

Preferences as Heuristics*

Erik Mohlin[†]

Alexandros Rigos[‡]

This version: November 19, 2025

First version: November 19, 2024

(For the latest version please click [here](#))

Abstract

We present a model of how an individual's preference for a specific behaviour can evolve as an adaptation to her social environment. Our key assumptions are that the decision maker (DM) (i) is boundedly rational in her ability to process information and evaluate which of two actions yields the highest payoff and (ii) can develop a subjective preference for one of the actions. A stronger preference for an action reduces the probability of not taking it when it is optimal (omission error), but it also increases the probability of taking the action when it is suboptimal (commission error). The optimally adapted preference strikes a balance between these errors and depends on the DM's environment. DMs are typically better off (in objective terms) if they are endowed with subjective preferences that deviate from maximisation of expected objective payoff. Importantly, if DMs can evaluate actions perfectly, they do not develop such biases. The results extend in an intuitive manner to n -player, two-strategy supermodular games. Our framework can be used to interpret preferences for specific actions or strategies (e.g., conditional cooperation, truth-telling, or norm-following) as psychological and cultural consequences of material incentives and social organisation.

Keywords: Preference evolution; Bounded rationality; Culture; Cooperation; Logit choice; Monotone comparative statics; Best response dynamics; Supermodular games

JEL classification: C72, C73, D01, D83, D91

*This paper has benefited from comments by audiences at the universities of Aarhus, Lund, Maastricht, and Nottingham, as well as the Stockholm School of Economics, the Nordic Economic Theory Group online seminar, ESA Europe Helsinki 2024, the 2024 Nordic Economic Theory Conference, the 2023 Nordic Behavioural Economics Conference, LGTC2024, EWBEE 2025, DETC2025 the 2024 Swedish Conference in Economics, LEG2022, and the 2021 Congress of the Game Theory Society. Mohlin gratefully acknowledges financial support from the Swedish Research Council (2015-01751 and 2019-02612), Handelsbankens forskningsstiftelser (P19-0204), and the Knut and Alice Wallenberg Foundation (Wallenberg Academy Fellowship 2016-0156). Rigos gratefully acknowledges financial support from Handelsbankens forskningsstiftelser (project P21-0244 and Tommy Andersson's project P22-0087).

[†]Swedish Defence University, Lund University, and Institute for Futures Studies. Email: erik.mohlin@nek.lu.se.

[‡]University of Copenhagen and Lund University. Email: alexandros.rigos@gmail.com.

1 Introduction

Substantial cross-cultural differences have been documented for a number of preferences and behavioural tendencies. These differences in preferences or behaviours can often be attributed to differences in material and social circumstances across societies. For example, Gächter and Schulz (2016) find that preferences against lying are weaker in countries with high prevalence of rule violations (in terms of political fraud, tax evasion, and corruption). Henrich et al. (2004) find that ultimatum game offers are positively correlated with the extent to which economic life depends on cooperation with non-immediate kin and market exposure. Bohnet, Herrmann, and Zeckhauser (2010) document differences in betrayal aversion between, on the one hand countries on the Arabian gulf and, on the other hand Switzerland and the USA, and relate this to different costs and risk of being betrayed under kin-based vs formal-legal enforcement systems.¹

In this paper we present a modelling framework that explains people’s preferences over actions and strategies as adaptations to the material incentives that prevail in their physical and social environment. Cultural differences in preferences, and the resulting behaviours, are thereby explained in terms of differences in incentive structures across societies.² In very general terms, our model implies that people will develop intrinsic preferences for actions or strategies that tend to be optimal in their environment.³

Our theory distinguishes between the subjective decision utility that underlies the individual’s choice behaviour, and the objective payoff that we use to evaluate the quality of her choices. We interpret subjective utility as a cardinal measure of value implemented in the brain, not merely a representation of behaviour. Objective payoff corresponds to whatever drives adaptation or evolution of preferences: if a preference trait confers a higher objective payoff the trait is more likely to be acquired and spread through the population. More on the mechanism involved below.

Our proposed mechanism builds on the observation that the cognitive processes underlying belief formation, deliberation, and choice are imperfect and noisy, implying that mistakes are unavoidable (Alós-Ferrer 2018; Gold and Shadlen 2007; Krajbich, Armel, and Rangel 2010; Woodford 2020). We say that an objective mistake occurs if an individual

¹Cross-country variation has also been documented for, e.g., dictator game donations, ultimatum game rejections, and anti-social punishment (Falk et al. 2018; Henrich et al. 2004; Herrmann, Thöni, and Gächter 2008)

²Naturally, we do not claim that all cross-cultural differences that have been documented by experimental economists are due to differences in preferences. Some differences are demonstrably generated by differences in beliefs (T. O. Weber et al. 2023) rather than preferences, or some combination thereof (e.g. Bigoni et al. 2019).

³This is in sharp contrast with theories of crowding-out which imply a negative relationship between intrinsic preferences and external rewards, at least for low levels of such rewards.

takes an action that does not maximise expected objective payoff, an a subjective mistake occurs if the action does not maximise expected subjective utility.

While eliminating all mistakes is impossible, adjusting subjective preferences can change the likelihood of different objective mistakes. Strengthening a preference for an action decreases the chance of wrongly avoiding it but raises the risk of wrongly choosing it. The optimal preference balances these mistakes. We suggest that subjective preferences for actions and strategies may be at least partially adapted to implement such an optimal solution. In this sense, preferences function as heuristics that guide behaviour away from actions associated with costly mistakes in the direction of actions associated with less costly mistakes.

We consider three interpretations of why subjective utility adapts to raise objective payoff. Biologically, evolution selects on subjective preferences that enhance fitness. Culturally, subjective preferences reflect internalized norms transmitted through social learning, where markers of success guide imitation. Psychologically, a dual-self mechanism allows a long-run planner to adjust the short-run doer's preferences to improve long-term outcomes. While biological adaptation may be relevant in some environments, our analysis emphasizes cultural and psychological mechanisms, which permit more rapid adjustment. We discuss these interpretations in more detail in Section 3.3.

In our model, we first consider a non-strategic setting where a decision maker (DM) faces a choice between two actions denoted 0 and 1. For simplicity one may think of action 0 as representing inaction and 1 as an active choice. The DM receives an *objective payoff* (π) which depends on her choice and the state of the world, the latter being a random variable drawn from a distribution, which we call the DM's *environment*. The DM tries to maximise her *utility* (u), but makes occasional mistakes, whose frequency and direction depend on the utility difference between acting (action 1) and not acting (action 0) at the realised state. The DM's behaviour is summarised by her *stochastic choice rule*, which maps subjective utility differences into choice probabilities.

The higher the utility difference is in favour of acting (action 1), the higher the probability that the DM acts. For example, utility differences may be transformed into choice probabilities via a binary logit choice rule. Ideally the stochastic choice rule should be derived from an explicit model of information-processing and decision making.⁴ For tractability reasons we refrain from doing this in our main analysis, but see Section 7.1 and Appendix G for a discussion of the relation to imperfect Bayesian decision making.

The DM can develop a *taste* for or against acting (action 1), which can bias her preferences. While the DM's utility from inaction (action 0) is equal to the objective payoff

⁴Such a model could incorporate a number of sources of mistakes, e.g. incorrect (but ex-ante unbiased) priors, base-rate neglect (Kahneman and Tversky 1973), signal neglect (Phillips and Edwards 1966), or noise in computation of expected values (Benjamin 2019; Enke and Graeber 2023; Gabaix 2014).

obtained from it, the utility she experiences from acting (action 1) is equal to the objective payoff plus the (positive or negative) taste for acting. As mentioned above, We let the DM's taste adjust due to (social or biological) evolutionary pressure or due to psychological adaptation. We make the idealising assumption that the DM develops a taste that maximises her expected objective payoff in her environment. This optimal taste is our object of interest.

We show that if the DM makes mistakes with positive probability, she typically develops an optimal taste that will bias her towards or against action (relative to inaction). In contrast, if the DM makes no mistakes, having unbiased preferences is always optimal. Crucially, the optimal taste varies with the DM's environment in predictable ways. If the distribution of objective payoff differences between acting and not acting increases in a likelihood-ratio-dominance sense, then the optimal taste becomes more favourable towards acting. The result extends to any first-order-stochastic-dominance increase of the environment if the DM's stochastic choice rule is *moderate*, meaning that it satisfies a condition that implies that the DM's probability of acting does not change too fast with her perceived utility difference.

We then turn our focus on two-action strategic environments, where several players interact. In this framework, each state is a supermodular game. States are ordered in such a way that for each player—holding others' behaviour constant—the payoff advantage of action (action 1) relative to inaction (action 0) is increasing in the state. As in the basic model, each player can develop a taste for one of her actions, and her probability of acting is governed by a (strictly increasing) stochastic choice rule. We find that, if all players' stochastic choice rules are moderate, their extremal equilibrium tastes are more in favour of acting as the environment increases in a first-order-stochastic-dominance sense.

We complement our static equilibrium analysis with a dynamic analysis. For tractability we assume that behaviour adjusts rapidly via a best response dynamic for fixed tastes, and that tastes also adjust slowly also via a best response dynamic. We show that the comparative statics described above are consistent with our explicit dynamic adjustment process (see Section 5).

Throughout our analysis we assume that evolution and adaptation only takes place in the domain of subjective utility and not in the domain of cognition and choice precision. In Section 7.2 we relax this assumption and show that the rationale for preference adaptation remains even in a model where both preferences and cognition can be adjusted at a (convex) cost.

Our theory formalises and elucidates a mechanism through which cross-cultural differences in preferences and values can be traced to cross-cultural differences in everyday practices and social organisation.⁵ We proceed to apply our theory to a selection of scenar-

⁵We will not make any attempt at summarising the vast literature on culture across the humanities and

ios with well-documented cross-cultural differences. We briefly mention some of them here (see Section 6 for complete treatment).

Consider the case of preferences against lying. Whether or not a person is willing to lie depends in part on the expected consequences of lying (Fischbacher and Föllmi-Heusi 2013; Gneezy 2005). In addition, there is evidence that people suffer a utility cost from the mere act of lying (Abeler, Nosenzo, and Raymond 2019). Our theory explains how the strength of the material consequences of lying may influence the development of subjective preferences against the act of lying. Specifically, our model implies that people will have a stronger preference against lying in environments where being caught lying is more likely. Moreover, we note that a higher prevalence of rule violations is plausibly associated with lower probability that violations are detected. Our model then provides a possible explanation for why higher prevalence of rule violations is associated with weaker preferences against lying, as evidenced by Gächter and Schulz (2016)

More generally, our model suggests that the strength of preferences for complying with social norms in a society may be related to the probability that norm violations are detected and the severity of the accompanying sanctions by society (Bicchieri 2005; Ellingsen and Mohlin 2024; Henrich et al. 2004; López-Pérez 2008).

When a social dilemma is indefinitely repeated with a sufficiently high repetition probability, entirely selfish players can be incentivised to cooperate. Cooperation can be similarly incentivised by community enforcement based on reputations. Our model suggests that preferences for conditionally cooperative strategies (Fischbacher, Gächter, and Fehr 2001; Kocher et al. 2008; Rustagi, Engel, and Kosfeld 2010) can evolve in such environments, and should be stronger in environments where there is a higher probability of repeated interactions with the same partner, or more reliable reputation transmission between partners. Moreover, our model predicts that, for a given environment, preferences for cooperation will be stronger in a society that has coordinated on a more cooperative equilibrium. This is in line with the experimental findings of Peysakhovich and Rand (2016), under the psychological interpretation of objective payoff. It is also consistent with the evidence of Rustagi (2023), that exposure to interactions in a market with community enforcement promotes conditionally cooperative behaviour outside the market.

Our model can also speak to the effect of kinship structures on preferences (Enke 2019; Greif and Tabellini 2017; Schulz 2022). We focus on one example, betrayal aversion, as

social sciences. We note that values and value systems are an important component in classical accounts of culture (Parsons 1951) and in recent economic models of culture (Alesina and Giuliano 2015; Guiso, Sapienza, and Zingales 2006). Other accounts (Acemoglu and Robinson 2021; Geertz 1973) stress the cognitive side of culture in the form of transmission and development of concepts and symbols. Still, it seems plausible that these in turn affect the formation of values and attitudes, corresponding to the subjective preferences in our model.

documented by Bohnet, Herrmann, and Zeckhauser (2010) and show how our model can explain why there would be a stronger preference against trusting strangers in a kin-based society.

In general, on our account, prosocial behaviours elicited in simple experimental games may reflect genuine prosocial preferences. Still, these preferences are not adapted to anonymous one-shot interactions but may be viewed as spillovers from adaptations to environments where the presence of non-anonymous and open-ended interactions allow prosocial preferences to evolve.

2 Relation to the literature

Our work is related to several strands of literature. The idea that cognitive processes are inherently noisy goes back to the work of Fechner in the nineteenth century (Fechner [1860] 1948; Woodford 2020). A successful modern incarnation of the idea is the drift-diffusion model (Gold and Shadlen 2007; Ratcliff 1978; Ratcliff et al. 2016) which has also recently entered economics (Alós-Ferrer, Fehr, and Netzer 2021; Fudenberg, Strack, and Strzalecki 2018). For reasons of tractability we do not provide an explicit model of the noise-generating process. We merely assume that choice is stochastic and that mistakes are decreasing in (perceived) utility differences, but we conjecture that one can rephrase our model in a drift-diffusion framework.

The mechanism we describe can generate non-materialistic preferences even in non-strategic individual decision problems.⁶ In this, we follow the literature that analyses preferences and other traits as constrained optimal solutions to evolutionary maximisation problems (Netzer 2009; Rayo and Becker 2007; Robson 2001; Robson and Samuelson 2011; Steiner and Stewart 2016). The closest precursor is Samuelson and Swinkels (2006) who point out that evolution cannot equip us with correct priors and beliefs for all situations and hence we are unable to avoid mistakes. Like us, they show that adjusting subjective preferences can help alleviate the cost of such mistakes. Their analysis is confined to non-strategic situations, whereas we also apply our framework to games. Moreover, in contrast to their work, we focus on conducting comparative statics, which allows us to discuss the directed effects that environments have on preference evolution. The presence of cognitive noise has also been used to explain the presence of biases in probabilistic reasoning (Enke and Graeber 2023; Steiner and Stewart 2016).

The idea that humans rely on adaptive heuristics is widespread in psychology and related disciplines, variously conceptualised as ecological rationality (Gigerenzer and Todd

⁶Other models of endogenous preferences that do not depend on the strategic context include models of habits as a function of past consumption (Becker 1996; von Weizsäcker 2014).

1999), rational analysis (Anderson 1991), or resource rational analysis (Lieder and Griffiths 2020). Within economics similar phenomena tend to fall under the broad label of bounded rationality (Simon 1957).⁷ The rational inattention paradigm provides particularly successful perspective on optimal utilisation of limited cognitive resources (Matějka and McKay 2015; Sims 2003). Our main difference compared to this literature is that our locus of adaptations is preferences, rather than different aspects of belief formation. We provide a detailed comparison with models of rational inattention in Section 8.3 and Appendix E.

There is a large literature on the evolutionary foundations of preferences in strategic settings, often referred to as the indirect evolutionary approach (Frank 1987; Güth and Yaari 1992; Sandholm 2001). In these models, non-materialistic preferences typically evolve because preferences are observable (Dekel, Ely, and Yilankaya 2007; Heifetz, Shannon, and Spiegel 2007; Heller and Mohlin 2019) or because there is assortative matching (Alger and Weibull 2013; Bergstrom 1995). In our model, preferences are unobservable and non-materialistic preferences can evolve even when matching is not assortative.⁸

There are also models of vertical social transmission (Bisin and Verdier 2001; Tabellini 2008; see also Bowles 1998) where parents choose to socialise their children in a way that puts at least some weight on the children sharing values with their parents. This means that if the parents have preferences that deviate from material self-interest, then they will to some extent install similar preferences in their children. For example, in Tabellini (2008) parents socialise children to a preference for cooperation. However, if parents only cared about the material success of their children, then they would not provide the children with preferences for cooperation. By contrast, in our model, selection is entirely payoff-based, and preferences that deviate from material self-interest may evolve even if everyone is initially entirely selfish and materialistic.

Beyond economics, there is a rich literature on cultural evolution (Cavalli-Sforza and Feldman 1981; Boyd and Richerson 1988).⁹ Typically in these models, selection occurs at the level strategies rather than preferences. We are not aware of any work in this area that explores cognitive noise as a source of prosocial preferences.

⁷See for example Samuelson (2001) on automata and Dow (1991), Jehiel and Mohlin (2023), and Rubinstein (1998) on coarse reasoning.

⁸Another mechanism works via the importance of relative payoffs in (small) finite populations, which may give rise to spiteful preferences (Huck and Oechssler 1999; Schaffer 1988). We work with a finite population, but our results do not depend on this and are intact if we let the population grow to infinity.

⁹For overviews see Boyd and Richerson (2009) and Chudek, Muthukrishna, and Henrich (2015).

3 Individual choice model

3.1 Model layout

3.1.1 Subjective utility and objective payoff

A decision maker (DM) chooses an action a from the binary set $\{0, 1\}$ to maximise her *subjective utility*, which depends on the action taken and a state of the world $\omega \in \Omega$ as described by the utility function $u : \{0, 1\} \times \Omega \rightarrow \mathbb{R}$. We find it helpful to think of option 1 as representing taking a specific action and option 0 representing abstaining from this action. While u expresses the DM's *subjective* preferences, the *objective payoff* that the DM receives from a particular action-state combination is expressed by the function $\pi : \{0, 1\} \times \Omega \rightarrow \mathbb{R}$. The objective payoff drives the adaptation or evolution of taste in the model (see Section 3.1.5). Importantly, subjective utility u need not be identical to the objective payoff π . The subjective-utility difference between actions 1 and 0 at state ω is denoted by

$$y(\omega) := u(1, \omega) - u(0, \omega),$$

and the objective-payoff difference is denoted by

$$x(\omega) := \pi(1, \omega) - \pi(0, \omega).$$

We assume that u and π are bounded, therefore y and x are also bounded. The smallest closed interval that contains $x(\Omega)$ (the set of values that objective-payoff differences take) is denoted by $X \subseteq \mathbb{R}$.

3.1.2 Stochastic decision making

We assume that the DM is boundedly rational, leading to her making occasional mistakes. Her cognitive ability and decision-making ability is captured by a *stochastic choice rule* (SCR), which is an exogenous function $q : \mathbb{R} \rightarrow [0, 1]$ that maps subjective utility differences to choice probabilities. In particular, $q(y)$ gives the probability that the DM takes action $a = 1$ when her utility difference between action 1 and action 0 is y . We assume that q is strictly increasing and continuous, so that costlier mistakes are less frequent. We denote the set of such SCRs by $\mathcal{Q}$. An example of an SCR is the logit rule with precision parameter $\beta > 0$, so that $q(y) = (1 + \exp(-\beta y))^{-1}$.

3.1.3 Tastes

Our analysis takes the cognitive capacity of the DM, as expressed by her SCR, as given. However, we allow the DM to develop a *taste* $\kappa \in K \subseteq \mathbb{R}$ that biases her towards one of

the two actions. In particular, at any action-state combination, we assume that the DM perceives a subjective utility equal to the objective payoff, but she additionally receives a utility of κ whenever she takes action 1. So, using $\mathbb{1}$ to denote the indicator function, her subjective utility is given by

$$u(a, \omega) = \pi(a, \omega) + \kappa \cdot \mathbb{1}_{a=1}. \quad (1)$$

The taste space K is assumed to be a closed interval around zero $K = [\underline{\kappa}, \bar{\kappa}]$ with $-\infty < \underline{\kappa} < 0 < \bar{\kappa} < +\infty$. If $\kappa > 0$, the DM has an inherent preference towards acting (action 1), whereas if $\kappa < 0$, the preference is towards not acting (action 0). Since $0 \in K$, it is always possible for the DM to choose without any taste biases.

With the above, the DM's subjective utility differences relate to objective payoff differences through

$$y(\omega) = x(\omega) + \kappa$$

and the probability that a DM with taste κ and SCR q takes action 1 at a state where the objective-payoff difference is x is given by

$$\Pr[a = 1 | x(\omega) = x] = q(y(\omega)) = q(x + \kappa).$$

As the taste κ increases, this stochastic choice is shifted to the left (see Figure 1) and—since q is strictly increasing—the probability of action 1 being chosen at any state increases.¹⁰

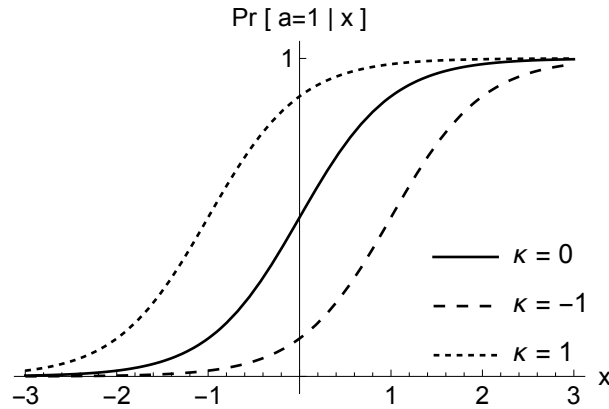


Figure 1: The effect of a taste κ to the stochastic choice rule.

3.1.4 Choice without mistakes

At a state ω , the DM's utility-maximising action is determined by the sign of $y(\omega)$. If $y(\omega) > 0$ the optimal action is 1, if $y(\omega) < 0$ the optimal action is 0, and if $y(\omega) = 0$ both

¹⁰Note that this can be easily generalised to a setting in which the DM's SCR is $q : X \times K \rightarrow \mathbb{R}$ with $(K, \geq_K)$ being a totally ordered space. In that setting, $q(x; \kappa)$ is strictly increasing in x for a fixed κ and strictly increasing (in $\geq_K$) in κ for a fixed x . We use the simpler setting as it is more intuitive and to ease exposition.

actions are optimal. Therefore, if the DM would be able to optimise exactly (i.e., if her SCR was the step function $q^*(y) = \mathbb{1}_{y \geq 0}$), she would maximise subjective utility at all states of the world. As a consequence, a DM with subjective utility equal to objective payoff ($u = \pi$ or equivalently $\kappa = 0$) and step-function SCR q^* maximises objective payoff at all states of the world ω . By assuming that the SCR is continuous (and strictly increasing) we rule out this case.

3.1.5 Objective payoff maximisation

While the DM's stochastic-choice rule q is assumed to be exogenous and constant, her tastes κ can *adapt* to her environment. Instead of explicitly modelling the evolutionary dynamics, we define objective payoff to be whatever matters for evolution or adaptation (whether interpreted biologically, culturally, or psychologically) and focus on the idealised case where selection is strong enough to eventually produce a DM who acts in a way that maximises objective payoff.

More specifically, the state of the world ω is a random variable, which is drawn according to a probability measure μ over the state space ($\mu \in \Delta(\Omega)$). We call μ the DM's *environment*. In this setup, evolution will equip the DM with a taste that induces actions that maximise her expected objective payoff, i.e., her tastes adapt to maximise

$$\begin{aligned} & \int \Pr[a = 1|\omega]\pi(1, \omega) + \Pr[a = 0|\omega]\pi(0, \omega)\mu(d\omega) = \\ &= \int q(y(\omega))\pi(1, \omega) + (1 - q(y(\omega)))\pi(0, \omega)\mu(d\omega) = \\ &= \int q(x(\omega) + \kappa)x(\omega)\mu(d\omega) + \int \pi(0, \omega)\mu(d\omega). \end{aligned} \quad (2)$$

Since the second term in the above expression does not depend on κ , the optimal taste in environment μ is the $\kappa \in K$ that maximises

$$U(\kappa; \mu) := \int q(x(\omega) + \kappa)x(\omega)\mu(d\omega). \quad (3)$$

This optimal taste is our central outcome of interest. More specifically, we are interested in exploring how the optimal taste is affected by the environment.

3.2 Effect of environments on tastes

3.2.1 Environments as distributions over objective-payoff differences

Notice that the DM's optimal taste depends on the environment μ only to the extent that μ affects the distribution of objective-payoff differences. Assuming that $x(\cdot)$ is μ -measurable,

it is straightforward to calculate the cumulative density function (CDF) of objective-payoff differences:

$$F_\mu(x) = \mu(\{\omega \in \Omega : x(\omega) \leq x\}).$$

We also assume that F_μ has a density function f_μ . A formal description of how we get from distributions over states to distributions over objective-payoff differences is given in Appendix A.

In what follows, we will often refer to CDFs F of objective-payoff differences as environments, suppressing the underlying dependence on the distribution over states (and the index μ). We denote the set of environments (i.e., the set of CDFs over X) by $\mathcal{F}$. In a similar spirit, we will occasionally refer to the objective-payoff difference x as the *state*, meaning that the state of the world ω is such that $x(\omega) = x$. With the above in mind and with a slight abuse of notation, we can write the maximand (3) in the simpler form:

$$U(\kappa; F) := \int xq(x + \kappa) dF(x). \quad (4)$$

We will also denote the set of maximising tastes in environment F by

$$\psi(F) := \arg \max_{\kappa \in K} U(\kappa; F). \quad (5)$$

Note that $U(\kappa; F)$ is continuous in κ and that K is a compact set. Therefore, for any environment F , $\psi(F)$ is a nonempty closed subset of K . From equation 5, it is clear that the optimal taste depends on the environment F .

3.2.2 Main result: comparative statics

Our goal is to conduct comparative statics exercises with respect to the environment. For this, we use the notions of first-order-stochastic-dominance and likelihood-ratio dominance orders.

Definition 1 (First-order stochastic dominance (FOSD)). *Let F_1, F_2 be CDFs over $X \subseteq \mathbb{R}$. We say that F_2 first-order stochastically dominates F_1 and write $F_2 \succeq_{FOSD} F_1$ if $F_2(x) \leq F_1(x)$ for all $x \in X$.*

Definition 2 (Likelihood-ratio (LR) dominance). *Let F_1 and F_2 be CDFs over $X \subseteq \mathbb{R}$ with respective density functions f_1 and f_2 . We say that F_2 LR dominates F_1 and write $F_2 \succeq_{LR} F_1$ if $f_2(x_2)f_1(x_1) \geq f_1(x_2)f_2(x_1)$ for any $x_2 > x_1$ in X .*

Note that $F_2 \succeq_{LR} F_1$ implies $F_2 \succeq_{FOSD} F_1$.¹¹

¹¹Letting $F|S$ denote the restriction of F on S , an equivalent definition of $F_2 \succeq_{LR} F_1$ is that $F_2|S \succeq_{FOSD} F_1|S$ for any $S \subseteq X$ (Wang and Lehrer 2024, Proposition 1). So, if $F_2 \succeq_{LR} F_1$, using $S = X$ this definition implies that $F_2 \succeq_{FOSD} F_1$.

For some of our results we need to assume that the SCR is *moderate* in the following sense.

Definition 3 (Moderate SCR). *Given $K \subseteq \mathbb{R}$ and $X \subseteq \mathbb{R}$, a stochastic choice rule $q \in \mathcal{Q}$ is called moderate if $xq(x + \kappa)$ exhibits increasing differences in $(x; \kappa)$ over $X \times K$. That is, if*

$$x_2q(x_2 + \kappa_2) - x_1q(x_1 + \kappa_2) \geq x_2q(x_2 + \kappa_1) - x_1q(x_1 + \kappa_1)$$

for any $x_2 > x_1$ in X and any $\kappa_2 > \kappa_1$ in K . The SCR q will be called strictly moderate if the inequality is always strict.

With this, our main technical result is the following.

Proposition 1. *Let $F_1, F_2 \in \mathcal{F}$ be environments and $q \in \mathcal{Q}$ be a stochastic choice rule.*

1. *If $F_2 >_{LR} F_1$, then $\min \psi(F_2) \geq \max \psi(F_1)$.*
2. *If q is moderate on $X \times K$ and $F_2 >_{FOSD} F_1$, then $\max \psi(F_2) \geq \max \psi(F_1)$ and $\min \psi(F_2) \geq \min \psi(F_1)$.*
3. *If q is strictly moderate on $X \times K$ and $F_2 >_{FOSD} F_1$, then $\min \psi(F_2) \geq \max \psi(F_1)$.*

Proof. See Appendix C.1.

Proposition 1 says that the tastes that individuals develop move in the direction of the objective incentives that their environments provide. If we compare two environments in the likelihood-ratio order, then part 1 states that the optimal tastes under the higher environment are at least as high as the optimal tastes under the lower environment. Part 3 states that the same is true under the FOSD order, provided that the SCR is strictly moderate. When we use the FOSD order but only moderate (not necessarily strictly moderate) part 2 states that the highest of the optimal tastes under the higher environment are at least as high as the highest of the optimal tastes under the lower environment, and the lowest of the optimal tastes under the higher environment are at least as high as the lowest of the optimal tastes under the lower environment.

The results in Proposition 1 rely critically on the assumption that mistakes are unavoidable. The following remark clarifies that if the SCR is a step function then $\kappa = 0$ is always optimal (see also Section 3.1.4).

Remark 1. *If $q(x) = \mathbb{1}_{x>0}$, then $0 \in \psi(F)$ for any environment $F \in \mathcal{F}$.*

The comparative statics of Proposition 1 are “robust” in the sense that they are detail-free. That is, the particular details of the q function don’t need to be known to establish

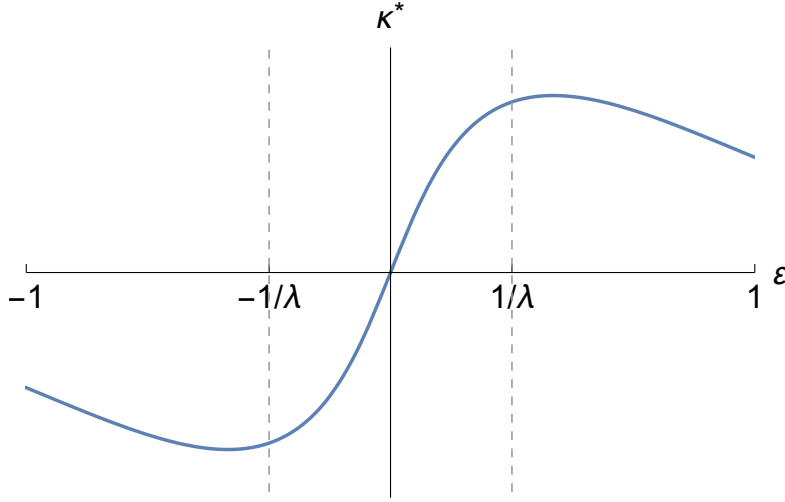


Figure 2: Optimal taste for three-state environments. Each environment has three equiprobable states that yield objective payoff differences $x \in \{-1, \varepsilon, 1\}$. The DM has a logistic SCR with imprecision parameter λ . For environments with $\varepsilon \in [-1/\lambda, 1/\lambda]$, monotone comparative statics are obtained, even though the SCR is not moderate over the whole $X \times K$.

the direction in which the optimal taste shifts when the environment changes. This is important, as it is very difficult to precisely measure such functions in the field, or even in the lab. As long as it is reasonable to assume that decision makers' q is strictly increasing in the application under consideration, then the taste they develop should follow monotone comparative statics in LR environment increases. If we further think that it is reasonable for q to be moderate, then monotone comparative statics are established *for all* FOSD increases.

3.2.3 Moderateness is Sufficient but not Necessary

We want to stress that the moderateness condition in Proposition 1 item 3 is sufficient but not necessary. That is, there are cases where higher environments (in the FOSD sense) lead to higher optimal tastes even with a q that is not moderate.

As an example, compare two environments with objective-payoff differences in $X = [0, 1]$. In the first environment F_1 , there are three equiprobable states that yield objective-payoff differences $x \in \{-1, \varepsilon_1, -1\}$, while in the second environment F_2 , the three equiprobable states yield objective-payoff differences $x \in \{-1, \varepsilon_2, 1\}$ for some $-1 < \varepsilon_1 < \varepsilon_2 < 1$. Notice that the change from F_1 to F_2 is an FOSD increase, but not an LR increase.

Now, consider a DM whose behaviour is represented by a logistic SCR with imprecision parameter λ , i.e., $q(y) = (1 + \exp(-y/\lambda))^{-1}$ and let the taste space be $K \supseteq [-1, 1]$. Notice that if $\lambda > 1$, then q is not moderate over $X \times K$. Nevertheless, for any small enough epsilon ($\varepsilon < 1/\lambda$), $\psi(F_2) > \psi(F_1)$. So, the increase in the environment $F_1 \rightarrow F_2$ leads to an increase in the optimal taste. Figure 2 illustrates this.

In Appendix B, we provide more details about the intuition behind the result of Proposition 1 and the relevance of the moderateness assumption.

3.3 Biological, Cultural, and Psychological Adaptation Mechanisms

We distinguish three broad potential interpretations of why subjective utility adapts to increase objective payoff. First, from the perspective of biological evolution, subjective preferences would have a genetic basis. Individuals with tastes that induce higher objective payoffs (fitness) enjoy more reproductive success. On this interpretation subjective utility evolves slowly over generations whereas behaviour can adjust much more quickly within a generation.

Second, from the perspective of cultural evolution or social learning, subjective preferences for actions may correspond e.g. to social norms which are internalised to different degrees by different individuals. Vertical (parent-to-child) and horizontal (peer-to-peer) transmission may be partly based on social status or wealth, or other culturally relevant markers of success, which would then correspond to objective payoff.

Third and finally, we may take an individualistic psychological perspective, assuming that both subjective preferences and the objective payoffs of our model are internal to an individual. We may think of this in terms of a dual self-model (Thaler and Shefrin 1981), such that the short-run self (doer or agent) maximises subjective utility whereas the long-run self (planner or principal) maximises objective payoff. The short-run doer makes all decisions and the long-run planner only occasionally steps in and adjusts the subjective preferences in order to better achieve the goal of maximising objective payoff.

We do not preclude the possibility that there are contexts in which the first (biological) interpretation is the most appropriate one. However, for our main applications we believe that the second perspective (cultural evolution and social learning) is most relevant. For interpreting laboratory evidence on preference adaptation, the third (psychological) interpretation, which is consistent with the quickest adaptation, is required.

4 Strategic interaction model

4.1 Setup

We extend the model to analyse the behaviour and taste formation of individuals in games. In this setting, a finite population N of $n \geq 2$ players play games drawn from a family of n -player, 2-strategy games, which are parametrised by a state $\theta \in \Theta \subseteq \mathbb{R}$. We will occasionally refer to state θ as “game θ ”. The state space $\Theta \subseteq \mathbb{R}$ is endowed with the Borel σ -algebra. We use $\bar{\theta} := \sup \Theta$ and $\underline{\theta} := \inf \Theta$ to denote the highest and the lowest state, respectively.

Each player $i \in N$ chooses an action a_i from the set $\{0, 1\}$ and the action profile space is denoted by $A := \{0, 1\}^n$ and is endowed with the coordinatewise order. When interacting in the games, players get *objective payoff*. Player i 's objective payoff from action a_i when her opponents play $\mathbf{a}_{-i} \in \{0, 1\}^{n-1}$ and the state is θ is denoted by $\pi_i(a_i, \mathbf{a}_{-i}; \theta)$.¹² We assume that for all action profiles $\mathbf{a} \in A$ and all players $i \in N$, $\pi_i(a_i, \mathbf{a}_{-i}; \cdot)$ is bounded and measurable. Finally, we denote player i 's objective-payoff difference between taking action 1 and action 0 at state θ and opponents' action profile $\mathbf{a}_{-i}$ by

$$x_i(\theta; \mathbf{a}_{-i}) := \pi_i(1, \mathbf{a}_{-i}; \theta) - \pi_i(0, \mathbf{a}_{-i}; \theta)$$

and use $\mathbf{x}(\theta; \mathbf{a})$ to denote the vector of objective-payoff differences of all players at state θ and action profile $\mathbf{a}$.

4.1.1 The games

We assume that all the games $\theta \in \Theta$ exhibit *strategic complementarities*. So, fixing a state θ , any player i 's objective-payoff difference increases as her opponents switch from action 0 to action 1. Moreover, as the state θ increases, all players' objective-payoff differences increase. Together, the above are summarised in the following assumption.

Assumption 1. *The payoff differences $\mathbf{x} : \Theta \times A \rightarrow \mathbb{R}^n$ are nondecreasing in both θ and $\mathbf{a}$.*

4.1.2 Behaviours

Consider a player i and denote by $p_i(\theta)$ the probability with which she takes action 1 when the state is θ . We will refer to $\theta \mapsto p_i(\theta)$ as player i 's *behaviour*. We let the behaviour space to be the set of nondecreasing measurable mappings $\Theta \mapsto [0, 1]$ and denote it by P .¹³ The *behaviour profile space* is denoted by $\mathcal{P} := P^n$ with typical element $\mathbf{p}$.

We extend the definition of objective payoffs to include (expected) objective payoffs from mixed-strategy profiles:

$$\pi_i(p_i(\theta), \mathbf{p}_{-i}(\theta); \theta) := \sum_{\mathbf{a} \in A} \pi_i(a_i, \mathbf{a}_{-i}, \theta) \prod_{j \in N} (p_j(\theta))^{a_j} (1 - p_j(\theta))^{1-a_j}$$

(where $0^0 := 1$, by convention). Doing the same for objective-payoff differences gives

$$x_i(\theta; \mathbf{p}_{-i}) := \pi_i(1, \mathbf{p}_{-i}(\theta); \theta) - \pi_i(0, \mathbf{p}_{-i}(\theta); \theta). \quad (6)$$

¹²The formulation of the payoffs here is very broad. In particular, it allows for π_i to express player i 's expected payoff when being matched to some other player j to play a 2×2 game. Of course, this matching can be uniformly random, assortative, or anything else (Jensen and Rigos 2018). Several of the applications of Section 6 make heavy use of this type of payoffs, which we think is more realistic for the respective cases.

¹³Even though somewhat limiting, the assumption that the various p_i are nondecreasing is natural in our setting: as θ increases, action 1 becomes more attractive and it should be expected that individuals are more likely to take this action.

Note that Assumption 1 implies that $x_i(\theta; \mathbf{p}_{-i})$ is nondecreasing in θ and in $\mathbf{p}_{-i}$.

As in Section 3, X_i denotes the smallest closed interval that contains all potential realisations of x_i , i.e.,

$$X_i := \left[\inf_{\theta \in \Theta, \mathbf{p}_{-i} \in P^{n-1}} x_i(\theta; \mathbf{p}_{-i}), \sup_{\theta \in \Theta, \mathbf{p}_{-i} \in P^{n-1}} x_i(\theta; \mathbf{p}_{-i}) \right] = [x_i(\underline{\theta}; \mathbf{0}), x_i(\bar{\theta}; \mathbf{1})]$$

and $X := \prod_{i \in N} X_i$.

4.1.3 Stochastic choice rules and tastes

As in Section 3, each player $i \in N$ makes decisions according to her stochastic choice rule $q_i \in \mathcal{Q}$ and subjective utility u_i which equals objective payoff with an added taste $\kappa_i \in K_i$ for action 1 (see eq. (1)). The taste space for player i is $K_i = [\underline{\kappa}_i, \bar{\kappa}_i]$ with $-\infty < \underline{\kappa}_i < 0 < \bar{\kappa}_i < +\infty$ and we use $K := \prod_{i \in N} K_i$ to denote the taste profile space. With that, when faced with objective payoff difference x_i , player i takes action 1 with probability $q_i(x_i + \kappa_i)$.

4.1.4 Strategic environments

A *strategic environment* ν is a probability measure over Θ . This is the parameter with respect to which will conduct comparative statics. The set of strategic environments is denoted by $\mathcal{N} := \Delta(\Theta)$ and the CDF of strategic environment ν by F_ν^θ . We use $F_{\nu; \mathbf{p}_{-i}}^{x_i}$ to denote the distribution of objective-payoff differences that player i faces in strategic environment ν under her opponents' behaviour profile $\mathbf{p}_{-i} \in P^{n-1}$. This can be calculated from the change of variables formula and eq. (6).

Now, since $p_j(\theta)$ is nondecreasing in θ for all $j \neq i$, Assumption 1 implies that $x_i(\theta; \mathbf{p}_{-i})$ is nondecreasing in θ . Therefore, $F_{\nu''}^\theta >_{FOSD} F_\nu^\theta$ implies that $F_{\nu''; \mathbf{p}_{-i}}^{x_i} \geq_{FOSD} F_{\nu; \mathbf{p}_{-i}}^{x_i}$ for any behaviours $\mathbf{p}_{-i} \in P^{n-1}$ and any player i . That is, as the environment increases in the FOSD sense, the distribution of objective-payoff differences that any player faces does not decrease.

4.2 Behaviour equilibrium

In this section we identify behaviour profiles that are self-consistent under a fixed taste profile κ and study their properties.

4.2.1 Noisy best responses

Fix the players' taste profile to be $\kappa \in K$ and consider a focal player i 's best response behaviour. When her opponents' behaviour profile is $\mathbf{p}_{-i}$ and the state is θ , player i experiences a subjective-utility difference given by $y_i(\theta) := x_i(\theta; \mathbf{p}_{-i}) + \kappa_i$. The SCR q_i , which

represents player i 's cognitive limitation, together with her taste κ_i imply player i 's noisy best response:

Definition 4 (Noisy best response). *Fix player i 's taste $\kappa_i \in K_i$. Player i 's noisy best response is the function $\phi_i(\cdot; \kappa_i) : P^{n-1} \rightarrow P$ defined through*

$$\phi_i(\mathbf{p}_{-i}; \kappa_i)(\theta) := q_i(x_i(\theta; \mathbf{p}_{-i}) + \kappa_i). \quad (7)$$

The joint noisy-best-response function under taste profile $\boldsymbol{\kappa}$ is denoted by $\boldsymbol{\phi}(\cdot; \boldsymbol{\kappa}) : \mathcal{P} \rightarrow ([0, 1]^\Theta)^n$.

Since (i) $x_i(\theta, \mathbf{p}_{-i})$ is nondecreasing in θ and $\mathbf{p}_{-i}$ (Assumption 1), (ii) $p_j(\theta)$ is nondecreasing in θ for all $j \neq i$, and (iii) q_i is strictly increasing, $\phi_i(\mathbf{p}_{-i}; \kappa_i)(\theta)$ is a nondecreasing function of θ . So, best responses to nondecreasing behaviour profiles are also nondecreasing, meaning that the set of nondecreasing behaviour profiles is both the domain and range of the joint noisy best response function, i.e., $\boldsymbol{\phi}(\cdot; \boldsymbol{\kappa}) : \mathcal{P} \rightarrow \mathcal{P}$ for any $\boldsymbol{\kappa}$.

The definition of noisy best responses gives rise to the following notion of *behaviour equilibrium* for a given taste profile $\boldsymbol{\kappa} \in K$.

Definition 5 (Behaviour equilibrium). *Fix a taste profile $\boldsymbol{\kappa} \in K$. A behaviour profile $\mathbf{p}^* \in \mathcal{P}$ is a behaviour equilibrium under tastes $\boldsymbol{\kappa}$ if it is a fixed point of $\boldsymbol{\phi}(\cdot; \boldsymbol{\kappa})$, i.e., if $\mathbf{p}^* = \boldsymbol{\phi}(\mathbf{p}^*; \boldsymbol{\kappa})$.*

For any taste profile $\boldsymbol{\kappa} \in K$, we denote the set of behaviour equilibria by $\mathcal{B}(\boldsymbol{\kappa})$. Proposition 2 establishes the existence of behaviour equilibria and that extremum equilibria follow monotone comparative statics in $\boldsymbol{\kappa}$.¹⁴

Proposition 2. *For any $\boldsymbol{\kappa} \in K$, $\mathcal{B}(\boldsymbol{\kappa})$ is nonempty and has a largest and a smallest element. Moreover, $\min \mathcal{B}(\boldsymbol{\kappa})$ and $\max \mathcal{B}(\boldsymbol{\kappa})$ are strictly increasing in $\boldsymbol{\kappa}$.*

Proof. First, $\boldsymbol{\phi}$ is a nondecreasing function of $\mathbf{p}$. Second, because each q_i is strictly increasing, $\phi_i(\mathbf{p}_{-i}; \kappa_i)$ is strictly increasing in κ_i and therefore $\boldsymbol{\phi}$ is a strictly increasing function of $\boldsymbol{\kappa}$. Third, $\mathcal{P}$ is a complete lattice. Finally, $(K, \geq)$ is a poset. With the above, the result follows from Milgrom and Roberts (1994), Theorem 3 (iii). $\square$

4.3 Taste equilibrium

Now consider the case where the players interact in strategic environment $\nu \in \mathcal{N}$. When the players' behaviour profile is $\mathbf{p}$, player i 's expected objective payoff is given by

$$\pi_i(\mathbf{p}, \nu) = \int_{\Theta} p_i(\theta) x_i(\theta; \mathbf{p}_{-i}) \nu(d\theta) + \int_{\Theta} \pi(0, \mathbf{p}_{-i}; \theta) \nu(d\theta). \quad (8)$$

¹⁴We use the componentwise order to compare behaviour profiles, i.e., $\mathbf{p}'' \geq \mathbf{p}'$ if $p_i''(\theta) \geq p_i'(\theta)$ for all $i \in N$ and all $\theta \in \Theta$.

As in Section 3, tastes are *adaptive*. So, any focal player i 's taste κ_i adjusts to her environment in order to deliver the highest expected objective payoff. As the last term in expression (8) above does not depend on player i 's own behaviour, her optimal taste results from the maximisation of the objective

$$\int_{\Theta} p_i(\theta) x_i(\theta; \mathbf{p}_{-i}) \nu(d\theta).$$

Since player i uses q_i to noisily best respond, her behaviour is given by $p_i = q(x_i + \kappa_i)$ and her objective when her taste is κ_i and the behaviours of her opponents are $\mathbf{p}_{-i}$ is given by

$$U_i(\kappa_i; \mathbf{p}_{-i}, \nu) = \int_{\Theta} [q_i(x_i(\theta; \mathbf{p}_{-i}) + \kappa_i) x_i(\theta; \mathbf{p}_{-i})] \nu(d\theta).$$

Thus, her optimal taste is

$$\psi_i(\mathbf{p}_{-i}; \nu) := \arg \max_{\kappa_i \in K_i} U_i(\kappa_i; \mathbf{p}_{-i}, \nu). \quad (9)$$

The function ψ_i gives the optimal taste of player i as a response to the behaviours of her opponents. In this sense, it is player i 's *taste best response*. The joint taste best response to behaviours $\mathbf{p}$ in environment ν is denoted by $\psi(\mathbf{p}; \nu)$.

An important distinction between our current setting and the individual-decision situations of Section 3 is that in strategic situations, a player's environment depends not only on the distribution of states of nature (games), but also on the behaviour of the player's opponents. Proposition 3 shows that an individual's preferences adjust in response to others' behaviour.

Proposition 3. Fix a strategic environment $\nu \in \mathcal{N}$ and let $\mathbf{p}'_{-i}, \mathbf{p}''_{-i} \in P^{n-1}$ be behaviour profiles of player i 's opponents with $\mathbf{p}''_{-i} \geq \mathbf{p}'_{-i}$. Consider player i 's taste best response $\psi_i(\cdot; \nu)$.

1. If player i 's SCR q_i is moderate, then

$$\max \psi_i(\mathbf{p}''_{-i}; \nu) \geq \max \psi_i(\mathbf{p}'_{-i}; \nu) \text{ and } \min \psi_i(\mathbf{p}''_{-i}; \nu) \geq \min \psi_i(\mathbf{p}'_{-i}; \nu).$$

2. If player i 's SCR q_i is strictly moderate, then

$$\min \psi_i(\mathbf{p}''_{-i}; \nu) \geq \max \psi_i(\mathbf{p}'_{-i}; \nu).$$

Proof. Because of Assumption 1, $\mathbf{p}''_{-i} \geq \mathbf{p}'_{-i}$ implies $x_i(\theta; \mathbf{p}''_{-i}) \geq x_i(\theta; \mathbf{p}'_{-i})$ for any $\theta \in \Theta$. It follows that $F_{\nu; \mathbf{p}''_{-i}}^{x_i} \geq_{FOSD} F_{\nu; \mathbf{p}'_{-i}}^{x_i}$. Finally, using Proposition 1 (items 2 and 3), we get items 1 and 2, respectively. $\square$

As seen in the above Proposition, one's optimal taste can vary with the behaviour of one's opponents. The resulting notion of a *taste equilibrium* requires that behaviours are in behaviour equilibrium given players' tastes and that tastes are optimal (i.e., taste best responses) given these behaviours and the strategic environment.

Definition 6 (Taste equilibrium). A pair $(\tilde{p}, \tilde{\kappa}) \in \mathcal{P} \times K$ is a taste equilibrium of strategic environment $v \in \mathcal{N}$ if

$$\tilde{p} = \phi(\tilde{p}; \tilde{\kappa}) \quad \text{and} \quad \tilde{\kappa} \in \psi(\tilde{p}; v).$$

We denote the set of taste equilibria of strategic environment v by $\mathcal{T}(v)$. Our main result for strategic-interaction settings is that taste equilibria are nondecreasing as the strategic environment increases. So, as action 1 becomes more objectively valuable, in equilibrium, individuals follow action 1 more frequently in all states θ and develop stronger tastes towards it.

Proposition 4. Let q_i be moderate for all players $i \in N$. Then, for any $v \in \mathcal{N}$, $\mathcal{T}(v)$ has a largest and a smallest element. Moreover, if $v', v'' \in \mathcal{N}$ and $v'' \geq_{FOSD} v'$, then $\min \mathcal{T}(v'') \geq \min \mathcal{T}(v')$ and $\max \mathcal{T}(v'') \geq \max \mathcal{T}(v')$.

Proof. See Appendix C.2.

5 Dynamics

In this section, we provide a dynamic justification for the static results presented above, in particular the monotone comparative statics of Proposition 4.

Our dynamic model has two key features: it is based on best responses and it takes place over two different time scales.¹⁵ In the short run, tastes are assumed to be constant and behaviour is allowed to equilibrate. During this process, a behaviour equilibrium is reached based on players' taste profile. At a slower time scale, tastes are allowed to adapt, so that the players whose tastes adapt are taste best responses to their opponents' behaviours. In this sense, behaviour equilibria can be viewed as *short-run equilibria*, whereas taste equilibria are *long-run equilibria*.

On a literal reading of the model, a psychological dual-self interpretation may be most natural. The short-run doer adjusts behaviour rapidly by best-responding to current behaviour, given fixed tastes, while the long-run planner occasionally adjusts the tastes by best-responding to current behaviour (as induced by current tastes). When adopting an evolutionary interpretation (either biological or cultural), the best response dynamic should be viewed as a tractable alternative to other more realistic, more noisy, payoff-monotonic selection dynamics.

¹⁵A similar dynamic is used by Sandholm (2001).

5.1 Short-run dynamics

Let κ be a taste profile and let the players' behaviours be in the behaviour equilibrium $\mathbf{p} \in \mathcal{B}(\kappa)$. Now, consider a change in the taste profile so that the new profile is κ' . How will behaviours adjust to the new profile? We consider the following dynamic process.

In the short run, tastes are fixed, so only behaviours are allowed to adjust. We denote the initial behaviour profile by $\mathbf{p}^0 := \mathbf{p}$. At each period t in discrete time, a player $r(t) \in N$ gets an opportunity to revise her behaviour. At each period $t > 0$, one player $r(t) \in N$ is prompted to revise her behaviour. We assume that the *revision sequence* r is such that all players receive opportunities to revise their behaviour over an infinite horizon. More formally, we give the following definition.

Definition 7 (Revision sequence). *A function $r : \mathbb{N}_+ \rightarrow N$ is a revision sequence if for any player $i \in N$ and any period $T \in \mathbb{N}_+$, there is some other period $t > T$ with $r(t) = i$.*

Upon receiving a revision opportunity, a player i revises her behaviour by noisily best responding to the previous-period behaviour profile of her opponents $\mathbf{p}_{-i}^{t-1}$, i.e., she adjusts to the behaviour $p_{r(t)}^t := \phi_{r(t)}(\mathbf{p}_{-i}^{t-1}; \kappa'_i)$, leading to the behaviour profile at time t being $\mathbf{p}^t = (p_1^{t-1}, \dots, p_{r(t)-1}^{t-1}, p_{r(t)}^t, p_{r(t)+1}^{t-1}, \dots, p_n^{t-1})$. We denote the sequence of behaviour profiles $(\mathbf{p}^0, \mathbf{p}^1, \mathbf{p}^2, \dots)$ generated in this way by $\mathcal{S}(\kappa, \kappa', \mathbf{p}; r)$.

For the short-run dynamic process, we provide Proposition 5, which shows that an increase in tastes leads to higher equilibrium behaviours.¹⁶

Proposition 5. *Let κ', κ be taste profiles such that $\kappa' > \kappa$, $\mathbf{p} \in \mathcal{B}(\kappa)$, and r be a revision sequence. Then $\mathcal{S}(\kappa, \kappa', \mathbf{p}; r)$ converges upwards to a behaviour equilibrium of κ' .*

Proof. See Appendix C.3.

That is, suppose the initial taste profile is κ and behaviour is at equilibrium given κ , and suppose there is an upward shift in tastes, to the taste profile κ' . The proposition then states that a sequence of best-response adjustments of behaviour will converge at a behavioural equilibrium given κ' . Moreover, this equilibrium is higher than the initial equilibrium.

5.2 Long-run dynamics

Let ν be an environment and the behaviour and tastes of players be in the taste equilibrium $(\mathbf{p}, \kappa) \in \mathcal{T}(\nu)$. Now, consider that the environment changes so that the new environment is ν' . How will the population's tastes adjust to the new environment?

To answer this, we consider the following dynamic process. We denote the initial behaviour-taste profile by $(\mathbf{p}^0, \kappa^0) := (\mathbf{p}, \kappa)$. At each period t in discrete time, a player

¹⁶See also Vives (1990), Theorem 5.1 for a similar result.

$r(t) \in N$ gets an opportunity to revise her taste. As in the short-run dynamics, we assume that players receive opportunities to revise their tastes according to a revision sequence r . When given a revision opportunity, the player revises her taste by best responding to the previous period's behaviour profile $\mathbf{p}^{t-1}$, i.e., she adjusts to some taste $\kappa_{r(t)}^t \in \psi_{r(t)}(\mathbf{p}_{-r(t)}^{t-1}; \nu')$, leading to the taste profile at time t being $\boldsymbol{\kappa}^t = (\kappa_1^{t-1}, \dots, \kappa_{r(t)-1}^{t-1}, \kappa_{r(t)}^t, \kappa_{r(t)+1}^{t-1}, \dots, \kappa_N^{t-1})$.

Once the new taste profile $\boldsymbol{\kappa}^t$ has been reached, a short-run adjustment process begins from $(\boldsymbol{\kappa}^t, \mathbf{p}^{t-1})$, as explained in Section 5.1, until a behaviour profile $\mathbf{p}^t \in \mathcal{B}(\boldsymbol{\kappa}^t)$ has been established (the sequence converges for the changes we consider, see Proposition 5). Then, a new period $t + 1$ begins. We denote the sequence of behaviour-taste profiles $((\mathbf{p}^0, \boldsymbol{\kappa}^0), (\mathbf{p}^1, \boldsymbol{\kappa}^1), (\mathbf{p}^2, \boldsymbol{\kappa}^2), \dots)$ generated in this way by $\mathcal{L}(\nu, \nu', (\mathbf{p}, \boldsymbol{\kappa}); r)$.

Following similar arguments as those used in Proposition 5, we arrive at Proposition 6, which states that the dynamic adjustment process leads to changes in tastes and behaviours in the direction of the environment change. Specifically, suppose we are initially at a taste equilibrium of an environment ν , and suppose there is an upward shift in environment, to ν' . The proposition then states that a sequence of best-response adjustments of behaviour and tastes will converge upwards to a taste equilibrium of ν' .

Proposition 6. *Let ν', ν be strategic environments such that $\nu' >_{FOSD} \nu$, $(\mathbf{p}, \boldsymbol{\kappa}) \in \mathcal{T}(\nu)$, and r be a revision sequence. Let also the SCR q_i of any player $i \in N$ be moderate. Then $\mathcal{L}(\nu, \nu', (\mathbf{p}, \boldsymbol{\kappa}); r)$ converges upwards to a taste equilibrium of ν' .*

6 Applications

We will now apply our theoretical framework and results to a number of empirically documented cases of cross-cultural variation in preferences.

6.1 Preferences against lying and other norm violations

6.1.1 Lying aversion and lie-detection

Whether or not a person is willing to lie depends in part on the expected consequences of lying, as influenced by the probability that a lie is detected, the severity of associated sanctions, and the material benefits of a successful lie (Fischbacher and Föllmi-Heusi 2013; Gneezy 2005). In addition, people appear to suffer a utility cost from the mere act of lying (Abeler, Nosenzo, and Raymond 2019). Our theory explains how the strength of the material consequences of lying (playing the role of objective payoffs in our model) may influence the development of subjective preferences against the act of lying.

We consider both the individual choice framework and the strategic choice framework. In either case, let strategy 1 correspond to truthfulness and let strategy 0 correspond to

lying. Furthermore, let the state of the world, denoted τ , correspond to the probability that a lie is detected. We assume $\tau \in [\underline{\tau}, \bar{\tau}] \subset [0, 1)$. An environment is a probability measure μ on $[\underline{\tau}, \bar{\tau}]$, so that a higher environment implies that lying is more likely to be detected.

6.1.2 Individual choice

To apply the individual-choice framework, suppose that truthfulness yields a payoff of 1. Lying yields an immediate added benefit of $g > 0$, but if the lie is detected, there is also a cost $c > 0$ of being sanctioned. Thus, the expected objective payoff from lying is $1 + g - \tau c$. The expected objective-payoff difference between truthfulness (strategy 1) and lying (strategy 0) is $x_i(\tau) = \tau c - g$.

An environment μ induces a CDF F_μ on the set of payoff differences. Consider two environments μ'' and μ' . If $F_{\mu''}$ is higher than $F_{\mu'}$ in the LR-order, then Proposition 1, part 1 implies that the objective-payoff-maximising tastes under the higher environment are at least as high as the objective-payoff-maximising tastes under the lower environment. Parts 2 and 3 deliver similar results for moderate stochastic choice rules and the FOSD-order. Thus, an environment where lying is more often detected generates a stronger preference against lying.

6.1.3 Strategic choice

Let us now expand our analysis to the strategic choice framework. As before, truthfulness yields a payoff of 1 and lying yields an immediate added benefit of $g > 0$. However, we now assume that if the lie is detected then the associated cost $c(\mathbf{p}_{-i}) > 0$ is a function of how prevalent lying is in the population. Thus, the expected objective-payoff difference between truthfulness and lying is now written $x_i(\tau; \mathbf{p}_{-i}) = \tau c(\mathbf{p}_{-i}) - g$. Note that this is strictly increasing in τ . It also seems reasonable to assume that the cost of being caught lying is non-increasing in the share of the population that lies. If many people lie, then lies will be punished less harshly and less frequently. It follows that $x_i(\tau; \mathbf{p})$ is non-decreasing in both $\mathbf{p}_{-i}$ and τ .

Suppose that the stochastic choice rules of all players are strictly moderate.¹⁷ Our results for the strategic setting now imply that people will have a stronger preferences against lying in environments where being caught lying is more likely (Proposition 4). More precisely, consider two environments μ'' and μ' . If $F_{\mu''}$ is higher than $F_{\mu'}$ in the FOSD-order, then Proposition 4 implies that the highest and the lowest of the equilibrium tastes and behaviours under the higher environment μ'' are at least as high the respective tastes and behaviours under the lower environment μ' . Even if the whole population is not in equilib-

¹⁷We also require that the function c is bounded such that the payoff difference x_i is contained in an interval.

rium, our Proposition 3 (part 2) suggests that each individual will develop stronger lying aversion when those that surround her are less likely to lie.

In line with these theoretical results, Gächter and Schulz (2016) document a negative cross-country correlation between preferences against lying and a measure of the prevalence of rule violations (based on measures of political fraud, tax evasion, and corruption). Since a higher prevalence of rule violations is plausibly associated with lower probability that violations are detected, our model provides a possible explanation for why higher prevalence of rule violations is associated with weaker preferences against lying.¹⁸

6.1.4 Social norms in general

Our explanation of lying aversion can be extended to other norm violations. In general, actions that violate a social norm may result in punishment, shunning, and loss of status. Whether the costs of violating a norm exceeds the benefits will depend on the probability that transgressions are detected, as well as on the accompanying sanctions. It seems plausible that the cost of violating a norm is decreasing in the share of other people violating the norm, since more wide-spread violations will probably dilute the associated stigma and make effective punishment more difficult to maintain. As in the case of lying aversion, our model would then predict that in environments where the costs of norm violations tend to outweigh the benefits, people will develop a preference against actions that count as norm violations (Bicchieri 2005; Ellingsen and Johannesson 2004; Ellingsen and Mohlin 2024; Krupka and R. A. Weber 2013; López-Pérez 2008).

For example, consider fairness norms in bargaining. Henrich et al. (2004) find that among the people of Lamalera in Indonesia almost everyone offers at least 50% in the Ultimatum game. The authors relate this to the importance of collective whale hunt, and the accompanying division of whale meat according to strict rules, where making an unfair offer would be costly in terms of a bad reputation or sanctions.

At a higher level, we conjecture that the identified mechanism could help explain the human capacity to identify and internalise social norms. Consider a society with many different social norms. In each of the many situations with a sufficiently strictly enforced social norm, our theory would predict a preference against the proscribed actions. It seems that a more economical way of handling all the situations together would be to combine (i) a general capacity to detect what the norms in a situation are with (ii) a general preference against actions identified as violating social norms. We plan to revisit this idea in future work.

¹⁸We do not wish to deny that there are other plausible explanations. It may be that weaker preferences against lying causes higher levels of political fraud, tax evasion, and corruption. Alternatively, it may be that a high prevalence of rule violations causes people to believe that lying is not immoral.

6.2 Conditional cooperation

6.2.1 Conditionally Cooperative Preferences

Conditional cooperation is a well-documented behavioural regularity, which has been invoked to explain behaviour in finitely repeated social dilemmas (Fischbacher, Gächter, and Fehr 2001; Kocher et al. 2008; Rustagi, Engel, and Kosfeld 2010). Experimentally, conditional cooperation is typically elicited via the strategy method in a sequential public goods game (Fischbacher, Gächter, and Fehr 2001; Kocher et al. 2008). Subjects are asked about both their unconditional contribution and their contributions conditional on average contributions by the other players. The correlation between own contributions and the others' average contributions serves as a measure of conditional cooperation. Rustagi, Engel, and Kosfeld (2010) find that a group's share of conditional cooperators in a lab experiment is positively correlated with the group's ability to cooperate in a field setting (forest management and monitoring).

Many different motivations can explain conditionally cooperative behaviour, including inequity aversion and different forms of preferences for reciprocity (Fischbacher, Gächter, and Fehr 2001). However, we will treat conditionally cooperative behaviour as motivated by preference specifically for conditionally cooperative strategies.

6.2.2 Conditional Cooperation via Repetition

When a social dilemma is indefinitely repeated with a sufficiently high repetition probability, entirely selfish players can be incentivised to cooperate by using repeated-game strategies that condition future cooperation on current behaviour. Our model explains how a preference for conditionally cooperative strategies may develop in such a setting. To see this, consider the following general description of the one-shot PD. Objective payoffs are normalised so that mutual cooperation yields 1 and mutual defection yields 0. The parameter $g > 0$ represents the gain from defecting against a cooperator and $l > 0$ is the loss suffered by the one who cooperates against defection.

	C	D
C	1	$-l$
D	$1 + g$	0

Now consider a setting where this game is indefinitely repeated. Let the state of the world correspond to the continuation probability $\delta \in [\underline{\delta}, \bar{\delta}] \subset [0, 1)$. An environment is a probability measure μ on $[\underline{\delta}, \bar{\delta}]$, meaning that a higher environment implies that games tend to last longer. We restrict attention to two salient strategies. Let action 1 correspond to a conditionally cooperative strategy (denoted CC), such as Tit-for-tat or Grim Trigger.

Let action 0 correspond to a strategy, such as Always Defect (denoted AD), that induces defection when pitted against the conditionally cooperative strategy. The objective payoffs of the repeated game are given by the discounted stage-game payoffs as follows.

	CC	AD
CC	1	$-l(1 - \delta)$
AD	$(1 - \delta)(1 + g)$	0

We assume that individuals from the population are randomly matched into pairs who play the indefinitely repeated game. Let $\bar{p}_{-i} := \sum_{j \neq i} p_j / (n - 1)$ denote the average behaviour of a player i 's opponents. The expected objective payoff difference between action 1 (CC) and action 0 (AD) is

$$x_i(\delta; \mathbf{p}_{-i}) = \bar{p}_{-i}(\delta)(1 - (1 - \delta)(1 + g)) + (1 - \bar{p}_{-i}(\delta))(-l)(1 - \delta).$$

This difference is increasing in δ . Moreover, if $l > g$ (i.e., if the underlying stage-game PD exhibits strategic complementarity), then the objective-payoff difference is nondecreasing in $\bar{p}_{-i}$ and hence in $\mathbf{p}_{-i}$ (c.f. Takahashi 2010, Heller and Mohlin 2018). Also note that the payoff difference x_i is contained in an interval.

Suppose that the stochastic choice rules of all players are moderate. Thus, if $l > g$, then the conditions of Proposition 4 are satisfied and we can conclude that people will have stronger preferences for conditional cooperation in environments where interactions tend to last longer (comparing environments with the FOSD-order). Moreover, Proposition 3 implies that for a given environment, preferences for cooperation will be stronger in a society that has coordinated on a more cooperative equilibrium (an equilibrium with a higher $\mathbf{p}$).

Our theoretical results are in line with the experimental evidence from Peysakhovich and Rand (2016), assuming a psychological interpretation of adaptation. They let subjects play indefinitely repeated games under parameters that either support cooperative equilibria or not. Subjects who experience the cooperative-conducive environment subsequently behave more prosocially in a range of games. They do not explicitly elicit preferences, but their results from the dictator game can be taken as evidence for a shift in preferences that is caused by the difference in experienced environments. Cooperative preferences can also develop if there is external, third-party, enforcement of cooperation in repeated games.¹⁹ Engl, Riedl, and R. Weber (2021) demonstrate experimentally that subjects who play a finitely repeated public goods game where low contributions are externally punished subsequently exhibit stronger preferences for conditional cooperation, compared with subjects who did not experience the external punishment institution.²⁰

¹⁹External punishment institutions may themselves build on repeated game strategies, as in Mohlin, Rigos, and Weidenholzer (2023).

²⁰Cassar, d'Adda, and Grosjean (2014) let subjects play series of one-shot Prisoner's dilemma games with

6.2.3 Conditional Cooperation via Reputation

The objective incentives to cooperate need not be provided in the form of a repeated interaction in a fixed constellation of players. It may also be enforced by reputations in an environment where individuals are randomly matched to play with different opponents but have some information about their partner's past behaviour, as in models of community enforcement (Heller and Mohlin 2018; Kandori 1992; Sugaya and Wolitzky 2021).

To represent this in our framework, suppose that agents from the population interact recurrently over an indefinite number of rounds, with continuation probability δ . In each round agents are randomly matched in pairs to play a one-shot PD (with the same parameterisation as above). Following Kandori (1992) we assume that there is a reputation system in place which operates as follows: (i) Everyone starts the first round in good standing. (ii) A player enters bad standing if she plays *D* against someone who is in good standing. (iii) A player regains good standing by playing *C* against someone who is in good standing.

The state of nature $\rho \in [\underline{\rho}, \bar{\rho}] \subset (0, 1]$ is the probability that a player who is actually in bad standing appears to her current opponent to be in bad standing. An agent in good standing never appears to be in bad standing. An environment is a probability measure on $[\underline{\rho}, \bar{\rho}]$. Thus, a higher environment implies that reputation or gossip is more reliable.

We reduce complexity and consider only two strategies. Strategy 1 is a reputation-conditional strategy, denoted *CC*. It starts out playing *C* and in later rounds it plays *C* if and only if the current opponent appears to be in good standing.²¹ Strategy 0 is always defect (*AD*).

Note that an *AD*-player enters bad standing at the end of the first round and then remains in bad standing forever. A *CC*-player may sometimes misperceive someone in bad standing as having good standing, leading her to cooperate against someone in bad standing. However, she will never defect against someone in good standing. Thus, she remains in good standing forever. Hence, the expected average discounted payoff of a *CC*-player is

$$\pi_i(1, \mathbf{p}_{-i}; \rho) = \bar{\mathbf{p}}_{-i}(\rho) - (1 - \bar{\mathbf{p}}_{-i}(\rho))(1 - \rho)l,$$

and the expected average discounted payoff of a *AD*-player is

$$\pi_i(0, \mathbf{p}_{-i}; \rho) = (1 - \delta)\bar{\mathbf{p}}_{-i}(\rho)(1 + g) + \delta\bar{\mathbf{p}}_{-i}(\rho)(1 + g)(1 - \rho).$$

and outside option. Those who play under a regime of impartial enforcement of cooperative behaviour later display higher trust and trustworthiness in a one-shot trust game, compared to those who play the series of one-shot Prisoner's dilemma games under less cooperation-inducing institutions. Note that trustworthiness can be interpreted as revealing a shift in preferences, whereas trusting behaviour also depends on beliefs.

²¹The *CC* strategy and the reputation system defined above corresponds to a Kandori's community-enforcement system with a single punishment period and the added possibility of misperception. Alternatively, it can be seen as a community version of Contrite Tit-For-Tat.

It can be verified that the difference $x_i(\rho; \mathbf{p}_{-i}) = \pi_i(1, \mathbf{p}_{-i}; \rho) - \pi_i(0, \mathbf{p}_{-i}; \rho)$ is increasing in ρ . Furthermore it can be shown that if the underlying game is strictly supermodular ($l > g$) then x_i is increasing in $\bar{\mathbf{p}}_{-i}$, hence increasing in $\mathbf{p}_{-i}$, for high enough δ (and for all values of ρ). Thus, assuming that stochastic choice rules are strictly moderate, we find that if the underlying PD exhibits strategic complementarity, and the continuation probability δ is sufficiently high, then preferences for conditional cooperation are stronger in environments where there is more reliable reputation transmission.²²

6.2.4 Market integration and conditional cooperation

Rustagi (2023) documents a plausibly causal effect of market participation (measured as distance from markets) on conditional cooperation (measured in a non-market setting).²³ Interestingly, the effect is only present for markets that feature asymmetric information about quality (specifically cattle markets), where the seller might be tempted to sell substandard goods, and hope to get away with it because of limited regulation and enforcement. From the perspective of our theory, the fact that a preference for conditional cooperation develops among those participating in markets with asymmetric information is explained by the fact that opportunistic behaviour, despite being tempting, typically does not pay off (since it may ruin one's good standing in a market where community enforcement is at work, as documented by Rustagi). Those who live further from markets instead rely on exchange within a close network of people they know. Hence the temptation to cheat that is present among strangers on a market is effectively quelled. This means there is less need to develop a preference that mitigates cheating, (again in line with documented by Rustagi's findings).

6.3 Kin-based Enforcement Systems and Betrayal aversion

Anthropologists point to the general importance of kinship systems for understanding differences in social organisation and culture (Parkin 1997). There is a long line of research in the social sciences, going back to M. Weber (2001) and Banfield (1967), suggesting that strong kin networks hamper the development of impartial formal institutions.²⁴ In a society

²²Reputation-based incentives may also explain why preferences for rejection of unfair offers develop in real-life bargaining situations, where interactions are open-ended and information about acceptance/rejection decisions may reach future bargaining partners. These preferences will then be expressed also in one-shot interactions in lab, such as the ultimatum game (c.f. Nowak, Page, and Sigmund 2000). We provide a formal argument in Appendix D.

²³Henrich et al. (2004) find positive correlation between market integration and size of offers in the ultimatum game.

²⁴Khaldun (2004) is an early precursor (c.f. Gellner 2000).

with strong kinship networks, transactions mainly occur between members of the same kin group and enforcement of cooperation is based on reciprocity and reputation. In a society with looser kinship structure, transactions will frequently occur between members of different kin groups, and hence enforcement needs to rely more on impersonal formal-legal institutions. Economists have recently become interested in this distinction in relation to the development of formal and informal institutions (Enke 2019; Greif and Tabellini 2017; Schulz 2022). We will focus on one example, betrayal aversion, as documented by Bohnet, Herrmann, and Zeckhauser (2010).

Betrayal aversion refers to the tendency to be less willing take on risk in the form of trusting another human than risk that is generated by a non-human randomisation device, despite the induced payoff distributions being identical. Experimentally this can be examined by comparing two versions of the trust game with binary actions. In the standard trust game, the first mover is asked for the threshold probability of the second mover being trustworthy (i.e., returning money) that would make the first mover indifferent between trusting and not trusting. This is compared with a version of the game where the second mover's decision whether to return money or not is made by a computer. If a person requires a higher probability of money being returned in the standard trust game than in the version with a computerised return decision, then she displays betrayal aversion.

Bohnet, Herrmann, and Zeckhauser (2010) find higher levels of betrayal aversion in three countries on the Arabian gulf (Kuwait, Oman and the United Arab Emirates) than in two western countries (Switzerland and the USA). The authors argue (building on Rosen 2000) that in Gulf countries the probability of facing untrustworthy behaviour is lower than in the West, since interactions to a large extent occur within clan and kinship networks, where reputations and repeated interactions discipline behaviour. Still, interactions outside of these networks may be more risky than in the west, and there is limited compensation to be had in case one is betrayed. Thus, it makes sense to behave differently towards kin and non-kin. Importantly, laboratory experiments take place outside clan and kin networks. Consequently behaviour in experiments should reflect tastes and behavioural tendencies that have developed in interactions with non-kin, even if interactions with kin are more common in real life. By contrast, enforcement in the West is more based on formal institutions, including the law and law enforcement, which will typically require payment of damages to betrayed parties. Hence, kin and non-kin can be treated in more or less the same way.²⁵ Our model clarifies why these differences in incentives surrounding trust and betrayal may induce differences in preferences that, in turn, induce differences in behaviour

²⁵Bigoni et al. (2019) similarly find higher levels of betrayal aversion in northern Italy than in southern Italy. It is possible that similar differences in kin-based enforcement mechanisms may matter for these differences as well, as suggested by Banfield (1967).

in anonymous one-shot laboratory settings.

6.3.1 Individual choice

We first analyse a single DM's decision whether to trust or not. We identify action 1 with trusting (T) and action 0 with not trusting (NT). Let the state of the world $\gamma \in [\underline{\gamma}, \bar{\gamma}] \subset (0, 1]$ measure the probability that trust is rewarded. An environment is a distribution on this set, so that a higher environment means that trust is more likely to be rewarded. Not trusting yields an outside option of 0. The payoff from trusting is 1 if trust is rewarded and $1 - S < 1$ if trust is not rewarded. The expected objective-payoff difference between trusting and not trusting is $x_i(\gamma) = 1 - (1 - \gamma)S$. Comparing two environments that are LR-ordered, Proposition 1 implies that the objective-payoff-maximising taste for trust are at least as low in the low trustworthiness environment as in the high trustworthiness environment. Again, parts 2 and 3 deliver similar conclusions for moderate stochastic choice rules and the FOSD order.

6.3.2 Strategic choice

Consider the following stag Hunt game, where players choose between Stag (action 1), which represents trusting and rewarding trust, and Hare (action 0), which represents not trusting and not rewarding trust.

	Stag	Hare
Stag	1, 1	$S(\chi), T(\chi)$
Hare	$T(\chi), S(\chi)$	0, 0

The payoff when one does not reward the opponent's trust is $T(\chi) \in (0, 1)$, and the payoff when the opponent does not reward one's trust is $S(\chi) < 0$. The state of the world $\chi \in [\underline{\chi}, \bar{\chi}] \subset \mathbb{R}$ measures the strength of enforcement so that S is increasing in χ and T is decreasing in χ . We assume that S and T are bounded on $[\underline{\chi}, \bar{\chi}]$. The payoff difference between Stag and Hare is

$$x_i(\chi; \mathbf{p}_{-i}) = \bar{p}_{-i}(1 - T(\chi)) + (1 - \bar{p}_{-i})S(\chi).$$

This is increasing in χ (since S is increasing and T decreasing in χ). It is also increasing in $\bar{p}_{-i}$ (hence increasing in $\mathbf{p}_{-i}$) since $T(\chi) \in (0, 1)$ and $S(\chi) < 0$. Under the assumption that the stochastic choice rules of all players are moderate, Proposition 4 now implies that people will have a stronger preference against the act of trusting (i.e., stronger betrayal aversion) in environments where being betrayed tends to be more costly because enforcement tends to be weaker, such as in non-kin interactions in the Gulf countries.

7 Extensions and Robustness

7.1 Imperfect Bayesian Foundations

The stochastic choice rule q captures the idea that the DM's cognitive processes, and hence decisions, are noisy. So far, our model has been silent about what causes this noise. In this section, we examine this question from the perspective of (imperfect) Bayesian decision making.

We note that an expected-utility maximiser who has a correct prior and uses Bayes's rule to update her beliefs receives maximum fitness when her tastes are unbiased, i.e., when $\kappa = 0$ (see Appendix F). However, there are many realistic departures from such an ideal Bayesian model, each of which may potentially create scope for $\kappa \neq 0$.

The decision maker may have incorrect priors (Samuelson and Swinkels 2006), be imperfectly able to compute expected values (Benjamin 2019; Enke and Graeber 2023; Gabaix 2014), or suffer from various forms of imperfect Bayesian updating, such as signal neglect (Phillips and Edwards 1966), and base-rate neglect (Kahneman and Tversky 1973). We now examine the last of these imperfections in some detail.

Consider a DM with a correct prior belief about ω (i.e., her prior is equal to the environment μ) who receives a signal about the state of the world, updates her belief, and then picks the optimal action a based on her posterior belief. If the DM were Bayesian, then her posterior belief would depend on her prior and, therefore, so would her SCR. So, in this case, the probability with which the DM takes action 1 should not be only a function of the state ω , but also of the environment μ . Clearly, as we assume that q does not depend on μ , this is not the case for the DMs in our model.

Since the probability with which our DMs take action 1 depends only on the objective-payoff difference and not on the environment, these DMs react only to the information they receive from the signal, but not on their prior. This means that DMs in our model suffer from an extreme case of base-rate neglect (Grether 1980; Kahneman and Tversky 1973) where the prior does not affect their choice probabilities whatsoever.

The assumption (implicit in our SCR) that the DM completely neglects her prior, is made for tractability, not realism. Still, moderate levels of base-rate neglect are quite common (Benjamin 2019). In appendix G, we show numerically that for simple environments with a binary state space and normally-distributed signals, our results still hold for any amount of base-rate neglect. First, generically, expected objective payoff is attained by some non-zero taste κ . Second, as the environment increases in the LR order, the optimal taste increases.

It should also be noted that higher levels of base-rate neglect lead to larger optimal tastes (in absolute value). This is intuitive: as the DM makes fewer mistakes in her assessments, the rectifying role of adapted preferences diminishes. This leads to a zero optimal taste for

the case of a Bayesian decision maker (who doesn't suffer from base-rate neglect).

7.2 Adaptation at the cognitive level

In the analysis of our model, we assumed that the cognitive capacity of the DM—summarised by q —is not subject to adaptation. That is, we assumed that evolution only acts on tastes, but not on cognition. Here, we explore the possibility that evolution and adaptation can affect both the DM's tastes and cognition. If such adaptation at the cognitive level is costless, then the first-best can be achieved by the DM using the step-function $\mathbb{1}_{x(\omega) \geq 0}$ as her stochastic choice rule. However, costless adjustment is unrealistic: we assume that both taste and cognition can be improved at a cost.

To operationalise the question, we assume that the DM has a logit stochastic choice rule given by $q_\beta(y) = (1 + \exp(-\beta y))^{-1}$ for some parameter $\beta \geq 0$, which can be increased at a cost. Similarly, we assume that developing a taste κ is also costly and that both costs are convex, specifically quadratic. Making explicit the dependence of the objective function on the SCR, we write $U(\kappa; F, q_\beta)$. With this, in each environment F a process of evolution or adaptation selects β and κ so as to maximise

$$U(\kappa; F, q_\beta) - \gamma\beta^2 - \xi\kappa^2. \quad (10)$$

We present the optimisation results for a family of binary environments (where $\Omega = \{\omega_1, \omega_2\}$ and $x(\omega_1) = 1$, $x(\omega_2) = -1$) and different parameters γ, ξ in fig. 3. The figures show that even when the SCR q can be improved at a cost, there is still room for a taste adjustment to offer further improvement. Moreover, at least in our specific example, the comparative statics follow the same direction as in Proposition 1: as the probability $p(\omega_1)$ increases, the optimal taste increases, even though the optimal precision parameter changes with the environment and doesn't follow a consistent pattern.

8 Discussion

8.1 Lab and Field

Our model provides a perspective on the relationship between behaviour that is observed in the field and behaviour in laboratory settings. According to one line of reasoning, which we may call the maladaptation-selfishness hypothesis, prosocial behaviour in the lab should not be taken as an indication of social preferences, but be interpreted as mistakes made by agents who are selfishly motivated (Binmore 2006; Hoffman, McCabe, and Smith 1996;

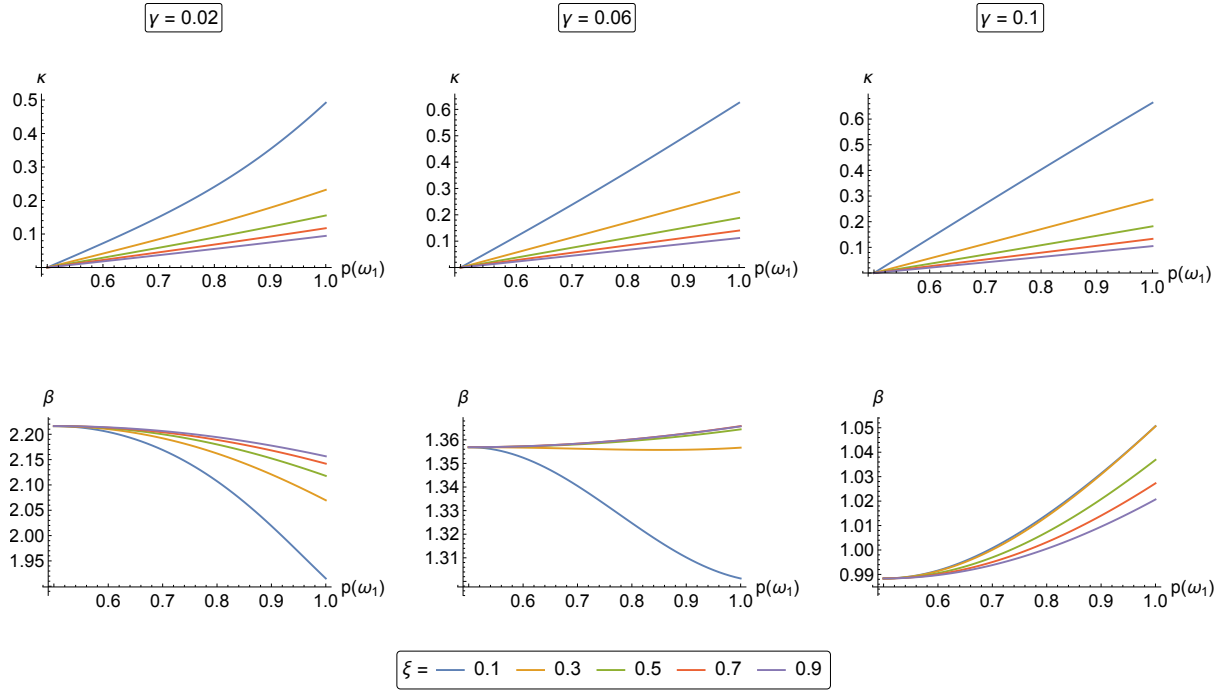


Figure 3: Optimal taste κ and logit precision β in various binary environments when the DM can adjust taste and logit precision at cost $\gamma\beta^2 + \xi|\kappa|^2$. The two states of the world $\{\omega_1, \omega_2\}$ have objective payoff differences $x(\omega_1) = 1$ and $x(\omega_2) = -1$. Environments vary in their probability $p(\omega_1)$.

Samuelson 2005).²⁶ The argument goes as follows: real-world interactions are typically non-anonymous and open-ended, making prosocial behaviours optimal even from a selfish point of view (due to reputation, contagion, and other repeated-game strategies). When people find themselves in the artificial environment of an anonymous one-shot experimental game, they search for analogies from the real world and hence end up behaving as if they were in a repeated non-anonymous interaction.²⁷

According to a contrasting view, which we may call the adaptation-sociality hypothesis, prosocial behaviour in experimental games reflects genuine social preferences. Moreover, it is argued that our evolutionarily relevant past did contain many short-term interactions where selfishness would dictate a lack of prosocial behaviours, but prosocial preferences still evolved, e.g., due to some gene-culture co-evolutionary version of group selection (Boyd,

²⁶A related objection that applies specifically to the ultimatum game is that there is a continuum of Nash equilibria featuring rejection, and it may be difficult to learn to play the unique subgame perfect equilibrium in which there is no rejection, see Andreoni and Blanchard (2006)

²⁷Binmore (2006) even questions the relevance of ascribing preferences to subjects who haven't had time to learn: "My own view is that we waste our time trying to work out what utility function inexperienced subjects are maximizing, because it is not useful to model them as maximizing anything at all. Economics is not the answer to everything, because we do not automatically behave rationally when confronted with novel problems. Insofar as we ever behave rationally, it is largely because we have the capacity to learn."

Gintis, et al. 2003; Fehr and Henrich 2003; Henrich 2004).

Our model takes the middle road by formalising a maladaptation-sociality hypothesis: prosocial behaviour in simple experimental games reflects genuine pro-social preferences, at least partly. Still, these preferences are not adapted to anonymous one-shot interactions but may be viewed as spillovers from preference adaptations to environments where the presence of non-anonymous and open-ended interactions allow prosocial preferences to evolve. The preferences in question function as useful heuristics in repeated interactions but not in one-shot lab experiments.

8.2 Preferences over Actions

We focus on preferences over actions and strategies rather than preferences over consequences of actions, i.e., outcomes. First, this allows us to construct a tractable model. Second, preferences over actions seem realistic in many cases, particularly in the context of social norms. We care about the consequences of our actions. But sometimes we also care about the actions per se regardless of their consequences, as discussed in Section 6.1 in relation to lying aversion and social norm violations. Explaining an action in terms of preferences for the action may invite the worry that “anything goes”. We believe that this worry should be handled by rigorous empirical testing of different preference assumptions (as in Abeler, Nosenzo, and Raymond 2019), rather than by an outright rejection of a whole class of plausible preferences.

8.3 Relation to Rational Inattention Models

In models of rational inattention (RI) (Matějka and McKay 2015; Sims 2003), a DM faces a distribution μ over states and chooses an information structure in order to maximise her expected payoff minus information cost. Each such information structure induces a joint distribution over states and actions. Its cost is typically assumed to be linear in the Shannon mutual information between the action and the state.

Let λ be the linear-cost parameter and $u(a, \omega)$ be the utility of choice alternative a at state ω . With two choice alternatives ($A = \{0, 1\}$), the solution implies that the DM optimally adopts a logistic choice rule (Matějka and McKay 2015):

$$r^{\text{RI}}(\omega) := \Pr(a = 1|\omega) = \left(1 + \exp\left(\frac{-(u(1, \omega) - u(0, \omega) + \alpha_1 - \alpha_0)}{\lambda}\right)\right)^{-1}, \quad (11)$$

where

$$\alpha_i := \lambda \log \left(\int_{\omega} \Pr(a_i|\omega) \mu(d\omega) \right) \quad \text{for } i = 1, 2. \quad (12)$$

Notice that the term in eq. (11) includes the state-dependent utility difference between the two actions ($u(1, \omega) - u(0, \omega) = y(\omega)$; see section 3). So, if the DM has materialistic preferences (i.e., $u(a, \omega) = \pi(a, \omega)$ and $y(\omega) = x(\omega)$), the solution can be re-written as

$$r^{\text{RI}}(\omega) = \left(1 + \exp\left(-\frac{x(\omega) + \alpha_1 - \alpha_0}{\lambda}\right) \right)^{-1}. \quad (13)$$

Now, compare the behaviour in the RI model to that in our adapted-preference model for a DM who has a logistic SCR with imprecision λ , i.e., $q(y) = (1 + \exp(-y/\lambda))^{-1}$. Since in our model the DM's utility is $u(a, \omega) = \pi(a, \omega) + \kappa$ for some κ , behaviour in our model is

$$r^{\text{AP}}(\omega) = \left(1 + \exp\left(-\frac{x(\omega) + \kappa}{\lambda}\right) \right)^{-1}, \quad (14)$$

with κ maximising $U(\kappa; \mu)$ (see eq. (3)).

Both models imply stochastic choice that is logistic with the same imprecision parameter λ . However, the “shifting” factors κ and $\alpha_1 - \alpha_0$ are solutions to different problems: in RI, $\alpha_1 - \alpha_0$ results from a fixed-point calculation, whereas in our model, κ results from objective payoff maximisation. Therefore, behaviour in the two models is generically different. We provide a longer discussion on this point in appendix E.

9 Conclusion

It is a natural idea that people prefer actions and strategies than tend to be beneficial for them. We cannot not claim any originality in this respect. However, we provide a precise formalisation of a particular account of why this is the case. Our explanation is centred on the role of preferences in mitigating the cost of mistakes. As we have shown, this allows us to build a framework that understands cross-cultural differences in preferences as the result of underlying differences in incentives, which are generated by different forms of social organisation. This sets us apart from theories of cultural variation that are based on models of vertical transmission, according to which parents inculcate children to prefer what the parents want them to prefer (Bisin and Verdier 2001; Tabellini 2008). It also distinguishes us from the literature on preference evolution in strategic interactions where preferences may prevail because of their effect on others, e.g., because they provide a commitment advantage (Frank 1987). In future work we wish to test the empirical implications of our theory and relate them to alternative theories.

Appendices

A Environments: from distributions over states to distributions over objective-payoff differences

Our aim is to conduct comparative statics with respect to probability distributions over the state space Ω . For this, we impose some additional structure on our spaces. We equip Ω with an arbitrary σ -algebra, forming the measurable space (Ω, Σ) . Similarly, the real line forms the measurable space $(\mathbb{R}, B_{\mathbb{R}})$ with the Borel σ -algebra. We now assume that the objective-payoff-difference function $x : (\Omega, \Sigma) \rightarrow (\mathbb{R}, B_{\mathbb{R}})$ is measurable. Finally, we equip (Ω, Σ) with the σ -finite reference measure $\tilde{\mu}$ and $(\mathbb{R}, B_{\mathbb{R}})$ with the pushforward reference measure $\tilde{\mu} \circ x^{-1}$.

An *environment* is a probability measure μ on Ω , absolutely continuous with respect to $\tilde{\mu}$. For any environment μ , the resulting distribution of objective-payoff differences will have a cumulative distribution function (CDF) that can be calculated in a straightforward manner

$$F_{\mu}(x) = \mu(\{\omega \in \Omega : x(\omega) \leq x\})$$

and a density function f_{μ} (with respect to $\tilde{\mu} \circ x^{-1}$). In our individual-choice framework, the only feature of states that matters for the DM's optimal taste κ^* is the objective-payoff difference they induce. To simplify the language and the analysis, we refer to these distributions F as environments, rather than the underlying distributions over states.

We use a similar construction for the strategic-interaction part. The state space $\Theta \subseteq \mathbb{R}$ is endowed with the Borel σ -algebra. Then, all functions defined on Θ (e.g., all behaviours p_i) are assumed to be measurable.

B Intuition behind Proposition 1 and the moderate- q condition

Here, we provide some intuition behind Proposition 1 and on why the moderate- q condition is needed to obtain monotone comparative statics for FOSD environment increases that are not LR increases. To better illustrate the intuition, we use binary environments in which state $x^+ > 0$ occurs with probability r and state $-x^- < 0$ occurs with probability $1-r$. Such environments are depicted in the examples of fig. 4, to which we will refer throughout.

The DM makes suboptimal choices with positive probability. These decision errors are of two types (see fig. 4). First, *commission errors* happen if the DM acts (takes action 1)

when she should not. They occur with probability $\Pr[a = 1 | -x^-]$ when $x = -x^-$ and each of them costs the DM x^- . Second, *omission errors* happen if the DM does not act when she should. They occur with probability $1 - \Pr[a = 1 | x^+]$ when the state is x^+ and each costs the DM x^+ .

B.1 LR increases

Let the DM be equipped with the optimal taste κ_1 for her initial environment F_1 ($\kappa_1 \in \psi(F_1)$). This taste equates the marginal benefits from a κ increase to the marginal losses.²⁸ Now, consider a change of the environment to $F_2 \geq_{LR} F_1$, for example an increase of the probability of the high state to $r_2 > r_1$ (fig. 4a middle panel). In the new environment, lower objective-payoff differences occur less frequently and higher objective-payoff differences more frequently, leading to an increase of (expected) losses due to omission mistakes and a decrease of losses due to commission mistakes (compared to F_1). A marginal increase of the DM's taste towards acting (action 1) leads to fewer omission errors and more commission errors, which is profitable at the margin. Therefore, to maximise payoff, optimal tastes in F_2 should be more favourable towards action 1 than those in F_1 ($\kappa_2 > \kappa_1$). This is in accordance with the Le Châtelier-Samuelson principle, which states that, after a change in an exogenous variable, (equilibrium) adjustments should be in the direction that counteract the effects of that change.

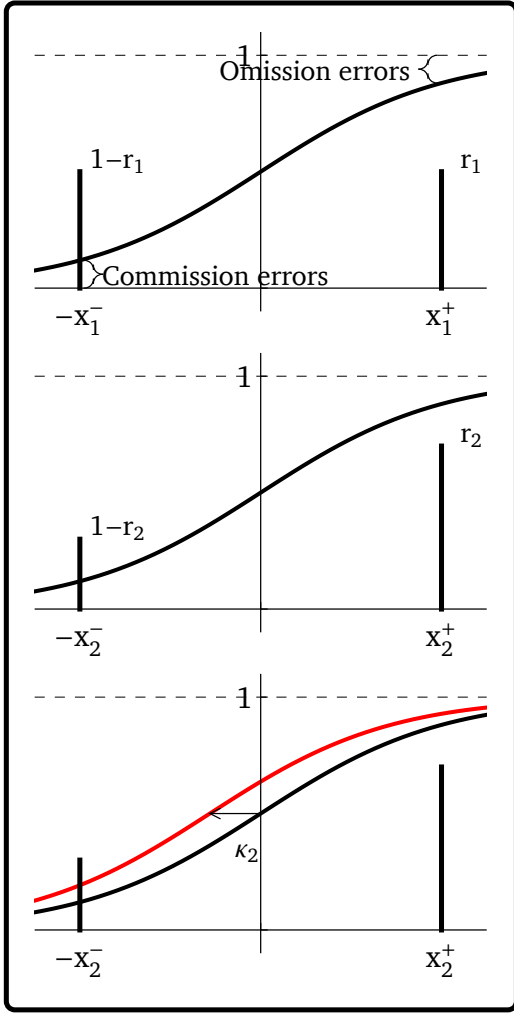
B.2 Non-LR FOSD increases

The generality of this intuition does not carry through if the environment only increases in the FOSD sense and not in the LR sense. Consider such an example where in the higher environment F_2 the cost of the low state decreases to $x_2^- < x_1^-$ (see fig. 4b middle panel). Clearly, in the new environment F_2 , each commission mistake is less costly compared to F_1 (because $x_2^- < x_1^-$). At the same time, though, a marginal increase of the taste κ leads to a much larger increase in the frequency of commission mistakes in the new environment F_2 compared to the lower environment F_1 (because $q'(-x_2^-)$ is much larger than $q'(-x_1^-)$).²⁹

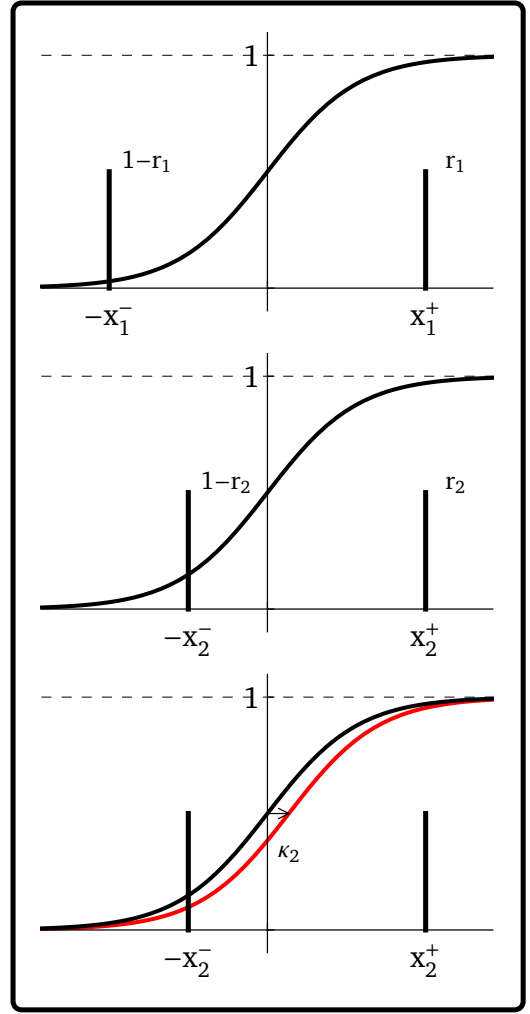
In the example of fig. 4b, the increase in frequency outcales the decrease in the cost of each commission mistake and the combined effect is that a marginal increase of the taste κ leads to an *increase* of losses due to commission mistakes and to a *decrease* of the optimal

²⁸To avoid unnecessary complications, in the examples of this section the optimal taste is unique, in the interior of K , and a continuous function of both r and x^- .

²⁹Of course, an increase in κ also leads to a reduction of omission mistakes. Since, in our example, losses (and also marginal losses) from omission mistakes are very small in the first place, we treat any gains from reducing them as negligible to gain sharper intuition.



(a) Top panel: The DM has optimal taste ($\kappa_1 = 0$) for her environment (F_1). Middle panel: The environment LR-increases to F_2 . States where omission errors occur become more frequent and states where commission errors occur become less frequent. So, marginal expected losses from omission errors are higher. Bottom panel: The new optimal taste κ_2 is higher (i.e., it “shifts” the SCR to the left), as this reduces the probability of omission errors and increases that of commission errors.



(b) Top panel: The DM has optimal taste ($\kappa_1 = 0$) for her environment (F_1). Middle panel: The environment FOSD- (but not LR-) increases to F_2 . Commission errors become less costly, but an increase in κ also makes them considerably more likely. Overall, marginal expected losses from commission errors are higher. Bottom panel: As the stochastic choice rule q is not moderate, monotone comparative statics are not guaranteed. Here, the optimal taste decreases ($\kappa_2 < \kappa_1$).

Figure 4: Two examples of taste adjustment. Environments are binary random variables where $x = x^+ > 0$ with probability r and $x = -x^- < 0$ with probability $1 - r$.

taste ($\kappa_2 < \kappa_1$). So, in this example, we are not able to obtain monotone comparative statics for an FOSD increase in the environment. This is because the SCR q was not moderate.

B.3 The moderate- q condition

If q is moderate (Definition 3), then any changes in mistake frequency caused by the shifting of probability mass from lower to higher states cannot be too large. In particular, these changes are guaranteed to be moderate enough so that a marginal increase in κ always leads to an overall decrease of the losses due to omission errors (or an overall increase of the losses due to commission errors). So, under the moderate- q condition, the intuition of Appendix B.1 is restored for any FOSD environment increase and we obtain monotone comparative statics (items 2–3 of Proposition 1).

C Proofs

C.1 Proof of Proposition 1

Item 1: First, we show that $h : X \times K \rightarrow \mathbb{R}$ defined through $h(x, \kappa) := xq(x + \kappa)$ has single-crossing differences in $(\kappa; x)$, i.e., that for any $\kappa_2 > \kappa_1$, and any $x_2 > x_1$, we have

$$h(x_1, \kappa_2) - h(x_1, \kappa_1) > (\geq) 0 \Rightarrow h(x_2, \kappa_2) - h(x_2, \kappa_1) > (\geq) 0.$$

Indeed, pick $x_1, x_2 \in X$ and $\kappa_1, \kappa_2 \in K$ such that $\kappa_2 > \kappa_1$ and $x_2 > x_1$ and assume that $h(x_1, \kappa_2) - h(x_1, \kappa_1) > (\geq) 0$, i.e., that $x_1(q(x_1 + \kappa_2) - q(x_1 + \kappa_1)) > (\geq) 0$. This means that $x_1 > (\geq) 0$, since q is strictly increasing. Also because q is strictly increasing, we have that $q(x_2 + \kappa_2) - q(x_2 + \kappa_1) > 0$ and, since $x_2 > x_1 > (\geq) 0$, we get that $h(x_2, \kappa_2) - h(x_2, \kappa_1) = x_2(q(x_2 + \kappa_2) - q(x_2 + \kappa_1)) > 0$.

Since h exhibits single-crossing differences and $F_2 \geq_{LR} F_1$, it follows from Theorem 2 of Athey (2002) that

$$\psi(F_2) = \arg \max_{\kappa \in K} \int h(x, \kappa) dF_2(x) \geq_{SSO} \arg \max_{\kappa \in K} \int h(x, \kappa) dF_1(x) = \psi(F_1),$$

where the ranking is in the strong set order. That is, for any $\kappa_1 \in \psi(F_1)$ and any $\kappa_2 \in \psi(F_2)$, $\kappa_2 \geq \kappa_1$.

Item 2: Let $\kappa_1, \kappa_2 \in K$ with $\kappa_2 > \kappa_1$. As q is moderate, $xq(x + \kappa)$ exhibits increasing differences in $(x; \kappa)$ and, therefore, $x(q(x + \kappa_2) - q(x + \kappa_1))$ is nondecreasing in x . Now, take two environments F_1, F_2 with $F_2 \geq_{FOSD} F_1$. We have that $F_2 >_{FOSD} F_1$ and, as $x(q(x +$

$\kappa_2) - q(x + \kappa_1))$ is an nondecreasing function of x , we get that

$$\begin{aligned} U(\kappa_2; F_2) - U(\kappa_1; F_2) &= \int x(q(x + \kappa_2) - q(x + \kappa_1)) dF_2(x) \\ &\geq \int x(q(x + \kappa_2) - q(x + \kappa_1)) dF_1(x) \\ &= U(\kappa_2; F_1) - U(\kappa_1; F_1). \end{aligned}$$

That is, $U(\kappa; F)$ exhibits increasing differences in $(\kappa; F)$ (in the $\geq_{FOSD}$ order). Since $\psi(F)$ is nonempty and closed, it has a maximum and a minimum element for any environment F . So, the conditions of Theorem 3.1 (ii) of Vives (1990) are satisfied and both $\min \psi(F)$ and $\max \psi(F)$ are nondecreasing functions of F .

Item 3: Following a similar procedure as in Item 2, with *strictly* moderate q , we can conclude that $U(\kappa; F)$ exhibits *strictly* increasing differences in $(\kappa; F)$. Additionally, as K is a closed interval of $\mathbb{R}$, it is a lattice under the usual $\geq$ order and so is $(\mathcal{F}, \geq_{FOSD})$. Finally, $U(\cdot; F)$ is supermodular on K for each $F \in \mathcal{F}$, since it is modular. It follows from Theorem 3.1 (iv) of Vives (1990) (see also Topkis 1978) that ψ is nondecreasing in the strong set order, i.e., that $\min \psi(F_2) \geq \max \psi(F_1)$. $\square$

C.2 Proof of Proposition 4

Existence of taste equilibria Fix an environment $\nu \in \mathcal{N}$. From Proposition 2, $\mathcal{B}(\kappa)$ has a maximum element for any $\kappa \in K$. Moreover, for any player $i \in N$ and any profile $\mathbf{p}_{-i} \in P^{n-1}$, $\psi_i(\mathbf{p}_{-i}; \nu)$ is the set of maximisers over the compact set K_i of $U_i(\kappa_i; \mathbf{p}_{-i}; \nu)$, which is continuous in κ_i . So, $\psi_i(\mathbf{p}_{-i}; \nu)$ is a compact subset of $\mathbb{R}$ and has a maximum element. Consequently, $\psi(\mathbf{p}; \nu)$ has a maximum element (the coordinate-wise maximum) for any $\mathbf{p} \in \mathcal{P}$.

Now, consider the function $\bar{\chi} : K \times \mathcal{N} \rightarrow K$ defined through

$$\bar{\chi}(\kappa; \nu) := \max \psi(\max \mathcal{B}(\kappa); \nu).$$

Because of the preceding analysis, $\bar{\chi}(\kappa; \nu)$ is well defined. Additionally, $\max \psi(\mathbf{p}; \nu)$ is nondecreasing in $\mathbf{p}$ (Proposition 3, item 1) and $\max \mathcal{B}(\kappa)$ is strictly increasing in κ (Proposition 2). So, $\bar{\chi}(\kappa; \nu)$ is nondecreasing in κ over K , which is a complete lattice. Therefore, Tarski's fixed-point theorem applies and the set of fixed points of $\bar{\chi}(\cdot; \nu)$ is nonempty (as it is a complete lattice) for any given environment ν .

Let κ^* be a fixed point of $\bar{\chi}(\cdot; \nu)$. Define $\mathbf{p}^* := \max \mathcal{B}(\kappa^*)$ (i.e., $\mathbf{p}^*$ is the largest behaviour equilibrium of κ^*), which implies that $\mathbf{p}^* = \phi(\mathbf{p}^*; \kappa^*)$. Moreover, since κ^* is a fixed point of $\bar{\chi}(\cdot; \nu)$, we have that: $\kappa^* = \max \psi(\mathbf{p}^*; \nu)$ and, thus, $\kappa^* \in \psi(\mathbf{p}^*; \nu)$. Therefore, $(\mathbf{p}^*, \kappa^*) \in \mathcal{T}(\nu)$ and $\mathcal{T}(\nu)$ is nonempty.

Existence of extremal taste equilibria As the set of fixed points of $\overline{\chi}(\cdot; \nu)$ is a complete lattice, it has a largest element $\overline{\kappa}(\nu)$. From the definition of $\overline{\chi}$, it is clear that $(\max \mathcal{B}(\overline{\kappa}(\nu)), \overline{\kappa}(\nu)) \in \mathcal{T}(\nu)$. We will show that $(\max \mathcal{B}(\overline{\kappa}(\nu)), \overline{\kappa}(\nu)) = \max \mathcal{T}(\nu)$.

Take any $(\tilde{p}, \tilde{\kappa}) \in \mathcal{T}(\nu)$. By definition of taste equilibrium, we have that $\tilde{p} \in \mathcal{B}(\tilde{\kappa})$ and $\tilde{\kappa} \in \psi(\tilde{p}; \nu)$. So, $\max \mathcal{B}(\tilde{\kappa}) \geq \tilde{p}$. Further, each q_i is moderate and therefore $\max \psi(\cdot; \nu)$ is nondecreasing (Proposition 3, item 1). So, $\overline{\chi}(\tilde{\kappa}; \nu) = \max \psi(\max \mathcal{B}(\tilde{\kappa}); \nu) \geq \max \psi(\tilde{p}; \nu) \geq \tilde{\kappa}$. From Theorem 3 of Milgrom and Roberts (1994), we get that $\overline{\kappa}(\nu) = \sup\{\kappa | \overline{\chi}(\kappa; \nu) \geq \kappa\}$. Therefore, $\overline{\kappa}(\nu) \geq \tilde{\kappa}$. Finally, because $\max \mathcal{B}(\kappa)$ is strictly increasing in κ , $\max \mathcal{B}(\overline{\kappa}(\nu)) \geq \max \mathcal{B}(\tilde{\kappa}) \geq \tilde{p}$. So, $(\max \mathcal{B}(\overline{\kappa}(\nu)), \overline{\kappa}(\nu)) \geq (\tilde{p}, \tilde{\kappa})$ for any $(\tilde{p}, \tilde{\kappa}) \in \mathcal{T}(\nu)$.

A similar analysis holds for the smallest taste equilibrium, if, instead of $\overline{\chi}$, one uses the function $\underline{\chi} : K \times \mathcal{N} \rightarrow K$, defined through $\underline{\chi}(\kappa; \nu) := \min \psi(\min \mathcal{B}(\kappa); \nu)$.

Nondecreasing extremal taste equilibria We note that $\max \psi(p; \nu)$ is nondecreasing in ν (Proposition 1) and that $\mathcal{B}(\kappa)$ does not depend on the environment ν . Therefore, $\overline{\chi}(\kappa; \nu)$ is nondecreasing in ν . We again use Theorem 3 of Milgrom and Roberts (1994), which guarantees that $\overline{\kappa}(\nu)$ is nondecreasing in ν . Similarly for the smallest taste equilibrium. $\square$

C.3 Proof of Proposition 5

We first show that $\mathcal{S}(\kappa, \kappa', p; r)$ is an nondecreasing sequence. That is, we will show that for any period $T > 0$ and all periods $t \in \{1, 2, \dots, T\}$, we have $p^t \geq p^{t-1}$. We do this by induction.

Step 1: Base case

We verify that the statement holds for $T = 1$, i.e., that $p^1 \geq p^0$.

Indeed, the only player whose behaviour may change in period 1 is player $r(1)$. Since $p \in \mathcal{B}(\kappa)$, we know that

$$p_{r(1)}^0 = \phi_{r(1)}(p_{-r(1)}^0; \kappa_{r(1)}).$$

In period 1, player $r(1)$ revises her behaviour to

$$p_{r(1)}^1 = \phi_{r(1)}(p_{-r(1)}^0; \kappa'_{r(1)}).$$

Note that $\phi_{r(1)}(p_{-r(1)}; \kappa_{r(1)})$ is strictly increasing in $\kappa_{r(1)}$ and that $\kappa'_{r(1)} \geq \kappa_{r(1)}$. It follows that $p_{r(1)}^1 \geq p_{r(1)}^0$ and so, $p^1 \geq p^0$.

Step 2: Induction step

We assume that the statement holds for $T = k$, i.e., that

$$p^t \geq p^{t-1} \quad \text{for all } t \in \{1, 2, \dots, k\}. \quad (15)$$

We will show that the statement holds for $T = k + 1$, i.e., that $\mathbf{p}^t \geq \mathbf{p}^{t-1}$ for all $t \in \{1, 2, \dots, k + 1\}$.

The only player whose behaviour may change in period $k + 1$ is player $r(k + 1)$, who revises her behaviour to

$$p_{r(k+1)}^{k+1} = \phi_{r(k+1)}(\mathbf{p}_{-r(k+1)}^k; \kappa'_{r(k+1)}).$$

Let

$$t' = \max(\{t \in \mathbb{N}_+ : (t < k + 1) \wedge (r(t) = r(k + 1))\} \cup \{0\})$$

be the last period in which player $r(k + 1)$ revised her behaviour (or period 0, if $k + 1$ is the first period player $r(k + 1)$ gets a revision opportunity). Since player $r(k + 1)$ hasn't revised her behaviour since t' , we know that

$$p_{r(k+1)}^k = p_{r(k+1)}^{t'} = \begin{cases} \phi_{r(k+1)}(\mathbf{p}_{-r(k+1)}^{t'-1}; \kappa'_{r(k+1)}) & \text{if } t' > 0 \\ \phi_{r(k+1)}(\mathbf{p}_{-r(k+1)}^0; \kappa_{r(k+1)}) & \text{if } t' = 0 \end{cases}.$$

Additionally, $\phi_{r(k+1)}(\mathbf{p}_{-r(k+1)}; \kappa_{r(k+1)})$ is nondecreasing in $\mathbf{p}_{-r(k+1)}$ and strictly increasing in $\kappa_{r(k+1)}$. Now, from the premise of the proposition we get that $\kappa' \geq \kappa$ and from the assumption in Step 2 (eq. (15)), we have that $\mathbf{p}^k \geq \mathbf{p}^0$ and if $t' > 0$ that $\mathbf{p}^k \geq \mathbf{p}^{t'-1}$. Thus, $p_{r(k+1)}^k \geq p_{r(k+1)}^{k-1}$ and, since no other player gets a revision opportunity in period $k + 1$, we also get that $\mathbf{p}^k \geq \mathbf{p}^{k-1}$. This together with eq. (15), yields

$$\mathbf{p}^t \geq \mathbf{p}^{t-1} \quad \text{for all } t \leq k + 1.$$

So, we have shown that for any $T > 0$ and for all $t \leq T$, $\mathbf{p}^t \geq \mathbf{p}^{t-1}$. That is, the sequence $\mathcal{S}(\kappa, \kappa', \mathbf{p}; r)$ is nondecreasing.

Now we show that the sequence $\mathcal{S}(\kappa, \kappa', \mathbf{p}; r)$ converges. Fix a state $\theta \in \Theta$ and a player $i \in N$ and consider the sequence $(p_i^t(\theta))_{t \in \mathbb{N}_+}$. Since this sequence is nondecreasing and because $[0, 1]$ is compact, the sequence converges to a unique limit $p_i^*(\theta)$. So, for each player i $p_i^t \rightarrow p_i^*$ and the whole sequence of behaviour profile converges, i.e., $\mathbf{p}^t \rightarrow \mathbf{p}^*$.

Finally, we show that the limit point $\mathbf{p}^*$ is a behaviour equilibrium of κ' . Fix a player i and define the sequence $R(i) := (t \in \mathbb{N}_+ : r(t) = i)$. Due to the properties of $r(\cdot)$, $R(i)$ is an infinite sequence for any $i \in N$. Now, consider the sequence of player i 's behaviours $(p_i^\tau)_{\tau \in R(i)}$. For this, we have that $p_i^\tau = \phi_i(\mathbf{p}_{-i}^{\tau-1}; \kappa_i)$ and so, $p_i^* = \lim_{\tau \rightarrow \infty} p_i^\tau = \lim_{\tau \rightarrow \infty} \phi_i(\mathbf{p}_{-i}^{\tau-1}; \kappa_i)$. Since q_i is continuous, we get that $p_i^* = \phi_i(\mathbf{p}_{-i}^*; \kappa_i')$. As the player i was picked arbitrarily, we get that $p_i^* = \phi_i(\mathbf{p}_{-i}^*; \kappa_i')$ for all $i \in N$, i.e., that $\mathbf{p}^* \in \mathcal{B}(\kappa')$. $\square$

D Evolution of responders' taste for rejection in ultimatum games

Here we provide a simple illustration of how preference for rejection in the ultimatum game may develop in an environment where there is a positive probability that information about the rejection decisions are transmitted to future bargaining partners. Consider the following (simultaneous-move) mini-Ultimatum Game (UG),

	<i>A</i>	<i>R</i>
<i>U</i>	$\frac{3}{4}, \frac{1}{4}$	0, 0
<i>F</i>	$\frac{1}{2}, \frac{1}{2}$	$\frac{1}{2}, \frac{1}{2}$

Agents from the population interact recurrently over an indefinite number of rounds. In each round agents are randomly matched to play the mini-UG twice, once as responder and once as proposer, each time with a different opponent. Before playing the game, the proposer receives a signal about the responder's choice as responder facing an unfair offer in a previous interaction. With probability $\rho \in (0, 1]$ the signal is equal to the action that the current responder actually took in the previous interaction, and with probability $1 - \rho$ the signal is instead equal to the action that the responder did not take. The probability of a correct signal ρ corresponds to the state of nature in our framework. Thus, a higher environment implies that information about previous interactions is more accurate.

As before, we consider only two strategies. In the proposer role both strategies prescribe the same behaviour, namely to play *U* in the first round and in later rounds play *U* if and only if the signal about past behaviour is *A*. The strategies differ for the responder role. Strategy 1 always rejects (AR) and strategy 0 always accepts (AA). Using the fact that there is no difference in proposer behaviour, it can be verified that the expected objective payoff difference between AR and AA is

$$x_i(\rho; \mathbf{p}_{-i}) = \frac{1}{2}\delta\rho - \frac{1}{4}(1 - \delta) - \delta \left[\frac{1}{4}\rho + \frac{1}{2}(1 - \rho) \right] = \frac{1}{4}((3\rho - 1)\delta - 1).$$

This is increasing in ρ . According to proposition Proposition 1, preferences for rejecting unfair offers should be stronger when decisions to reject are more clearly visible to other people.

E Comparing Adapted Preferences to Rational Inattention

Consider two decision makers. One has a logistic stochastic choice rule $q(y) = (1 + \exp(-y/\lambda))^{-1}$ for a given parameter $\lambda > 0$ and develops an adapted preference κ , as in the

main text of this paper. The second decision maker is rationally inattentive and faces information costs that are linear in Shannon-entropy reduction with the same parameter λ . Let the two decision makers be in the same environment μ and make choices between the two actions $a = 1, 0$. Finally, let us denote the set of all stochastic choice rules as $\mathcal{R} := [0, 1]^\Omega$, the logistic choice rule with imprecision parameter λ and shift parameter κ as

$$r_{\kappa, \lambda}(\omega) := \frac{1}{1 + \exp\left(-\frac{x(\omega) + \kappa}{\lambda}\right)},$$

the set of all logistic choice rules with imprecision $\lambda > 0$ as $\mathcal{R}_\lambda := \{r_{\kappa, \lambda} : \kappa \in \mathbb{R}\}$, and the expected objective payoff for rule r in environment μ as $U(r, \mu) := \int x(\omega)r(x(\omega))\mu(d\omega)$.

The DM with adapted preferences develops a taste κ^* that solves $\max_{\kappa \in \mathbb{K}} U(r_{\kappa, \lambda}, \mu)$, leading to the stochastic choice rule $r^{\text{AP}} := r_{\kappa^*, \lambda}$. For simplicity, let's assume that there is a unique such taste. The rationally inattentive DM adopts the stochastic choice rule r^{RI} that solves

$$\max_{r \in \mathcal{R}} U(r, \mu) - I(r, \mu),$$

where $I(r, \mu)$ denotes the mutual information between the action a and payoff difference x in environment μ when the DM is using choice rule r .

We observe that, generically, if κ^* is interior, the DM with adapted preferences achieves higher expected objective payoff than the rationally inattentive DM in the same environment μ . This is true even without taking the information cost of the RI agent into account. The reason is the following.

First, we note that the stochastic choice rule $r^{\text{RI}} \in \mathcal{R}_\lambda$ (Matějka and McKay 2015, see also section 8.3 in the main text). At the same time, as r^{AP} is the choice rule among $\mathcal{R}_\lambda$ that yields the highest expected objective payoff, we have that $U(r^{\text{AP}}, \mu) \geq U(r^{\text{RI}}, \mu)$, i.e., the highest expected objective payoff the rationally inattentive DM can achieve is no higher than that of the DM with adapted preferences.

Second, we will show that the rationally inattentive DM's shifting parameter is generically different from κ^* . As the adapted-preference DM's optimal taste κ^* is interior, we have that

$$\frac{d}{d\kappa} U(r_{\kappa, \lambda}, \mu) \Big|_{\kappa=\kappa^*} = 0.$$

Thus, making a small change $\Delta\kappa$ in the taste from κ^* to $\kappa^* + \Delta\kappa$ has no first-order effect on the expected objective payoff for the DM with adapted preferences. The same change, though, generically has a non-zero effect on the rationally inattentive DM's information costs (generically, $\frac{d}{d\kappa} I(r_{\kappa, \lambda}, \mu) \Big|_{\kappa=\kappa^*} \neq 0$), meaning that the rationally inattentive DM is generically strictly better off using another choice rule within the class $\mathcal{R}_\lambda$ with a slightly different shifting parameter $\tilde{\kappa}$ than using κ^* :

$$U(r_{\tilde{\kappa}, \lambda}, \mu) - I(r_{\tilde{\kappa}, \lambda}, \mu) > U(r_{\kappa^*, \lambda}, \mu) - I(r_{\kappa^*, \lambda}, \mu).$$

Therefore, the rationally inattentive DM's optimal shifting parameter is different than κ^* , meaning that $U(r^{\text{AP}}, \mu) > U(r^{\text{RI}}, \mu)$.

F Bayesians don't Need Tastes

Let the payoff difference x be distributed according to the CDF F and consider an expected-utility-maximising DM who has a correct prior belief and unbiased preferences, i.e., $\kappa = 0$. Before making a decision, the DM receives a message s from an arbitrary message space S and updates her prior, forming the posterior belief $b(\cdot|s) \in \Delta(X)$ about the payoff difference. As the DM maximises expected payoff, she takes action 1 whenever the expected utility difference (which is equal to the expected objective-payoff difference) according to her posterior is non-negative:

$$a^* = \begin{cases} 1 & \text{if } \int_X x b(x|s) dx \geq 0 \\ 0 & \text{otherwise} \end{cases}.$$

It is clear that if $b(\cdot|s)$ is given by Bayes's rule, then the DM receives maximum expected fitness, given the information she has available *after any received message* s . Therefore, no bias $\kappa \neq 0$ can increase her expected fitness and $\kappa = 0$ is a fitness-maximising taste.

This is true for any environment, i.e., any distribution F of the objective-payoff difference. Notice that the DM's stochastic choice rule q in this case will depend on the prior belief F .

G Adapted Tastes and Base-rate Neglect

In the environments that we are exploring in this section, there are two states of the world $\Omega = \{\omega_1, \omega_2\}$. At each state ω_i , the objective-payoff is $x(\omega_i)$, while the state occurs with probability $p(\omega_i)$. After the state realises, the DM receives a signal s that is normally distributed around the objective payoff of the realised state, i.e., $s \sim N(x(\omega_i), \sigma^2)$ for some $\sigma > 0$ when the state is ω_i . We denote the respective probability density of signal realisation s as $f(s|\omega_i)$.

We assume that the DM has a correct prior (i.e., assigns a probability p that the state is ω_1 and a probability p_2 that it is ω_2), is aware of the data-generating process, but has base-rate neglect. This means that her posterior belief $h(\omega_i|s)$ about state ω_i after receiving signal realisation s satisfies (see Bohren and Hauser 2024):

$$\frac{h(\omega_2|s)}{h(\omega_1|s)} = \left(\frac{p(\omega_2)}{p(\omega_1)} \right)^\alpha \frac{f(s|\omega_2)}{f(s|\omega_1)}.$$

The parameter $\alpha \in [0, 1]$ captures the DM's degree of rationality, in terms of complying with the Bayesian ideal: when $\alpha = 1$, the DM is Bayesian and when $\alpha = 0$ she completely neglects her prior, just like the DMs in the main model. Upon observing her signal s , the DM picks the action $a^*(s)$ with the highest expected utility according to her posterior. So,

$$a^*(s) = \begin{cases} 1 & \text{if } \kappa + \sum_{\omega_i \in \Omega} x(\omega_i)h(\omega_i|s) \geq 0 \\ 0 & \text{otherwise} \end{cases}.$$

As in the main text, we are interested in finding the taste κ that maximises expected objective payoff.

For this exercise, we use $x(\omega_1) = 1$, $x(\omega_2) = -1$ and we vary the probabilities p and the rationality parameter α . In Figure 5, we plot the optimal taste κ in the different environments and for different values of α . We make three observations. First, the DM typically

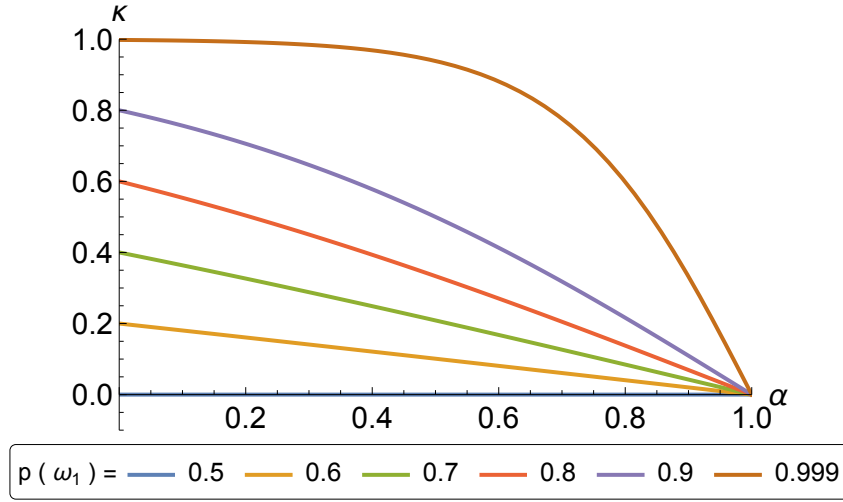


Figure 5: Optimal taste for different levels of base-rate neglect α in binary environments. The two states of the world $\{\omega_1, \omega_2\}$ have objective payoff differences $x(\omega_1) = 1$ and $x(\omega_2) = -1$. Plots for environments with different $p(\omega_1)$.

develops some nonzero taste κ to correct her mistakes. Second, the strength of the optimal taste decreases as the DM becomes more rational (as α increases) and eventually vanishes when she is perfectly Bayesian ($\alpha = 1$). Third, higher environments (those where $p(\omega_1)$ is higher), lead to higher optimal tastes, which is in line with the comparative statics results of Proposition 2. Of course, this latter observation is likely to be dependent on the specific type of environments we have considered.

References

Abeler, J., D. Nosenzo, and C. Raymond (2019). “Preferences for Truth-Telling”. *Econometrica* 87.4, pp. 1115–1153.

- Acemoglu, D. and J. A. Robinson (2021). *Culture, institutions and social equilibria: A framework*. Tech. rep. w28832. National Bureau of Economic Research.
- Alesina, A. and P. Giuliano (2015). “Culture and institutions”. *Journal of Economic Literature* 53.4, pp. 898–944.
- Alger, I. and J. W. Weibull (2013). “Homo Moralis, Preference Evolution under Incomplete Information and Assortative Matching”. *Econometrica* 81.6, pp. 2269–2302.
- Alós-Ferrer, C. (2018). “A review essay on Social neuroscience: Can research on the social brain and economics inform each other?” *Journal of Economic Literature* 56.1, pp. 234–264.
- Alós-Ferrer, C., E. Fehr, and N. Netzer (2021). “Time will tell: Recovering preferences when choices are noisy”. *Journal of Political Economy* 129.6, pp. 1828–1877.
- Anderson, J. R. (1991). “The adaptive nature of human categorisation”. *Psychological Review* 98.3, pp. 409–429.
- Andreoni, J. and E. Blanchard (2006). “Testing subgame perfection apart from fairness in ultimatum games”. *Experimental Economics* 9, pp. 307–321.
- Athey, S. (2002). “Monotone Comparative Statics under Uncertainty”. *The Quarterly Journal of Economics* 117.1, pp. 187–223.
- Banfield, E. C. (1967). *The moral basis of a backward society*. New York: Free Press.
- Becker, G. S. (1996). *Accounting for tastes*. Harvard University Press.
- Benjamin, D. J. (2019). “Chapter 2 - Errors in probabilistic reasoning and judgment biases”. In: *Handbook of Behavioral Economics - Foundations and Applications 2*. Ed. by B. D. Bernheim, S. DellaVigna, and D. Laibson. Vol. 2. Handbook of Behavioral Economics: Applications and Foundations 1. North-Holland, pp. 69–186.
- Bergstrom, T. C. (1995). “On the Evolution of Altruistic Ethical Rules for Siblings”. *American Economic Review* 85.1, pp. 58–81.
- Bicchieri, C. (2005). *The Grammar of Society: The Nature and Dynamics of Social Norms*. Cambridge, MA: Cambridge University Press.
- Bigoni, M. et al. (2019). “At the root of the North-South cooperation gap in Italy: Preferences or beliefs?” *The Economic Journal* 129.619, pp. 1139–1152.
- Binmore, K. (2006). “Why do people cooperate?” *Politics, Philosophy and Economics* 5.1, pp. 81–96.
- Bisin, A. and T. Verdier (2001). “The economics of cultural transmission and the dynamics of preferences”. *Journal of Economic Theory* 97.2, pp. 298–319.
- Bohnet, I., B. Herrmann, and R. Zeckhauser (2010). “Trust and the reference points for trustworthiness in Gulf and Western countries”. *The Quarterly Journal of Economics* 125.2, pp. 811–828.

- Bohren, J. A. and D. Hauser (2024). *Behavioral Foundations of Model Misspecification*. Tech. rep. SSRN.
- Bowles, S. (1998). “Endogenous preferences: The cultural consequences of markets and other economic institutions”. *Journal of Economic Literature* 36.1, pp. 75–111.
- Boyd, R., H. Gintis, et al. (2003). “The evolution of altruistic punishment”. *Proceedings of the National Academy of Sciences* 100.6, pp. 3531–3535.
- Boyd, R. and P. J. Richerson (1988). *Culture and the evolutionary process*. University of Chicago Press.
- Boyd, R. and P. J. Richerson (2009). “Culture and the evolution of human cooperation”. *Philos Trans R Soc B* 364.1533.
- Cassar, A., G. d’Adda, and P. Grosjean (2014). “Institutional quality, culture, and norms of cooperation: Evidence from behavioral field experiments”. *The Journal of Law and Economics* 57.3, pp. 821–863.
- Cavalli-Sforza, L. L. and M. W. Feldman (1981). *Cultural transmission and evolution: A quantitative approach*. 16. Princeton University Press.
- Chudek, M., M. Muthukrishna, and J. Henrich (2015). “Cultural evolution”. In: *The Handbook of Evolutionary Psychology, Volume 2*. Ed. by D. M. Buss. John Wiley & Sons, Ltd. Chap. 30, pp. 749–769.
- Dekel, E., J. C. Ely, and O. Yilankaya (2007). “Evolution of Preferences”. *Review of Economic Studies* 74, pp. 685–704.
- Dow, J. (1991). “Search decisions with limited memory”. *The Review of Economic Studies* 58.1, pp. 1–14.
- Ellingsen, T. and M. Johannesson (2004). “Promises, threats and fairness”. *The Economic Journal* 114.495, pp. 397–420.
- Ellingsen, T. and E. Mohlin (2024). *A model of social duties*. Tech. rep. Working Paper.
- Engl, F., A. Riedl, and R. Weber (2021). “Spillover effects of institutions on cooperative behavior, preferences, and beliefs”. *American Economic Journal: Microeconomics* 13.4, pp. 261–299.
- Enke, B. (2019). “Kinship, cooperation, and the evolution of moral systems”. *The Quarterly Journal of Economics* 134.2, pp. 953–1019.
- Enke, B. and T. Graeber (May 2023). “Cognitive Uncertainty*”. *The Quarterly Journal of Economics* 138.4, pp. 2021–2067.
- Falk, A. et al. (2018). “Global evidence on economic preferences”. *The Quarterly Journal of Economics* 133.4, pp. 1645–1692.
- Fechner, G. T. ([1860] 1948). “Elements of psychophysics”. In: *Readings in the history of psychology*. Appleton-Century-Crofts, pp. 206–213.

- Fehr, E. and J. Henrich (2003). "Is strong reciprocity a maladaptation? On the evolutionary foundations of human altruism". In: *Genetic and cultural evolution of cooperation*. Ed. by P. Hammerstein. MIT Press, pp. 55–82.
- Fischbacher, U. and F. Föllmi-Heusi (2013). "Lies in disguise-an experimental study on cheating". *Journal of the European Economic Association* 11.3, pp. 525–547.
- Fischbacher, U., S. Gächter, and E. Fehr (2001). "Are people conditionally cooperative? Evidence from a public goods experiment". *Economics Letters* 71.3, pp. 397–404.
- Frank, R. H. (1987). "If Homo Economicus Could Choose His Own Utility Function, Would He Want One with a Conscience?" *The American Economic Review* 77.4, pp. 593–604.
- Fudenberg, D., P. Strack, and T. Strzalecki (2018). "Speed, Accuracy, and the Optimal Timing of Choices". *American Economic Review* 108.12, pp. 3651–84.
- Gabaix, X. (2014). "A sparsity-based model of bounded rationality". *The Quarterly Journal of Economics* 129.4, pp. 1661–1710.
- Gächter, S. and J. F. Schulz (2016). "Intrinsic honesty and the prevalence of rule violations across societies". *Nature* 531.7595, pp. 496–499.
- Geertz, C. (1973). *The interpretation of cultures*. Vol. 5019. New York: Basic books.
- Gellner, E. (2000). "Trust, cohesion, and the social order". In: *Trust: Making and breaking cooperative relations*. Ed. by D. Gambetta. Oxford: Department of Sociology, University of Oxford. Chap. 9, pp. 142–157.
- Gigerenzer, G. and P. M. Todd (1999). *Simple heuristics that make us smart*. USA: Oxford University Press.
- Gneezy, U. (2005). "Deception: The role of consequences". *American Economic Review* 95.1, pp. 384–394.
- Gold, J. I. and M. N. Shadlen (2007). "The neural basis of decision making". *Annu. Rev. Neurosci.* 30, pp. 535–574.
- Greif, A. and G. Tabellini (2017). "The clan and the corporation: Sustaining cooperation in China and Europe". *Journal of Comparative Economics* 45.1, pp. 1–35.
- Grether, D. M. (1980). "Bayes rule as a descriptive model: The representativeness heuristic". *The Quarterly journal of economics* 95.3, pp. 537–557.
- Guiso, L., P. Sapienza, and L. Zingales (2006). "Does culture affect economic outcomes?" *Journal of Economic Perspectives* 20.2, pp. 23–48.
- Güth, W. and M. E. Yaari (1992). "Explaining Reciprocal Behavior in Simple Strategic Games: An Evolutionary Approach". In: *Explaining Process and Change*. Ed. by U. Witt. Ann Arbor, MI: University of Michigan Press, pp. 22–34.
- Heifetz, A., C. Shannon, and Y. Spiegel (2007). "What to Maximize if You Must". *Journal of Economic Theory* 133.1, pp. 31–57.

- Heller, Y. and E. Mohlin (2018). “Observations on cooperation”. *Review of Economic Studies* 85.4, pp. 2253–2282.
- Heller, Y. and E. Mohlin (2019). “Coevolution of deception and preferences: Darwin and Nash meet Machiavelli”. *Games and Economic Behavior* 113, pp. 223–247.
- Henrich, J. (2004). “Cultural group selection, coevolutionary processes and large-scale cooperation”. *Journal of Economic Behavior and Organization* 53.1, pp. 3–35.
- Henrich, J. et al., eds. (2004). *Foundations of human sociality: Economic experiments and ethnographic evidence from fifteen small-scale societies*. Oxford University Press.
- Herrmann, B., C. Thöni, and S. Gächter (2008). “Antisocial punishment across societies”. *Science* 319.5868, pp. 1362–1367.
- Hoffman, E., K. McCabe, and V. L. Smith (1996). “Social distance and other-regarding behavior in dictator games”. *American Economic Review* 86.3, pp. 653–660.
- Huck, S. and J. Oechssler (1999). “The Indirect Evolutionary Approach to Explaining Fair Allocations”. *Games and Economic Behavior* 28, pp. 13–24.
- Jehiel, P. and E. Mohlin (2023). *Categorization in Games: A Bias-Variance Perspective*. Tech. rep. Working Paper.
- Jensen, M. K. and A. Rigos (2018). “Evolutionary games and matching rules”. *International Journal of Game Theory* 47.3, pp. 707–735.
- Kahneman, D. and A. Tversky (1973). “On the psychology of prediction”. *Psychological Review* 80.4, p. 237.
- Kandori, M. (1992). “Social Norms and Community Enforcement”. *Review of Economic Studies* 59.1, pp. 63–80.
- Khaldun, I. (2004). *The Muqaddimah: Volume II*. Trans. by F. Rosenthal. Princeton: Princeton University Press.
- Kocher, M. G. et al. (2008). “Conditional cooperation on three continents”. *Economics Letters* 101.3, pp. 175–178.
- Krajbich, I., C. Armel, and A. Rangel (2010). “Visual fixations and the computation and comparison of value in simple choice”. *Nature Neuroscience* 13.10, pp. 1292–1298.
- Krupka, E. L. and R. A. Weber (2013). “Identifying social norms using coordination games: Why does dictator game sharing vary?” *Journal of the European Economic Association* 11.3, pp. 495–524.
- Lieder, F. and T. L. Griffiths (2020). “Resource-rational analysis: Understanding human cognition as the optimal use of limited computational resources”. *Behavioral and Brain Sciences* 40, pp. 1–60.
- López-Pérez, R. (2008). “Aversion to Norm-breaking: A Model”. *Games and Economic Behavior* 64, pp. 237–267.

- Matějka, F. and A. McKay (2015). “Rational inattention to discrete choices: A new foundation for the multinomial logit model”. *American Economic Review* 105.1, pp. 272–298.
- Milgrom, P. and J. Roberts (1994). “Comparing Equilibria”. *The American Economic Review* 84.3, pp. 441–459.
- Mohlin, E., A. Rigos, and S. Weidenholzer (2023). “Emergence of specialized third-party enforcement”. *Proceedings of the National Academy of Sciences* 120.24, e2207029120.
- Netzer, N. (2009). “Evolution of time preferences and attitudes toward risk”. *American Economic Review* 99.3, pp. 937–955.
- Nowak, M. A., K. M. Page, and K. Sigmund (2000). “Fairness versus reason in the ultimatum game”. *Science* 289.5485, pp. 1773–1775.
- Parkin, R. (1997). *Kinship: An introduction to basic concepts*. Oxford: Blackwell.
- Parsons, T. (1951). *The Social System*. New York: Free Press.
- Peysakhovich, A. and D. G. Rand (2016). “Habits of virtue: Creating norms of cooperation and defection in the laboratory”. *Management Science* 62.3, pp. 631–647.
- Phillips, L. D. and W. Edwards (1966). “Conservatism in a simple probability inference task”. *Journal of Experimental Psychology* 72.3, p. 346.
- Ratcliff, R. (1978). “A theory of memory retrieval”. *Psychological Review* 85.2, p. 59.
- Ratcliff, R. et al. (2016). “Diffusion decision model: Current issues and history”. *Trends in Cognitive Sciences* 20.4, pp. 260–281.
- Rayo, L. and G. S. Becker (2007). “Evolutionary Efficiency and Happiness”. *Journal of Political Economy* 115.2, pp. 302–337.
- Robson, A. J. (2001). “The biological basis of economic behavior”. *Journal of Economic Literature* 39.1, pp. 11–33.
- Robson, A. J. and L. Samuelson (2011). “The Evolutionary Foundations of Preferences”. In: *The Social Economics Handbook*. Ed. by J. Benhabib, A. Bisin, and M. O. Jackson. Amsterdam: North Holland, pp. 221–310.
- Rosen, L. (2000). *The justice of Islam: comparative perspectives on Islamic law and society*. USA: Oxford University Press.
- Rubinstein, A. (1998). *Modeling bounded rationality*. MIT Press.
- Rustagi, D. (2023). *Market exposure, civic values, and rules*. Tech. rep. CeDEx Discussion Paper Series.
- Rustagi, D., S. Engel, and M. Kosfeld (2010). “Conditional cooperation and costly monitoring explain success in forest commons management”. *Science* 330.6006, pp. 961–965.
- Samuelson, L. (2001). “Analogies, Adaptation, and Anomalies”. *Journal of Economic Theory* 97.2, pp. 320–366.
- Samuelson, L. (2005). “Foundations of human sociality: A review essay”. *Journal of Economic Literature* 43.2, pp. 488–497.

- Samuelson, L. and J. M. Swinkels (2006). “Information, evolution and utility”. *Theoretical Economics* 1.1, pp. 119–142.
- Sandholm, W. H. (2001). “Preference Evolution, Two-Speed Dynamics, and Rapid Social Change”. *Review of Economic Dynamics* 4, pp. 637–679.
- Schaffer, M. E. (1988). “Evolutionarily Stable Strategies for a Finite Population and a Variable Contest Size”. *Journal of Theoretical Biology* 132, pp. 469–478.
- Schulz, J. F. (2022). “Kin networks and institutional development”. *The Economic Journal* 132.647, pp. 2578–2613.
- Simon, H. A. (1957). *Models of man; social and rational*. Wiley.
- Sims, C. A. (2003). “Implications of rational inattention”. *Journal of Monetary Economics* 50.3, pp. 665–690.
- Steiner, J. and C. Stewart (2016). “Perceiving prospects properly”. *American Economic Review* 106.7, pp. 1601–1631.
- Sugaya, T. and A. Wolitzky (2021). “Communication and Community Enforcement”. *Journal of Political Economy* 129.9, pp. 2595–2628.
- Tabellini, G. (2008). “The scope of cooperation: Values and incentives”. *The Quarterly Journal of Economics* 123.3, pp. 905–950.
- Takahashi, S. (2010). “Community enforcement when players observe partners’ past play”. *Journal of Economic Theory* 145.1, pp. 42–62.
- Thaler, R. H. and H. M. Shefrin (1981). “An economic theory of self-control”. *Journal of political Economy* 89.2, pp. 392–406.
- Topkis, D. M. (1978). “Minimizing a Submodular Function on a Lattice”. *Operations Research* 26.2, pp. 305–321.
- Vives, X. (1990). “Nash equilibrium with strategic complementarities”. *Journal of Mathematical Economics* 19.3, pp. 305–321.
- Von Weizsäcker, C. C. (2014). “Adaptive preferences and institutional stability”. *Journal of Institutional and Theoretical Economics: JITE*, pp. 27–36.
- Wang, T. and E. Lehrer (2024). “Weighted Utility and Optimism/Pessimism: A Decision-Theoretic Foundation of Various Stochastic Dominance Orders”. *American Economic Journal: Microeconomics* 16.1, pp. 210–23.
- Weber, M. (2001). *The Protestant Ethic and the Spirit of Capitalism*. [1920]. London: Routledge.
- Weber, T. O. et al. (2023). “The behavioral mechanisms of voluntary cooperation across culturally diverse societies: Evidence from the US, the UK, Morocco, and Turkey”. *Journal of Economic Behavior and Organization* 215, pp. 134–152.
- Woodford, M. (2020). “Modeling imprecision in perception, valuation, and choice”. *Annual Review of Economics* 12, pp. 579–601.