

Risk Perception: Measurement and Aggregation*

Nick Netzer

University of Zurich

Arthur Robson

Simon Fraser University

Jakub Steiner

University of Zurich, Cerge-Ei and CTS

Pavel Kocourek

University of Duisburg-Essen

October 16, 2024

Abstract

In a model inspired by neuroscience, we study choice between lotteries as a process of encoding and decoding noisy perceptual signals. The implications of this process for behavior depend on the decision-maker's understanding of risk. When the aggregation of perceptual signals is coarse, encoding and decoding generate behavioral risk attitudes even for vanishing perceptual noise. We show that the optimal encoding of lottery rewards is S-shaped and that low-probability events are optimally oversampled. Taken together, the model can explain adaptive risk attitudes and probability weighting, as in prospect theory. Furthermore, it predicts that risk attitudes are influenced by the anticipation of risk, time pressure, experience, salience, and availability heuristics.

*We have benefited from comments of Olivier Compte, Cary Frydman, Philippe Jehiel, Franz Ostrizek, Antonio Rosato, Ryan Webb, and various seminar audiences. Giorgi Chavchanidze and Pavel Ilinov provided excellent research assistance. Earlier versions of this paper were circulated under the titles “Endogenous Risk Attitudes” and “Risk Perception.” This work was supported by the Alfred P. Sloan Foundation and the NOMIS Foundation (Netzer), by SSHRC grant 435-2013-0426 (Robson), by ERC grant 770652 and GAČR 24-10145S (Steiner), and by DFG grant 450175676 (E. Kováč, Kocourek).

1 Introduction

Choice under risk exhibits inconsistencies with expected utility theory. Some of these revolve around the inability to represent behavior with fixed risk preferences. We have known at least since Kahneman and Tversky (1979) that risk attitudes adapt to the environment. In a similar vein, Rabin’s (2000) paradox implies that choices involving small and large risks are represented by distinct utility functions. Additionally, risk attitudes are further modulated by factors such as the salience of rewards (Bordalo et al., 2012), time pressure (Kirchler et al., 2017), and experience (Ert and Haruvy, 2017). Additional inconsistencies include the overweighting of small-probability events relative to large-probability events (Kahneman and Tversky, 1979) and the relevance of the availability of an event in memory for assessing its probability (Tversky and Kahneman, 1973).

We present a decision-making model grounded in neuroscience and psychophysics. When evaluating a lottery, our decision-maker (DM) first *measures* the rewards that the lottery pays in different states of the world and then *aggregates* the collected information to assess the expected value of the lottery. Frictions occur at both stages of the process. The measurement stage involves encoding the rewards into perceptual signals, which are noisy. During the aggregation stage, the DM acts as a statistician, decoding the perceptual data to estimate the lottery’s value, possibly using a misspecified statistical model. The interaction of these two frictions generates the aforementioned behavioral phenomena.

For illustration, consider insurance against a small-stakes risk, such as the baggage insurance often offered during a flight booking. Assume that the DM’s genuine preferences are risk-neutral regarding this risk. When evaluating the option of traveling uninsured (referred to as “the lottery”), the DM contemplates various possible contingencies, some involving no damage to the baggage and others involving varying degrees of damage or complete loss. We envision the DM collecting impressions of these various contingencies, which we refer to as signals. The DM then aggregates all the signals into an estimate of the expected value of traveling uninsured and compares it to the value of buying insurance (referred to as “the safe option”).

First, we focus on the aggregation stage. For this, we take the measurement stage to be exogenous, modeled as a sampling process governed by an *encoding function* and *sampling frequencies*. The non-linear encoding function maps lottery rewards into subjective impressions (e.g., represented by neural firing rates), which are perturbed by additive Gaussian noise. Depending on the application, these perturbed signals can be one’s own experiences retrieved from memory, experiences provided by others (e.g., the insurance company), or cues extracted from the visual description of the decision problem. The sampling frequen-

cies specify the proportion of signals dedicated to each state of the world, which do not necessarily coincide with the true probabilities of the states. One possible and particularly simple aggregation procedure then takes the arithmetic average of all the sampled signals and applies the inverse of the encoding function to compute an estimate of the expected value of the lottery.

When the number of signals which the DM collects becomes large, this particular decision procedure generates choices represented by expected utility maximization with a utility function equal to the encoding function and subjective probabilities equal to the sampling frequencies. We thus establish a tight connection between the elements of the measurement stage (the encoding function and the sampling frequencies) and the elements of expected utility theory (the utility function and subjective probabilities). In this case, however, the representation does not reflect the DM’s true preferences and welfare but rather the DM’s errors in aggregating the sampled information.

This simple model makes several additional predictions that are consistent with empirical regularities. First, because behavioral risk attitudes depend on the properties of encoding, they will adapt if the DM’s encoding strategy adjusts to the environment (Kahneman and Tversky, 1979; Frydman and Jin, 2022). Since behavior depends on sampling frequencies rather than true probabilities, subjective probability weights from the representation may differ from objective probabilities (Kahneman and Tversky, 1979) and can furthermore be affected by the salience of states (Bordalo et al., 2012) or availability in memory (Tversky and Kahneman, 1973).

The simple aggregation procedure described above is a special case of misspecified maximum likelihood estimation using a coarse partitional model of the true state space, similar to Savage’s (1954) decision-maker employing a small-world model of the grand world. If the DM anticipates a riskless lottery that pays the same reward in all states, averaging all the signals and applying the inverted encoding function gives the maximum likelihood estimate of the reward. The estimation is misspecified if the rewards differ across the states, resulting in a biased estimate and behavioral risk attitudes. If, on the other extreme, a fully sophisticated DM anticipates that the lottery pays different rewards in different states, her maximum likelihood estimate averages the signals and applies the inverted encoding function state-by-state, leading to an unbiased estimate and to risk-neutral behavior in the limit of rich perceptual data. For intermediate levels of sophistication, captured by a general coarse partitional model of the state space, maximum likelihood estimation leads to risk-neutral choices whenever the DM faces risk that she comprehends (i.e., is measurable with respect to her partition), but implies perception-driven risk attitudes for uncomprehended elements of risk.

For instance, our traveler may simplify her reasoning by bundling all states with distinct levels of damage into one element of the partition, distinguishing only between damage and no damage. If the planned travel indeed involves only two possible contingencies (say, no damage or complete loss of the baggage), then the traveler’s partition is correct and she makes a risk-neutral choice, regardless of the details of the encoding function and the sampling frequencies that govern her reasoning. If, however, the travel involves possible contingencies with various degrees of damage, then the DM’s coarse partition induces estimation bias and behavioral risk attitudes. The estimated expected damage will be influenced by the proportions of sampling of each possible contingency and by the curvature of the encoding function. This makes the DM manipulable. For instance, an advertisement for baggage insurance may emphasize scenarios involving total baggage loss, causing the traveler to oversample such contingencies, resulting in an exaggerated estimate of expected damage.

To illustrate the applicability of our results across settings, we also discuss a financial investor who omits a variable from her model of the risk and consequently displays behavioral risk attitudes. In general, these behavioral risk attitudes become less pronounced with a finer partition. To the extent that more experience with a decision environment implies a better understanding of the true state space and thus a more appropriate partition, behavior is predicted to become more risk-neutral with experience, in line with Ert and Haruvy (2017).

In addition to the above partitional model, we also study a Bayesian variant in which the DM is endowed with a prior over lotteries and forms Bayesian posteriors upon collecting perceptual signals. To study how prior information and perceptual information interact in the aggregation process, we develop an asymptotic approach in which both the amounts of prior and perceptual information diverge at comparable rates. By adjusting the relative rates of divergence, we can vary the influence of these two information sources on choice. Choice becomes more risk-neutral when the DM anticipates large risk a priori, consistent with the finding that framing a decision problem as one featuring high risk dampens risk attitudes (Rabin, 2000). Choice also becomes more risk-neutral when the DM collects a lot of data, which predicts that risk attitudes are intensified under time pressure when decisions rely more on prior information and less on perceptual signals (Kirchler et al., 2017).

The preceding arguments treated the encoding strategy as exogenous. In Section 3, we develop an evolutionary argument to complement the empirical evidence that perception of rewards and probabilities has properties similar to prospect theory. We assume that these patterns evolved in a stable ancestral environment, with correct decoding occurring at the time when optimization of the encoding strategy took place. We find that combining an S-shaped encoding function with oversampling of small-probability states is jointly optimal for a large class of environments, regardless of their distributional specifics. The argument,

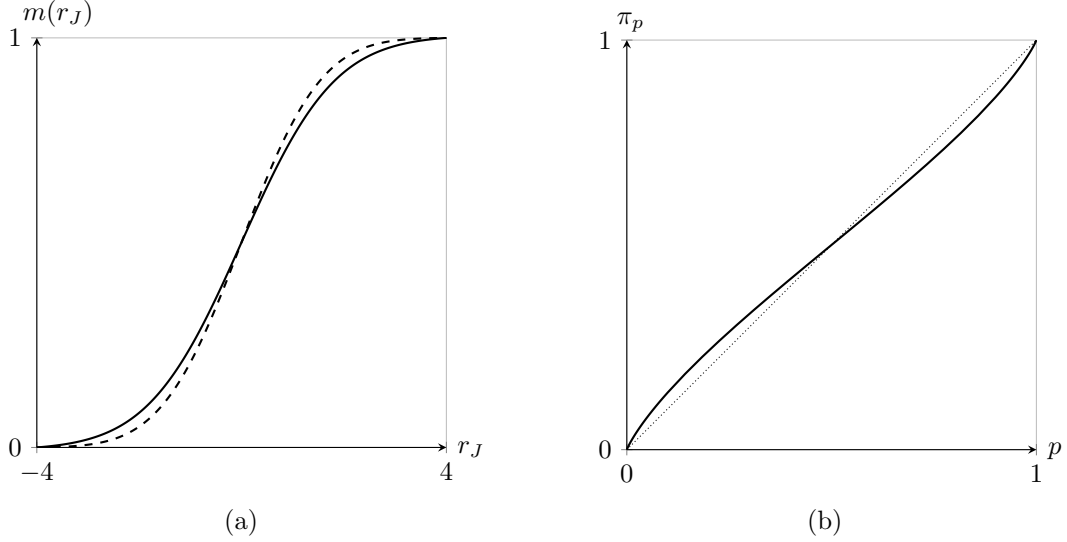


Figure 1: (a) The encoding function maps the reward r_J to the perception signal $m(r_J)$. Dashed line: Optimal encoding function for a benchmark riskless choice. Solid line: Optimal encoding function for the choice between a non-trivial lottery and a safe option. (b) Optimal sampling frequency π_p as a function of the objective probability p , shown in comparison to the diagonal line. See Example 3 for details.

therefore, does not require detailed knowledge of the ancestral environment. Since these perception patterns are optimal across diverse environments, it is plausible that we are adapted to use prospect-theory-like perception heuristics. When such heuristics are combined with aggregation frictions that plausibly arise in complex environments, they result in behavioral patterns consistent with prospect theory.

The optimization of perception is an attention allocation problem. The DM is akin to an engineer who measures a physical input by reading the stochastic position of a needle on a meter (Robson, 2001). Similar to the engineer, our DM can enhance the precision of reward perception within a specific range by steepening the encoding function in that range, though this comes at the expense of precision in other ranges. Furthermore, the DM can allocate attention to a specific state by frequently sampling it, at the expense of attention to other states. We again analyze the limit of rich perceptual data, which serves as an approximation for the more realistic case of non-vanishing small noise and allows for tractable results. We prove that the expected loss from noisy perception approximates the mean squared error in the estimate of the lottery value, integrated over all decision problems where the lottery value ties with the safe option. This conditioning on ties arises because choice is trivial except when the values of two options are nearly equal, given that information is nearly complete in the limit. Since information is instrumental for choice, it must be precise only in decision problems where precision matters.

We then prove that the optimal encoding strategy, which minimizes the mean squared error conditional on ties, exhibits the typical prospect theory properties. It is illustrated in Figure 1, which plots the optimal encoding function and sampling frequencies for a specific example (to be found in Section 3.2). The DM selects the encoding function to be steep near the modal rewards and flatter toward the tails of the reward distribution. Thus, she perceives the reward values that are typical for her environment with relative precision, sacrificing precision for tail rewards. Jointly with this encoding function, the optimal sampling frequencies oversample low-probability states. Our derivation of oversampling hinges on a subtle argument. Conditioning on ties establishes a statistical association between tail rewards and low-probability states. Since tail rewards from high-probability states typically result in very attractive or unattractive lotteries rather than ties, tail rewards arise relatively often in low-probability states when conditioning on a tie. Thus, a DM with an S-shaped encoding function frequently struggles to estimate rewards in low-probability states when at ties, making it optimal to compensate by oversampling these states.

Our results on optimal perception substantially generalize earlier findings in the literature that studied choice between simple, one-dimensional objects (Netzer, 2009, see the dashed line in Figure 1(a) for this benchmark). Here, we provide a microfoundation for an objective rooted in choice that involves a lottery—a complex, multi-dimensional object—which allows us to study the perception of rewards and probabilities in the context of risk.

2 Aggregation of Perceptual Data

We introduce a model of measurement and aggregation of lottery rewards that rests on the noisy encoding of these rewards and subsequent decoding of the perceptual data. Throughout this section, we treat the measurement stage as fixed and focus on the behavioral implications of aggregation frictions. We make predictions consistent with the findings of Oprea (2023) and the literature on prospect theory (Kahneman and Tversky, 1979), Rabin’s paradox (Rabin, 2000), salience (Bordalo et al., 2012), time pressure (Kirchler et al., 2017), experience (Ert and Haruvy, 2017), and availability heuristics (Tversky and Kahneman, 1973).

2.1 Maximum Likelihood Decoding

There is a set of states of the world $i \in \{1, \dots, I\}$, where each state i has a fixed positive probability p_i . The DM chooses between a safe option of value s and a lottery that pays a reward $r_i \in [\underline{r}, \bar{r}]$ in each state i , where $\underline{r} < \bar{r}$ are arbitrary bounds. We let $\mathbf{r} = (r_i)_i$ denote the tuple of rewards and refer to it as the *lottery*. The pair $(\mathbf{r}, s)$ is the *decision problem*.

We focus on the DM’s frictions in the perception of the lottery rewards. For simplicity, we assume that she knows the probability tuple and observes the value of the safe option without frictions, thus abstracting from possible interactions between learning the rewards and the probabilities.

The DM receives n signals, where each signal is a monotonic transformation of one of the rewards, perturbed with additive noise. That is, signals are given by $x_k = (\hat{m}_k, i_k)$, $k = 1, \dots, n$, with exogenous sample size n . We refer to the first component, $\hat{m}_k$, as the *perturbed message*. The second component, i_k , indicates the state that the k ’th signal pertains to. Each perturbed message is generated by encoding the reward r_{i_k} in state i_k into *unperturbed message* $m(r_{i_k})$ and then perturbing it to $\hat{m}_k = m(r_{i_k}) + \hat{\varepsilon}_k$, where the noise term $\hat{\varepsilon}_k$ is independently and identically distributed (iid) standard normal.¹ The sampled state i_k is one of the states $i = 1, \dots, I$, iid with positive probabilities π_i . The function $m : [\underline{r}, \bar{r}] \rightarrow [\underline{m}, \bar{m}]$ is assumed throughout the paper to be continuously differentiable and strictly increasing, with $m' > 0$. The assumption that the DM encodes rewards into messages from a finite range and that this encoding is noisy follows a long tradition in the psychometric literature (see e.g., the discussion in Frydman and Jin, 2022). We refer to m as the *encoding function* and to $(\pi_i)_i$ as the *sampling frequencies*. The pair $(m, (\pi_i)_i)$ of encoding function and sampling frequencies is the *encoding strategy*, kept exogenous in this section.

In the insurance example from the introduction, the sure payoff s corresponds to the traveler’s wealth minus the insurance price, while the lottery corresponds to the traveler’s uncertain wealth if she travels uninsured. The traveler repeatedly contemplates all possible travel contingencies, with the proportion of thoughts on each contingency governed by sampling frequencies that may not reflect true probabilities (Tversky and Kahneman, 1973). Consistent with the insights from psychophysics, the non-linear shape of the encoding function determines how the financial consequences of any contemplated contingency map to subjective perceptions.²

The DM forms a maximum likelihood estimate of the lottery rewards. She is endowed with a compact set $\mathcal{A} \subseteq [\underline{r}, \bar{r}]^I$ of lotteries she deems possible and, given the n signals,

¹We use the Gaussian assumption mainly because it will yield a tractable form of the Kullback-Leibler divergence in the next subsection.

²Even when lottery rewards are presented as numbers, noisy perception occurs in the process of visual inspection of the options (Schaffner et al., 2023). In this case, and in line with a large literature (e.g., Robson, 2001; Netzer, 2009; Frydman and Jin, 2022), our approach assumes that the DM processes stated numbers like any other percept (see Khaw et al., 2021, for a careful discussion and motivation of this assumption). Furthermore, memory is involved when processing the consequences of different events such as winning or losing (Ludvig et al., 2015).

concludes that she has encountered the lottery

$$\mathbf{q}_n \in \arg \max_{\mathbf{r}' \in \mathcal{A}} \prod_{k=1}^n \varphi(\hat{m}_k - m(r'_{i_k}))$$

that maximizes the likelihood of the observed signals, where φ denotes the standard normal density.³ Given the estimate $\mathbf{q}_n = (q_{1n}, \dots, q_{In})$, she estimates the value of the lottery to be $q_n = \sum_i p_i q_{in}$ and chooses the lottery if $q_n > s$ and the outside option otherwise. Risk neutrality with respect to rewards embodied in this rule is an implicit assumption about the units of measurement in which the rewards are expressed. For instance, the rewards might be an appropriate concave function of monetary prizes if the DM chooses among monetary lotteries and money has diminishing returns.⁴

The set $\mathcal{A}$ of anticipated lotteries is as follows. The DM employs a possibly coarse model of the risk in the spirit of the *small world* of Savage (1954). That is, she anticipates, rightly or wrongly, distinctions among some of the states of the world to be payoff-irrelevant, similarly to Jehiel (2005). Let $\mathcal{P}$ be a partition of the set of all the states $\{1, \dots, I\}$. The DM anticipates that $r_i = r_j$ for all pairs of states $i, j \in J$ from the same element J of the partition $\mathcal{P}$. That is, she anticipates lotteries from a set

$$\mathcal{A}_{\mathcal{P}} = \left\{ \mathbf{r} \in [\underline{r}, \bar{r}]^I : r_i = r_{i'} \text{ for all } i, i', J \text{ such that } i, i' \in J, J \in \mathcal{P} \right\}. \quad (1)$$

If $\mathcal{P} = \{\{1, \dots, I\}\}$ is the coarsest partition, then the DM anticipates only degenerate lotteries that pay the same reward in all states. If, on the other extreme, $\mathcal{P} = \{\{1\}, \dots, \{I\}\}$ is the finest partition, then the DM anticipates any reward tuple.

We treat the partition $\mathcal{P}$ as exogenous. There are various paths that could lead the DM to adopt a partition that is too coarse to measure the lotteries she actually encounters. Real-world DMs are sometimes unaware of all contingencies that affect their payoffs, thus effectively omitting relevant variables from their risk models. Alternatively, the DM may have been incorrectly assured (possibly by a strategically interested party) that her next lottery would be relatively riskless. Finally, the DM may know that she is encountering a risky lottery but applies a coarse estimation procedure for simplicity, as it requires estimating fewer distinct rewards. Our approach is also consistent with the idea that the DM's partition

³The maximum-likelihood estimate exists since $\mathcal{A}$ is compact. It is unique for the specifications that we consider in the following.

⁴For example, rewards can be $r_i = f(w_i)$, where w_i is a monetary prize and f is a fitness function, and the DM measures the rewards r_i by applying a non-linear encoding function $m(r_i)$. Then, from the perspective of an analyst who measures the rewards in monetary units, the DM's measurement $m(f(w))$ compounds f and m . For monetary lotteries with relatively small stakes, where a linear function f appears particularly plausible, all non-degenerate risk attitudes are entirely driven by the aggregation frictions.

changes over time. An inexperienced agent may initially use a coarse partition but refine it in subsequent choices.

The following result characterizes the behavior of the DM given a partition and encoding strategy.

Proposition 1. *With maximum likelihood decoding, the probability that the DM chooses the lottery in problem $(\mathbf{r}, s)$ converges almost surely to 1 (0) as $n \rightarrow \infty$ if*

$$\sum_{J \in \mathcal{P}} p_J r_J^* > (<) s,$$

where $p_J = \sum_{i \in J} p_i$ and r_J^* is the certainty equivalent defined by

$$m(r_J^*) = \sum_{i \in J} \frac{\pi_i}{\sum_{j \in J} \pi_j} m(r_i)$$

for each element of the partition $J \in \mathcal{P}$.

In the limit, the DM makes choices as if she were treating the lottery $\mathbf{r}$ as a compound lottery, where each element J of the partition constitutes a sub-lottery that occurs with probability p_J . She behaves as if she first reduces each sub-lottery to its certainty equivalent under the utility function m and subjective probabilities equal to the (renormalized) sampling frequencies π_i . After this reduction, she evaluates the overall lottery in a risk-neutral manner using the true probabilities of each J .

Let us return to the traveler from the introduction and assume she uses a binary partition that categorizes states as either having no damage or some damage. She first conditions on the event of damage, computing the certainty equivalent for this conditional lottery under the utility function m and subjective probabilities equal to the renormalized sampling frequencies. She then aggregates this certainty equivalent with her wealth in the event of no damage, using the objective probability of this favorable outcome. Consequently, while this traveler's insurance choice cannot be influenced by making the general event of positive damages more salient, drawing her attention to events with high damages at the expense of those with low damages increases her willingness to pay for insurance. In general, sampling frequencies determine subjective probabilities across states that the DM does not distinguish, while objective probabilities matter whenever the DM does distinguish between states.

The proof of the proposition (in Appendix A.1) relies on a result about misspecified maximum-likelihood estimation by White (1982). White lets an agent observe iid signals from a signal density and form the maximum likelihood estimate from a set of hypothesized signal densities, which may fail to include the true density. He proves that the estimate

almost surely converges to the minimizer of the Kullback-Leibler divergence from the true signal density as the number of signals diverges. In our setting, the DM observes the empirical distribution of approximately $\pi_i n$ perturbed messages drawn iid from $\mathcal{N}(m(r_i), 1)$ for each state i . Since the DM anticipates the same reward in all states $i \in J$, she forms an estimate of a single unperturbed message for each $J \in \mathcal{P}$, a perturbation of which has generated the observed data. For Gaussian errors, this estimate is the arithmetic average of the perturbed messages for J , which almost surely converges to $\sum_{i \in J} \frac{\pi_i}{\sum_{j \in J} \pi_j} m(r_i)$. Thus, the DM’s estimate of the reward in J converges to the certainty equivalent r_J^* defined in the proposition. Across elements J of the partition, the DM’s anticipation of distinct rewards implies that she aggregates the values r_J^* in a risk-neutral manner.

This result has implications that differ from those of expected utility theory and align with observed empirical behavioral patterns. First, the behavior in Proposition 1 arises not due to risk preferences but as a consequence of the friction in aggregating rewards. This basis for behavioral risk attitudes has empirical support in the recent experiment by Oprea (2023). To isolate risk preferences from aggregation frictions, Oprea compares choices over lotteries to choices over their ‘deterministic mirrors.’ A deterministic mirror pays the lottery’s expected value without any risk but is still described in terms of an arithmetic operation that involves the same aggregation. The subjects’ behavior is essentially identical for the lotteries and their deterministic mirrors, suggesting that the aggregation frictions, rather than genuine risk preferences, drive the choice. Our approach indeed implies that behavioral risk attitudes would persist if the lottery were replaced by its expected value presented arithmetically, when evaluating the latter involves the same aggregation frictions.

Second, behavior is influenced by sampling frequencies π_i . Tversky and Kahneman (1973) point out that subjective probability weights are influenced by the availability of the respective events in memory and do not necessarily correspond to true probabilities. In Bordalo et al. (2012), the salience of a state shifts subjective probabilities in favor of that state. Our proposition provides a microfoundation for salience-driven behavior based on non-representative sampling. Note, however, that we treat the sampling frequencies as exogenous in this section and thus abstain from predicting how salience weights depend on the menu; see Bordalo et al. for such predictions.

Third, since behavior depends on the encoding strategy, choice will be influenced by both the adaptation of the encoding function (see Frydman and Jin, 2022; Schaffner et al., 2023) and variations in the sampling frequencies, such as through marketing interventions that manipulate focus (see Koszegi and Szeidl, 2013).

An interesting special case arises for the DM who anticipates no risk and uses the coarsest partition. She follows a simple decoding procedure that averages all the perturbed mes-

sages and then applies the inverted encoding function to obtain an estimated lottery value $m^{-1}(\sum_{k=1}^n \hat{m}_k/n)$. According to the following corollary, this DM behaves as if the encoding function m serves as her utility function and the sampling frequencies π_i , rather than the objective probabilities p_i , represent her subjective probabilities.

Corollary 1. *With the coarsest partition, the probability that the DM chooses the lottery in problem $(\mathbf{r}, s)$ converges almost surely to 1 (0) as $n \rightarrow \infty$ if $\sum_{i=1}^I \pi_i m(r_i) > (<) m(s)$.*

At the other extreme, a DM who uses the finest partition always behaves in a risk-neutral manner based on correct probabilities. More generally, whenever the DM encounters a lottery $\mathbf{r} \in \mathcal{A}_{\mathcal{P}}$ that she has anticipated, she learns within a well-specified model. Asymptotic results for well-specified maximum-likelihood estimation by Wald (1949) imply that she correctly learns the encountered lottery as the number of signals diverges. This implies risk-neutral choice regardless of the encoding strategy.

Between these two extremes, our model generates interesting comparative statics. The behavior is modulated by the degree of understanding of the risk and thus by the DM's experience with the situation. More experience (conceptualized here as a finer partition) generates more risk-neutral and less manipulable behavior. This is consistent with experimental results showing that experience with a decision problem tends to shift risk preferences toward risk neutrality (Bradbury et al., 2015; Ert and Haruvy, 2017; Charness et al., 2023).⁵

We conclude this section with two examples illustrating Corollary 1 and Proposition 1 in the context of an investor who omits a risk factor in her asset evaluation.

Example 1 (investor unaware of risk). An investor chooses between a safe return s and a risky asset with return $r_{x,i}$, which depends on variables x and i drawn from a joint distribution $p_{x,i}$. She learns the return function $r_{x,i}$ in a misspecified model that omits the variable i , believing the return to be solely a function of x , and forms its estimate q_x . For instance, the investor may be aware that a firm's profit depends on the interest rate (x) but is unaware of the firm's trade exposure and thus neglects the impact of the exchange rate (i). She estimates q_x for each value of x separately from a large sample of signals $m(r_{x,i_k}) + \hat{\varepsilon}_k$, $k = 1, \dots, n$, using maximum likelihood estimation, where the shares of signals about the return $r_{x,i}$ are determined by her sampling frequencies π_i .

Corollary 1 applies to each value x separately. For each x , the asset return is a lottery with risk induced by random $i \mid x$. The investor, who is unaware of the risk, attributes the randomness observed in the data to the encoding noise $\hat{\varepsilon}_k$, and forms the estimate q_x equal

⁵To be precise, our model predicts that refining a partition eliminates perception-driven risk attitudes and generates risk-neutral behavior for those lotteries that become measurable with respect to the refined partition.

to the certainty equivalent of the lottery that pays $r_{x,i}$ with probabilities π_i under the utility function coinciding with the encoding function m . $\triangle$

Example 2 (investor with partial risk awareness). Now suppose the asset return r_{i_1,i_2} depends on a pair of unobserved risk factors $i = (i_1, i_2)$, and the investor observes a variable y that correlates with i . Thus, conditional on the observed value of y , the return is a lottery that attains values r_{i_1,i_2} with probabilities $p_{i_1,i_2}(y)$. The investor is aware of the payoff relevance of the risk factor i_1 but not i_2 , so her estimate q_{i_1} of the return is based solely on i_1 . That is, for each value of y , the investor forms a coarse model of the actual lottery in which the estimated return q_{i_1} arises with probability $p_{i_1}(y) = \sum_{i_2} p_{i_1,i_2}(y)$. She samples the states (i_1, i_2) with frequencies π_{i_1,i_2} and forms a coarse estimate q_{i_1} of the return function from a large sample of signals $m(r_{i_{1,k},i_{2,k}}) + \hat{\varepsilon}_k$, $k = 1, \dots, n$.

The behavior of the investor is now characterized by Proposition 1 for a partition $\mathcal{J}$ that is measurable with respect to i_1 . She behaves as if she viewed the risky asset return as a compound lottery. Given each i_1 , she first computes the certainty equivalent $r_{i_1}^*$ for a lottery that pays r_{i_1,i_2} with probabilities $\pi_{i_1,i_2} / \sum_{i_2} \pi_{i_1,i_2}$ under the utility function equal to the encoding function m . Then, given each observed value of y , she proceeds to evaluate the risky asset in the risk neutral manner as if the asset paid $r_{i_1}^*$ in each state i_1 arising with probabilities $p_{i_1}(y)$. The certainty equivalents $r_{i_1}^*$ correspond to the estimates q_{i_1} , while the risk premia reflect the estimation biases arising from the misspecified econometric model. $\triangle$

2.2 Decoding with Additional Prior Information

The distinction between anticipated and unanticipated lotteries was dichotomous in the previous section. We now adopt a more continuous approach to deliver additional comparative statics. In a specific example, we study a Bayesian DM who possesses a comparable amount of both a priori and perceptual information. Both sources influence the choice. We consider a sequence of information structures where the prior simultaneously concentrates on the set of riskless lotteries and the quantity of perceptual data diverges. By varying the relative rates of divergence, we can modulate the relative influence of a priori versus perceptual information. This considered limit is a modeling tool that facilitates analytical results.

We set the prior belief of the DM over lotteries $\mathbf{r}$, indexed by n , to the density

$$\varrho_n(\mathbf{r}) = \varrho_n^0 \exp\left(-\frac{n}{2\Delta}\sigma^2(\mathbf{r})\right) \quad (2)$$

with support $[r, \bar{r}]^I$, where $\sigma^2(\mathbf{r}) = \sum_{i=1}^I p_i(r_i - \sum_j p_j r_j)^2$ is the variance of the states' rewards and ϱ_n^0 is the normalization factor. For large n , this prior approximates a prior over a riskless value uniformly distributed on $[r, \bar{r}]$. The parameter $\Delta > 0$ specifies the

level of a priori anticipated risk, with a larger Δ indicating that the DM anticipates more risk for any given n .⁶ In addition to indexing the prior, we let n control the volume of the DM's perceptual data. Similar to our model in the previous subsection, for each state i , the DM observes a sequence of approximately $a\pi_i n$ messages equal to $m(r_i)$ perturbed with iid additive standard normal noise, where m and $(\pi_i)_i$ continue to denote the exogenous encoding strategy.⁷ The parameter $a > 0$ represents attention span, with a larger a implying that the DM observes more signals for every fixed n . Finally, the DM chooses the lottery over the safe option if and only if the posterior expectation of the lottery value exceeds s .

The new parameters Δ and a jointly determine the influence of the DM's prior and the perceptual data on her posterior. If $a\Delta$ is small, the DM's prior anticipation of a relatively riskless lottery dominates the perceptual data. Conversely, if $a\Delta$ is large, the rich perceptual data dominates the relatively dispersed prior.

The result of Berk (1966), which characterizes the Bayesian posterior asymptotically for an expanding sample while keeping the prior fixed, does not apply directly to our setting where the prior varies. However, by adapting this result, we demonstrate that the Bayesian posterior converges to an atom concentrated on a lottery that reconciles the unexpected elements observed in the data with the deviations from our increasingly focused prior.

To formulate the result, we define the function $\mathbf{q}^* : [\underline{r}, \bar{r}]^I \rightarrow [\underline{r}, \bar{r}]^I$ as

$$\mathbf{q}^*(\mathbf{r}) = \arg \min_{\mathbf{r}' \in [\underline{r}, \bar{r}]^I} \left\{ \frac{\sigma^2(\mathbf{r}')}{a\Delta} + \sum_{i=1}^I \pi_i (m(r'_i) - m(r_i))^2 \right\}. \quad (3)$$

We impose the regularity condition that the minimizer is unique for the given $\mathbf{r}$ of interest, which holds generically. We will show that the DM's posterior expected lottery converges to $\mathbf{q}^*(\mathbf{r})$ almost surely as $n \rightarrow \infty$. This asymptotic estimate is a compromise lottery that is not too risky and does not generate messages too far from the true messages. When $a\Delta$ is small, the primary concern in (3) is the minimization of $\sigma^2(\mathbf{r}')$, and hence $\mathbf{q}^*(\mathbf{r})$ will involve little risk. In the limit as $a\Delta \rightarrow 0$, the estimate minimizes the Kullback-Leibler divergence from the true lottery (the last term on the right of (3)) among the riskless lotteries. When $a\Delta$ is large, the main concern in (3) is the minimization of the discrepancy between the estimate and the messages. In the limit as $a\Delta \rightarrow \infty$, this yields the correct lottery $\mathbf{q}^*(\mathbf{r}) = \mathbf{r}$.

Proposition 2. *With Bayesian decoding and concentrated prior, the probability that the DM*

⁶For large n , one can think of this prior as drawing a common value for the rewards in all states uniformly from $[\underline{r}, \bar{r}]$ and then perturbing each reward with a Gaussian shock with variance proportional to Δ^2 .

⁷Since we take the number n of signals to be large, from now on we abstract from uncertainty over the number of perturbed messages sampled for each state and from divisibility issues. Thus, we assume that the average of the sampled messages for state i is drawn from $\mathcal{N}(m(r_i), 1/(a\pi_i n))$.

chooses the lottery in problem $(\mathbf{r}, s)$ converges almost surely to 1 (0) as $n \rightarrow \infty$ if

$$\sum_{i=1}^I p_i q_i^*(\mathbf{r}) > (<) s.$$

See Appendix A.2 for the proof. If the DM anticipates relatively large risk and/or collects a lot of perceptual data ($a\Delta$ large), the relevance of the prior information is diminished, and the DM learns the true value of the lottery, becoming risk-neutral like the well-specified DM in Subsection 2.1. If the DM anticipates relatively little risk and/or collects little perceptual data ($a\Delta$ small), prior information remains influential and the DM's posterior mirrors that of the coarse DM. This leads her to maximize the expectation of utility, which is equal to her encoding function, and to use sampling frequencies as subjective probabilities.

Furthermore, the result yields additional comparative statics between these two extremes. As the volume of perceptual data grows, choice eventually shifts from the risk attitudes shaped by the encoding strategy to risk neutrality. Similarly, the anticipation of high risk attenuates risk attitudes.⁸ To illustrate these comparative statics further, we compute the risk premium, defined in the standard manner as the excess return needed to compensate the DM for accepting risk. As is common in expected utility theory, we compute the risk premium for a lottery $\mathbf{r}$ with expected value $r = \sum_J p_J r_J$ that involves little risk, in that the variance $\sigma^2(\mathbf{r})$ of the rewards across states is small. To further facilitate comparison with the standard risk premium, we set the sampling frequencies equal to the actual probabilities.⁹

Proposition 3. *Let the encoding function m be twice differentiable, and assume $\pi_i = p_i$. Given a lottery $\mathbf{r}$, the expected value of its asymptotic Bayesian estimate is*

$$r + \frac{1}{2} \frac{m''(r)}{m'(r)} \cdot \frac{1 + 4z(r)}{(1 + z(r))^2} \cdot \sigma^2(\mathbf{r}) + o(\sigma^2(\mathbf{r})), \quad (4)$$

where $z(r) = a\Delta m'^2(r)$.

See Appendix A.3 for the proof. To interpret the result, recall that the risk premium of an expected utility maximizer is given by $\frac{1}{2} \frac{u''(r)}{u'(r)} \sigma^2(\mathbf{r})$. The risk premium of our DM is the

⁸Our prediction that anticipation of risk diminishes the DM's risk aversion is reminiscent of Köszegi and Rabin (2007) but is based on a different mechanism. In Köszegi and Rabin (2007), anticipation is modeled by a reference lottery and behavioral risk attitudes are generated by aversion to losses relative to that reference lottery. In our model, anticipation is modeled by prior beliefs about the decision problem and behavioral risk attitudes are generated by biases in the estimation.

⁹Otherwise, the effect of distorted sampling dominates the effect of the curvature of the encoding function, because the first effect is of the order of $\sigma(\mathbf{r})$, while the second effect is only of the order of $\sigma^2(\mathbf{r})$. We use the expression $o(\cdot)$ for 'terms of smaller order than.' That is, as is standard, we say that the function $f(\mathbf{r})$ is $o(g(\mathbf{r}))$ if $f(\mathbf{r}_k)/g(\mathbf{r}_k) \rightarrow 0$ for any sequence $\mathbf{r}_k$ such that $g(\mathbf{r}_k) \rightarrow 0$.

same for $u(\cdot) = m(\cdot)$, but it is scaled by the positive factor $\frac{1+4z(r)}{(1+z(r))^2}$, which depends on the parameters Δ and a and approaches 1 and 0 as $a\Delta \rightarrow 0$ and $a\Delta \rightarrow \infty$, respectively.

The bias in the lottery valuation, interpreted here as the risk premium, arises as follows. When facing a risky lottery, the decision maker (DM) experiences a conflict between perceptual data and prior information, resolving it by determining that the lottery is riskier than anticipated a priori but less risky than suggested by the perceptual data. This underestimation of the reward variance leads to a mismatch with the perceptual data. To minimize this mismatch, the curvature of the encoding function necessitates that the DM biases the estimated average reward by an amount equal to the risk premium specified in (4).

The dependence of the risk premium on the parameter a sheds light on the instability of risk preferences pointed out by Kahneman (2011), who distinguishes between fast and slow decision-making. In his work, the fast mode employs an a priori heuristics and favors the risk attitudes described in prospect theory, while the slow mode mitigates the biases associated with prospect theory by analytical reasoning. Our results formalize his intuition. We observe encoding-based risk attitudes when $a \rightarrow 0$, which corresponds to a fast decision that is based primarily on prior beliefs. For slow decisions, corresponding to $a \rightarrow \infty$, the collected data make the prior irrelevant, leading to a risk-neutral choice.¹⁰ Therefore, the seller of the baggage insurance in our traveler’s example might impose time pressure hoping that behavioral risk aversion induces insurance purchase.

Rabin (2000) argues that risk-averse choices observed for small risks imply implausibly high risk-aversion for large risks under a stable utility function. The risk attitudes of our DM depend on the level of a priori anticipated risk. Anticipation of low risk, indicated by small Δ , induces strong risk attitudes because it makes risky lotteries surprising, and this leads to distortion of the posteriors when a risky lottery is encountered. If the DM anticipates relatively large risk, indicated by large Δ , then her risk attitudes are attenuated. Risky lotteries become relatively unsurprising, and the DM’s posterior expectation approaches the lottery’s true expected value, leading to risk neutrality. For illustration, consider again the seller of baggage insurance. The seller benefits when the traveler makes the purchase decision based on a prior with little anticipated risk, which then paints the potential baggage losses as relatively large and induces behavioral risk aversion. For example, it may be more profitable to sell baggage insurance in isolation rather than jointly with other insurance contracts

¹⁰Fast versus slow here refers to exogenous changes in the decision environment, for example due to time pressure. This differs from a literature where response times are endogenous and may vary systematically with the difficulty of the decision problem (e.g., Fudenberg et al., 2018; Alós-Ferrer et al., 2021; Liu and Netzer, 2023). Kirchler et al. (2017) show experimentally that time pressure increases risk aversion for gains and risk loving for losses. Relatedly, Porcelli and Delgado (2009) and Cahliková and Cingl (2017) find that stress accentuates risk attitudes in lab choices. But see also Kocher, Pahlke, and Trautmann (2013) who do not find an increase of risk aversion due to time pressure in their design.

covering more substantial losses, such as health or liability insurance, which would frame the buyer to contemplate much larger risks from the outset.

3 Stylized Properties of the Encoding Strategy

In the previous section, we argued that the properties of the encoding strategy have behavioral implications for a DM facing decoding frictions, but we have not yet specified these properties or considered their possible basis. In this section, we complement the above behavioral link with arguments supporting the stylized facts about perception of rewards and probabilities from prospect theory. These stylized facts—the S-shaped evaluation of rewards and the overweighting of small probabilities—have been documented in both human and animal studies, suggesting an evolutionary basis for the perception of rewards and probabilities.¹¹

Accordingly, we assume that the perception of lotteries evolved in a stable and simple environment where both encoding and decoding are jointly optimized. We show that both perception features of prospect theory are jointly optimal adaptations to natural distributions of decision problems, regardless of the fine details of these distributions. This robustness to the details of distributions is important for two reasons. First, it allows us to study the evolutionary process without needing detailed knowledge of the ancestral environment. Second, it suggests that these perception patterns have evolved as heuristics that both human and non-human species apply rigidly without continuous readaptation.

Our foundation for prospect-theory-like behavior is ultimately based on a maladaptation argument. We argue that prospect theory perception patterns, evolved for a transparent ancestral environment where decoding was optimal, are also applied as heuristics in more complex environments with finer distinctions and potentially coarse decoding. Some form of maladaptation underlies the economic literature that derives choice predictions from psychophysics arguments. The papers by Robson (2001) and Netzer (2009) derive optimal perception of riskless rewards but are motivated by prospect theory for risky prospects. Similarly, Frydman and Jin (2022) test the adaptation of perception of riskless rewards in an experiment involving risky prospects. Our approach assumes a somewhat distinct form of maladaptation. Instead of assuming that the evolutionary process ignored the lottery structure, we posit that the decision process does not adapt to novel, finer distinctions in the rewards that arise subsequently.

¹¹Marsh and Kacelnik (2002), Chen et al. (2006), and Lakshminarayanan et al. (2011) document the S-shaped evaluation of rewards in monkeys and birds. Stauffer et al. (2015), Ferrari-Toniolo et al. (2019), and Ferrari-Toniolo et al. (2022) show that monkeys’ choices over risky prospects are best explained by a probability weighting function that overweights small probabilities.

Our evolutionary optimization takes the form of an attention allocation problem. Increasing the sampling frequency of a state enhances the DM’s attention to its reward but reduces attention to other states. Similarly, steepening the encoding function near a reward value increases attention to this area but results in increased perception error elsewhere. We show below that, in the limit of many signals, the DM’s loss equals the mean squared error of her estimate of the lottery value, averaged over the pivotal decision problems where perception matters, even when the data volume is large. We then optimize the encoding strategy. The derived loss-minimizing perception generalizes results in the literature (Robson, 2001; Netzer, 2009; Woodford, 2012; Payzan-LeNestour and Woodford, 2021) and, under reasonable assumptions, entails an S-shaped encoding function and oversampling of low-probability states.

3.1 Objective

We fix a partition $\mathcal{P}$ of the state space and assume that it accurately describes the environment when the encoding strategy is optimized, i.e., all lotteries are measurable with respect to $\mathcal{P}$. As a result, the decoding process utilizes a well-specified statistical model of the perceptual data. We interpret this optimization process as a form of evolutionary selection.

Since the distinction between states within each $J \in \mathcal{P}$ is redundant, we treat J as an index of a state, refer to the rewards in states $i \in J$ simply as r_J , and model the entire lottery $\mathbf{r} = (r_J)_{J \in \mathcal{P}}$ as having $|\mathcal{P}|$ rewards, each with probability $p_J = \sum_{i \in J} p_i$. Recall that an encoding strategy consists of the encoding function $m(\cdot)$ and positive sampling frequencies $(\pi_J)_J$.¹² The strategy is optimized ex ante for a given distribution of decision problems. Specifically, the rewards r_J are assumed to be iid with a continuous density h , and the safe option s is drawn from a continuous density h_s independently of the lottery rewards. Both densities have support on $[\underline{r}, \bar{r}]$.

Let $r = \sum_J p_J r_J$ denote the lottery value, as before, and $q_n = \sum_J p_J q_{Jn}$ its maximum likelihood estimate, where each $q_{Jn} = m^{-1}(\hat{m}_{Jn})$ is the maximum-likelihood estimate of r_J given the average perturbed message $\hat{m}_{Jn}$ for state J . The DM’s ex ante expected loss relative to choice under complete information is

$$L(n) = \mathbb{E} [\max\{r, s\} - \mathbb{1}_{q_n > s} r - \mathbb{1}_{q_n \leq s} s],$$

where the expectation is over the estimate q_n and the decision problem $(\mathbf{r}, s)$.

¹²Assuming positive sampling frequencies is without loss, as it becomes optimal to gather at least some information about the rewards in every state as the number of signals increases.

Proposition 4. *The expected loss is given by*

$$\lim_{n \rightarrow \infty} nL(n) = \frac{1}{2} \mathbb{E} \left[h_s(r) \sum_J \frac{p_J^2}{\pi_J m'^2(r_J)} \right], \quad (5)$$

where the expectation is taken with respect to $\mathbf{r}$.

See Appendix A.4 for the proof. The limit loss characterization in (5) has an intuitive interpretation. It is the mean squared error (MSE) in the perception of the lottery value, integrated over all decision problems in which the true lottery value r ties with the outside option value s (multiplied by $n/2$). This conditioning on ties arises because the likelihood of large perception errors vanishes quickly as n increases. Asymptotically, only small perception errors contribute significantly to the loss, distorting choice only in decision problems in which an approximate tie arises. In the limit, the set of decision problems that contribute nontrivially to the expected loss approaches the set of problems with exact ties.¹³ To understand the relevance of the MSE for loss, fix the true and perceived lottery values. The perception error distorts choice and causes loss if and only if the safe option s falls between these two values. Hence, a mistaken choice arises with a probability proportional to the size of the perception error. Since the loss caused by such mistakes is also proportional to the error size, the expected loss becomes proportional to the MSE.

To understand the details of the characterization in (5), note that the precision of the DM's perception varies across lotteries. For large n , the DM approximately measures each reward r_J with an MSE of $1/(\pi_J n m'^2(r_J))$. This is because she observes the encoded value $m(r_J)$ with an MSE of $1/(\pi_J n)$, and the encoding function can be locally linearized for large n , with its local slope $m'(r_J)$ determining the precision of the estimate. Therefore, the sum in (5) represents the MSE of the perception of the lottery value, scaled up by n .

Motivated by the asymptotic loss characterization, we now fix a large n , use (5) to approximate the expected loss, and minimize this loss by controlling the encoding strategy. We define the *information-processing problem* as the minimization of the expected MSE, conditional on ties:

$$\min_{m'(\cdot) > 0, (\pi_J)_J > 0} \mathbb{E} \left[\sum_J \frac{p_J^2}{\pi_J m'^2(r_J)} \mid r = s \right] \quad (6)$$

¹³Steiner and Stewart (2016) use a related marginal argument in their analysis of the optimal perception of probabilities. In Herold and Netzer (2023), the negative association between rewards and probabilities in choice problems with a tie implies that a DM with an S-shaped value function makes systematic and large mistakes that can be corrected by probability weighting.

$$\text{s.t.:} \quad \int_{\underline{r}}^{\bar{r}} m'(\tilde{r}) d\tilde{r} \leq \bar{m} - \underline{m} \quad (7)$$

$$\sum_J \pi_J = 1. \quad (8)$$

The objective in (6) equals the asymptotic loss characterized in (5), up to a factor independent of the encoding strategy.¹⁴ The DM controls the derivative $m'(\cdot)$. Constraint (7) is implied by the finite range of the encoding function—it cannot be steep everywhere. Constraint (8) requires $(\pi_J)_J$ to be a probability distribution, implying that the sampling frequencies are also a scarce resource. We say that an encoding strategy $(m(\cdot), (\pi_J)_J)$ is optimal if $(m'(\cdot), (\pi_J)_J)$ solves the information-processing problem.

3.2 Optimization

We say that a density $f(x)$ on $[\underline{r}, \bar{r}]$ is *unimodal* and *symmetric* around the mode $r_m = (\underline{r} + \bar{r})/2$ if it is strictly decreasing on $(r_m, \bar{r}]$ and $f(r_m + y) = f(r_m - y)$ for all y .¹⁵

Proposition 5. *If the densities h and h_s are unimodal and symmetric, then the optimal encoding strategy has the following properties:*

1. *The encoding function is S-shaped: strictly convex below and strictly concave above r_m . Additionally, it is continuously differentiable and symmetric: $m'(r_m + y) = m'(r_m - y)$ for all y .*
2. *The DM oversamples low-probability states: for any two states J, J' such that $p_J < p_{J'}$, it holds that $\frac{\pi_J}{p_J} > \frac{\pi_{J'}}{p_{J'}}$.*

The proof of Proposition 5 in Appendix A.5 derives the first-order conditions for the

¹⁴The conditional expectation in (6),

$$\mathbb{E} \left[\sum_J \frac{p_J^2}{\pi_J m'^2(r_J)} \mid r = s \right] = \frac{\mathbb{E} \left[h_s(r) \sum_J \frac{p_J^2}{\pi_J m'^2(r_J)} \right]}{\mathbb{E}[h_s(r)]},$$

coincides with the asymptotic loss from (5) up to the ex ante likelihood of a tie, $\mathbb{E}[h_s(r)]$, and the scaling factor $2/n$. Conditioning on ties can be ignored in the special case when s is uniformly distributed, because conditional and unconditional MSE then coincide up to a constant. Earlier work has assumed the minimization of unconditional MSE (e.g., Woodford, 2012).

¹⁵Symmetry is sufficient but not necessary for the statement in Proposition 5. Our proof exploits that symmetry combined with unimodality is preserved by summation. Then, the distribution of each reward r_J conditional on a tie is unimodal. Unimodality in the absence of symmetry is generally not preserved by summation. Note that if the safe option is the value of an alternative lottery with rewards drawn from h , then unimodality and symmetry of h imply unimodality and symmetry of h_s .

information-processing problem for general distributions and then leverages unimodality and symmetry.

The proposition does not specify how the agent allocates sampling frequencies across states she does not distinguish during evolutionary adaptation. However, since the results of the proposition are robust to environmental details, it is plausible that decision-makers develop perceptual heuristics involving S-shaped reward encoding and oversampling of rare contingencies in modern times. When combined with coarseness frictions in decoding, these heuristics imply behavioral patterns consistent with prospect theory.

The intuition for the first statement of the proposition is that, to minimize the loss, the optimal encoding function is steep in the range of rewards that frequently occur at ties. The novel challenge in our setting is to prove that the reward in each state, conditional on a tie, inherits the unimodality property of its unconditional distribution, which then implies the S-shape of the encoding function. Our solution, when restricted to riskless lotteries with a single state, coincides with the optimal encoding from Netzer (2009). We thus extend his result to the perception of nontrivial risk (see Lemma 4 in Appendix A.5 for details).

The second statement of the proposition depends crucially on our conditioning on a tie. While rewards are assumed to be iid across the states unconditionally, they are no longer identically distributed conditional on a tie. The tie condition $\sum_J p_J r_J = s$ is relatively uninformative about rewards in low-probability states; hence the conditional reward distributions for these states are more spread out compared to those for high-probability states. In simple terms, a tail reward in a high-probability state makes the lottery value extreme and unlikely to tie with the outside option. Once conditioned on a tie, more tail rewards are encountered in low- than in high-probability states. Because the optimal encoding function is relatively flat at the tails, the DM expects to measure the rewards for low-probability states relatively poorly. As a consequence, she finds it optimal to oversample low-probability states.¹⁶ In particular, for binary lotteries, $\pi_J > p_J$ for the state with probability $p_J < 1/2$ and vice versa for the high-probability state. Had the DM minimized the unconditional MSE, this effect would not have arisen. By focusing on the consequences of perception errors instrumental for payoff maximization, we obtain an objective that conditions on ties and induces nontrivial sampling frequencies as an optimal adaptation.

The following example provides a numerical illustration of Proposition 5.

Example 3. The DM chooses between a binary lottery $(p, r_1; 1 - p, r_2)$ and a safe option s ,

¹⁶In Woodford (2012), a decision-maker collects information about multi-attribute objects through an information channel and allocates channel capacity across different attributes. It is optimal to allocate more capacity to attributes with exogenously higher prior variance. In our analysis, the more spread-out reward distributions for low-probability states arise endogenously in pivotal decision problems, and the (jointly optimal) S-shaped value function is crucial for making oversampling optimal.

with r_1 , r_2 , and s independently drawn from the standard normal distribution truncated to the interval $[-4, 4]$. The case with $p = 1$, serving as a benchmark here, corresponds to the choice over two iid-generated riskless values and was studied in a related model by Netzer (2009). See Appendix A.6 for details on the numerical optimization used to obtain the following results.

Figure 1a in the Introduction depicts the optimal encoding functions $m_p(r)$ for the benchmark case with $p = 1$ and for non-degenerate risk with $p = 0.2$. Consistent with Proposition 5, $m_p(r)$ is S-shaped regardless of p . Thus, the existing results on optimal reward perception for riskless choices extend to the perception of lotteries.

Figure 1b depicts the optimal sampling frequency π_p of a state as a function of the state’s objective probability p . For this, we vary the objective probability $p \in [0, 1]$ of the first state (and the complementary probability $1 - p$). For each p , we jointly optimize the sampling frequency π_p of the first state (which determines the complementary sampling frequency $1 - \pi_p$) and the encoding function $m_p(r)$. Again consistent with Proposition 5, the DM oversamples the low-probability state.

The example reveals a qualitative difference between the encoding functions $m_p(r)$ for the riskless benchmark ($p = 1$) and for non-degenerate lotteries ($0 < p < 1$). The encoding derived for risky lotteries is less S-shaped than that of the benchmark. Specifically, $m_1(r) < m_p(r)$ for $r \in (-4, 0)$ and $m_1(r) > m_p(r)$ for $r \in (0, 4)$.¹⁷ Thus, the optimal perception of rewards is context-dependent, which may inform empirical quantitative neuroeconomic studies focusing on reward perception, such as that of Schaffner et al. (2023). The comparison of the optimal encoding functions $m_p(r)$ across values of p is intuitive. The encoding function is steep at reward values that are relatively likely to appear conditional on a tie. In the benchmark case of a degenerate lottery, the event of a tie provides more information about the rewards than with a non-degenerate lottery. This is because, in non-degenerate lotteries, a tail reward value in one state can be offset by an opposite tail reward in the other state. This results in relatively more dispersed rewards conditional on a tie and, consequently, a relatively flatter encoding function. $\triangle$

4 Literature

We build on a rich literature in neuroscience and economics, making two distinct contributions. First, we clarify the role of misspecification in the behavioral consequences of the perception strategy when the stakes are large relative to perception frictions. Second, we jointly optimize both encoding of the lottery rewards and their sampling frequencies.

¹⁷We verified this claim numerically for all (p, r) from a discrete 1000×1000 grid over the set $(0, 1) \times [-4, 4]$.

For the first contribution, we apply the statistical results of Berk (1966) and White (1982) regarding the asymptotic outcomes of misspecified Bayesian and maximum-likelihood estimation. The concept of Berk-Nash equilibrium in Esponda and Pouzo (2016), defined as a fixed point of misspecified learning, has motivated renewed interest in misspecification across economics. Heidhues et al. (2018) characterize a vicious circle of overconfident learning, Molavi (2019) studies the macroeconomic consequences of misspecification, Frick et al. (2024) rank the short- and long-run costs of various forms of misspecification, and Eliaz and Spiegler (2020) focus on the political-economy consequences of misspecification. We study the interplay of encoding and misspecified decoding of rewards, using the classical results from the misspecification literature to derive behavioral risk attitudes.

Our results contrast with Savage (1954). In his discussion of small world models, Savage argues that a coarse representation of the complex grand world does not necessarily distort behavior. Savage’s argument assumes that the DM learns the correct average reward for each element of the coarse state space partition. Our approach departs from Savage in that we explicitly model the process of learning. We argue that the DM is unlikely to learn the correct average rewards for each element of her partition. If she learns within the small world model, then, instead of the average reward, her estimate converges to the certainty equivalent under her encoding function and subjective probabilities equal to her sampling frequencies.

Salant and Rubinstein (2008) and Bernheim and Rangel (2009) provide a revealed-preference theory of the behavioral and welfare implications of frames—payoff-irrelevant aspects of decision problems. We provide an account of how a specific frame—anticipation of the risk structure—affects choice and welfare. As in Kahneman, Wakker, and Sarin (1997), our model implies a distinction between decision and welfare utilities. In the case of the misspecified DM, the gap between the decision utility that she anticipates from the lottery and the welfare utility—the true expected lottery reward—may be large. Our model facilitates an analysis of systematic mistakes in decision making as outlined in Koszegi and Rabin (2008) and, specifically for framing effects, in Benkert and Netzer (2018).

Our second contribution ultimately stems from psychophysics, a field that originated in Fechner’s (1860) study of stochastic perceptual comparisons based on Weber’s data.¹⁸ A large literature in brain sciences and psychology regards perception as information processing via a limited channel and studies the optimal encoding of stimuli for a given channel capacity. Laughlin (1981) derives and tests the hypothesis that optimal encoding under an information-theoretic objective encodes random stimuli with neural activities proportional

¹⁸Woodford (2020) provides a review of psychophysics from an economics perspective.

to their cumulative distribution value, implying S-shaped encoding for unimodal densities.¹⁹

Neuroscience studies encoding adaptations under various optimization objectives, such as maximizing mutual information between the stimulus and its perception, maximizing Fisher information, or minimizing the mean squared error of perception.²⁰ Economics can contribute by providing microfoundations for the most appropriate optimization objective for perceptions related to choice. Robson (2001) studied reward encoding that minimizes the probability of a wrong choice and has shown that, in the limit of vanishing perception frictions, the optimal encoding function likewise coincides with the cumulative distribution function of rewards in the decision environment. Netzer (2009) studied maximization of the expected chosen reward, an objective rooted in the instrumental approach of economics to information. The optimal encoding function still follows the cumulative distribution function but is flattened. Schaffner et al. (2023) report that Netzer’s optimal encoding function provides a better fit to neural data than encodings derived under competing objectives.

These models examine choices involving riskless prizes; therefore, the derived encoding functions do not directly apply to choices involving gambles. Indeed, in that literature, encoding functions are often interpreted as hedonic anticipatory utilities rather than decision utilities (see Rayo and Becker, 2007).²¹ We extend Netzer’s instrumental approach to choices involving gambles and find a connection to one of the aforementioned reduced-form objectives.²² That is, in the limit of rich perceptual data, maximization of the expected chosen reward is equivalent to minimization of the expected mean squared error in the perceived lottery value, conditioned on a tie. This conditioning not only generates a better fit for the optimal encoding function documented by Schaffner et al. (2023) but it is also crucial for the result of optimal oversampling of low-probability contingencies. This oversampling would not arise under reduced-form objectives that maximize unconditional measures of precision.²³

Several recent papers examine risk attitudes arising from noisy reward encoding. Khaw

¹⁹See Attneave (1954) and Barlow (1961) for early contributions and Heng et al. (2020) for recent work.

²⁰See e.g., Bethge et al. (2002) and Wang et al. (2016).

²¹The optimal hedonic utility function of Rayo and Becker (2007) is a step function. They provide an extension in which this function becomes S-shaped. Robson et al. (2023) is a dynamic version of Robson (2001) and Netzer (2009) that captures low-rationality, real-time adaptation of a hedonic utility function over riskless rewards. Friedman (1989) is an early approach dealing with gambles.

²²Our model differs from Robson (2001) and Netzer (2009) in regard to perception frictions. Those papers model frictions as minimal just noticeable differences, while here we rely on the modeling framework of Thurstone (1927) who hypothesized that perception is a Gaussian perturbation of an encoded stimulus. Payzan-LeNestour and Woodford (2021) have shown that the Gaussian approach yields the same limiting results as in Robson (2001) and Netzer (2009).

²³Steiner and Stewart (2016) find probability weighting to be an optimal correction for naive noisy information processing. Herold and Netzer (2023) derive probability weighting as the optimal correction for an exogenous, distortive S-shaped value function. Lieder et al. (2017) argue that a contingency should be oversampled if it has extreme payoff consequences and decisions are based on a small sample. The present paper derives both S-shaped encoding and low-probability oversampling in a joint optimization.

et al. (2021) show theoretically and verify experimentally that exogenous logarithmic stochastic encoding and Bayesian decoding generate risk attitudes compatible with the paradox of Rabin (2000) stemming from an effect akin to reversion to the mean. In Vieider (2024), probabilities are also exogenously encoded in a logarithmic manner, establishing a link to prospect theory. Frydman and Jin (2022) and Juechems et al. (2021) optimize the encoding of lottery rewards and demonstrate both theoretically and experimentally that this encoding adapts to distribution of the decision problems, affecting choice. Compared to these papers, we analyze the optimal encoding of rewards alongside the optimization of sampling frequencies. We also differ in the proposed source of behavioral distortions. The discussed models assume well-specified learning, approximating the frictionless benchmark as noise diminishes. We focus on the limit of small encoding noise right away. This focus uncovers a novel connection between coding and behavior. While the impact of coding on behavior necessarily vanishes when the decoding model is well-specified, as in previous literature, the implications for behavior remain substantial if the cognitive model used for decoding oversimplifies the risk structure.

5 Conclusion

We develop a model of perception and aggregation of lotteries. Our model includes an *encoding* stage, during which the DM generates signals about lottery rewards. Following this, there is a *decoding* stage where the DM estimates the lottery’s value based on the signals, employing a more or less sophisticated estimation procedure.

The behavioral impact of encoding vanishes for rich perceptual data if the DM encounters a lottery that she has anticipated and that she therefore decodes using a well-specified estimation procedure. Conversely, encoding-induced behavioral risk attitudes arise for lotteries that the DM has not anticipated and therefore decodes using a misspecified estimation procedure. In the latter case, our model provides a unified explanation for multiple well-documented empirical patterns: adaptive risk preferences, varying risk attitudes for small and large risks, probability weighting, availability heuristics, the role of salience, time pressure, experience, and behavioral risk attitudes toward safe prospects arising due to aggregation frictions. Additionally, we derive the properties of optimal encoding in a model with well-specified decoding and show that the optimal encoding strategy exhibits S-shaped reward encoding and oversampling of low-probability states.

A natural and feasible extension of our model would replace the safe option with a lottery, letting the DM choose between two risky prospects. Our results from Propositions 1 and 2 generalize to this setting, with the friction in the decoding stage again specified either by

a partitional model or a concentrated prior over the two lotteries. Only if the partitional model or the prior did not treat the two lotteries independently would new interaction effects between the perceptions of the two lotteries arise. Our result on the evolutionary optimal encoding strategy can also accommodate an extension to two lotteries. This is because the safe option s may represent, in reduced form, the value of the alternative lottery. If the lottery rewards are unimodally distributed, the resulting s is also unimodal, as required for Proposition 5.

An interesting question left open by our paper concerns the properties of optimal encoding when decoding is anticipated to be misspecified. Even though evolutionary processes are typically not forward-looking, there may be benefits to a robust encoding solution that performs well when the DM acts in changing environments for which she is misspecified. When considering the limit of rich perceptual data, as we did in this paper, large mistakes from misspecification will always outweigh small mistakes from perception errors, constituting a force toward linear encoding functions (Rustichini et al., 2017) that mitigate misspecification bias. A more suitable framework for studying optimal encoding with misspecified decoding may be a Bayesian model with prior information, as in our Subsection 2.2, where perception errors and misspecification errors remain of comparable size. Optimal perception of lotteries in the presence of large encoding noise and with the anticipation of one’s own decoding frictions remains a challenging open problem that may require new conceptual breakthroughs.

References

- Alós-Ferrer, C., E. Fehr, and N. Netzer (2021). Time will tell: Recovering preferences when choices are noisy. *Journal of Political Economy* 129(6), 1828–1877.
- Attneave, F. (1954). Some informational aspects of visual perception. *Psychological review* 61(3), 183.
- Barlow, H. B. (1961). Possible principles underlying the transformation of sensory messages. *Sensory communication* 1(01).
- Benkert, J.-M. and N. Netzer (2018). Informational requirements of nudging. *Journal of Political Economy* 126(6), 2323–2355.
- Berk, R. H. (1966). Limiting behavior of posterior distributions when the model is incorrect. *The Annals of Mathematical Statistics*, 51–58.

- Bernheim, B. D. and A. Rangel (2009). Beyond revealed preference: choice-theoretic foundations for behavioral welfare economics. *The Quarterly Journal of Economics* 124(1), 51–104.
- Bethge, M., D. Rotermund, and K. Pawelzik (2002). Optimal short-term population coding: when fisher information fails. *Neural computation* 14(10), 2317–2351.
- Birnbaum, Z. W. (1948). On random variables with comparable peakedness. *The Annals of Mathematical Statistics* 19(1), 76–81.
- Bordalo, P., N. Gennaioli, and A. Shleifer (2012). Salience theory of choice under risk. *The Quarterly Journal of Economics* 127(3), 1243–1285.
- Bradbury, M., T. Hens, and S. Zeisberger (2015). Improving investment decisions with simulated experience. *Review of Finance* 19, 1019–1052.
- Cahlíková, J. and L. Cingl (2017). Risk preferences under acute stress. *Experimental Economics* 20(1), 209–236.
- Charness, G., N. Chemaya, and D. Trujano-Ochoa (2023). Learning your own risk preferences. *Journal of Risk and Uncertainty* 67, 1–19.
- Chen, M. K., V. Lakshminarayanan, and L. R. Santos (2006). How basic are behavioral biases? evidence from capuchin monkey trading behavior. *Journal of Political Economy* 114(3), 517–537.
- Eliaz, K. and R. Spiegel (2020). A model of competing narratives. *American Economic Review* 110(12), 3786–3816.
- Ert, E. and E. Haruvy (2017). Revisiting risk aversion: Can risk preferences change with experience? *Economics Letters* 151, 91–95.
- Esponda, I. and D. Pouzo (2016). Berk–nash equilibrium: A framework for modeling agents with misspecified models. *Econometrica* 84(3), 1093–1130.
- Fechner, G. T. (1860). Elements of psychophysics.
- Ferrari-Toniolo, S., P. M. Bujold, and W. Schultz (2019). Probability distortion depends on choice sequence in rhesus monkeys. *Journal of Neuroscience* 39(15), 2915–2929.
- Ferrari-Toniolo, S., L. C. U. Seak, and W. Schultz (2022). Risky choice: probability weighting explains independence axiom violations in monkeys. *Journal of Risk and Uncertainty* 65(3), 319–351.

- Frick, M., R. Iijima, and Y. Ishii (2024). Welfare comparisons for biased learning. *American Economic Review* 114, 1612–1649.
- Friedman, D. (1989). The S-shaped value function as a constrained optimum. *American Economic Review* 79(5), 1243–1248.
- Frydman, C. and L. J. Jin (2022). Efficient coding and risky choice. *Quarterly Journal of Economics* 137(1), 161–213.
- Fudenberg, D., P. Strack, and T. Strzalecki (2018). Speed, accuracy, and the optimal timing of choices. *American Economic Review* 108(12), 3651–3684.
- Heidhues, P., B. Köszegi, and P. Strack (2018). Unrealistic expectations and misguided learning. *Econometrica* 86(4), 1159–1214.
- Heng, J. A., M. Woodford, and R. Polania (2020). Efficient sampling and noisy decisions. *Elife* 9, e54962.
- Herold, F. and N. Netzer (2023). Second-best probability weighting. *Games and Economic Behavior* (138), 112–125.
- Jehiel, P. (2005). Analogy-based expectation equilibrium. *Journal of Economic Theory* 123(2), 81–104.
- Johnson, D. H. and G. C. Orsak (1993). Relation of signal set choice to the performance of optimal non-gaussian detectors. *IEEE Transactions on Communications* 41(9), 1319–1328.
- Juechems, K., J. Balaguer, B. Spitzer, and C. Summerfield (2021). Optimal utility and probability functions for agents with finite computational precision. *Proceedings of the National Academy of Sciences* 118(2).
- Kahneman, D. (2011). *Thinking, fast and slow*. Macmillan.
- Kahneman, D. and A. Tversky (1979). Prospect theory: An analysis of decision under risk. *Econometrica* 47(2), 263–291.
- Kahneman, D., P. P. Wakker, and R. Sarin (1997). Back to Bentham? Explorations of experienced utility. *The Quarterly Journal of Economics* 112(2), 375–406.
- Khaw, M. W., Z. Li, and M. Woodford (2021). Cognitive imprecision and small-stakes risk aversion. *The Review of Economic Studies* 88, 1979–2013.

- Kirchler, M., D. Andersson, C. Bonn, M. Johannesson, E. Ø. Sørensen, M. Stefan, G. Tinghög, and D. Västfjäll (2017). The effect of fast and slow decisions on risk taking. *Journal of Risk and Uncertainty* 54(1), 37–59.
- Kocher, M. G., J. Pahlke, and S. T. Trautmann (2013). Tempus fugit: time pressure in risky decisions. *Management Science* 59(10), 2380–2391.
- Köszegi, B. and M. Rabin (2007). Reference-dependent risk attitudes. *American Economic Review* 97, 1047–1073.
- Koszegi, B. and M. Rabin (2008). Revealed mistakes and revealed preferences. In A. Caplin and A. Schotter (Eds.), *The Foundations of Positive and Normative Economics: A Handbook*, pp. 193–209. Oxford University Press.
- Koszegi, B. and A. Szeidl (2013). A model of focusing in economic choice. *Quarterly Journal of Economics* 128, 53–104.
- Lakshminarayanan, V. R., M. K. Chen, and L. R. Santos (2011). The evolution of decision-making under risk: Framing effects in monkey risk preferences. *Journal of Experimental Social Psychology* 47(3), 689–693.
- Laughlin, S. (1981). A simple coding procedure enhances a neuron’s information capacity. *Zeitschrift für Naturforschung c* 36(9-10), 910–912.
- Lieder, F., T. Griffiths, and M. Hsu (2017). Overrepresentation of extreme events in decision making reflects rational use of cognitive resources. *Psychological Review* 125, 1–32.
- Lindvall, T. (2002). *Lectures on the coupling method*. Courier Corporation.
- Liu, S. and N. Netzer (2023). Happy times: Measuring happiness using response times. *American Economic Review* 113, 3289–3322.
- Ludvig, E., C. Madan, and M. Spetch (2015). Priming memories of past wins induces risk seeking. *Journal of Experimental Psychology: General* 144, 24–29.
- Marsh, B. and A. Kacelnik (2002). Framing effects and risky decisions in starlings. *Proceedings of the National Academy of Sciences* 99(5), 3352–3355.
- Molavi, P. (2019). Macroeconomics with learning and misspecification: A general theory and applications. *Mimeo*.
- Netzer, N. (2009). Evolution of time preferences and attitudes toward risk. *The American Economic Review* 99(3), 937–955.

- Ok, E. A. (2011). *Real analysis with economic applications*. Princeton University Press.
- Oprea, R. (2023). Simplicity equivalents. Mimeo.
- Payzan-LeNestour, E. and M. Woodford (2021). Outlier blindness: A neurobiological foundation for neglect of financial risk. *Journal of Financial Economics* (in press).
- Porcelli, A. J. and M. R. Delgado (2009). Acute stress modulates risk taking in financial decision making. *Psychological Science* 20(3), 278–283.
- Purkayastha, S. (1998). Simple proofs of two results on convolutions of unimodal distributions. *Statistics & probability letters* 39(2), 97–100.
- Rabin, M. (2000). Risk aversion and expected-utility theory: A calibration theorem. *Econometrica* 68(5), 1281–1292.
- Rayo, L. and G. S. Becker (2007). Evolutionary efficiency and happiness. *Journal of Political Economy* 115(2), 302–337.
- Robson, A. J. (2001). The biological basis of economic behavior. *Journal of Economic Literature* 39(1), 11–33.
- Robson, A. J., L. A. Whitehead, and N. Robalino (2023). Adaptive utility. *Journal of Economic Behavior and Organization* 211, 60–81.
- Rustichini, A., K. Conen, X. Cai, and C. Padoa-Schioppa (2017). Optimal coding and neuronal adaptation in economic decisions. *Nature Communications* 8, 1–14.
- Salant, Y. and A. Rubinstein (2008). (a, f): choice with frames. *The Review of Economic Studies* 75(4), 1287–1296.
- Savage, L. J. (1954). *The Foundations of Statistics*. John Wiley & Sons.
- Schaffner, J., S. D. Bao, P. N. Tobler, T. A. Hare, and R. Polania (2023). Sensory perception relies on fitness-maximizing codes. *Nature Human Behavior* 7, 1135–1151.
- Stauffer, W. R., A. Lak, P. Bossaerts, and W. Schultz (2015). Economic choices reveal probability distortion in macaque monkeys. *Journal of Neuroscience* 35(7), 3146–3154.
- Steiner, J. and C. Stewart (2016). Perceiving prospects properly. *The American Economic Review* 106(7), 1601–31.
- Thurstone, L. L. (1927). A law of comparative judgment. *Psychological review* 34(4), 273.

- Tversky, A. and D. Kahneman (1973). Availability: A heuristic for judging frequency and probability. *Cognitive Psychology* 5, 207–232.
- Vieider, F. (2024). Decisions under uncertainty as bayesian inference on choice options. *Management Science*, forthcoming.
- Wald, A. (1949). Note on the consistency of the maximum likelihood estimate. *The Annals of Mathematical Statistics* 20(4), 595–601.
- Wang, Z., A. A. Stocker, and D. D. Lee (2016). Efficient neural codes that minimize lp reconstruction error. *Neural computation* 28(12), 2656–2686.
- White, H. (1982). Maximum likelihood estimation of misspecified models. *Econometrica: Journal of the Econometric Society*, 1–25.
- Woodford, M. (2012). Prospect theory as efficient perceptual distortion. *American Economic Review, Papers & Proceedings* 102(3), 41–46.
- Woodford, M. (2020). Modeling imprecision in perception, valuation, and choice. *Annual Review of Economics* 12, 579–601.

A Proofs

A.1 Proof of Proposition 1

Let $f_{\mathbf{r}}(x)$ be the signal density conditional on the encountered lottery $\mathbf{r}$. That is, for signal $x = (\hat{m}, i)$, $f_{\mathbf{r}}(x) = \pi_i \varphi(\hat{m} - m(r_i))$ where φ is the standard normal density. The Kullback-Leibler divergence of the signal densities for any two lotteries $\mathbf{r}, \mathbf{r}'$ is

$$\begin{aligned}
D_{\text{KL}}(f_{\mathbf{r}} \parallel f_{\mathbf{r}'}) &= \int_{\mathbb{R} \times \{1, \dots, I\}} f_{\mathbf{r}}(x) \ln \frac{f_{\mathbf{r}}(x)}{f_{\mathbf{r}'}(x)} dx \\
&= \sum_{i=1}^I \int_{\mathbb{R}} \pi_i \varphi(\hat{m} - m(r_i)) \ln \frac{\pi_i \varphi(\hat{m} - m(r_i))}{\pi_i \varphi(\hat{m} - m(r'_i))} d\hat{m} \\
&= \sum_{i=1}^I \pi_i \int_{\mathbb{R}} \varphi(\hat{m} - m(r_i)) \ln \frac{\varphi(\hat{m} - m(r_i))}{\varphi(\hat{m} - m(r'_i))} d\hat{m} \\
&= \sum_{i=1}^I \pi_i D_{\text{KL}}(\varphi_{m(r_i)} \parallel \varphi_{m(r'_i)})
\end{aligned}$$

$$= \frac{1}{2} \sum_{i=1}^I \pi_i (m(r_i) - m(r'_i))^2,$$

where $\varphi_{\tilde{m}}(\hat{m}) = \varphi(\hat{m} - \tilde{m})$ is the density of the perturbed message $\hat{m}$ conditional on the unperturbed message $\tilde{m}$. The last equality follows from the fact that the Kullback-Leibler divergence of two Gaussian densities with means μ_1, μ_2 and variances equal to 1 is $(\mu_1 - \mu_2)^2/2$ (see e.g., Johnson and Orsak, 1993).

Let

$$\mathbf{r}^* = \arg \min_{\mathbf{r}' \in \mathcal{A}_{\mathcal{P}}} D_{\text{KL}}(f_{\mathbf{r}} \parallel f_{\mathbf{r}'}) = \arg \min_{\mathbf{r}' \in \mathcal{A}_{\mathcal{P}}} \sum_{i=1}^I \pi_i (m(r_i) - m(r'_i))^2.$$

This minimizer $\mathbf{r}^* = (r_i^*)_i$ is unique and satisfies, for each state $i = 1, \dots, I$,

$$\begin{aligned} m(r_i^*) &= \arg \min_{m \in [\underline{m}, \bar{m}]} \sum_{j \in J(i)} \pi_j (m(r_j) - m)^2 \\ &= \sum_{j \in J(i)} \frac{\pi_j}{\pi_{J(i)}} m(r_j), \end{aligned}$$

where $J(i)$ is the element of the partition $\mathcal{P}$ that contains i and $\pi_{J(i)} = \sum_{j \in J(i)} \pi_j$. The estimated lottery value q_n almost surely converges to $\sum_{i=1}^I p_i r_i^*$, which follows from White (1982) who proves that $\mathbf{q}_n$ almost surely converges to the minimizer of the Kullback-Leibler divergence (provided the minimizer is unique).

A.2 Proof of Proposition 2

We first state a lemma that will be useful for proving Proposition 2.

Lemma 1. *Let $\psi_n(\mathbf{x}) : [\underline{r}, \bar{r}]^I \rightarrow \mathbb{R}$ be a sequence of continuous functions uniformly converging to a function $\psi(\mathbf{x})$ which has a unique minimizer $\mathbf{x}^*$. Then, the random variable X_n with PDF equal to $\alpha_n \exp(-n\psi_n(\mathbf{x}))$, where α_n is the normalization factor, converges to $\mathbf{x}^*$ in probability as $n \rightarrow \infty$.*

Proof. We need to prove that for every $\delta > 0$, the probability $P(X_n \in B_\delta) \rightarrow 1$ as $n \rightarrow \infty$, where B_δ is the open Euclidean δ -ball centered at $\mathbf{x}^*$. Fix $\delta > 0$ and define

$$d = \min_{\mathbf{x} \in [\underline{r}, \bar{r}]^I \setminus B_\delta} \{\psi(\mathbf{x}) - \psi(\mathbf{x}^*)\}.$$

The minimum exists as ψ is continuous and the set $[\underline{r}, \bar{r}]^I \setminus B_\delta$ is closed. Additionally, $d > 0$ since $\mathbf{x}^*$ is the unique minimizer of ψ on $[\underline{r}, \bar{r}]^I$.

Because the convergence $\psi_n \rightarrow \psi$ is uniform, for any $d' > 0$ there exists $n_{d'} \in \mathbb{N}$ such that $|\psi_n(\mathbf{x}) - \psi(\mathbf{x})| < d'$ for all $\mathbf{x} \in [\underline{r}, \bar{r}]^I$ and $n \geq n_{d'}$. Consider $n \geq n_{d/4}$. Because $\psi_n(\mathbf{x}) \geq \psi(\mathbf{x}) - \frac{d}{4} \geq \psi(\mathbf{x}^*) + \frac{3d}{4}$ for $\mathbf{x}$ outside of the ball B_δ , the probability density of X_n is at most $\alpha_n \exp(-n\psi(\mathbf{x}^*) - \frac{3d}{4}n)$. This implies,

$$P(X_n \notin B_\delta) \leq \tilde{\alpha}_n \exp\left(-\frac{3d}{4}n\right) (\bar{r} - \underline{r})^I, \quad \text{where } \tilde{\alpha}_n := \alpha_n \exp(-n\psi(\mathbf{x}^*)). \quad (9)$$

We conclude by establishing an upper bound for $\tilde{\alpha}_n$. Given $\delta > 0$, let $\delta' > 0$ be such that $\psi(\mathbf{x}) \leq \psi(\mathbf{x}^*) + d/4$ for all $\mathbf{x} \in B_{\delta'} \cap [\underline{r}, \bar{r}]^I$. Existence of such δ' follows from the continuity of ψ . Then, $\psi_n(\mathbf{x}) \leq \psi(\mathbf{x}) + \frac{d}{4} \leq \psi(\mathbf{x}^*) + \frac{d}{2}$ for all $\mathbf{x} \in B_{\delta'} \cap [\underline{r}, \bar{r}]^I$ and $n > n_{d/4}$. Thus the probability density of X_n is at least $\tilde{\alpha}_n \exp(-\frac{d}{2}n)$ on this set. It follows that,

$$1 \geq P(X_n \in B_{\delta'}) \geq \tilde{\alpha}_n \exp\left(-\frac{d}{2}n\right) b',$$

where $b' > 0$ is the volume of the set $B_{\delta'} \cap [\underline{r}, \bar{r}]^I$. Substituting the implied upper bound on $\tilde{\alpha}_n$ into (9) gives

$$P(X_n \notin B_\delta) \leq \exp\left(-\frac{d}{4}n\right) \frac{(\bar{r} - \underline{r})^I}{b'}.$$

The right side vanishes as $n \rightarrow \infty$, so the claim follows. $\square$

We now prove Proposition 2. Let $\hat{\mathbf{m}}_n = (\hat{m}_{in})_{i=1}^I$ be the tuple of the averages of $a\pi_i n$ perturbed messages received for each state i . Since the encoding errors are standard normal, this tuple of averages is a sufficient statistic for the Bayesian estimation, and we have $\hat{m}_{in} \sim \mathcal{N}\left(m(r_i), \frac{1}{a\pi_i n}\right)$. By Bayes' Rule, the posterior density of each lottery $\mathbf{r}' \in [\underline{r}, \bar{r}]^I$ is, for given $\hat{\mathbf{m}}_n$, proportional to

$$\varrho_n(\mathbf{r}') \prod_{i=1}^I \varphi\left((\hat{m}_{in} - m(r'_i))\sqrt{a\pi_i n}\right) \propto \exp\left(-n\psi(\mathbf{r}'; \hat{\mathbf{m}}_n)\right),$$

where $\propto$ denotes equality modulo normalization and

$$\psi(\mathbf{r}; \hat{\mathbf{m}}) := \frac{1}{2} \sum_{i=1}^I \left(\frac{\sigma^2(\mathbf{r})}{\Delta} + a\pi_i (m(r_i) - \hat{m}_i)^2 \right).$$

Throughout this paragraph, consider a fixed realization of the sequence $(\hat{\mathbf{m}}_n)_n$ such that $\hat{m}_{in} \rightarrow m(r_i)$ for all i . Then, $\psi(\mathbf{r}'; \hat{\mathbf{m}}_n)$ converges to $\psi(\mathbf{r}'; (m(r_i))_i)$, uniformly in $\mathbf{r}'$. Additionally, $\psi(\mathbf{r}'; (m(r_i))_i)$ as a function of $\mathbf{r}'$ has the unique minimizer $\mathbf{q}^*(\mathbf{r})$ by assumption. Lemma 1 implies that the posterior formed given $\hat{\mathbf{m}}_n$ converges in probability to $\mathbf{q}^*(\mathbf{r})$.

Since the support of the rewards is bounded, convergence in probability implies convergence in expected value, and thus the Bayesian estimate $\mathbb{E}[\hat{\mathbf{r}} \mid \hat{\mathbf{m}}_n] \in [\underline{r}, \bar{r}]^I$ converges to $\mathbf{q}^*(\mathbf{r})$. Since $\hat{m}_{in} \rightarrow m(r_i)$ almost surely, we conclude that $\mathbb{E}[\sum_{i=1}^I p_i \hat{r}_i \mid \hat{\mathbf{m}}_n] \in [\underline{r}, \bar{r}]$ converges to $\sum_{i=1}^I p_i q_i^*(\mathbf{r})$ almost surely. Here, $\hat{\mathbf{r}}$ and $\hat{r}_i$ stand for random variables and $\mathbf{r}$ and r_i are their realizations.

A.3 Proof of Proposition 3

By Proposition 2, the Bayesian estimate of $\mathbf{r}$ converges to $\mathbf{q}^*(\mathbf{r})$ almost surely. We write $\mathbf{q}^* = (q_i^*)_{i=1}^I$ as an abbreviation for $\mathbf{q}^*(\mathbf{r})$ and let $q^* = \sum_{i=1}^I p_i q_i^*$. The first-order condition of the minimization in (3) implies

$$(q_i^* - q^*) + a\Delta(m(q_i^*) - m(r_i))m'(q_i^*) = 0 \quad (10)$$

for all $i = 1, \dots, I$, where we have used that $\pi_i = p_i$ and $\sum_{i=1}^I p_i (q_i^* - q^*) = q^* - q^* = 0$. We write σ^2 for the true reward variance $\sigma^2(\mathbf{r})$ and $\sigma^{*2} := \sum_{i=1}^I p_i (q_i^* - q^*)^2$ for the estimated variance. We will prove the following claims (see Footnote 9 for the definition of the “order smaller than” convention $o(\cdot)$):

Claim 1: Any function that is $o(r_i - r)$ or $o(q_i^* - r)$ is also $o(\sigma)$.

Claim 2: $q^* = r + o(\sigma)$.

Claim 3: $\sigma^{*2} = \frac{z(r)^2}{(1+z(r))^2} \sigma^2 + o(\sigma^2)$.

Claim 4: $q^* = r + \frac{1}{2} \frac{m''(r)}{m'(r)} \left(\sigma^2 + \left(\frac{2}{z(r)} - 1 \right) \sigma^{*2} \right) + o(\sigma^2)$.

To prove Claim 1, we provide a bound on the distance of r_i and q_i^* from r . It follows from definition of σ^2 that $(r_i - r)^2 \leq \sigma^2/p_i$, and thus $|r_i - r| \leq \sigma/\sqrt{p_i}$. Therefore, any function that is $o(r_i - r)$ is also $o(\sigma)$. Bounding $|q_i^* - r|$ is complicated by the fact that $\mathbf{q}^*$ is defined implicitly. We first establish a bound on $|q^* - r|$. Define $\underline{m}'$ and $\overline{m}'$ to be the minimum and the maximum of $m'(\cdot)$ on $[\underline{r}, \bar{r}]$, respectively, and let $\underline{z} = a\Delta\underline{m}'^2$, $\bar{z} = a\Delta\overline{m}'^2$. We have $0 < \underline{m}' \leq \overline{m}' < +\infty$ and $0 < \underline{z} \leq \bar{z} < +\infty$ since $m'(\cdot)$ is continuous and strictly positive on the closed interval $[\underline{r}, \bar{r}]$.

For fixed values of $\mathbf{r}$ and $\mathbf{q}^*$ define $z_i \in \mathbb{R}$ by

$$a\Delta m'(q_i^*) (m(q_i^*) - m(r_i)) = (q_i^* - r_i) z_i$$

whenever $q_i^* \neq r_i$, and $z_i := a\Delta m'^2(r_i)$ otherwise. It follows from its definition that $z_i \geq \underline{z}$

for all i . Then, equation (10) can be written as

$$0 = (q_i^* - q^*) + (q_i^* - r_i)z_i = (1 + z_i)(q_i^* - q^*) - (r_i - q^*)z_i,$$

and thus,

$$q_i^* - q^* = \frac{z_i}{1+z_i}(r_i - q^*) = \frac{z_i}{1+z_i}(r_i - r) + \frac{z_i}{1+z_i}(r - q^*). \quad (11)$$

Summing up the last equation weighted by p_i over i gives

$$0 = \sum_{i=1}^I \left(p_i \frac{z_i}{1+z_i} (r_i - r) \right) + (r - q^*) \sum_{i=1}^I \left(p_i \frac{z_i}{1+z_i} \right),$$

in which $0 < \frac{\underline{z}}{1+\underline{z}} \leq \frac{z_i}{1+z_i} < 1$. The triangle inequality implies that

$$|q^* - r| \leq \frac{1+\underline{z}}{\underline{z}} \sum_{i=1}^I p_i |r_i - r| \leq \frac{1+\underline{z}}{\underline{z}} \sigma \sum_{i=1}^I \sqrt{p_i} \leq \frac{1+\underline{z}}{\underline{z}} I \sigma.$$

Returning to equation (11), we have

$$|q_i^* - r| \leq \frac{z_i}{1+z_i} |r_i - r| + \frac{z_i}{1+z_i} |r - q^*| + |q^* - r| < |r_i - r| + 2|r - q^*| \leq \left(p_i^{-1/2} + 2 \frac{1+\underline{z}}{\underline{z}} I \right) \sigma.$$

We conclude that $|q_i^* - r| \leq \left(p_i^{-1/2} + 2 \frac{1+\underline{z}}{\underline{z}} I \right) \sigma$ for any $\mathbf{r} \in [\underline{r}, \bar{r}]^I$, and thus any function that is $o(q_i^* - r)$ is also $o(\sigma)$. This establishes Claim 1.

We will prove the remaining claims by taking first- and second-order approximations of the first-order condition (10) for $\sigma > 0$ small. Since $m(\cdot)$ is twice differentiable, the functions m and m' can be expressed using first-order Taylor approximations around r :

$$\begin{aligned} m(r_i) &= m(r) + m'(r)(r_i - r) + o(\sigma), \\ m(q_i^*) &= m(r) + m'(r)(q_i^* - r) + o(\sigma), \\ m'(q_i^*) &= m'(r) + m''(r)(q_i^* - r) + o(\sigma), \end{aligned}$$

where we used Claim 1 to replace $o(r_i - r)$ and $o(q_i^* - r)$ by $o(\sigma)$. Equation (10) then implies

$$\begin{aligned} 0 &= (q_i^* - q^*) + a\Delta \left(m'(r)(q_i^* - r_i) + o(\sigma) \right) \left(m'(r) + m''(r)(q_i^* - r) + o(\sigma) \right) \\ &= (q_i^* - q^*) + a\Delta m'^2(r)(q_i^* - r_i) + o(\sigma), \end{aligned}$$

where we used that $(q_i^* - r_i)(q_i^* - r) = o(\sigma)$. The last inline equation can be written as

$$0 = (q_i^* - q^*) + z(r)(q_i^* - r_i) + o(\sigma). \quad (12)$$

Summing up these equations weighted by p_i , we get $0 = z(r)(q^* - r) + o(\sigma)$. Thus $|q^* - r| \leq \frac{1}{z} o(\sigma)$, as needed for Claim 2.

We rewrite (12) as

$$(1 + z(r))(q_i^* - q^*) = z(r)(r_i - r) + z(r)(r - q^*) + o(\sigma) = z(r)(r_i - r) + o(\sigma),$$

where the second equality follows from Claim 2. Squaring both sides of the equation and summing up the equations weighted by p_i , we get

$$(1 + z(r))^2 \sigma^{*2} = z^2(r) \sigma^2 + o(\sigma^2),$$

where we used that $z(r) \leq \bar{z}$ and thus $z(r)(r_i - r)o(\sigma)$ is $o(\sigma^2)$. Claim 3 follows.

To prove Claim 4, we use the second-order Taylor approximation of $m(\cdot)$ around r :

$$\begin{aligned} m(q_i^*) &= m(r) + m'(r)(q_i^* - r) + \frac{1}{2}m''(r)(q_i^* - r)^2 + o(\sigma^2), \\ m(r_i) &= m(r) + m'(r)(r_i - r) + \frac{1}{2}m''(r)(r_i - r)^2 + o(\sigma^2). \end{aligned}$$

This implies the second-order approximation of equation (10),

$$\begin{aligned} 0 &= (q_i^* - q^*) + a\Delta \left(m'(r)(q_i^* - r_i) + \frac{1}{2}m''(r)((q_i^* - r)^2 - (r_i - r)^2) + o(\sigma^2) \right) \\ &\quad \cdot \left(m'(r) + m''(r)(q_i^* - r) + o(\sigma) \right), \end{aligned}$$

which we rewrite as

$$0 = (q_i^* - q^*) + z(r) \left((q_i^* - r_i) + \frac{1}{2} \frac{m''(r)}{m'(r)} ((q_i^* - r)^2 - (r_i - r)^2) \right) \left(1 + \frac{m''(r)}{m'(r)} (q_i^* - r) \right) + o(\sigma^2).$$

Summing up these equations weighted by p_i and dividing by $z(r)$, we arrive at

$$0 = (q^* - r) - \frac{1}{2} \frac{m''(r)}{m'(r)} \left(\sigma^2 - \sigma^{*2} + 2 \sum_{i=1}^I p_i (r_i - q_i^*) (q_i^* - r) \right) + o(\sigma^2). \quad (13)$$

Expressing $q_i^* - r_i$ from (12) allows us to write

$$\sum_{i=1}^I p_i (r_i - q_i^*) (q_i^* - r) = \frac{1}{z(r)} \sum_{i=1}^I p_i (q_i^* - r)^2 + o(\sigma^2) = \frac{1}{z(r)} \sigma^{*2} + o(\sigma^2),$$

where we used that $r = q^* + o(\sigma)$ for the second equality. Substituting the last inline equation back into (13) completes the proof of Claim 4.

Finally, substituting for σ^{*2} from Claim 3 into the expression from Claim 4 gives

$$\begin{aligned} q^* &= r + \frac{1}{2} \frac{m''(r)}{m'(r)} \left(1 + \left(\frac{2}{z(r)} - 1 \right) \frac{z(r)^2}{(1+z(r))^2} \right) \sigma^2 + o(\sigma^2) \\ &= r + \frac{1}{2} \frac{m''(r)}{m'(r)} \left(1 + \frac{2z(r) - z(r)^2}{(1+z(r))^2} \right) \sigma^2 + o(\sigma^2), \end{aligned}$$

and using $1 + \frac{2z(r) - z(r)^2}{(1+z(r))^2} = \frac{1+4z(r)}{(1+z(r))^2}$, we obtain (4), concluding the proof.

A.4 Proof of Proposition 4

The encoding error $\hat{m}_{Jn} - m(r_J)$ is drawn from $\mathcal{N}(0, 1/(\pi_J n))$. For each n , we set $\hat{m}_{Jn} - m(r_J) := \varepsilon_J / \sqrt{\pi_J n}$, where $\varepsilon_J \sim \mathcal{N}(0, 1)$ is an error factor common across all n , and independent across J . This choice of the correlation of the errors across n is without loss of generality because it does not affect the expected loss for each n (see e.g. Lindvall, 2002, for this technique known in probability theory as coupling).

We extend the inverse encoding function m^{-1} outside of the interval $(\underline{m}, \bar{m})$ by setting $m^{-1}(\hat{m}_{Jn}) = \underline{r}$ for $\hat{m}_{Jn} \leq \underline{m}$ and $m^{-1}(\hat{m}_{Jn}) = \bar{r}$ for $\hat{m}_{Jn} \geq \bar{m}$. This allows us to express the ML estimate of the lottery value as

$$\begin{aligned} q_n &= \sum_J p_J q_{Jn} \\ &= \sum_J p_J m^{-1} \left(m(r_J) + \frac{\varepsilon_J}{\sqrt{\pi_J n}} \right). \end{aligned}$$

We start the proof of Proposition 4 with a lemma that we will use below for an application of the Dominated Convergence theorem. The lemma establishes an integrable bound on the rescaled error of the estimated lottery value. Let $\varepsilon := (\varepsilon_J)_J$.

Lemma 2. *There exists a function $\bar{e}(\varepsilon)$ such that $|\sqrt{n}(q_n - r)| \leq \bar{e}(\varepsilon)$ for all $\mathbf{r}$, ε and n , and $\mathbb{E} \bar{e}^2(\varepsilon) < \infty$.*

Proof. Let $\underline{m}' > 0$ be a lower bound for $m'(r)$ on $[\underline{r}, \bar{r}]$, which exists since m' is positive and

continuous. Observe the bound on the estimation error for the reward r_J ,

$$|q_{Jn} - r_J| \leq \frac{|\varepsilon_J|}{\underline{m}'\sqrt{\pi_J n}}, \quad (14)$$

which holds uniformly for all J , r_J and ε_J . To see that (14) holds, note that $q_{Jn} - r_J$ and ε_J have the same sign since m^{-1} is monotone and $q_{Jn} = r_J$ if $\varepsilon_J = 0$. Consider a positive $q_{Jn} - r_J$ (the negative case is analogous). If $m(r_J) + \frac{\varepsilon_J}{\sqrt{\pi_J n}} \leq \bar{m}$, then (14) follows from $\partial m^{-1}(\cdot) \leq 1/\underline{m}'$. If $m(r_J) + \frac{\varepsilon_J}{\sqrt{\pi_J n}} > \bar{m}$, then $q_{Jn} = \bar{r}$ and thus $q_{Jn}(r_J, \varepsilon_J) - r_J = q_{Jn}(r_J, \varepsilon'_J) - r_J$, where $\varepsilon'_J \in (0, \varepsilon_J)$ is defined by $m(r_J) + \frac{\varepsilon'_J}{\sqrt{\pi_J n}} = \bar{m}$. Then, $q_{Jn}(r_J, \varepsilon_J) - r_J = q_{Jn}(r_J, \varepsilon'_J) - r_J \leq \frac{|\varepsilon'_J|}{\underline{m}'\sqrt{\pi_J n}} \leq \frac{|\varepsilon_J|}{\underline{m}'\sqrt{\pi_J n}}$, as needed. We can thus define $\bar{e}(\varepsilon) := \sum_J p_J \frac{|\varepsilon_J|}{\underline{m}'\sqrt{\pi_J}}$. $\square$

For a given $\mathbf{r}$, ε and s , we denote the DM's loss by

$$\tilde{\ell}_n(\mathbf{r}, \varepsilon, s) = \max\{r, s\} - \mathbb{1}_{q_n > s} r - \mathbb{1}_{q_n \leq s} s,$$

which depends on $\mathbf{r}$ and ε via q_n and r . We introduce substitution $s = r + \frac{\sigma}{\sqrt{n}}$ and denote the rescaled loss as $\ell_n(\mathbf{r}, \varepsilon, \sigma) := \sqrt{n} \tilde{\ell}_n\left(\mathbf{r}, \varepsilon, r + \frac{\sigma}{\sqrt{n}}\right)$. Let

$$\ell^*(\mathbf{r}, \varepsilon, \sigma) = \begin{cases} \sigma & \text{if } 0 \leq \sigma \leq \sum_J p_J \frac{\varepsilon_J}{\sqrt{\pi_J m'(r_J)}}, \\ -\sigma & \text{if } 0 \geq \sigma \geq \sum_J p_J \frac{\varepsilon_J}{\sqrt{\pi_J m'(r_J)}}, \\ 0 & \text{otherwise.} \end{cases}$$

Lemma 3. $\lim_{n \rightarrow \infty} \ell_n(\mathbf{r}, \varepsilon, \sigma) = \ell^*(\mathbf{r}, \varepsilon, \sigma)$ *almost everywhere*.

Proof. The choice of the DM differs from the optimal choice under complete information if and only if s attains a value in between r and q_n . In such cases, the loss of the DM relative to the complete-information choice is $|s - r|$. Therefore,

$$\ell_n(\mathbf{r}, \varepsilon, \sigma) = \begin{cases} \sigma & \text{if } 0 \leq \sigma \leq \sum_J p_J \left(m^{-1}\left(m(r_J) + \frac{\varepsilon_J}{\sqrt{\pi_J n}}\right) - r_J \right) \sqrt{n}, \\ -\sigma & \text{if } 0 \geq \sigma \geq \sum_J p_J \left(m^{-1}\left(m(r_J) + \frac{\varepsilon_J}{\sqrt{\pi_J n}}\right) - r_J \right) \sqrt{n}, \\ 0 & \text{otherwise.} \end{cases}$$

The right side converges pointwise to $\ell^*(\mathbf{r}, \varepsilon, \sigma)$, because

$$\lim_{n \rightarrow \infty} \left(m^{-1}\left(m(r_J) + \frac{\varepsilon_J}{\sqrt{\pi_J n}}\right) - r_J \right) \sqrt{n} = \partial m^{-1}(m(r_J)) \frac{\varepsilon_J}{\sqrt{\pi_J}}$$

$$= \frac{\varepsilon_J}{m'(r_J)\sqrt{\pi_J}}.$$

□

To prove Proposition 4, observe that

$$\begin{aligned} nL(n) &= \int_{[\underline{r}, \bar{r}]^{|\mathcal{P}|+1} \times \mathbb{R}^{|\mathcal{P}|}} n\tilde{\ell}_n(\mathbf{r}, \varepsilon, s) h_s(s) \prod_J (h(r_J)\varphi(\varepsilon_J)) ds \prod_J (dr_J d\varepsilon_J) \\ &= \int \ell_n(\mathbf{r}, \varepsilon, \sigma) h_s\left(r + \frac{\sigma}{\sqrt{n}}\right) \prod_J (h(r_J)\varphi(\varepsilon_J)) d\sigma \prod_J (dr_J d\varepsilon_J), \end{aligned}$$

where we applied the substitution $s = r + \frac{\sigma}{\sqrt{n}}$. To apply the Dominated Convergence Theorem, we note that the last integrand is bounded as follows:

$$0 \leq \ell_n(\mathbf{r}, \varepsilon, \sigma) h_s\left(r + \frac{\sigma}{\sqrt{n}}\right) \prod_J (h(r_J)\varphi(\varepsilon_J)) \leq \bar{\ell}(\mathbf{r}, \varepsilon, \sigma) \bar{h}_s \prod_J (h(r_J)\varphi(\varepsilon_J)), \quad (15)$$

where $\bar{h}_s$ is an upper bound on the density h_s ,²⁴

$$\bar{\ell}(\mathbf{r}, \varepsilon, \sigma) = \begin{cases} |\sigma| & \text{if } |\sigma| \leq \bar{e}(\varepsilon) \text{ and } r + \frac{\sigma}{\sqrt{n}} \in [\underline{r}, \bar{r}], \\ 0 & \text{otherwise,} \end{cases}$$

and $\bar{e}(\varepsilon)$ is the bound from Lemma 2. Since $\bar{e}(\varepsilon)$ is an upper bound on the size of the error of the DM's estimate, $\bar{\ell}$ expands, relative to ℓ_n , the set of $(\mathbf{r}, \varepsilon, \sigma)$ for which the erroneous choice occurs.

The upper bound in (15) is integrable as needed for the use of the Dominated Convergence Theorem:

$$\begin{aligned} & \int \bar{\ell}(\mathbf{r}, \varepsilon, \sigma) \bar{h}_s \prod_J (h(r_J)\varphi(\varepsilon_J)) d\sigma \prod_J (dr_J d\varepsilon_J) \\ & \leq \bar{h}_s \int \int_{-\bar{e}(\varepsilon)}^{\bar{e}(\varepsilon)} |\sigma| d\sigma \prod_J (h(r_J)\varphi(\varepsilon_J)) \prod_J (dr_J d\varepsilon_J) \\ & = \bar{h}_s \int \bar{e}^2(\varepsilon) \prod_J (h(r_J)\varphi(\varepsilon_J)) \prod_J (dr_J d\varepsilon_J) \\ & = \bar{h}_s \mathbb{E} \bar{e}^2(\varepsilon), \end{aligned}$$

²⁴The bound on h_s exists since h_s is continuous and the support is compact.

where the expectation in the last line is finite, as needed, by Lemma 2.

Hence, the Dominated Convergence Theorem, Lemma 3, and continuity of h_s imply

$$\begin{aligned}
\lim_{n \rightarrow \infty} nL(n) &= \int \ell^*(\mathbf{r}, \varepsilon, \sigma) h_s(r) \prod_J (h(r_J) \varphi(\varepsilon_J)) d\sigma \prod_J (dr_J d\varepsilon_J) \\
&= \mathbb{E} \left[\int_0^{\sum_J p_J \frac{\varepsilon_J}{\sqrt{\pi_J m'(r_J)}}} \sigma h_s(r) d\sigma \right] \\
&= \mathbb{E} \left[\frac{1}{2} \left(\sum_J p_J \frac{\varepsilon_J}{\sqrt{\pi_J m'(r_J)}} \right)^2 h_s(r) \right] \\
&= \mathbb{E} \left[\frac{1}{2} \sum_J \frac{p_J^2}{\pi_J m'^2(r_J)} h_s(r) \right],
\end{aligned}$$

where the first two expectations are over $\mathbf{r}$ and ε and the last expectation is over $\mathbf{r}$. The last step follows from the fact that ε_J are iid standard normal and thus $\mathbb{E} \varepsilon_J^2 = 1$ and $\mathbb{E} [\varepsilon_J \varepsilon_{J'}] = 0$ for all $J \neq J'$.

A.5 Proof of Proposition 5

The next lemma states the first-order conditions of the information-processing problem without imposing unimodality and symmetry on the reward densities. To state the result, define

$$h_J(\tilde{r}) = \frac{h(\tilde{r}) \mathbb{E}[h_s(r) | r_J = \tilde{r}]}{\mathbb{E}[h_s(r)]} = \frac{\int h_s(r) h(\tilde{r}) \prod_{J' \neq J} h(r_{J'}) dr_{J'}}{\int h_s(r) \prod_{J'} h(r_{J'}) dr_{J'}}, \quad (16)$$

which is the density of reward r_J in state J conditional on a tie $r = s$.

Lemma 4. *The information-processing problem has a unique optimal encoding strategy. This optimal strategy has the following properties:*

1. *The encoding function satisfies, for all $\tilde{r} \in [\underline{r}, \bar{r}]$,*

$$m'(\tilde{r}) = m_0 \cdot \left(\sum_J \frac{p_J^2}{\pi_J} h_J(\tilde{r}) \right)^{\frac{1}{3}}, \quad (17)$$

where $m_0 \in \mathbb{R}_+$ is a normalization factor chosen such that $\int_{\underline{r}}^{\bar{r}} m'(\tilde{r}) d\tilde{r} = \bar{m} - \underline{m}$.

2. The sampling frequencies satisfy, for all $J, J' \in \mathcal{P}$,

$$\left(\frac{p_J}{\pi_J}\right)^2 \mathbb{E} \left[\frac{1}{m'^2(r_J)} \mid r = s \right] = \left(\frac{p_{J'}}{\pi_{J'}}\right)^2 \mathbb{E} \left[\frac{1}{m'^2(r_{J'})} \mid r = s \right], \quad (18)$$

where the expectations are over r_J and $r_{J'}$ with respect to the densities h_J and $h_{J'}$, respectively.

Proof. The objective of the information-processing problem is a functional

$$\mathcal{L}(m'(\cdot), (\pi_J)_J) = \mathbb{E} \left[\sum_J \frac{p_J^2}{\pi_J m'^2(r_J)} \mid r = s \right].$$

Since $\frac{p_J^2}{\pi_J m'^2(r_J)}$ is convex with respect to $(m'(r_J), \pi_J)$, the functional $\mathcal{L}$ is convex. Thus, considering that the constraints are linear, the first-order conditions are sufficient for a global minimum of the information-processing problem. Since the objective (6) is strictly decreasing in $m'(\cdot)$, the constraint (7) is binding. The Lagrangian of the constrained optimization problem (6)-(8) is

$$\begin{aligned} & \sum_J \mathbb{E} \left[\frac{p_J^2}{\pi_J m'^2(r_J)} \mid r = s \right] + \lambda \left(\int_{\underline{r}}^{\bar{r}} m'(\tilde{r}) d\tilde{r} - (\bar{m} - \underline{m}) \right) + \mu \left(\sum_J \pi_J - 1 \right) = \\ & \sum_J \int_{\underline{r}}^{\bar{r}} \frac{p_J^2}{\pi_J m'^2(\tilde{r}_J)} h_J(\tilde{r}_J) d\tilde{r}_J + \lambda \left(\int_{\underline{r}}^{\bar{r}} m'(\tilde{r}) d\tilde{r} - (\bar{m} - \underline{m}) \right) + \mu \left(\sum_J \pi_J - 1 \right), \end{aligned}$$

where λ and μ are the Lagrange multipliers for (7) and (8), respectively. For any $\tilde{r} \in [\underline{r}, \bar{r}]$, summing the derivatives w.r.t. $m'(\tilde{r})$ of all the integrands in the last inline expression gives the first-order condition

$$2 \sum_J \frac{p_J^2}{\pi_J m'^3(\tilde{r})} h_J(\tilde{r}) = \lambda. \quad (19)$$

Expressing $m'(\tilde{r})$ from (19) gives (17). The first-order condition with respect to each π_J is

$$\left(\frac{p_J}{\pi_J}\right)^2 \mathbb{E} \left[\frac{1}{m'^2(r_J)} \mid r = s \right] = \mu,$$

which implies (18). □

Observe that the optimal m' is continuous since each h_J as defined above is continuous: h is continuous and, since h_s is continuous on a compact interval, it is uniformly continuous, thus, the function $\tilde{r} \mapsto \mathbb{E}[h_s(r) \mid r_J = \tilde{r}]$ is continuous as well.

Observe also that, when the DM compares two riskless rewards drawn independently from the same density h , the first statement of the lemma replicates the optimal encoding result from Netzer (2009). In this case, (17) implies that $m'(r) \propto h^{\frac{2}{3}}(r)$.

We next state three auxiliary lemmas about unimodal and symmetric random variables.

Definition 1. *A real-valued continuous random variable is unimodal and symmetric around 0 if its density function $h(x)$ is strictly decreasing on the positive part of its domain and $h(x) = h(-x)$ for all $x \in \mathbb{R}$.*

This property is preserved by summation: the sum of unimodal and symmetric random variables is unimodal and symmetric, see e.g., Purkayastha (1998).

Definition 2 (Birnbaum (1948)). *Let X and Y be two unimodal random variables symmetric around 0. We say that X is more peaked than Y if $P(|X| < \alpha) > P(|Y| < \alpha)$ for all $\alpha > 0$ (unless the right side is 1).*

Equivalently, for two unimodal symmetric random variables, X is more peaked than Y whenever the CDF of X is greater than the CDF of Y at any $\alpha > 0$ from the support of Y .

For the next two lemmas, let $X_0, X_1, \dots, X_I$ be independent real-valued continuous random variables that are unimodal and symmetric around 0, where $X_1, \dots, X_I$ are identically distributed while the distribution of X_0 may be distinct. Denote by h the density of each of the variables $X_1, \dots, X_I$ and let h_0 be the density of X_0 . Let $(p_1, \dots, p_I) \in \Delta(\{1, \dots, I\})$ and $X := \sum_{i=1}^I p_i X_i$. We define the density of $(X_1, \dots, X_I) | (X = X_0)$ to be $h(X_1) \times \dots \times h(X_I) \times h_0(X)$, up to normalization, and we define $X_i | (X = X_0)$ by marginalizing it.

Lemma 5. *The random variable $X_i | (X = X_0)$, $i = 1, \dots, I$, is unimodal and symmetric around 0.*

Proof. Unimodality together with symmetry is preserved by multiplication by a constant and by summation, so the variable $X_{-i} := \frac{1}{p_i}(X_0 - \sum_{k \neq i} p_k X_k)$ is unimodal and symmetric around 0. Denote by h_{-i} the density of X_{-i} . Then $X_i | (X = X_0)$ is identical to $X_i | (X_i = X_{-i})$, and so its density is, up to a normalization constant, $h(x_i)h_{-i}(x_i)$. This function is unimodal and symmetric around 0, as needed. $\square$

Lemma 6. *The random variable $X_i | (X = X_0)$ is more peaked than $X_j | (X = X_0)$ if and only if $p_i > p_j$.*

Proof. Without loss of generality, assume $\{i, j\} = \{1, 2\}$ (that is, either $i = 1$ and $j = 2$ or $i = 2$ and $j = 1$). Define $X_{-12} := X_0 - \sum_{k=3}^I p_k X_k$ (if $I = 2$, then $X_{-12} = X_0$) and let

h_{-12} be its density. This is a unimodal random variable symmetric around 0. The random variable $X_i \mid (X = X_0)$ is identical to $X_i \mid (p_i X_i + p_j X_j = X_{-12})$ and so its density equals

$$h_i(x_i) = \frac{\int_{\mathbb{R}} h_{-12}(p_1 x_1 + p_2 x_2) h(x_1) h(x_2) dx_j}{\mathbb{E}[h_{-12}(p_1 X_1 + p_2 X_2)]},$$

where the expectation, which is with respect to X_1 and X_2 , is independent of i . Thus, for any $\alpha > 0$,

$$P(|X_1| < \alpha \mid X = X_0) = \frac{\iint_{(-\alpha, \alpha) \times \mathbb{R}} h_{-12}(p_1 x_1 + p_2 x_2) h(x_1) h(x_2) dx_1 dx_2}{\mathbb{E}[h_{-12}(p_1 X_1 + p_2 X_2)]},$$

and

$$\begin{aligned} P(|X_2| < \alpha \mid X = X_0) &= \frac{\iint_{\mathbb{R} \times (-\alpha, \alpha)} h_{-12}(p_1 x_1 + p_2 x_2) h(x_1) h(x_2) dx_1 dx_2}{\mathbb{E}[h_{-12}(p_1 X_1 + p_2 X_2)]} \\ &= \frac{\iint_{(-\alpha, \alpha) \times \mathbb{R}} h_{-12}(p_1 x_2 + p_2 x_1) h(x_1) h(x_2) dx_1 dx_2}{\mathbb{E}[h_{-12}(p_1 X_1 + p_2 X_2)]}, \end{aligned}$$

where we used for the last equation that $P(|X_1| < \alpha \mid X = X_0)$ and $P(|X_2| < \alpha \mid X = X_0)$ are both (up to the same normalization constant) integrals of the same function $(x_1, x_2) \mapsto h_{-12}(p_1 x_1 + p_2 x_2) h(x_1) h(x_2)$, but the first is over the region $[-\alpha, \alpha] \times \mathbb{R}$ and the second is over $\mathbb{R} \times [-\alpha, \alpha]$. This is equivalent to integrating both over the same region while switching the roles of x_1 and x_2 . Then,

$$\begin{aligned} & (P(|X_1| < \alpha \mid X = X_0) - P(|X_2| < \alpha \mid X = X_0)) \cdot \mathbb{E}[h_{-12}(p_1 X_1 + p_2 X_2)] = \\ & \iint_{(-\alpha, \alpha) \times \mathbb{R}} \left(h_{-12}(p_1 x_1 + p_2 x_2) - h_{-12}(p_1 x_2 + p_2 x_1) \right) h(x_1) h(x_2) dx_1 dx_2 = \\ & \iint_{(-\alpha, \alpha) \times (\mathbb{R} \setminus (-\alpha, \alpha))} \left(h_{-12}(p_1 x_1 + p_2 x_2) - h_{-12}(p_1 x_2 + p_2 x_1) \right) h(x_1) h(x_2) dx_1 dx_2 = \\ & 2 \iint_{(-\alpha, \alpha) \times [\alpha, +\infty)} \left(h_{-12}(p_1 x_1 + p_2 x_2) - h_{-12}(p_1 x_2 + p_2 x_1) \right) h(x_1) h(x_2) dx_1 dx_2, \end{aligned}$$

where we used that the integral is 0 on the region $(-\alpha, \alpha) \times (-\alpha, \alpha)$ and that h and h_{-12} are symmetric around 0.

Suppose that $p_2 > p_1$, and consider any $(x_1, x_2) \in (-\alpha, \alpha) \times [\alpha, +\infty)$. It follows from the identity

$$p_1 x_1 + p_2 x_2 = (p_1 x_2 + p_2 x_1) + (p_2 - p_1)(x_2 - x_1)$$

that

$$p_1x_1 + p_2x_2 > p_1x_2 + p_2x_1,$$

where the last left side (LS) is always positive. The right side (RS) is either positive or negative, but is smaller in absolute value than the LS. Indeed, if the RS is negative, then $x_1 < 0$, and

$$|p_1x_2 + p_2x_1| = -p_1x_2 + p_2|x_1| = -p_1|x_1| + p_2x_2 - (p_1 + p_2)(x_2 - |x_1|) < -p_1|x_1| + p_2x_2 = |p_1x_1 + p_2x_2|,$$

and due to the symmetry and unimodality of h_{-12} ,

$$h_{-12}(p_1x_1 + p_2x_2) < h_{-12}(p_1x_2 + p_2x_1),$$

unless both are zero. It follows that $X_2 \mid (X = X_0)$ is more peaked than $X_1 \mid (X = X_0)$, as needed. $\square$

Lemma 7. *Let a function f be continuous, symmetric ($f(x) = f(-x)$), and increasing on $\mathbb{R}_+$, and let X_1, X_2 be unimodal continuous random variables that are symmetric around 0 and have bounded support. Then $\mathbb{E}[f(X_1)] < \mathbb{E}[f(X_2)]$ whenever X_1 is more peaked than X_2 .*

Proof. Denote by $h_i(x)$ and $H_i(x)$ the PDF and CDF of X_i , $i = 1, 2$. Then,

$$\begin{aligned} \frac{1}{2} \mathbb{E}[f(X_i)] &= \int_0^\infty f(x)h_i(x)dx \\ &= \left[f(x)(H_i(x) - 1) \right]_0^{+\infty} - \int_0^\infty (H_i(x) - 1)df(x) \\ &= \frac{1}{2}f(0) + \int_0^\infty (1 - H_i(x))df(x), \end{aligned}$$

where we have used integration by parts for the Stieltjes integral (see e.g., Ok, 2011). If X_1 is more peaked than X_2 , then $1 - H_1(x) < 1 - H_2(x)$ unless both are zero for all $x > 0$. It follows that $\mathbb{E}[f(X_1)] < \mathbb{E}[f(X_2)]$. $\square$

We now prove Proposition 5. Statement 1 follows from (17) because, by Lemma 5, each conditional reward density h_J is unimodal with the same mode as that of the unconditional reward density h . Additionally, m' is symmetric around r_m since each h_J is symmetric around r_m . Now consider Statement 2. Suppose $p_J < p_{J'}$. By (18) it suffices to show that

$$\mathbb{E} \left[\frac{1}{m'^2(r_J)} \mid r = s \right] > \mathbb{E} \left[\frac{1}{m'^2(r_{J'})} \mid r = s \right]. \quad (20)$$

This indeed holds since, by Lemma 6, $r_{J'} \mid (r = s)$ is more peaked than $r_J \mid (r = s)$, and the inequality (20) then follows from Lemma 7 and the fact that $1/m'^2(r)$ is continuous and symmetric around r_m and increasing above r_m .

A.6 Numerical Computation

The numerical optimization of the encoding function $m(r)$ and the sampling frequencies π involves the following steps:

1. Discretize the interval $[-4, 4]$ using a grid to calculate function values.
2. For $J = 1, 2$, calculate the distribution h_J of $r_J \mid (r = s)$ as specified in equation (16).
3. Initialize the sampling frequencies as $\pi_J = p_J$, $J = 1, 2$.
4. With the current guess of $(\pi_J)_J$, calculate the function $m'(\cdot)$ based on the first-order condition in equation (17).
5. Update the sampling frequencies $(\pi_J)_J$ using the first-order condition in equation (18) and the current approximation of the function $m'(\cdot)$.
6. Repeat steps 4 and 5 until $(\pi_J)_J$ converges.

The computations were performed in Python. The script is available upon request.