

A Measure of Decision-Based Payoff Uncertainty

Thomas A. Weber

Chair of Operations, Economics and Strategy
École Polytechnique Fédérale de Lausanne
Lausanne, Switzerland

Abstract—We introduce decision-based payoff (DBP) uncertainty as a novel measure of informational uncertainty in decision-making. It is defined over an observed sample of nonnegative payoffs from past decisions, evaluated as a fraction of an ex-post optimal payoff benchmark. The resulting distribution of relative payoffs is supported on a subset of the unit interval. DBP uncertainty, taking values between zero and one, quantifies the average deviation from optimality: It vanishes when optimal payoffs are always achieved and reaches one when observed payoffs are consistently zero despite the availability of strictly positive outcomes. The measure is compatible with first- and second-order stochastic dominance and enables meaningful comparisons across decision problems with nonnegative payoffs. It is equal to the average relative regret and vanishes for degenerate problems with singleton action sets, regardless of the observed outcomes. A numerical example involving investment in a call option under uncertain asset prices illustrates the concept's applicability and interpretability.

Keywords—Data-driven decision making, relative regret, robust optimization, uncertainty quantification

I. INTRODUCTION

We consider a decision problem in which a feasible action must be selected to maximize an objective that depends on an unknown state. In other words, the decision-maker must choose an action despite the objective function not being fully known. Such problems are common—for instance, when formulating a business strategy without a fully specified goal, or when investing in financial markets where key asset parameters are uncertain.

Despite this ambiguity, the resulting uncertainty may be limited if decisions can be found that consistently yield high payoffs—relative to what would be achievable under full information. This observation motivates the development of a measure of *decision-based payoff* (DBP) uncertainty, derived from imperfect observations of past decisions and outcomes. Our central idea is to transform these observations into a measure of relative performance, which then serves to quantify the uncertainty that is most relevant to the decision-making process.

A. Literature

The concept of decision-based payoff uncertainty builds on classical decision theory, beginning with Wald's maximin principle [1] and the notion of absolute regret introduced by Savage [2]. Absolute regret, later axiomatized by Gilboa and Schmeidler [3], has seen extensive use in economics [4] and operations research [5]. However, this robustness criterion suffers from a lack of normalization and is therefore unsuitable for reliably comparing *different* decisions.

Moreover, as a decision criterion, it typically fails to guarantee basic performance levels, such as, for example, positive profits in a pricing problem [6]. By instead adopting a relative measure of decision success, we obtain a natural normalization to unity. Our approach builds on Weber [7], who introduces a performance index based on relative regret, and applies this concept to pricing and engineering design problems [8, 9], as well as to general multicriteria decision-making [10].

B. Outline

The remainder of the paper is organized as follows. Section II presents the general decision problem and defines the corresponding performance index. Section III introduces the DBP uncertainty measure and relates it to relative performance. Section IV discusses the consequences of imperfect observations of payoff performance in constructing the empirical relative performance sample. Section V examines key properties of DBP uncertainty, including its behavior under stochastic ordering and the composition of decision problems. Section VI illustrates the use of the measure in a numerical investment example. Section VII concludes.

II. MODEL

A decision-maker aims to select an action x from a (nonempty) compact action set X in order to *maximize* a continuous objective function $f: X \times \Theta \rightarrow \mathbb{R}$, which also depends continuously on a state $\theta \in \Theta$, where the state space Θ is a (nonempty) compact finite-dimensional set. By the extreme value theorem, both the global minimum of the objective function,

$$f_0 = \min_{(x,\theta) \in X \times \Theta} f(x, \theta),$$

and the state-dependent decision-based optimum,

$$f^*(\theta) = \max_{x \in X} f(x, \theta) = f(x(\theta), \theta), \quad \theta \in \Theta,$$

exist, where $x(\theta) \in X(\theta) = \operatorname{argmax}_{x \in X} f(x, \theta)$ is an optimal decision. Since one can always consider the shifted function $\hat{f} = f - f_0$, which preserves all optimal decisions, it is possible—without any loss of generality—to impose the normalization

$$f_0 = 0. \quad (1)$$

Doing so restricts attention to nonnegative-valued objectives that attain a value of zero in at least one contingency. The resulting normalization eliminates ambiguity in the interpretation of gain magnitudes, which is important, as our approach to quantifying decision-based payoff uncertainty relies on measuring relative performance.

To avoid trivialities, we further assume that the decision-maker can always achieve a nonzero gain, i.e., that the decision-based optimum is strictly positive:

$$f^*(\theta) > 0, \quad \theta \in \Theta. \quad (2)$$

For any combination of decision and state, the *performance ratio* is defined as

$$\varphi(x, \theta) = \frac{f(x, \theta)}{f^*(\theta)}, \quad (x, \theta) \in \mathcal{X} \times \Theta.$$

This ratio is a well-defined number in the interval $[0, 1]$, so

$$\varphi(\mathcal{X}, \Theta) \subseteq [0, 1].$$

For any fixed decision x , the mapping $\varphi(x, \cdot)$ is continuous by the maximum theorem. Thus, by the extreme value theorem, the *performance index*

$$\rho(x) = \min_{\theta \in \Theta} \varphi(x, \theta), \quad x \in \mathcal{X},$$

is well-defined. Moreover, $\rho(x)$ is continuous (again by the maximum theorem) and its image over $\mathcal{X}$ is contained in a “minimal” interval, that is,

$$\{0, \rho^*\} \subset \rho(\mathcal{X}) \subseteq [0, \rho^*],$$

where the optimal performance ratio ρ^* is strictly positive and bounded below by the ratio of the worst to the best optimal payoff. Specifically,

$$\rho^* \in [\underline{\rho}, 1], \text{ where } \underline{\rho} = \frac{\min f^*(\Theta)}{\max f^*(\Theta)} > 0.$$

As the lower bound for the optimal performance index is strictly positive, it offers a nontrivial performance guarantee.

III. DECISION-BASED PAYOFF UNCERTAINTY

We assume that decision-based payoff uncertainty has been encountered in $N \geq 1$ past attempts to maximize the objective. From these attempts, a sequence of set-valued samples $\mathcal{S}_1, \dots, \mathcal{S}_N$ has been recorded, where $\mathcal{S}_i = (\mathcal{X}_i, \Theta_i)$, with $\mathcal{X}_i \subset \mathcal{X}$ denoting a compact set of possible decisions, and $\Theta_i \subset \Theta$ a compact set of possible states that could have given rise to the (possibly) experienced reward $y_i = f(x_i, \theta_i) \in \mathbb{R}_+$ in the i -th trial, for some $(x_i, \theta_i) \in \mathcal{X}_i \times \Theta_i$. By continuity of the objective function f , the observed reward must lie within the compact set

$$\mathcal{Y}_i = f(\mathcal{X}_i, \Theta_i), \quad i \in \{1, \dots, N\}.$$

This information structure reflects three intertwined types of uncertainty: decision uncertainty, state uncertainty, and reward uncertainty. These are not independent of one another but are intricately linked through the payoff function. In specific applications, knowledge of any two of the three components determines the third. For instance, if $\Theta \subset \mathbb{R}$ and f is strictly increasing in θ , then a precise reward observation $y_i \in \mathbb{R}_+$ (so $\mathcal{Y}_i = \{y_i\}$) and a precise record of the decision $x_i \in \mathcal{X}$ in the i -th trial (so $\mathcal{X}_i = \{x_i\}$) together imply the existence of a unique state $\theta_i \in \Theta$ such that $y_i = f(x_i, \theta_i)$, yielding $\Theta_i = \{\theta_i\}$. From this, we extract the (relative) performance samples

$$r_i = \min \varphi(\mathcal{X}_i, \Theta_i) \in [0, 1], \quad i \in \{1, \dots, N\},$$

each representing a tight lower bound on the actual relative performance achieved in trial i . Based on the collection $\mathcal{R} = \{r_1, \dots, r_N\}$, we define the empirical cumulative distribution function (CDF) of relative performance as

$$G(s|\mathcal{R}) = \left(\frac{1}{N}\right) \sum_{i=1}^N \mathbf{1}_{\{r_i \leq s\}}, \quad s \in [0, 1].$$

Our measure of (observed) *decision-based payoff uncertainty* (“DBP uncertainty”) is then defined as

$$U(\mathcal{R}) = \int_0^1 G(s|\mathcal{R}) ds = 1 - \int_0^1 s dG(s|\mathcal{R}), \quad (3)$$

where the second expression follows from integration by parts. The term

$$\int_0^1 s dG(s|\mathcal{R}) = \frac{1}{N} \sum_{i=1}^N r_i = r$$

is simply the arithmetic mean of the performance sample. Thus, DBP uncertainty is an affine function of the observed average performance:

$$U(\mathcal{R}) = 1 - r = u(r).$$

This value can be interpreted as the average *relative regret* experienced by the decision-maker, arising from uncertainty about the underlying state as well as from potential imprecision in recalling past actions or rewards.

Example 1 (Perfect Information). If the decision-maker observes θ_i , then it is possible to select $x_i = x(\theta_i)$ and attain the payoff $y_i = f^*(\theta_i)$, so that $r_i = 1$, assuming perfect recall of decisions and rewards. Accordingly, DBP uncertainty vanishes: $U(\{1, \dots, 1\}) = 0$.

Example 2 (Relatively Robust Decisions). By maximizing the performance index and thereby implementing a relatively robust decision,

$$\hat{x}^* \in \operatorname{argmax}_{x \in \mathcal{X}} \rho(x),$$

the decision-maker secures, in any given state $\theta \in \Theta$, a reward $f(\hat{x}^*, \theta)$ which, by construction, is bounded below by $\rho^* f^*(\theta)$. This guarantee further exceeds the uniform lower bound $\underline{\rho} \cdot \min f^*(\Theta) > 0$. In this case, the DBP uncertainty satisfies $U(\mathcal{R}) \leq 1 - \rho^*$, for any performance sample $\mathcal{R}$ (again assuming perfect recall).

IV. IMPERFECT OBSERVATIONS

A key reason for shortcomings in relative payoff performance is that the decision-maker cannot observe the state θ directly, but instead receives an informative random signal Z with realizations in a measurable signal space $\mathcal{Z}$. Accordingly, the decision-maker’s policy is to choose an action based on the signal, i.e., $x = \xi(z)$, where

$$\xi(z) \in \Xi(z) = \operatorname{argmax}_{x \in \mathcal{X}} \hat{f}(x, z) \subset \mathcal{X},$$

and $\hat{f}: \mathcal{X} \times \mathcal{Z} \rightarrow \mathbb{R}$ is a surrogate objective that reflects the decision-maker’s treatment of uncertainty. This surrogate is constructed so that, under perfect information, $\hat{f}$ -optimal decisions coincide with f -optimal decisions. Specifically, if there exists a map $\zeta: \Theta \rightarrow \mathcal{Z}$ such that, for any $\theta \neq \hat{\theta}$,

$$P(Z \in \zeta(\theta)|\theta) = 1 \quad \text{and} \quad P(Z \in \zeta(\theta) \cap \zeta(\hat{\theta})|\theta) = 0,$$

then perfect identifiability of the state via the signal is achieved. In that case, making $\hat{f}$ -optimal decisions implies f -optimality:

$$P(\Xi(Z) \triangle X(\theta) \neq \emptyset | \theta) = 0, \quad \theta \in \Theta,$$

where Δ denotes the symmetric difference and $X(\theta) = \operatorname{argmax}_{x \in \mathcal{X}} f(x, \theta)$.

Let $B(\cdot | z)$ denote the posterior belief over Θ after observing z , and let $\beta(z | \theta)$ denote the sampling density of Z conditional on θ . Standard examples of surrogate objectives include:

- *Expected payoff (with belief distribution $B(\cdot | z)$):*

$$\hat{f}(x, z) = \int_{\Theta} f(x, \theta) dB(\theta | z);$$

- *Worst-case payoff (with sampling density $\beta(\cdot | \theta)$):*

$$\hat{f}(x, z) = \inf \{f(x, \theta) : \theta \in \Theta, \beta(z | \theta) > 0\};$$

- *Negative worst-case absolute regret:*

$$\hat{f}(x, z) = -\sup \{f^*(\theta) - f(x, \theta) : \theta \in \Theta, \beta(z | \theta) > 0\};$$

- *Updated performance index:*

$$\hat{f}(x, z) = \inf \{\varphi(x, \theta) : \theta \in \Theta, \beta(z | \theta) > 0\}.$$

A. Relative-Performance Samples

Given a signal realization z_i , the decision-maker implements $x_i = \xi(z_i)$. Moreover, given a set-valued observation $\mathcal{Y}_i \subset \mathbb{R}_+$ of possible payoffs, the effectively observed subset of states is

$$\Theta_i = \{\theta \in \Theta : f(x_i, \theta) \in \mathcal{Y}_i\}.$$

Hence, the effectively observed (posterior-averaged) relative performance is

$$\tilde{r}_i = \frac{1}{P(\Theta_i | z_1, \dots, z_N)} \int_{\Theta_i} \varphi(x_i, \theta) dB(\theta | z_1, \dots, z_N).$$

The posterior distribution $B(\cdot | z_1, \dots, z_N)$ is obtained from a prior h and likelihood β via Bayes' rule:

$$B(\theta | z_1, \dots, z_N) = \frac{h(\theta) \prod_{i=1}^N \beta(z_i | \theta)}{\int_{\Theta} h(\theta) \prod_{i=1}^N \beta(z_i | \theta) d\theta}, \quad \theta \in \Theta.$$

The prior h can be informed by the set-valued state realizations Θ_i , e.g., by (parametrized) maximum likelihood:

$$\begin{aligned} \hat{\theta} &\in \operatorname{argmax}_{\theta \in \mathcal{T}} L(\vartheta | \Theta_1, \dots, \Theta_N) \\ &= \operatorname{argmax}_{\theta \in \mathcal{T}} \prod_{i=1}^N \left(\frac{1}{|\Theta_i|} \int_{\Theta_i} \hat{h}(\theta | \vartheta) d\theta \right), \end{aligned}$$

so that the estimated prior becomes $h(\cdot) = \hat{h}(\cdot | \hat{\theta})$ on Θ .

B. Problem Inputs and Outputs

In summary, the inputs to the decision problem are:

- $\hat{h}(\cdot | \vartheta)$, $\vartheta \in \mathcal{T}$: a parameterized family of priors with common support Θ ;
- $\beta(\cdot | \theta)$, $\theta \in \Theta$: the sampling distribution of Z characterizing the decision-maker's information;
- $z_1, \dots, z_N$; $\mathcal{Y}_1, \dots, \mathcal{Y}_N$: observations of signals and set-valued payoff realizations from N independent experiments;
- $f(\cdot, \theta)$, $\theta \in \Theta$: the state-dependent objective; and $\hat{f}(\cdot, z)$: the surrogate objective.

The outputs are:

- $h(\cdot) = \hat{h}(\cdot | \hat{\theta})$: the maximum-likelihood prior; and $B(\cdot | z_1, \dots, z_N)$: the posterior distribution;
- Actions $x_i = \xi(z_i)$ that maximize the surrogate objective, producing relative-performance lower bounds $\tilde{r}_i$;
- The empirical distribution $G(\cdot | \mathcal{R})$ of $\mathcal{R} = \{r_i\}$ as defined in Section III, along with the resulting uncertainty measure $U(\mathcal{R}) = 1 - r$.

V. PROPERTIES

We now turn our attention to the properties of the proposed DBP uncertainty measure, in terms of stochastic ordering and composition of decision problems.

A. Stochastic Ordering

Since DBP uncertainty in (3) is the integral of the empirical CDF of $\mathcal{R}$, second-order stochastic dominance (SOSD) [11] implies a weak ordering of DBP uncertainty. *A fortiori*, a stochastic increase via first-order stochastic dominance (FOSD) [12] also induces a weak ordering of DBP uncertainty (as FOSD implies SOSD). Because U is affine in r , higher-order distributional changes that keep the average relative performance fixed do not affect DBP uncertainty.

Proposition 1: Let $\mathcal{R}, \hat{\mathcal{R}}$ be two performance samples. Then:

- (i) $G(\cdot | \hat{\mathcal{R}}) \succ_{\text{FOSD}} G(\cdot | \mathcal{R}) \Rightarrow U(\hat{\mathcal{R}}) \leq U(\mathcal{R})$; and
- (ii) $G(\cdot | \hat{\mathcal{R}}) \succ_{\text{SOSD}} G(\cdot | \mathcal{R}) \Rightarrow U(\hat{\mathcal{R}}) \leq U(\mathcal{R})$.

Similarly, a better observation of the decision problem reduces DBP uncertainty. To see this, let $\mathcal{S}_N = \{\mathcal{S}_i\}_{i=1}^N$, with $\mathcal{S}_i = (X_i, \Theta_i) \neq \emptyset$, for all $i \in \{1, \dots, N\}$, an observed sample of past decisions. The sample $\hat{\mathcal{S}}_N = \{\hat{\mathcal{S}}_i\}_{i=1}^N$ is called *sharper* than $\mathcal{S}_N$ (denoted by $\hat{\mathcal{S}}_N \subseteq \mathcal{S}_N$) if $\hat{\mathcal{S}}_i \subset \mathcal{S}_i$, for all $i \in \{1, \dots, N\}$.

Proposition 2: Let $\mathcal{R}$ and $\hat{\mathcal{R}}$ be performance samples associated with $\mathcal{S}_N$ and $\hat{\mathcal{S}}_N$, respectively. Then: $\hat{\mathcal{S}}_N \subseteq \mathcal{S}_N \Rightarrow U(\hat{\mathcal{R}}) \leq U(\mathcal{R})$.

A sharper sample observation leads to a decrease in DBP uncertainty.

Proposition 3: For any given sample $\mathcal{S}_N = \{\mathcal{S}_i\}_{i=1}^N$, the DBP uncertainty vanishes if $\mathcal{S}_i = X(\theta_i) \times \{\theta_i\}$, for $i \in \{1, \dots, N\}$. If in addition $\theta \neq \hat{\theta} \Rightarrow X(\theta) \neq X(\hat{\theta})$, for all $\theta, \hat{\theta} \in \Theta$, then the preceding condition is also necessary for the DBP uncertainty to vanish.

When only optimal decisions are being taken, a performance sample with positive DBP uncertainty cannot exist. On the other hand, if the state is unknown to the decision-maker, then a zero DBP might only be observed if two states could potentially have the same optimal action set.

B. Composition of Decision Problems

Consider now the combination of independent decision problems. For this, assume that $\mathcal{X} = \mathcal{X}_1 \times \mathcal{X}_2$ and that $f(x, \theta) = f_1(x_1, \theta) + f_2(x_2, \theta)$, where $x_1 \in \mathcal{X}_1$ and $x_2 \in \mathcal{X}_2$. Then

$$\begin{aligned}
\varphi(x, \theta) &= \frac{f_1(x_1, \theta) + f_2(x_2, \theta)}{f_1^*(\theta) + f_2^*(\theta)} \\
&\geq \min \left\{ \frac{f_1(x_1, \theta)}{f_1^*(\theta)}, \frac{f_2(x_2, \theta)}{f_2^*(\theta)} \right\} \\
&= \min \{ \varphi_1(x_1, \theta), \varphi_2(x_2, \theta) \},
\end{aligned}$$

so the sample average r of the joint performance $r_i \in \mathcal{R}_{1+2}$ weakly exceeds the sample average of $\min \{r_{1,i}, r_{2,i}\}$, with $(r_{1,i}, r_{2,i}) \in \mathcal{R}_1 \times \mathcal{R}_2$. The following result formalizes the additive risk pooling of decisions in the presence of unknown states.

Proposition 4: $U(\mathcal{R}_{1+2}) \leq 1 - (\frac{1}{N}) \sum_{i=1}^N \min \{r_{1,i}, r_{2,i}\}$.

It is possible to obtain a similar result on multiplicative risk pooling, with $f(x, \theta) = f_1(x_1, \theta) f_2(x_2, \theta)$, for DBP uncertainty based on the joint performance $r_i \in \mathcal{R}_{1 \times 2}$ as follows.

Proposition 5: $U(\mathcal{R}_{1 \times 2}) \leq 1 - (\frac{1}{N}) \sum_{i=1}^N r_{1,i} r_{2,i}$.

In Prop. 5, for independent states, when $\theta = (\theta_1, \theta_2)$ and each f_j depends on the state only through θ_j , equality applies.

VI. APPLICATION

Consider the decision problem of investing in a European call option with strike price κ at market price $p > 0$, given that the decision-maker has investment capital $\alpha > 0$. Without loss of generality, we assume a zero risk-free interest rate, which simplifies the analysis by removing any dependence on the option's time to maturity $T > 0$.¹

Let $\Theta = [\theta_1, \theta_2]$ denote the decision-maker's subjective range of feasible stock prices at maturity, with $\theta_1 < \kappa < \theta_2 - p$ to avoid trivialities. The corresponding objective function is

$$f(x, \theta) = \alpha + (\max\{\theta - \kappa, 0\} - p)x, \quad (x, \theta) \in \mathcal{X} \times \Theta,$$

where $\mathcal{X} = [0, \bar{x}]$ denotes the set of feasible investment amounts, and $\bar{x} = \alpha/p$ ensures the normalization condition (1). If the state $\theta \in \Theta$ is known, the optimal decision is

$$x(\theta) = \bar{x} \cdot \mathbf{1}_{\{\theta - \kappa - p \geq 0\}}, \quad \theta \in \Theta,$$

yielding the optimal payoff

$$f^*(\theta) = \alpha + \bar{x} \max\{\theta - \kappa - p, 0\} > 0,$$

in accordance with assumption (2). The performance index then becomes

$$\rho(x) = \min_{\theta \in \Theta} \frac{\alpha + (\max\{\theta - \kappa, 0\} - p)x}{\alpha + \bar{x} \max\{\theta - \kappa - p, 0\}}, \quad x \in \mathcal{X}.$$

To interpret this ratio economically, we note that it measures the worst-case ex-post efficiency, quantifying how well a given investment x performs relative to the state-contingent optimum that could have been achieved if θ were known. The numerator captures realized gains from investment, whereas the denominator represents the ideal benchmark payoff that one could obtain by acting with full knowledge of θ .

The decision-maker's challenge is thus inherently epistemic: Any departure from full knowledge incurs potential performance loss, and the degree of this loss, aggregated across a sample of decisions, forms the basis for DBP

uncertainty. Importantly, the normalization to $[0, 1]$ ensures that this measure retains interpretability across different investment magnitudes or risk scales.

As shown in [7], due to the quasiconcavity of the minimand in θ , the performance index $\rho(x)$ is determined by the minimum of the boundary ratios $\varphi(x, \theta_1)$ and $\varphi(x, \theta_2)$:

$$\rho(x) = \min \left\{ \frac{\alpha - px}{\alpha}, \frac{\alpha + (\theta_2 - \kappa - p)x}{\alpha + (\theta_2 - \kappa - p)\bar{x}} \right\}, \quad x \in \mathcal{X}.$$

The relatively robust decision $\hat{x}^*$ maximizes the performance index, occurring when the two boundary ratios are equal:

$$\hat{x}^* = \frac{\bar{x}}{1 + \left(\frac{\bar{x}}{\alpha} + \frac{1}{\theta_2 - \kappa - p} \right) p} = \frac{\bar{x}}{2 + \frac{p}{\theta_2 - \kappa - p}} < \frac{\bar{x}}{2}.$$

The decision $\hat{x}^*$ balances caution and aggressiveness in an environment of irreducible uncertainty. It represents a strategic compromise between maximizing potential gains when the state is favorable and minimizing regret when it is not. From a policy perspective, it is a conservative decision, chosen not because it yields the best possible outcome, but because it guarantees the best worst-case relative performance.

This leads to the optimal performance ratio

$$\rho^* = \rho(\hat{x}^*) = \frac{\alpha}{\alpha + \left(1 - \frac{p}{\theta_2 - \kappa} \right) p \bar{x}} = \frac{\theta_2 - \kappa}{2(\theta_2 - \kappa) - p},$$

which lies in the open interval $(\frac{1}{2}, 1)$. The naïve lower bound $\underline{\rho} = \min f^*(\Theta) / \max f^*(\Theta) = p / (\theta_2 - \kappa)$ is strictly smaller than ρ^* .²—At $x = \hat{x}^*$, for any performance sample $\mathcal{R}$ from past investments, the realized DBP uncertainty satisfies

$$U(\mathcal{R}) \leq 1 - \rho^* = \frac{\theta_2 - \kappa - p}{2(\theta_2 - \kappa) - p} \in (0, \frac{1}{2}).$$

This upper bound is notable and reflects the payoff uncertainty embedded in the problem, as the decision has to be taken in the absence of state information. The conclusion is thus twofold: First, relatively robust decisions provide the best possible guarantees against relative regret, and second, the generic impossibility of attaining ex-post optimal payoffs means that DBP uncertainty is typically a nontrivial measure that in practice takes on strictly positive values.

Example. Consider now a numerical example in which the underlying stock price S_t (with $S_0 = 100$ at time $t = 0$) follows a geometric Brownian motion for $t \in [0, T]$. That is, S_t satisfies the stochastic differential equation

$$\frac{dS_t}{S_t} = \mu dt + \sigma dW_t,$$

where W_t is a Wiener process, $\mu \in \mathbb{R}$ is the drift, and $\sigma \in \mathbb{R}_+$ is the volatility of returns; see [13, Ch. 3]. The normalized stock price satisfies

$$\frac{S_t}{S_0} = \exp \left[\left(\mu - \frac{\sigma^2}{2} \right) t + \sigma W_t \right], \quad t \in [0, 1],$$

and is thus log-normally distributed with parameters $\hat{\mu}_t = (\mu - \sigma^2/2)t$ and $\hat{\sigma}_t^2 = \sigma^2 t$. The distribution has expectation $S_0 e^{\mu t}$ and standard deviation $S_0 e^{\mu t} \sqrt{e^{\sigma^2 t} - 1}$. Its density is given by

$$g_t(s) = \frac{1}{\sqrt{2\pi t} \sigma s} \exp \left[-\frac{(\ln(s/S_0) - (\mu - \sigma^2/2)t)^2}{2\sigma^2 t} \right],$$

¹ All results in this section extend to the case with a risk-free rate $r \geq 0$ by substituting α/δ for α and p/δ for p , where $\delta = \exp(-rT) \in (0, 1]$ is the discount factor.

² Using $\eta = \theta_2 - \kappa$, we find that $\rho^* - \underline{\rho} = \frac{(\eta - p)^2}{(2\eta - p)\eta} > 0$.

for all $(s, t) \in \mathbb{R}_{++} \times [0, T]$. Given a critical value ζ of the normal distribution, the corresponding quantile of the log-normal distribution is

$$\hat{S}_t(\zeta) = S_0 \exp(\hat{\mu}_t + \hat{\sigma}_t \zeta),$$

where $\zeta \in \{1.282, 1.645, 2.326\}$ corresponds to the 90th, 95th, and 99th percentiles, respectively.

For $T = 1$, $\mu = 0$, $\sigma = 0.3$, and $r = 0$, we obtain $\hat{\mu}_T = -0.045$ and $\hat{\sigma}_T^2 = 0.09$ (with higher volatility $\sigma = 0.5$, we get $\hat{\mu}_T = -0.125$ and $\hat{\sigma}_T^2 = 0.25$). The associated vector of critical values $(\hat{S}_T(1.282), \hat{S}_T(1.645), \hat{S}_T(2.326))$ is equal to $(140.4386, 156.5961, 192.0912)$, and it is equal to $(167.5313, 200.8725, 282.3564)$ for the larger volatility. The Black-Scholes price for an option at the strike price $\kappa = 110$ is $p = 8.1410$; at the higher volatility with the same strike it becomes $p = 16.0957$.

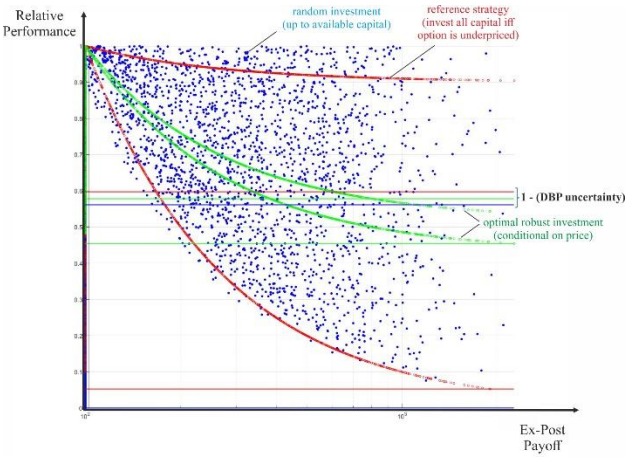


Fig. 1. DBP Uncertainty: $U(\mathcal{R}) = 1 - (\frac{1}{N}) \sum_{i=1}^N r_i$.

An example with capital $\alpha = 1000$ is illustrated in Fig. 1. The stock follows a geometric Brownian motion as described. As a signal, one may consider a perfect reading z of the stock price at time t . As t varies from 0 to 1, the informativeness of the signal increases from none to complete. This determines the posterior distribution $B(\theta|z)$. The DBP uncertainty decreases with increasing volatility (assuming prices adjust via Black-Scholes), as the relative size of the investment diminishes. One further observes that relatively robust decisions implicitly reflect risk aversion.

VII. CONCLUSION

This paper introduces *decision-based payoff uncertainty* (DBP uncertainty), a distribution-free measure for quantifying the informational quality of decisions in the presence of unknown states. The measure is rooted in evaluating the *relative* performance of past decisions and is fundamentally connected to relative robustness and mean relative regret. Five key insights emerge. First, DBP uncertainty departs from entropy-based approaches by focusing on observed payoffs rather than belief distributions. This distinction is crucial: Agnostic measures such as (relative) entropy may interpret repeated poor outcomes (e.g., consistent zero payoffs) as reflecting low uncertainty, whereas DBP uncertainty correctly

recognizes that, in such cases, the decision-maker has no actionable information about the underlying state. Second, the proposed measure exhibits robustness to limited sample contamination. Given a sufficiently large dataset, DBP uncertainty remains stable under the presence of a few outliers. Third, by maximizing the performance index, the decision-maker can systematically reduce DBP uncertainty. This identifies relatively robust strategies that ensure the highest guaranteed relative performance, even without precise state information. Fourth, DBP uncertainty is formally equivalent to mean relative regret. This equivalence not only grounds the measure in established decision-theoretic constructs but also lends interpretability: DBP uncertainty reflects the average opportunity lost due to lack of information. Finally, the measure is compatible with second-order stochastic dominance (SOSD): Improvements in the distribution of relative performance, whether through reduced risk or stochastic shifts, weakly reduce DBP uncertainty.

Taken together, these features suggest that DBP uncertainty offers a pragmatic reorientation in the assessment of uncertainty. Rather than placing primary emphasis on the internal coherence of beliefs, it shifts attention to the external adequacy of decisions, as revealed through observable outcomes. This is particularly valuable in environments where probabilistic calibration is either infeasible or conceptually ill-posed—for instance, in one-off strategic settings, opaque markets, or poorly understood systems. In such cases, traditional measures may understate the true difficulty of the decision problem by relying on assumptions about belief precision. DBP uncertainty, by contrast, gauges what the decision-maker can justify not in theory, but in hindsight, that is, how close their chosen actions came to what could have been achieved under full knowledge.

REFERENCES

- [1] A. Wald, "Statistical decision functions which minimize the maximum risk," *Ann. Math.*, vol. 46, no. 2, pp. 265–280, April 1945.
- [2] L. J. Savage, *The Foundations of Statistics*. New York: Wiley, 1954.
- [3] I. Gilboa and D. Schmeidler, "Maxmin expected utility with non-unique prior," *J. Math. Econ.*, vol. 18, no. 2, pp. 141–153, 1989.
- [4] G. Loomes and R. Sugden, "Regret theory: An alternative theory of rational choice under uncertainty," *Econ. J.*, vol. 92, no. 368, pp. 805–824, December 1982.
- [5] D. E. Bell, "Regret in decision making under uncertainty," *Oper. Res.*, vol. 30, no. 5, pp. 961–981, September/October 1982.
- [6] T. A. Weber, "Monopoly pricing with unknown demand," *Scand. J. Econ.*, vol. 127, no. 1, pp. 235–285, January 2025.
- [7] T. A. Weber, "Relatively robust decisions," *Theory and Decis.*, vol. 94, no. 1, pp. 35–62, January 2023.
- [8] J. Han and T. A. Weber, "Price discrimination with robust beliefs," *Eur. J. Oper. Res.*, vol. 306, no. 2, pp. 795–809, April 2023.
- [9] T. A. Weber, "Optimal depth of discharge for electric batteries with robust capacity shrinkage estimator," in *Proc. Int. Conf. Smart Grid Renew. Energy (SGRE)*, Doha, Qatar, pp. 1–5, January 2024.
- [10] T. A. Weber, "Relatively robust multicriteria decisions," *Manage. Sci.*, in press. Available: <https://doi.org/10.1287/mnsc.2025.00510>
- [11] J. Hadar and W. R. Russell, "Rules for ordering uncertain prospects," *Am. Econ. Rev.*, vol. 59, no. 1, pp. 25–34, March 1969.
- [12] E. L. Lehmann, "Ordered families of distributions," *Ann. Math. Stat.*, vol. 25, no. 3, pp. 399–419, September 1954.
- [13] S. E. Shreve, *Stochastic Calculus for Finance II: Continuous-Time Models*. New York: Springer, 2004.