human behavior, but are also of practical significance because many small errors can add up to large costs for both individuals and societies (Kahneman et al., 1982).

In an effort to reduce these costs, researchers have proposed using behavioral theory to redesign decision problems in order to help people make better choices and fewer costly errors (Thaler & Sunstein, 2008). These changes—often referred to as nudges—are an increasingly popular alternative to government initiatives such as educational programs, legislation, and incentives, and are often significantly less expensive to administer (Benartzi et al., 2017).

While promising, nudges are controversial. Many are uncomfortable with having their choices influenced by changes they are not aware of or cannot control, and there is often disagreement about how nudges should be evaluated (Goodwin, 2012). Furthermore, while inexpensive to administer, the development of nudges in new domains can involve a costly

search process due to a lack of rigorous theory about how choice architectures interact with people's decision-making processes.

To address these limitations, we propose a method of constructing *optimal nudges* based on a formal theory of resource-bounded decision making (Griffiths et al., 2015; Lieder & Griffiths, 2019). In this framework, a person's decision-making process is modeled as a sequential interaction with their own mental resources. Building on this idea, we formalize nudges as modifications to a decision maker's cognitive environment, for example, making new cognitive operations possible or reducing the cost of existing ones. Having formalized a specific kind of decision problem in this way, we can precisely specify the goal of a nudge with an objective function that can be programatically optimized. That is, we both provide a rigorous and transparent evaluation metric for nudges, and also eliminate the need to manually search for nudges that perform well on that metric.

In this paper, we begin by outlining the computational framework underlying our approach. We then present a concrete method to construct near-optimal nudges within a restricted set of nudges that reduce the cost of specific cognitive operations. As a proof of concept, we apply this method to the Mouselab process-tracing paradigm (Payne et al., 1988), where cognitive operations are externalized as information-gathering clicks. We find that our method both improves the quality of participants' decisions and also reduces the cognitive cost of those decisions. We conclude by discussing the limitations of the current approach and directions for future work.

Background

Nudging

In line with classical theories of rationality, public policy programs were traditionally guided by the assumption that people act optimally with respect to their self interest (Moseley & Stoker, 2013; Jackson, 2005). The goal of many public programs was thus to increase people's freedom of choice and remove government mandates and suggestions. When policy makers wanted to change people's behavior, they would often recommend new incentives, educational programs, or legislation (Benartzi et al., 2017).

In a landmark book, Thaler and Sunstein (2008) challenged

this approach, arguing that research in psychology and behavioral economics showed that people often do not act optimally with respect to their self interest. Instead, they argued, behavioral theories suggest an alternative framework for designing public programs. Specifically, they proposed that governments should use psychological theory to implement subtle changes to the structure of decision environments, *nudging* people towards making better choices without restricting their freedom of choice. These changes would be developed by “choice architects” who would leverage findings on people’s heuristics and biases to design effective nudges (Kahneman et al., 1982; Benartzi & Thaler, 2007).

Nudges have since been successfully applied to domains such as retirement savings, energy consumption, and personal health (Benartzi et al., 2017; Marteau et al., 2011; Newell & Siikamäki, 2014). For example, many companies have programs where workers can sign up to automatically save a proportion of their earnings in a tax-deferred investment account. Despite the obvious benefits of doing so, many undersave and underinvest – a recent study found that 68% of 401(k) participants thought their savings rate was too low (Choi et al., 2004). In response, proponents of nudge theory have suggested changing savings plans to “opt-out” programs, in which employees save a certain proportion of their pay by default but can choose not to (Madrian & Shea, 2001), a change that can lead to significantly higher savings rates at virtually no administrative cost (Chetty et al., 2014).

Despite these promising results, nudging can still be controversial. Many people are uncomfortable with having their decisions influenced by processes beyond their control or that they are unaware of. Even when developed openly, there is often disagreement about what behaviors nudges should aim to influence and optimize (Goodwin, 2012). Furthermore, the application of nudges has been limited by ad-hoc use of psychological theories and informal heuristics (Vlaev et al., 2016), as well as practical difficulties in adapting psychological models to real-world contexts (Moseley & Stoker, 2013).

Resource-rational analysis

Resource-rational analysis is a formal framework for deriving cognitive models based on the assumption that people act optimally with respect to their limited cognitive resources (Griffiths et al., 2015; Lieder & Griffiths, 2019). Within this approach, a cognitive process is understood as the solution to an optimization problem, where the objective function explicitly trades off external utility with internal computational cost. Critically, the theory predicts that people’s behavior will depend on both the structure of the external environment and also the computational actions they can execute and the costs of those actions (their internal computational architecture). As a result, resource-rational models often make behavioral predictions that differ dramatically from classical theories of rationality, and have been shown to account for a wide range of apparent biases and errors in human decision making (e.g. Lieder et al., 2012, 2018; Nobandegani & Shultz, 2020).

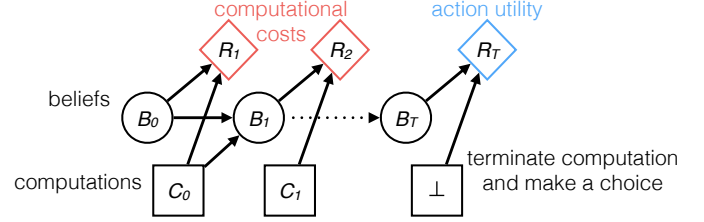


Figure 1: Metalevel Markov decision process.

Metalevel Markov decision processes

A key challenge in resource-rational models is appropriately specifying the relevant computational architecture for a given cognitive process and how these computations ultimately lead to a decision. Recent work has approached this challenge using formal tools developed in a subfield of artificial intelligence known as *rational metareasoning* (Matheson, 1968; Russell & Wefald, 1991), which studies the problem of computational resource allocation. In particular, computation—i.e., reasoning—is framed as a sequential decision problem in which an agent refines its beliefs about the quality of different possible physical actions by executing a series of computational actions. The goal of metareasoning is to select computations in a way that results in good decisions with minimal computation cost.

Concretely, this process is modeled as a metalevel Markov decision process (metalevel MDP; Hay et al., 2012). A metalevel MDP, graphically depicted in Figure 1, is formally identical to a standard Markov decision process (MDP), which is a well-established formalism for representing temporally extended interactions between an agent and its external environment (Puterman, 2014). A standard MDP, $(\mathcal{S}, \mathcal{A}, T, r)$, is defined by a set of states, $\mathcal{S}$, a set of actions, $\mathcal{A}$, a transition function, T , and a reward function, r . The transition function specifies the dynamics of the environment (i.e., how taking actions moves the agent from one state to another) and the reward function specifies the goal, giving a scalar state-dependent reward for each action that the agent takes. The agent chooses actions to maximize cumulative reward using a policy, π , which specifies which action to take based only on the current state.

While a standard MDP describes the interaction between an agent and its *external* environment, a metalevel MDP describes the interaction between an agent and its *internal* computational environment. A metalevel MDP is defined $(\mathcal{B}, \mathcal{C}, T_{\text{meta}}, r_{\text{meta}})$, where the states, $\mathcal{B}$, correspond to the agent’s beliefs and the actions, $\mathcal{C}$, correspond to computations (or cognitive operations). The transition function, T_{meta} , describes how computations update the agent’s beliefs. Finally, the reward function, r_{meta} , describes both the cost of computation and also the utility of the resulting decision. That is, $r_{\text{meta}}(b, c)$ is negative for all computations except $\perp$, for which it gives the expected utility of making a decision based on the final belief state. The metalevel policy, π_{meta} , chooses which computation to perform based on the current belief.

For example, consider the case of one-shot decision making, where a decision maker faces a decision between a set of possible actions, $\mathcal{A}$. Each of these actions, $a \in \mathcal{A}$, has some associated utility that depends on some initially unknown state, θ . The agent would like to maximize $U(a; \theta)$ but must choose a only on the basis of her belief, $b \in \mathcal{B}$, which we assume to be a distribution over θ . Before making a choice, she can execute any number of computational actions, $c \in \mathcal{C}$, which refine her belief. On average, the more computations she executes, the more accurate her belief and the better her decision. However, each of those computations incurs a cost, given by $r_{\text{meta}}(b, c)$. Thus, at some point the agent executes a special operation, $\perp$, which terminates the decision-making process. At this point, an action is (stochastically) selected by the action policy, $\pi_{\text{act}}(a | b_{\perp})$, which is typically assumed to be uniform over all actions with maximal expected utility given the final belief state, $b_{\perp}$. The final metalevel reward is the expected utility of that action: $r_{\text{meta}}(b, \perp) = \sum_{a \in \mathcal{A}} \pi_{\text{act}}(a | b) U(a; \theta)$. That is, the final metalevel reward is the *objective* expected utility of the action(s) with maximal *subjective* expected utility given the final belief state. The total metalevel reward is the difference between this expected utility and the sum of the incurred computational costs.

Constructing optimal nudges

The metalevel MDP formalism provides a computational foundation for understanding, predicting, and controlling the effect of nudges. With this lens, nudging is viewed as a method for modifying a person’s computational architecture, making some sequences of reasoning easier than others. These modifications are formalized as altering the components of the metalevel MDP describing a person’s decision-making process. Formalizing nudging in this way allows us to leverage computational optimization tools to automatically identify nudges that achieve precisely specified goals. Optimal nudging consists of four steps:

1. Model a decision problem as a metalevel MDP, M .
2. Specify a space of possible nudges as a set of possible modified metalevel MDPs, $\tilde{\mathcal{M}}$. Modifications might include adding computational actions or modifying the cost of existing actions.
3. Specify the goal of the nudge with an objective function, $f(\tilde{M}; \theta)$, that indicates how desirable the decision maker’s behavior will be given the modified metalevel MDP, $\tilde{M}$, and the true state of the world, θ .
4. Identify the optimal nudge as the modification that maximizes the objective function:

$$\tilde{M}_{\theta}^* = \operatorname{argmax}_{\tilde{M} \in \tilde{\mathcal{M}}} f(\tilde{M}; \theta)$$

Example: Reducing the cost of outcome evaluation

We now illustrate the general approach in the context of a simple example. Consider a decision problem in which you

must take an action that will have different consequences depending on the outcome of some random process in the environment. You are familiar with this environment and can easily call to mind the relevant probabilities, but you are facing a new set of possible actions and must carefully consider the consequences of each action depending on the random outcome in order to determine how desirable each (action, outcome) pair would be. Formally, your goal is to choose an action with high expected utility,

$$\text{EU}(a) = \sum_o p(o) U(a, o), \quad (1)$$

where the outcome probabilities, $p(o)$, are known but the outcome-dependent action utilities, $U(a, o)$, must be computed if they are to figure into your decision.

This describes the situation faced by participants of an experiment using the Mouselab paradigm (Figure 2; Payne et al., 1988), which has been used extensively in the study of human decision making (and is described in greater detail below). Previous work has shown that Mouselab can be formalized as a metalevel MDP (Gul et al., 2018), and it thus serves as a good initial test case for optimal nudging. In the following sections, we describe how we apply the proposed four-step process to create optimal nudges for Mouselab.

1. Metalevel MDP The metalevel MDP for Mouselab is defined by $(\mathcal{B}, \mathcal{C}, T_{\text{meta}}, r_{\text{meta}})$. The unknown state, θ , corresponds to the utilities, $U(a, o)$, for each combination of action and outcome. A belief, $b \in \mathcal{B}$, is thus a distribution over those utilities. The belief is initialized to a multivariate Gaussian with known mean and isotropic covariance corresponding to an i.i.d. prior over each utility. A computation, $c \in \mathcal{C}$, reveals one of these unknown utilities. As such, the transition function, T_{meta} , specifies that beliefs are updated by fixing the mean of the corresponding element of the belief to the true value and setting its variance to 10^{-10} . The reward function, r_{meta} , specifies the cost of revealing each utility (operationalized in our experiment as a number of clicks) as well as the decision utility. To compute the latter, we assume, as usual, that π_{act} selects the action with maximal expected utility given the final belief (i.e. maximizing Equation 1 with unknown utilities replaced by the mean of the corresponding element of the belief). $r_{\text{meta}}(b, \perp)$ is then the true expected utility of this action (or distribution over actions) given the true utilities.

2. Space of nudges Although endowing decision makers with novel computational actions is an exciting direction, we restrict ourselves here to a narrower class of nudges that simply reduce the cost of the existing computations.¹ Concretely, $\tilde{M}$ must be identical to the original metalevel MDP, except that the reward function is modified to be $\tilde{r}_{\text{meta}}(b, c) = r_{\text{meta}}(b, c) + \lambda_c$, where λ_c is the reduction in cost for computation c . The modification is subject to three constraints: $\lambda_c \leq -r_{\text{meta}}(b, c)$ (computational costs cannot become negative), $\lambda_c \geq 0$ (costs can only be reduced not increased), and $\sum_c \lambda_c \leq Z$ (there is a budget on total cost reduction). This

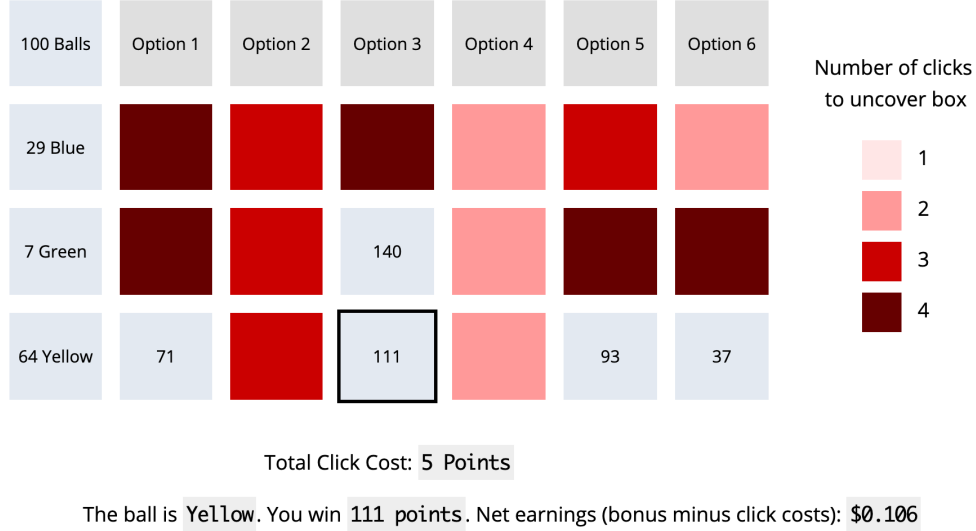


Figure 2: Experimental interface. On every trial, participants made a choice between six options. After choosing an option, a single ball color (Blue, Green, or Yellow) was selected with percentage probability equal to the number of balls of that color. The option then paid out with the value indicated by the corresponding cell. The values in some cells were initially shown, but others were hidden at trial onset. Participants could click on these cells to reveal the values, paying one point for each click. The number of clicks required to reveal each cell was indicated by its color.

final restriction prevents the trivial optimization $\tilde{r}_{\text{meta}}(b, c) = 0 \forall c \in \mathcal{C}$. The budget, Z , represents the degree to which the decision maker’s computational architecture can be modified.

3. Objective function The choice of a suitable objective function depends on a nudge’s goal. For example, many nudges aim to maximize the probability that people take a certain action, e.g. recycling or registering to become an organ donor. This kind of goal can be formalized as maximizing the probability of the decision maker choosing a specific action,

$$f_{\text{action}}(\tilde{M}; \theta, a^*) = \mathbb{E}_{b_{\perp}} [\pi_{\text{act}}(a^* | b_{\perp}) | \tilde{M}, \theta],$$

where the expectation is taken with respect to the belief state of the decision maker when she makes a choice, $b_{\perp}$. Although the distribution over $b_{\perp}$ depends on $\tilde{M}$ in complex ways, it will generally be the case that reducing the cost of a computation that measures some dimension of θ will increase the probability that the decision maker acquires an accurate belief about that dimension. Thus, this objective function will select modifications that make it inexpensive to measure dimensions of θ that make the desired action, a^* , appear desirable (or make competing actions seem undesirable).

Other nudges do not aim to make people choose a specific option, but rather to improve the overall quality of their decisions, encouraging them, for example, to make healthier eating choices or choose more diversified investment portfolios.

We can model this kind of goal as maximizing the expected utility of the decision maker’s choice,

$$f_{\text{utility}}(\tilde{M}; \theta) = \mathbb{E}_{b_{\perp}} \left[\sum_a \pi_{\text{act}}(a | b_{\perp}) \text{EU}(a; \theta) | \tilde{M}, \theta \right].$$

Finally, we might want to not only encourage people to make *better* decisions, but also to make it *easier* to make those decisions. We can formalize this goal as maximizing the cumulative metalevel reward which captures both decision quality and computational cost,

$$f_{\text{meta}}(\tilde{M}; \theta) = \mathbb{E}_{\mathcal{C}} \left[\sum_t^T r(B_t, C_t) | \tilde{M}, \theta \right]. \quad (2)$$

Here, the expectation is taken over all possible sequences of computations the decision maker could execute. This quantity is also called the *metalevel return* and it is the quantity that the optimal metalevel policy maximizes. We chose this as our objective function for the present study; however, any of the above objectives can be trivially implemented.

All the above objectives implicitly depend on assumptions about the metalevel policy, π_{meta} . One principled choice is to assume that the decision maker is metalevel optimal. However, computing the optimal metalevel policy is computationally intensive. We thus instead assume that the agent follows the *meta-greedy* policy (Russell & Wefald, 1991), which chooses each computation as if it were committed to making a decision on the following time step. This policy behaves similarly to the optimal policy, and explains human behavior in Mouselab nearly as well as the optimal policy (Gul et al.,

¹This excludes interventions that introduce novel deliberative strategies, such as “boosts” that make previously inaccessible information available (Hertwig & Grüne-Yanoff, 2017). Applying our framework to such cases is an important direction for future work

2018). More to the point, it is easy to compute; runtime is linear in the number of possible computations and independent of the dimensionality of the belief space, under the assumption of “independent actions” (Hay et al., 2012). This allows us to construct optimal nudges for complex decision problems where identifying an optimal solution is intractable.

4. Optimization Selecting a nudge that optimizes Equation 2 can be expressed as a k -dimensional maximization problem where $k = |C| - 1$ is the number of computations, excluding $\perp$. To reduce the complexity of this problem, we applied the additional constraint that the budget be divided equally and into integer amounts among one, two, three, or six options. We then selected a modification within this space using a simple greedy search procedure, first placing the full budget on one computation, then considering splitting it to reduce the cost of second computation, and so on. We found this approach to slightly outperform more sophisticated genetic programming techniques, suggesting that it is an effective optimization strategy.

Experiment: Testing optimal nudging

We evaluated the proposed optimal nudging method in a modified version of the Mouselab paradigm (Payne et al., 1988). In this setup, participants make choices between different options with known and unknown payoff values. The paradigm externalizes computations as information-gathering operations (clicks) that reveal these values, computational cost as an explicit monetary cost for clicking, and belief states as configurations of revealed and hidden values. By using such a paradigm, we can more easily make assumptions about the metalevel MDP underlying the participant’s decisions, thus allowing us to test the cost-modification approach directly.

Methods

An example of the experimental interface is given in Figure 2. Participants chose between six options (columns), each with three possible payoff values (rows). After making a choice, a ball was drawn from a simulated lottery machine with 100 balls, and the chosen option paid out depending on the color of the drawn ball. The percentage probability that a certain color ball was drawn was simply the number of balls indicated in the far left column. Different options paid different values depending on which color ball was drawn, and some of these values were revealed at trial onset, while others were hidden. To reveal a hidden value, participants had to click on the value they wished to reveal between one and four times (see Figure 2), paying one point for each click. The number of clicks necessary to reveal each cell was sampled uniformly from $\{0, 1, 2, 3, 4\}$ to mask the cost-reductions (described below). Cell values were sampled from a normal distribution with a mean of 75 points and a standard deviation 36 points (truncated at 0 points and discretized to integers).

On each trial, the cost structure was modified according to either the proposed optimal nudging method, or a random baseline. In both cases, the cost-modification budget was set

to $Z = 6$ points, which is a fairly modest budget in comparison to the average cost of 36 points for revealing every cell. Optimal costs were chosen to maximize the metalevel return of the meta-greedy policy (Equation 2), using the greedy search method described above. The random cost modification was determined by randomly sampling three costs and reducing each by 2 points.

We recruited 150 participants from Amazon’s Mechanical Turk, limiting our study to those living in the United States. Participants who failed an attention check were excluded from the experiment. Participants first completed a practice trial, and then 20 test trials in which 10 problems had random modifications and 10 had optimal modifications.² Each participant completed the same set of 21 problems, but problem order and each problem’s modification type varied randomly between participants. At the end of the game, participants’ total points were paid to them as a bonus with 10 points equal to 1 cent. Participants earned \$0.25 for participating in the study plus an average bonus of \$1.71.

Results

On average, participants earned 81.55 points on trials with random modifications and 89.66 points on trials with optimal modifications (see Figure 3). To test whether this difference was significant, we ran a crossed mixed-effects regression predicting total points earned on each trial with a fixed effect for the cost-modification condition (optimal vs. random) and random effects for both participant and problem. A likelihood ratio test of the mixed effects model with and without the condition fixed effect was significant ($\chi^2(1) = 37.711, p < 0.001$). Similar models predicting the click cost and decision quality also revealed significant effect, (click cost: 3.99 vs. 3.48, $\chi^2(1) = 8.6866, p = 0.003$; choice payout: 85.54 vs. 93.14, $\chi^2(1) = 33.821, p < 0.001$).

Discussion

In this paper, we have proposed a formal framework for developing, comparing, and evaluating nudges. The framework is based on theoretical work characterizing human decision making as making optimal use of limited computational resources (Griffiths et al., 2015; Lieder & Griffiths, 2019). Viewing error in decision making as the consequence of limited resources suggests that we can improve peoples’ decisions by alleviating those limitations. In particular, by modeling decision making as a metalevel Markov decision process (Hay et al., 2012), we formalized nudging as giving people access to more powerful or less costly computational actions, which they can deploy to make better decisions with less cognitive effort. Formalizing nudging in this way allows us to apply tools from artificial intelligence to design *optimal nudges*, i.e. nudges that most improve people’s decisions (subject to constraints). As a proof of concept, we applied the framework to the Mouselab process-tracing paradigm, finding that

²Due to a programming error, the first test problem was always the same as the practice problem. We thus exclude data from this problem from our analysis, leaving 19 trials per participant.

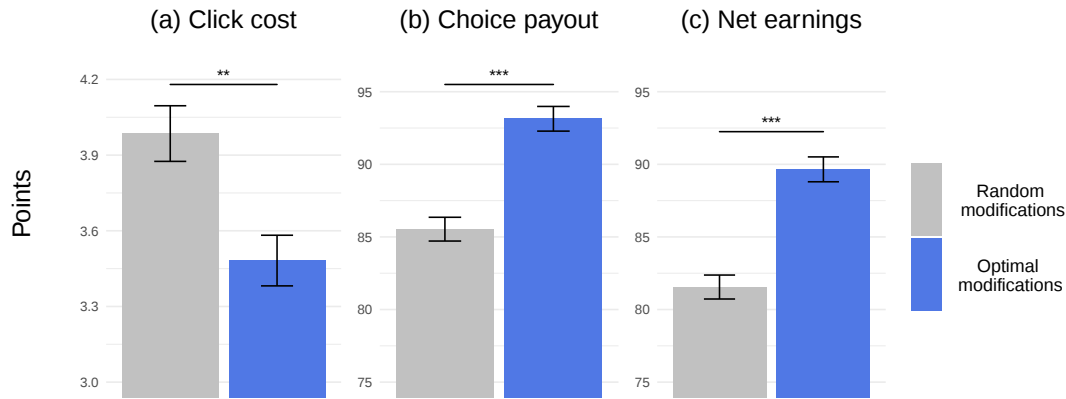


Figure 3: Average points per game spent and earned under random and optimal cost modifications. Plot (a) shows the average points spent uncovering values, (b) the average reward from participants' choices, and (c) their average net reward (choice payout minus click cost). Stars indicate significant differences in means with p-values estimated from a likelihood ratio test between an unrestricted crossed mixed-effects regression on the variable of interest with modification type as a fixed effect and a restricted regression with no fixed effect. Errors bars show standard error estimates derived from the residuals of the unrestricted crossed mixed-effects regressions.

optimal nudges both increased the quality of participants decisions and also reduced the effortfulness of making those decisions. This provides preliminary evidence that reducing computational costs can be an effective way to help people make decisions more effectively.

Our approach has a number of advantages over other approaches to nudging. First, it provides a theoretical foundation for understanding and predicting nudges' effects. Second, we explicitly specify the goal of a nudge using an objective function. This increases the transparency of nudges, provides a natural way to think about novel goals for nudges (e.g., making people's decisions easier without systematically changing their choices), and allows an end user to have control over how they are nudged. Third, given a model of the decision-making process and an objective, our method automatically discovers an optimal nudge using computational optimization techniques. This reduces the human labor involved in designing nudges, and can potentially identify better nudges than a person would be likely to discover.

Despite the advantages, optimal nudging presents several challenges. First, it requires a detailed model of the computational process underlying the decision we would like to intervene on. In the present work, we avoided this challenge by using a process tracing paradigm that externalizes these typically unobservable processes. Applying the method in the real world, however, requires one to infer this model from behavior. Nevertheless, even a heavily simplified decision-making model may be adequate to construct helpful, if not truly optimal, nudges. Second, the method makes strong assumptions about the decision maker's cognitive process, i.e. that it is near-optimal given the metalevel MDP. While this assumption may not be borne out in practice, it is not critical to the basic framework and could easily be modified. Third, we do not account for the potential communicative content

of nudges. If the decision maker can infer the nudger's goal, social reasoning may affect her decision.

Nudging is an increasingly popular method for improving people's choices and reducing the costs of their errors. However, nudges are both controversial and often difficult to implement and evaluate. In this paper we proposed optimal nudging as a way to leverage resource-rational modeling to address these limitations and ethical concerns. Priorities for future work include extending optimal nudging to more naturalistic tasks where deliberative processes are unobserved and applying the framework to improve existing nudges.

Acknowledgements

This research was supported by a grant from Facebook Reality Labs.

References

- Benartzi, S., Beshears, J., Milkman, K. L., Sunstein, C. R., Thaler, R. H., Shankar, M., ... Galing, S. (2017). Should governments invest more in nudging? *Psychological Science*, 28(8), 1041-1055.
- Benartzi, S., & Thaler, R. (2007). Heuristics and biases in retirement savings behavior. *Journal of Economic perspectives*, 21(3), 81-104.
- Chetty, R., Friedman, J. N., Leth-Petersen, S., Nielsen, T. H., & Olsen, T. (2014). Active vs. passive decisions and crowd-out in retirement savings accounts: Evidence from Denmark. *The Quarterly Journal of Economics*, 129(3), 1141-1219.
- Choi, J. J., Laibson, D., Madrian, B. C., & Metrick, A. (2004). For better or for worse: Default effects and 401(k) savings behavior. In *Perspectives on the economics of aging* (pp. 81-126). University of Chicago Press.

- Goodwin, T. (2012). Why we should reject Nudge. *Politics*, 32(2), 85–92.
- Griffiths, T. L., Lieder, F., & Goodman, N. D. (2015). Rational use of cognitive resources: Levels of analysis between the computational and the algorithmic. *Topics in cognitive science*, 7(2), 217–229.
- Gul, S., Krueger, P. M., Callaway, F., Griffiths, T. L., & Lieder, F. (2018). Discovering rational heuristics for risky choice. In *The 14th biannual conference of the german society for cognitive science*.
- Hay, N., Russell, S., Tolpin, D., & Shimony, S. E. (2012). Selecting computations: Theory and applications. In *Proceedings of the 28th conference on uncertainty in artificial intelligence*.
- Hertwig, R., & Grüne-Yanoff, T. (2017). Nudging and boosting: Steering or empowering good decisions. *Perspectives on Psychological Science*, 12(6), 973–986.
- Jackson, T. (2005). Motivating sustainable consumption. *Sustainable Development Research Network*, 29(1), 30–40.
- Kahneman, D., Slovic, P., & Tversky, A. e. (1982). *Judgment under uncertainty: heuristics and biases*. Cambridge, England: Cambridge University Press.
- Lieder, F., Griffiths, T., & Goodman, N. (2012). Burn-in, bias, and the rationality of anchoring. In *Advances in neural information processing systems* (pp. 2690–2798).
- Lieder, F., & Griffiths, T. L. (2019). Resource-rational analysis: understanding human cognition as the optimal use of limited computational resources. *Behavioral and Brain Sciences*, 1–85.
- Lieder, F., Griffiths, T. L., & Hsu, M. (2018). Overrepresentation of extreme events in decision making reflects rational use of cognitive resources. *Psychological review*, 125(1), 1.
- Madrian, B. C., & Shea, D. F. (2001). The power of suggestion: Inertia in 401(k) participation and savings behavior. *The Quarterly journal of economics*, 116(4), 1149–1187.
- Marteau, T. M., Ogilvie, D., Roland, M., Suhrcke, M., & Kelly, M. P. (2011). Judging nudging: can nudging improve population health? *Bmj*, 342, d228.
- Matheson, J. E. (1968). The Economic Value of Analysis and Computation. *IEEE Transactions on Systems Science and Cybernetics*, 4(3), 325–332. doi: 10.1109/TSSC.1968.300126
- Moseley, A., & Stoker, G. (2013). Nudging citizens? Prospects and pitfalls confronting a new heuristic. *Resources, Conservation and Recycling*, 79, 4–10.
- Newell, R. G., & Siikamäki, J. (2014). Nudging energy efficiency behavior: The role of information labels. *Journal of the Association of Environmental and Resource Economists*, 1(4), 555–598.
- Nobandegani, A. S., & Shultz, T. R. (2020). A resource-rational, process-level account of the St. Petersburg paradox. *Topics in Cognitive Science*.
- Payne, J. W., Bettman, J. R., & Johnson, E. J. (1988). Adaptive strategy selection in decision making. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 14(3), 534–552.
- Puterman, M. L. (2014). *Markov decision processes: discrete stochastic dynamic programming*. John Wiley & Sons.
- Russell, S., & Wefald, E. (1991). Principles of metareasoning. *Artificial Intelligence*, 49(1-3), 361–395.
- Thaler, R. H., & Sunstein, C. R. (2008). *Nudge: Improving decisions about health, wealth, and happiness*. New Haven: Yale University Press.
- Vlaev, I., King, D., Dolan, P., & Darzi, A. (2016). The theory and practice of “nudging”: changing health behaviors. *Public Administration Review*, 76(4), 550–561.