
Infra-Bayesian Reinforcement Learning Agents Outperform Classical RL For Worst-Case Robustness

Anonymous Author(s)

Affiliation

Address

email

Abstract

1 Classical reinforcement learning assumes the agent interacts with a fixed envi-
2 ronment whose behavior does not depend on the agent’s policy. This assumption
3 breaks down in non-realizable settings where other actors might anticipate the
4 agent’s behavior – most notably those environments crucial to AI safety, where
5 any given agent interacts with predictors, human and other AI agents, and insti-
6 tutions. In such environments, the agent’s model class fails to capture the world
7 in which it operates. Under such misspecification, Classical Bayesian methods
8 can produce confidently wrong posteriors, unreliable decisions, and unbounded
9 regret, as realizability fails to obtain. Infra-Bayesianism is a decision-theoretic
10 framework that addresses these failures by distinguishing ordinary probabilistic
11 uncertainty – where priors can be reasonably chosen – from Knightian uncertainty,
12 where no grounds exist for the construction of such a prior. Infra-Bayesianism
13 does so by evaluating actions on their worst-case outcomes in environments, rather
14 than from posterior expectations or weighted averaging.

15 We present the first proof-of-concept implementation of an infra-Bayesian rein-
16 forcement learning architecture for finite-outcome stateless decision problems.
17 Our agent maintains a set of imprecise hypotheses, updates them using infra-
18 Bayesian conditioning, and selects actions by maximizing worst-case expected
19 value. We apply this implementation of the infra-Bayesian maximin decision
20 process to an environment with Knightian uncertainty, and demonstrate a lower
21 worst-case regret as compared to classical reinforcement learning agents. We also
22 investigate Newcomb’s problem and show that the infra-Bayesian agent picks the
23 optimal strategy, outperforming classical decision theory agents. Our results pro-
24 vide a step towards reinforcement learning agents that remain robust under model
25 misspecification and policy-dependent uncertainty.

26 1 Introduction

27 Reinforcement learning (RL) is usually analyzed under the assumption that the agent interacts with
28 a fixed environment: a Markov decision process (MDP), a partially observable Markov decision
29 process (POMDP), or a Bayesian posterior over such models. This assumption is mathematically
30 convenient and has enabled a large body of convergence and regret guarantees. However, it is poorly
31 suited to settings in which the agent is embedded in an environment too complex to be represented by
32 the agent, or environments containing other actors, that model the agent’s behavior and thereby re-
33 spond to the agent’s policy rather than merely to its realized actions. Examples include autonomous
34 vehicles whose driving style is anticipated by human drivers, recommendation systems whose users
35 adapt to the system’s known behavior, and decision-theoretic problems such as Newcomb’s prob-
36 lem. In these settings, the environment is not merely unknown; it can be policy-dependent and
37 non-realizable from the agent’s point of view [Bell et al., 2021].

Classical Bayesian RL provides a principled treatment of ordinary uncertainty by placing a prior over possible environments. Its strongest guarantees, however, rely on a *realizability* or *grain of truth* assumption: the true environment, or a sufficiently exact model of it, must lie in the agent’s hypothesis class. This assumption is implausible for open-ended deployment. The real world contains the agent itself and other systems of comparable or greater complexity, so no tractable hypothesis class can be expected to contain a complete description of the environment. Under such misspecification, Bayesian agents can become confidently wrong, and value-based RL agents can fail to converge to optimal policies [Gelb, 2025b].

Infra-Bayesianism (IB) was developed to address this problem by replacing precise Bayesian hypotheses with imprecise, worst-case-evaluated hypotheses. Instead of representing uncertainty only by a probability distribution over worlds, an IB agent may represent *Knightian uncertainty* over possible worlds and evaluate actions by a lower expectation: the value guaranteed against the worst admissible member of the hypothesis class. It then selects the action that maximizes this worst-case expected value. This maximin structure is safety-relevant because it turns the agent’s objective into a lower-bound guarantee, rather than an average-case prediction. IB also supplies an update rule for these objects that is designed to preserve dynamic consistency across sequential observations [Diffraitor and Kosoy, 2020, Appel, 2020].

More broadly, IB belongs to the agent foundations and AI safety research agenda, which seeks mathematically precise models of embedded agency. From this perspective, non-realizability is not an edge case but a central feature of advanced agents operating in the real world: the agent is itself part of the environment it is trying to model.

Despite substantial theoretical work on IB and its role in the Learning-Theoretic Agenda for AI alignment, comparatively little work has translated its mathematical objects into concrete reinforcement learning agents. This leaves a gap between the formal promise of IB – robust reasoning under non-realizability, Knightian uncertainty, and policy-dependent environments – and the practical question of how an agent could actually represent, update, and plan within it. In this paper, we present a proof-of-concept IB RL architecture. The implementation-level object is not a posterior distribution over MDPs, but a finite set of extremal affine evaluators whose lower envelope defines the agent’s value estimate. This representation allows the agent to perform tractable lower-expectation evaluation, distinguish classical probabilistic mixtures from Knightian uncertainty, and apply dynamically consistent IB-style updates without enumerating full history trees.

Our contributions are:

1. We formulate a finite-outcome infra-Bayesian reinforcement learning (IBRL) architecture based on α -measures and infradistributions, including classical and Knightian mixtures, lower-expectation evaluation, and IB-style updates.
2. We show that the architecture recovers ordinary Bayesian behavior in the degenerate case where each infradistribution has a single minimal point and all uncertainty is classical.
3. We demonstrate empirically that an IB agent outperforms a Bayesian agent in terms of worst-case performance in a Knightian uncertain environment.
4. We show that IB decision theory obtains optimal rewards in simple policy-dependent environments, where classical decision theories fail to comprehensively model the causal structure.

2 Background

2.1 Classical and Bayesian reinforcement learning

In RL, an agent repeatedly chooses actions and receives observations from an environment. The usual formalism is a Markov decision process, which models the environment as a set of states through which the agent can transition by choosing from a set of actions, receiving a reward upon each transition. The agent aims to optimize its policy, to maximize the cumulative expected reward over future actions. Maximizing rewards is equivalent to minimizing *regret*, the difference between the expected reward obtainable by the optimal policy and the actually realized reward.

Classical Bayesian RL places a prior over environments and updates by Bayes’ rule, which is appropriate in realizable settings where the agent’s hypothesis class contains the truth, and is the basis for

many standard convergence arguments [Sutton and Barto, 2018, Ghavamzadeh et al., 2015]. This assumption is a poor fit for embedded or open-world agents, where no computationally feasible hypothesis class can contain the exact environment. In such non-realizable settings, Bayesian updating still produces a posterior, but the posterior need not converge to a useful model for decision-making. This is one of the core motivations for IB in the Learning-Theoretic Agenda [Gelb, 2025b, Kosoy, 2023].

2.2 Non-realizability and policy-dependent environments

The central motivation for this work is that difficulties in RL arise not merely from statistical uncertainty. They also include model misspecification, Knightian uncertainty, and environments whose behavior depends on the agent’s policy rather than only on its realized actions.

Policy-dependence is especially important for embedded agents. In ordinary RL, the environment is usually modeled as responding to the current state and action. In a policy-dependent environment, however, rewards or transitions may depend directly on the agent’s policy. This can occur when other agents, predictors, markets, users, or institutions form expectations about the agent’s behavior and respond to those expectations.

Newcomb-like decision processes formalize these environments by allowing rewards or transitions to depend on the policy itself. Bell et al. show that value-based RL agents in such environments can converge only to ratifiable policies, and may fail to converge to optimal policies at all [Bell et al., 2021]. This suggests that policy-dependent environments require a decision-theoretic treatment that goes beyond ordinary value learning in fixed MDPs.

2.3 Brief description of Newcomb’s Problem

Newcomb’s problem (named after physicist William Newcomb) is a classic decision theory dilemma [Nozick, 1969]. In this problem, the agent is presented with a transparent box and an opaque box, and may choose either to take only the opaque box (one-boxing) or to take both boxes (two-boxing). The agent can see that the transparent box contains one thousand dollars. The opaque box contains either nothing or one million dollars, depending on a prediction that has already been made. Crucially, the prediction was about the agent’s choice: the opaque box contains a million dollars if and only if the agent is predicted to take just the opaque box. The prediction is highly accurate, but not necessarily perfect. All aspects of the scenario are known to the agent. This scenario poses a dilemma because in it, the principle of dominance conflicts with the principle of expected-utility maximization. In particular, agents following causal or evidential decision theories come to different conclusions, and both fail to capture the causal structure of the environment accurately.

2.4 Infra-Bayesianism at a high level

Infra-Bayesianism generalizes Bayesian reasoning by allowing hypotheses to contain Knightian uncertainty. Instead of representing uncertainty only by a probability distribution over possible worlds, an IB agent may represent a set of admissible evaluators. The agent then evaluates a policy by its lower expectation: the value guaranteed under the least favorable admissible evaluator.

This distinction is important. A Bayesian agent averages over hypotheses using posterior probabilities. An IB agent can also represent ambiguity that should not be averaged away. In this sense, IB separates ordinary probabilistic uncertainty from Knightian uncertainty. The resulting decision rule is maximin: choose the policy whose worst admissible value is best.

This lower-bound perspective is especially relevant in safety-motivated settings. When the agent’s model class may be misspecified, or when the environment may respond to the agent’s policy, average-case posterior value may be the wrong object to optimize. IB instead provides a formalism for reasoning with partial information and worst-case guarantees.

2.5 Affine measures and infradistributions

The basic implementation-level object used in this paper is an affine measure, or a-measure. In the finite non-signed setting considered here, an a-measure is a pair

$$a = (\lambda\mu, b), \quad (1)$$

where μ is a probability measure over possible observation histories, $\lambda \geq 0$ is a scale factor, and $b \geq 0$ is an offset. Given a bounded return function f over histories, the a-measure evaluates f by

$$a(f) = \lambda \mathbb{E}_\mu[f] + b. \quad (2)$$

The probability component μ represents ordinary stochastic uncertainty over histories. The scale λ determines the weight assigned to that component. The offset b records value associated with branches of the history tree that have been ruled out by observation. Intuitively, when an observation eliminates some branches, IB does not simply delete the value associated with those branches. Instead, their contribution is carried forward into the affine offset, which helps preserve dynamic consistency between ex ante and ex post evaluation.

An infradistribution Ψ can be viewed, for our finite purposes, as a set of affine evaluators. Its lower expectation is

$$\mathbb{E}_\Psi[f] = \inf_{a \in \Psi} a(f). \quad (3)$$

Thus, a policy is evaluated by the worst admissible a-measure in the infradistribution. In reward-maximization language, the agent chooses a policy maximizing this lower expectation.

2.6 Classical mixtures and Knightian uncertainty

IB distinguishes between classical probabilistic uncertainty and Knightian uncertainty. Classical uncertainty is represented by a mixture. If Ψ_i are infradistributions and w_i are prior weights, their classical mixture is

$$\sum_i w_i \Psi_i = \left\{ \sum_i w_i a_i : a_i \in \Psi_i \right\}, \quad \sum_i w_i = 1. \quad (4)$$

This corresponds to uncertainty that can be averaged over using prior weights, as in Bayesian reasoning.

Knightian uncertainty is different. It represents ambiguity that remains exposed to the outer infimum rather than being averaged away. A singleton infradistribution recovers ordinary probabilistic evaluation. A non-singleton infradistribution represents ambiguity over several admissible evaluators, where the value of a policy is determined by the least favorable one.

This distinction is central to the behavior of an IB agent. A Bayesian agent facing several possible models assigns probabilities to them and optimizes posterior expected value. An IB agent may instead treat several models as admissible without assigning meaningful probabilities between them, and then choose the policy whose lower expectation is highest.

2.7 IB-style updating

The IB update rule can be understood as a dynamically consistent analogue of Bayesian conditioning. After observing an event L , each a-measure is restricted to the observed branch. However, unlike ordinary Bayesian conditioning, the value associated with the unobserved branch is not simply discarded. Instead, it is transferred into the offset term.

To formalize this procedure, we must define the return function g . Unlike classical RL, where reward is typically a stepwise scalar $R(s, a)$ evaluated at each transition, IB evaluates policies using a bounded return (or loss) function g defined over entire histories, formally, over infinite sequences of actions and observations (destinies). During evaluation, we will set $g \equiv f$. Let $\mu(g)$ denote the expected value of g under the measure μ . Viewing the observation L as an indicator function for the realized event, the raw update takes the following schematic form:

$$(\lambda\mu, b) \mapsto (\lambda\mu L, b + \lambda\mu((1-L)g)). \quad (5)$$

Here, μL denotes the restriction of μ to the observed event. The complementary indicator $(1-L)$ isolates the branches of the history tree that were ruled out by the observation. The term $\lambda\mu((1-L)g)$ calculates the exact expected return of those unrealized branches, which gets added to the offset b . Intuitively, the offset records value that has already been settled by the observation, so that future lower-expectation evaluation remains dynamically consistent with the original ex ante evaluation.

180 After this restriction step, the resulting infradistribution is renormalized, such that

$$\mathbb{E}_{\Psi}(0) = 0, \quad \mathbb{E}_{\Psi}(1) = 1. \quad (6)$$

181 This normalization plays a role analogous to posterior normalization in Bayesian conditioning. In
182 the special case where every infradistribution has exactly one minimal point and all uncertainty
183 is represented by classical mixtures, the IB update reduces to ordinary Bayesian updating. This
184 recovery of Bayes’ rule is an important sanity check for both the formalism and our implementation.

185 An important property of the raw IB update (5) is its linearity, implying that the IB update is a trans-
186 formation in the space of $\mathbf{a}$ -measures mapping straight lines to straight lines. Intuitively, this means
187 that the IB update does not produce new vertices. To reconstruct all relevant information contained
188 in the infradistribution, the implementation thus only needs to track and update the extremal minimal
189 points, which are the vertices of the set of minimal points.

190 3 Related Work

191 Existing IB work is mostly theoretical, covering inframeasures, update rules, dynamic consistency,
192 learnability, and regret-style guarantees [Diffractor and Kosoy, 2020, Appel, 2020, Gelb, 2025a].
193 Our work complements this by giving a finite RL implementation that explicitly represents infradis-
194 tributions, performs IB-style updates, and plans using lower expectations.

195 We situate this contribution relative to three nearby areas: robust RL, imprecise probability and
196 credal sets, and policy-dependent RL.

197 3.1 Robust reinforcement learning

198 Robust RL studies agents that optimize performance under worst-case uncertainty over environment
199 models. In robust MDPs, uncertainty is usually represented as a set of possible transition kernels
200 or reward functions, and the agent chooses a policy that performs well against the least favorable
201 model in that set [Iyengar, 2005, Nilim and El Ghaoui, 2005].

202 This literature is closely related to our work because both approaches replace average-case evalu-
203 ation with worst-case evaluation. The main difference is the representation of uncertainty. Robust
204 MDPs typically retain a fixed, policy-independent environment model with uncertainty over transi-
205 tion or reward parameters. Our approach instead represents uncertainty using infradistributions over
206 affine evaluators and updates those evaluators using the IB update rule.

207 3.2 Imprecise probability and credal sets

208 Imprecise probability theory represents uncertainty using sets of probability measures rather than
209 a single precise distribution. This includes credal sets, lower and upper probabilities, and lower
210 previsions. Walley’s framework of coherent lower previsions is one of the standard foundations for
211 reasoning with imprecise probabilities, while Levi’s work on credal probability helped establish sets
212 of probability functions as representations of belief states [Walley, 1991, Levi, 1980].

213 Credal sets are closed and convex sets of probability distributions. They provide a simple rep-
214 resentation of Knightian uncertainty, where uncertainty cannot be represented by a single precise
215 probability distribution. Credal-set methods are closely related to our work because they also evalu-
216 ate choices using lower expectations over a set of admissible probability models. However, IB is not
217 merely a direct application of credal sets: in the finite setting considered here, an $\mathbf{a}$ -measure contains
218 both a probabilistic component and an offset term.

219 Recent domain-theoretic work on imprecise probability further emphasizes credal sets as structured
220 objects supporting refinement and convergence [Edalat et al., 2026]; our work shares this finite-
221 representation perspective, but implements lower-expectation planning using IB affine evaluators
222 rather than ordinary probability measures.

223 3.3 Policy-dependent and Newcomb-like environments

224 Policy-dependent environments challenge the standard RL assumption that the environment is fixed
225 independently of the agent’s policy. Bell et al. formalize this issue for RL through Newcomb-like

226 decision processes, where rewards or transitions may depend directly on the policy. They show that
 227 value-based RL agents in such environments cannot converge to non-ratifiable policies and may fail
 228 to converge to optimal policies [Bell et al., 2021]. This motivates decision procedures that evaluate
 229 policies more directly, rather than only learning action values in a fixed environment. Our work
 230 provides a finite IB implementation for studying lower-expectation planning under non-realizability
 231 and policy-dependent uncertainty.

232 4 Methods

233 We implement a finite-outcome IB agent for stateless bandit and Newcomb-like decision problems.
 234 The code is available at <https://anonymous.4open.science/r/Project-Newcomb-ebdb/>.

235 The belief state of the agent is stored as a single infradistribution and a world model. The world
 236 model describes the type of environment in which the agent operates. In particular, the representation
 237 of the probability measure μ of each a-measure depends on the world model.

238 4.1 Representing infradistributions

239 General infradistributions are infinite sets of (signed) affine measures. Operationally, measures can
 240 be discarded from this set if they are never needed to determine the lower expectation of any relevant
 241 function through equation (3). For any infradistribution, only its minimal points contribute to the
 242 expectation values [Appel, 2020]. Minimal points that can be written as non-trivial convex combi-
 243 nations of other minimal points also cannot contribute. This means that only the extremal minimal
 244 points need to be stored, analogous to representing a convex polytope by its vertices. This is the key
 245 computational representation used in our implementation.

246 Infradistributions are represented by their extremal minimal points, which is a finite set of a-
 247 measures. Three constructions cover every belief state we use:

- 248 1. Singleton: a set containing a single a-measure encodes a hypothesis without uncertainty.
 249 These infradistributions correspond to knowing the true environment exactly.
- 250 2. Classical (Bayesian) mixture: a weighted combination of infradistributions, according to
 251 equation (4). These represent classical probabilistic uncertainty. If each component is a
 252 singleton, this mixture reduces to standard Bayesian averaging.
- 253 3. Knightian mixture: the set-union of the constituent infradistributions. This mixture carries
 254 no weights. It is evaluated by the worst-case across components.

255 The two mixture operations compose and nest freely, which allows representing the rich belief states
 256 on which an IB agent operates.

257 4.2 World models

258 Conceptually, measures are probability distributions over infinite histories. Making them compu-
 259 tationally tractable requires compressed representations. Each type of environment uses a separate
 260 world model that defines how measures are represented and how predictions are derived from them,
 261 how past observations are represented, and how weighted mixtures of a-measures are computed.
 262 This paper focuses on two world models: Bernoulli bandits and Newcomb-like problems.

263 In a one-armed Bernoulli bandit, the history can be represented as a pair of integers (N, R) , where N
 264 is the number of times the arm was pulled and R is the number of times a reward was obtained. This
 265 representation uses the fact that each round is independent, such that the order of observations does
 266 not matter. A non-mixed measure is represented by a fixed reward probability p . Mixed measures
 267 are represented as a set of pairs (c_i, p_i) , where p_i is the probability and c_i the weight of the i^{th}
 268 component of the mixture and $\sum_i c_i = 1$. Given such a history and measure, the probability of
 269 a certain branch is computed as $\sum_i c_i (1 - p_i)^{N-R} p_i^R$. Predictions for future outcomes can be
 270 computed similarly. Learning proceeds by updating the history (N, R) , implicitly recreating Bayes’
 271 rule. For a k -armed Bernoulli bandit, histories are represented by k pairs (N_j, R_j) and measures by
 272 k sets of pairs (c_{ji}, p_{ji}) . Expectation values are computed per-arm as for the one-armed bandit. This
 273 factorization is possible because arms are independent of each other, i.e. observing one arm does not
 274 give any information about the others.

275 We consider Newcomb-like environments with imperfect predictors. The world model itself con-
 276 tains the entire reward matrix and the predictor accuracy, i.e. the agent knows the entire structure of
 277 the environment. In this setting, there is nothing to be learned. Therefore, measures and histories do
 278 not have an internal state and are not updated upon observations.

279 4.3 Decision and update procedure

280 The aim of the agent is to select the optimal policy, which might be non-deterministic, based on its
 281 belief state. The agent achieves this by discretizing policy space and maintaining a set of candidate
 282 policies Π . At every step of the interaction, the agent will:

- 283 1. Compute the expected value of each policy $\pi \in \Pi$ as $\mathbb{E}_{\Psi(\pi)}[f] = \inf_{a \in \Psi(\pi)} a(f)$, where
 284 $\Psi(\pi)$ is the infradistribution and f is the reward function. The infradistribution can depend
 285 on the policy explicitly, such as in Newcomb-like environments, or implicitly, by sampling
 286 an action from the policy. With $|\Psi| = 1$ this operation reduces to an ordinary expected
 287 value. With $|\Psi| > 1$ it picks the worst-case environment.
- 288 2. Select the optimal policy $\pi^* = \operatorname{argmax}_{\pi \in \Pi} \mathbb{E}_{\Psi(\pi)}[f]$, and sample an action from π^* . Both
 289 the policy and the action are passed to the environment.
- 290 3. Update its belief state upon receiving an observation from the environment. This includes
 291 raw updates of a-measures according to equation (5), updates of the probability measures
 292 as specified by the world model, and the renormalization step. Because of linearity of the
 293 raw update, no new extremal points appear, so it suffices to update the existing set.

294 5 Results and Discussion

295 In the following sections, we introduce experiments that validate our implementation of an infra-
 296 Bayesian agent. First, we demonstrate proper behavior under Knightian uncertainty; next, we val-
 297 idate behavior in Newcomb’s problem. Additional results in the appendix show that an IB agent
 298 reduces to a Bayesian agent when initialized equivalently in a standard stochastic bandit setting (for
 299 which see A), and how an infra-Bayesian agent can meaningfully learn under Knightian uncertainty
 300 (for which see B).

301 5.1 Knightian Uncertainty

302 To demonstrate Knightian uncertainty, we consider a two-armed, adversarial Bernoulli bandit envi-
 303 ronment. Each arm yields reward 1 with a probability chosen anew at the beginning of each step, and
 304 otherwise reward 0. The mechanism by which probabilities are chosen is unknown to the agent and
 305 may potentially be time-dependent or even adversarial. This makes it impossible to learn reward
 306 probabilities over multiple episodes. Past observations do not provide useful data for assigning a
 307 prior to the current interaction.

308 The reward probabilities are constrained to the range $p_1 \in [0.3, 0.7]$ and $p_2 \in [0.4, 0.8]$. The left of
 309 figure 1 displays the space of possible environments. There are two regions: in the upper triangle of
 310 the diagram $p_2 > p_1$, and thus it is optimal to pull arm 2. In the lower triangle, the opposite is true.
 311 For decision making, it only matters whether the agent believes it is in the upper or lower region.
 312 Both regions are possible, within the constraint.

313 A classical Bayesian agent cannot act from the interval constraints alone. It must first replace them
 314 with a precise prior over environments. Different choices of this prior can lead to different policies.
 315 For example, a prior concentrated on an environment where $p_1 > p_2$ recommends arm 1, while a
 316 prior concentrated on an environment where $p_2 > p_1$ recommends arm 2. In this experiment, we
 317 illustrate this prior-dependence by considering classical agents whose priors are point masses at the
 318 corners of the allowed set; of course, this should not be read as the only possible classical choice.
 319 A uniform prior over the allowed set would recommend arm 2 in this example, matching the infra-
 320 Bayesian action. The point is instead that the classical agent’s behavior depends on an additional
 321 prior-selection assumption that is not specified by the interval constraints themselves. By contrast,
 322 the infra-Bayesian agent can represent the constraint directly as Knightian uncertainty, which will
 323 lead it to maximize the worst-case outcome. The worst allowed environment is $(p_1, p_2) = (0.3, 0.4)$,

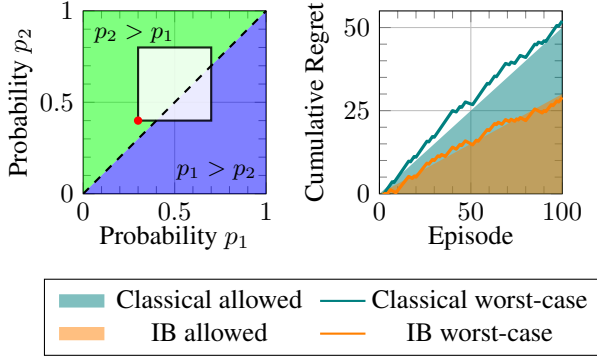


Figure 1: Left: Visualization of environment space (p_1, p_2) . In the blue (green) area, arm 1 (2) yields the higher expected reward. The white box indicates the region that is allowed by the constraint. The red dot indicates the worst allowed environment. Right: cumulative regret of classical and IB agents. The shaded areas show the theoretically allowed ranges. The lines show simulated results from a single roll-out of the worst-case configuration for each agent.

as indicated in the figure. In this environment $p_2 > p_1$, so the agent will always pull arm 2. It is guaranteed to achieve an average reward of 0.4, and does not risk getting 0.3 on arm 1.

The exact reward and regret realized depend on the true environment. The right panel of figure 1 shows the possible ranges of regret achieved by the two agents, for any true environment consistent with the constraint. Also shown are simulated regret curves for the worst-case configurations for each agent. We find that the worst-case outcome for the IB agent realizes a lower regret than the worst-case of the classical agent.

Note that the results here are not meant to demonstrate a meaningful form of learning, either by the infra-Bayesian or classical agents. Indeed, even if arm 1 repeatedly gives higher rewards, an infra-Bayesian agent will not update to exploit them, which is the correct behavior under Knightian uncertainty. Rewards could be controlled by an unknown or adversarial mechanism that makes arm 1 look attractive, only to give it a lower reward when the agent chooses it. Sticking to arm 2 is therefore the robust and safe strategy in this scenario. We present a stochastic bandit setting, in which the IB agent does learn despite Knightian uncertainty in appendix B.

5.2 Policy dependence

Policy dependence is investigated via Newcomb’s problem with an imperfect predictor. We consider a setting in which the transparent box always contains \$1. The opaque box contains \$10 if the agent is predicted to one-box, and \$0 otherwise. The resulting payoffs are shown in table 1. The agent chooses a policy that one-boxes with probability p . A predictor with accuracy $\alpha \in [0.5, 1]$ reads the agent’s policy and predicts one-boxing with probability $p \cdot (2\alpha - 1) + 0.5 \cdot (2 - 2\alpha)$. A perfect predictor ($\alpha = 1$) predicts one-boxing at the same rate p that the agent one-boxes. A random predictor ($\alpha = 0.5$) guesses at random, thus its prediction is independent of the agent’s policy.

	Predicted one-box	Predicted two-box
One-box	10	0
Two-box	11	1

Table 1: Reward matrix for Newcomb’s problem

We implement this setting using the Newcomb-like world model described in section 4.2. The optimal policy and reward are shown in figure 2. For a sufficiently accurate predictor ($\alpha > 0.55$), one-boxing is the optimal strategy. For a nearly random predictor ($\alpha < 0.55$), two-boxing is optimal. The figure also shows simulated results from our implementation. We find that the agent consistently selects the optimal policy and achieves the corresponding optimal reward.

An agent following causal decision theory would two-box in Newcomb’s problem, arguing that its decision cannot influence the already fixed box contents after the predictor’s move, thus missing out on the large reward when the predictor is sufficiently accurate. An agent following evidential decision theory also one-boxes in this scenario, but for a different reason. It fails to accurately model the causal structure of the problem and can misbehave in similar scenarios. This scenario demonstrates IB decision theory selecting the optimal action, and for the correct reason, given full information of the reward structure. The same approach also achieves optimal behavior in the (Asymmetric) Death

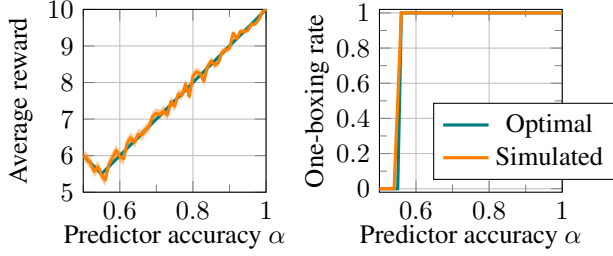


Figure 2: Average reward (left) and one-boxing rate (right) in Newcomb’s problem as a function of the predictor accuracy. Shown are both optimal and simulated values, averaged over 1000 episodes. For $\alpha = 0.55$, the reward is independent of the one-boxing rate and thus every rate is optimal.

in Damascus and Coordination Game scenarios, some of which require mixed strategies [Bell et al., 2021].

6 Conclusion, Limitations, and Future Work

This paper presents a proof-of-concept infra-Bayesian reinforcement learning architecture for finite-outcome stateless decision problems. Our design implements a-measures, infradistributions, classical and Knightian mixtures, and the IB conditioning rule, while remaining computationally tractable by storing only extremal minimal points and using optimized representations for histories and belief states. We find that the architecture recovers standard Bayesian behavior as a special case, confirming that IB generalizes rather than replaces classical reasoning. In a Knightian-uncertain bandits setting, the IB agent identifies and optimizes against the worst-case environment within the admissible set, while classical agents depend on arbitrary priors. This demonstrates concretely that the distinction between probabilistic and Knightian uncertainty, central to the IB formalism, produces meaningfully different agent behavior in settings where classical methods may be ambiguous. More broadly, real-world agents must operate under misspecification and policy-dependent environments. In such settings, optimizing expected value can lead to brittle behavior, while worst-case optimization offers robustness. Bridging IB theory to concrete agent implementations is a step toward RL systems that are robust by design.

Our implementation is restricted to finite outcomes, nonnegative a-measures, and small hypothesis spaces. Scaling to continuous state spaces, large hypothesis classes, and function approximation remains open. Future work – some of which is already underway – would seek to address these concerns, most notably permitting optimization of multi-step decision processes under Knightian uncertainty and fully leveraging IB’s advantages in keeping track of multiple possible environments. Nevertheless, the results show that IB reasoning can be translated into a working RL agent. Additionally, our regret bounds, while an improvement over those of classical RL, remain linear, though those regret bounds for classical RL presuppose that a classical agent can converge to a ratifiable policy at all.

References

- Alexander Appel. Basic inframeasure theory. *AI Alignment Forum*, September 2020. URL <https://www.alignmentforum.org/posts/basic-inframeasure-theory>.
- James Bell, Linda Linsefors, Caspar Oesterheld, and Joar Skalse. Reinforcement learning in new-comblike environments. In *Advances in Neural Information Processing Systems*, volume 34, pages 27284–27295, 2021.
- Diffraction and Vanessa Kosoy. Introduction To The Infra-Bayesianism Sequence — LessWrong. August 2020. URL <https://www.lesswrong.com/posts/zB4f7QqKhBH5b37a/introduction-to-the-infra-bayesianism-sequence>.
- Abbas Edalat, Pietro Di Gianantonio, and Amin Farjudian. A domain-theoretic foundation for imprecise probability and credal sets, 2026. URL <https://arxiv.org/abs/2604.09272>.
- Brittany Gelb. An introduction to credal sets and infra-bayes learnability. *AI Alignment Forum*, August 2025a. URL <https://www.alignmentforum.org/s/n7qFxakSnxGuvYAX/p/rkhaRnAc6dLzQT2sJ>.
- Brittany Gelb. What is inadequate about bayesianism for ai alignment: Motivating infra-bayesianism. *AI Alignment Forum*, May 2025b. URL <https://www.alignmentforum.org/posts/wzCtwYtojMabyEg2L/what-is-inadequate-about-bayesianism-for-ai-alignment>.
- Mohammad Ghavamzadeh, Shie Mannor, Joelle Pineau, and Aviv Tamar. Bayesian reinforcement learning: A survey. *Foundations and Trends® in Machine Learning*, 8(5–6):359–489, 2015.
- Garud N. Iyengar. Robust dynamic programming. *Mathematics of Operations Research*, 30(2):257–280, 2005.
- Vanessa Kosoy. The Learning-Theoretic Agenda: Status 2023 — LessWrong. April 2023. URL <https://www.lesswrong.com/posts/ZwshvqiqCvXPszEct/the-learning-theoretic-agenda-status-2023>.
- Isaac Levi. *The Enterprise of Knowledge: An Essay on Knowledge, Credal Probability, and Chance*. MIT Press, Cambridge, MA, 1980.
- Arnab Nilim and Laurent El Ghaoui. Robust control of markov decision processes with uncertain transition matrices. *Operations Research*, 53(5):780–798, 2005.
- Robert Nozick. Newcomb’s problem and two principles of choice. In *Essays in honor of carl g. hempel: A tribute on the occasion of his sixty-fifth birthday*, pages 114–146. Springer, 1969.
- Richard S. Sutton and Andrew G. Barto. *Reinforcement Learning: An Introduction*. MIT Press, 2018.
- Peter Walley. *Statistical Reasoning with Imprecise Probabilities*. Chapman and Hall, London, 1991.

423 Appendix

424 A Empirical validation that an infra-Bayesian reduces to a classical 425 Bayesian agent given identical hypotheses

426 To validate our implementation, we test whether our infra-Bayesian implementation reproduces ordi-
427 nary Bayesian behavior when the agents are initialized with identical hypothesis sets. We consider
428 standard two-armed Bernoulli bandits and compare a classical discrete Bayesian agent against an
429 infra-Bayesian agent with a single a -measure. This a -measure contains the same finite Bernoulli
430 hypothesis grid used by the classical agent, so the pessimistic minimum over a -measures is vac-
431 uous and the infra-Bayesian posterior predictive should reduce to the classical Bayesian posterior
432 predictive.

433 Figure 3 shows the resulting cumulative regret and action probabilities across four stochastic, two-
434 armed bandit environments. Solid, transparent curves are the classical Bayesian agent and opaque,
435 dashed curves are the corresponding infra-Bayesian agent; colors indicate different bandit environ-
436 ment. The curves coincide exactly under matched random seeds, confirming that the infra-Bayesian
437 agent recovers the classical Bayesian update and action rule in the single-measure setting.

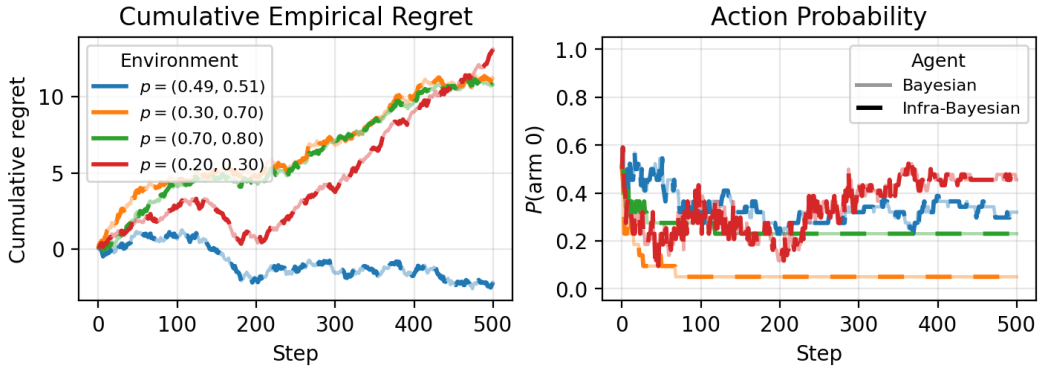


Figure 3: Validation that the infra-Bayesian agent reproduces classical Bayesian bandit behavior in the precise-prior, single- a -measure case. The legend gives the priors of each arm in four unique Bernoulli bandit settings.

438 B Trap Bandit Experiments

439 We consider a simple experiment illustrating how a robust infra-Bayesian learner can be useful even
440 in a stateless stochastic bandit setting. The environment has two arms, indexed by $i \in \{1, 2\}$. At the
441 start of each run, the arm reward probabilities (p_1, p_2) are sampled uniformly from

$$\{(0.3, 0.7), (0.7, 0.3)\},$$

442 so that one arm has significantly higher expected reward than the other. Independently, the world is
443 sampled to be either safe or risky, with probability α of being risky and probability $1 - \alpha$ of being
444 safe.

445 In a safe world, each arm is Bernoulli with reward probability p_i . In a risky world, the arm with
446 the larger reward probability is also the trapped arm. Pulling this arm yields a catastrophic reward
447 of -1000 with probability p_{cat} , yields reward 1 with probability p_i , and yields reward 0 otherwise.

α_{DGP}	α_{prior}	Agent	Catastrophe rate	p50, 95% CI	p95, 95% CI
0.99	n/a	infra_bayesian	0.045	9.60 [9.60, 9.60]	125.76 [105.60, 183.36]
0.99	0.99	bayes_greedy	0.045	9.60 [9.60, 9.60]	125.76 [96.48, 183.36]
0.99	0.01	bayes_greedy	0.645	624.00 [528.00, 748.80]	960.00 [950.88, 960.00]
0.99	0.99	bayes_thompson	0.080	38.40 [28.80, 38.40]	115.20 [96.00, 144.00]
0.99	0.01	bayes_thompson	0.645	580.80 [489.60, 720.00]	950.40 [940.80, 950.40]
0.01	n/a	infra_bayesian	0.000	39.60 [39.60, 39.60]	40.00 [40.00, 40.00]
0.01	0.01	bayes_greedy	0.015	0.40 [0.40, 0.80]	5.60 [4.02, 9.64]
0.01	0.01	bayes_thompson	0.015	1.60 [1.20, 1.60]	7.60 [5.22, 11.60]

Table 2: Final cumulative expected-regret percentiles with bootstrap confidence intervals.

448 The lower-probability arm remains Bernoulli. Thus each run is generated as follows:

sample $(p_1, p_2) \sim \text{Unif}\{(0.3, 0.7), (0.7, 0.3)\},$
sample world type $\sim \text{Bernoulli}(\alpha),$
if safe: $r_i \sim \text{Bernoulli}(p_i),$
if risky: $\begin{cases} \text{the trap arm } \arg \max_i p_i \text{ yields } -1000 \text{ with probability } p_{\text{cat}}, \\ \text{the trap arm yields } 1 \text{ with probability } p_i, \\ \text{the trap arm yields } 0 \text{ otherwise,} \\ \text{the other arm yields } r_i \sim \text{Bernoulli}(p_i). \end{cases}$

449 We compare classical Bayesian agents with an infra-Bayesian agent using the same joint hypothesis
450 machinery. Bayesian agents represent uncertainty using `Infradistribution.mix(...)`. The
451 infra-Bayesian agent instead uses Knightian uncertainty over the safe and risky world families via
452 `Infradistribution.mixKU(...)`, while retaining ordinary Bayesian uncertainty over (p_1, p_2)
453 within each family.

454 The experiments vary the relationship between the true risky-world probability, α_{DGP} , and the
455 Bayesian agent’s point prior, α_{prior} . In the main mostly-risky setting, we set $\alpha_{\text{DGP}} = 0.99$. We
456 first consider a correctly specified Bayesian prior, $\alpha_{\text{prior}} = 0.99$, and then a severely misspecified
457 prior, $\alpha_{\text{prior}} = 0.01$. Across these conditions, the infra-Bayesian agent uses the same classical prior
458 over (p_1, p_2) as the Bayesian agents, but maintains Knightian uncertainty over whether the world is
459 safe or risky.

460 For Bayesian agents, we evaluate two exploration strategies: greedy action selection and Thompson
461 sampling. The infra-Bayesian agent uses greedy action selection with respect to its robust lower
462 values, with uniform tie-breaking. Regret is measured relative to the best policy with full knowl-
463 edge of the true world. We report cumulative expected regret percentiles and trapped-arm pull-rate
464 percentiles. Results are shown in Figure 4.

465 When the Bayesian prior is correctly specified, greedy Bayes and the infra-Bayesian agent behave
466 nearly identically in risky worlds. This is expected: when the risky-world probability is high,
467 expected-value maximization already favors the conservative action. The infra-Bayesian agent nev-
468 ertheless obtains this behavior without committing to a point prior over the safe/risky model class.
469 Under severe misspecification, however, greedy Bayes initially treats the world as mostly safe, re-
470 peatedly pulls the high-reward trapped arm, and incurs much larger expected regret and a high
471 catastrophe rate. Thompson sampling slightly reduces regret in this misspecified condition, but does
472 not remove the failure mode.

473 These results show infra-Bayesian agents matching or outperforming classical Bayesian agents in
474 risky worlds, which raises a natural question: at what cost? Table 2 includes an alternative mostly-
475 safe scenario with $\alpha_{\text{DGP}} = 0.01$. In this setting, the infra-Bayesian agent incurs substantially higher
476 regret, reflecting the cost of maintaining Knightian uncertainty over the risky-world hypothesis.

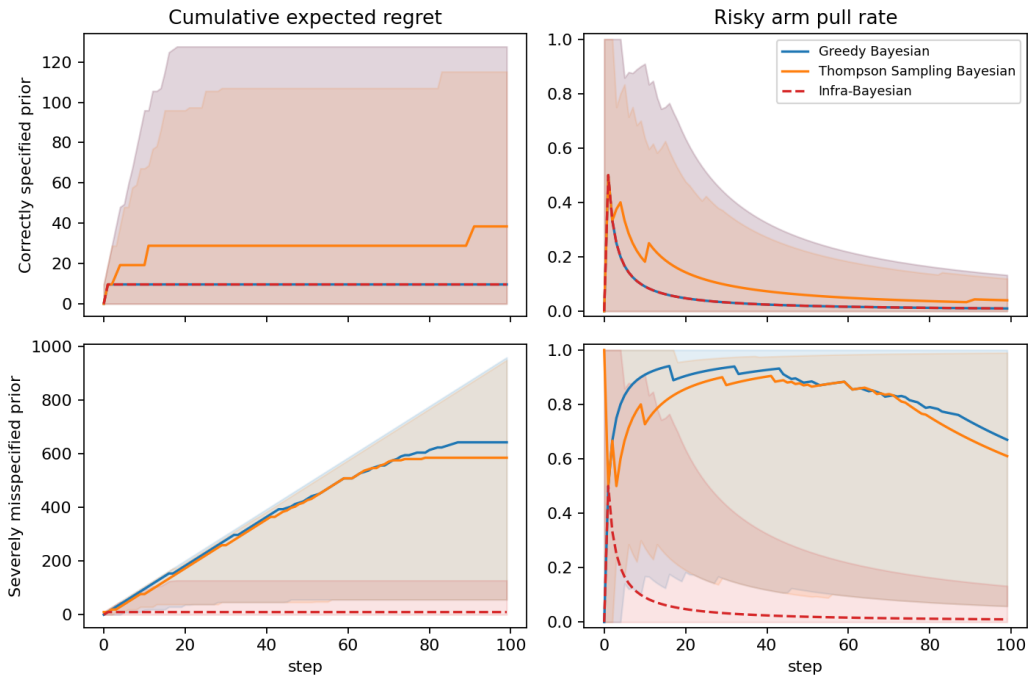


Figure 4: Comparing the performance of infra-Bayesian and classical Bayesian agents (with either greedy or Thompson Sampling exploration strategies) in the trap bandit setting. The first row shows results for a correctly specified Bayes prior condition; the second for a severely misspecified Bayes prior condition. The first column shows cumulative expected regret, and the second shows the average pull rate of the risky (trap) arm.

NeurIPS Paper Checklist

1. Claims

Question: Do the main claims made in the abstract and introduction accurately reflect the paper’s contributions and scope?

Answer: [Yes]

Justification: We claim to implement an infra-Bayesian RL agent and run some tests on it. We do precisely that.

Guidelines:

- The answer [N/A] means that the abstract and introduction do not include the claims made in the paper.
- The abstract and/or introduction should clearly state the claims made, including the contributions made in the paper and important assumptions and limitations. A [No] or [N/A] answer to this question will not be perceived well by the reviewers.
- The claims made should match theoretical and experimental results, and reflect how much the results can be expected to generalize to other settings.
- It is fine to include aspirational goals as motivation as long as it is clear that these goals are not attained by the paper.

2. Limitations

Question: Does the paper discuss the limitations of the work performed by the authors?

Answer: [Yes]

Justification: Limitations are mention in the "Limitations and Future Work" section.

Guidelines:

- The answer [N/A] means that the paper has no limitation while the answer [No] means that the paper has limitations, but those are not discussed in the paper.
- The authors are encouraged to create a separate "Limitations" section in their paper.
- The paper should point out any strong assumptions and how robust the results are to violations of these assumptions (e.g., independence assumptions, noiseless settings, model well-specification, asymptotic approximations only holding locally). The authors should reflect on how these assumptions might be violated in practice and what the implications would be.
- The authors should reflect on the scope of the claims made, e.g., if the approach was only tested on a few datasets or with a few runs. In general, empirical results often depend on implicit assumptions, which should be articulated.
- The authors should reflect on the factors that influence the performance of the approach. For example, a facial recognition algorithm may perform poorly when image resolution is low or images are taken in low lighting. Or a speech-to-text system might not be used reliably to provide closed captions for online lectures because it fails to handle technical jargon.
- The authors should discuss the computational efficiency of the proposed algorithms and how they scale with dataset size.
- If applicable, the authors should discuss possible limitations of their approach to address problems of privacy and fairness.
- While the authors might fear that complete honesty about limitations might be used by reviewers as grounds for rejection, a worse outcome might be that reviewers discover limitations that aren’t acknowledged in the paper. The authors should use their best judgment and recognize that individual actions in favor of transparency play an important role in developing norms that preserve the integrity of the community. Reviewers will be specifically instructed to not penalize honesty concerning limitations.

3. Theory assumptions and proofs

Question: For each theoretical result, does the paper provide the full set of assumptions and a complete (and correct) proof?

Answer: [Yes]

Justification: We only make use of theoretical results, which we cite extensively; this paper contains no new theoretical results.

Guidelines:

- The answer [N/A] means that the paper does not include theoretical results.
- All the theorems, formulas, and proofs in the paper should be numbered and cross-referenced.
- All assumptions should be clearly stated or referenced in the statement of any theorems.
- The proofs can either appear in the main paper or the supplemental material, but if they appear in the supplemental material, the authors are encouraged to provide a short proof sketch to provide intuition.
- Inversely, any informal proof provided in the core of the paper should be complemented by formal proofs provided in appendix or supplemental material.
- Theorems and Lemmas that the proof relies upon should be properly referenced.

4. Experimental result reproducibility

Question: Does the paper fully disclose all the information needed to reproduce the main experimental results of the paper to the extent that it affects the main claims and/or conclusions of the paper (regardless of whether the code and data are provided or not)?

Answer: [Yes]

Justification: Our linked Github repository contains all code needed to replicate our results, and we hope to develop it further in future.

Guidelines:

- The answer [N/A] means that the paper does not include experiments.
- If the paper includes experiments, a [No] answer to this question will not be perceived well by the reviewers: Making the paper reproducible is important, regardless of whether the code and data are provided or not.
- If the contribution is a dataset and/or model, the authors should describe the steps taken to make their results reproducible or verifiable.
- Depending on the contribution, reproducibility can be accomplished in various ways. For example, if the contribution is a novel architecture, describing the architecture fully might suffice, or if the contribution is a specific model and empirical evaluation, it may be necessary to either make it possible for others to replicate the model with the same dataset, or provide access to the model. In general, releasing code and data is often one good way to accomplish this, but reproducibility can also be provided via detailed instructions for how to replicate the results, access to a hosted model (e.g., in the case of a large language model), releasing of a model checkpoint, or other means that are appropriate to the research performed.
- While NeurIPS does not require releasing code, the conference does require all submissions to provide some reasonable avenue for reproducibility, which may depend on the nature of the contribution. For example
 - (a) If the contribution is primarily a new algorithm, the paper should make it clear how to reproduce that algorithm.
 - (b) If the contribution is primarily a new model architecture, the paper should describe the architecture clearly and fully.
 - (c) If the contribution is a new model (e.g., a large language model), then there should either be a way to access this model for reproducing the results or a way to reproduce the model (e.g., with an open-source dataset or instructions for how to construct the dataset).
 - (d) We recognize that reproducibility may be tricky in some cases, in which case authors are welcome to describe the particular way they provide for reproducibility. In the case of closed-source models, it may be that access to the model is limited in some way (e.g., to registered users), but it should be possible for other researchers to have some path to reproducing or verifying the results.

5. Open access to data and code

Question: Does the paper provide open access to the data and code, with sufficient instructions to faithfully reproduce the main experimental results, as described in supplemental material?

Answer: [Yes]

Justification: As for the previous question, all the necessary code is contained in the Github repository, which we hope that you will take a look at and which we hope to further improve in future.

Guidelines:

- The answer [N/A] means that paper does not include experiments requiring code.
- Please see the NeurIPS code and data submission guidelines (<https://neurips.cc/public/guides/CodeSubmissionPolicy>) for more details.
- While we encourage the release of code and data, we understand that this might not be possible, so [No] is an acceptable answer. Papers cannot be rejected simply for not including code, unless this is central to the contribution (e.g., for a new open-source benchmark).
- The instructions should contain the exact command and environment needed to run to reproduce the results. See the NeurIPS code and data submission guidelines (<https://neurips.cc/public/guides/CodeSubmissionPolicy>) for more details.
- The authors should provide instructions on data access and preparation, including how to access the raw data, preprocessed data, intermediate data, and generated data, etc.
- The authors should provide scripts to reproduce all experimental results for the new proposed method and baselines. If only a subset of experiments are reproducible, they should state which ones are omitted from the script and why.
- At submission time, to preserve anonymity, the authors should release anonymized versions (if applicable).
- Providing as much information as possible in supplemental material (appended to the paper) is recommended, but including URLs to data and code is permitted.

6. Experimental setting/details

Question: Does the paper specify all the training and test details (e.g., data splits, hyperparameters, how they were chosen, type of optimizer) necessary to understand the results?

Answer: [Yes]

Justification: Once again, our experiments were very simple and did not require extensive test details;

Guidelines:

- The answer [N/A] means that the paper does not include experiments.
- The experimental setting should be presented in the core of the paper to a level of detail that is necessary to appreciate the results and make sense of them.
- The full details can be provided either with the code, in appendix, or as supplemental material.

7. Experiment statistical significance

Question: Does the paper report error bars suitably and correctly defined or other appropriate information about the statistical significance of the experiments?

Answer: [Yes]

Justification: All of our plots include appropriate statistical significance information.

Guidelines:

- The answer [N/A] means that the paper does not include experiments.
- The authors should answer [Yes] if the results are accompanied by error bars, confidence intervals, or statistical significance tests, at least for the experiments that support the main claims of the paper.
- The factors of variability that the error bars are capturing should be clearly stated (for example, train/test split, initialization, random drawing of some parameter, or overall run with given experimental conditions).

- The method for calculating the error bars should be explained (closed form formula, call to a library function, bootstrap, etc.)
- The assumptions made should be given (e.g., Normally distributed errors).
- It should be clear whether the error bar is the standard deviation or the standard error of the mean.
- It is OK to report 1-sigma error bars, but one should state it. The authors should preferably report a 2-sigma error bar than state that they have a 96% CI, if the hypothesis of Normality of errors is not verified.
- For asymmetric distributions, the authors should be careful not to show in tables or figures symmetric error bars that would yield results that are out of range (e.g., negative error rates).
- If error bars are reported in tables or plots, the authors should explain in the text how they were calculated and reference the corresponding figures or tables in the text.

8. Experiments compute resources

Question: For each experiment, does the paper provide sufficient information on the computer resources (type of compute workers, memory, time of execution) needed to reproduce the experiments?

Answer: [No]

Justification: The experiments we ran to verify the performance of our IBRL agents can be straightforwardly run on a laptop.

Guidelines:

- The answer [N/A] means that the paper does not include experiments.
- The paper should indicate the type of compute workers CPU or GPU, internal cluster, or cloud provider, including relevant memory and storage.
- The paper should provide the amount of compute required for each of the individual experimental runs as well as estimate the total compute.
- The paper should disclose whether the full research project required more compute than the experiments reported in the paper (e.g., preliminary or failed experiments that didn't make it into the paper).

9. Code of ethics

Question: Does the research conducted in the paper conform, in every respect, with the NeurIPS Code of Ethics <https://neurips.cc/public/EthicsGuidelines>?

Answer: [Yes]

Justification: As our research concerns itself entirely with a seminal implementation of a type of RL agent with an improved decision theory, concerns about human subjects, data concerns, and societal impact are obviated.

Guidelines:

- The answer [N/A] means that the authors have not reviewed the NeurIPS Code of Ethics.
- If the authors answer [No], they should explain the special circumstances that require a deviation from the Code of Ethics.
- The authors should make sure to preserve anonymity (e.g., if there is a special consideration due to laws or regulations in their jurisdiction).

10. Broader impacts

Question: Does the paper discuss both potential positive societal impacts and negative societal impacts of the work performed?

Answer: [N/A]

Justification: We implemented a novel form of RL which specifically avoids worst-case outcomes and which is as yet only suitable for toy problems.

Guidelines:

- The answer [N/A] means that there is no societal impact of the work performed.

- If the authors answer [N/A] or [No], they should explain why their work has no societal impact or why the paper does not address societal impact.
- Examples of negative societal impacts include potential malicious or unintended uses (e.g., disinformation, generating fake profiles, surveillance), fairness considerations (e.g., deployment of technologies that could make decisions that unfairly impact specific groups), privacy considerations, and security considerations.
- The conference expects that many papers will be foundational research and not tied to particular applications, let alone deployments. However, if there is a direct path to any negative applications, the authors should point it out. For example, it is legitimate to point out that an improvement in the quality of generative models could be used to generate Deepfakes for disinformation. On the other hand, it is not needed to point out that a generic algorithm for optimizing neural networks could enable people to train models that generate Deepfakes faster.
- The authors should consider possible harms that could arise when the technology is being used as intended and functioning correctly, harms that could arise when the technology is being used as intended but gives incorrect results, and harms following from (intentional or unintentional) misuse of the technology.
- If there are negative societal impacts, the authors could also discuss possible mitigation strategies (e.g., gated release of models, providing defenses in addition to attacks, mechanisms for monitoring misuse, mechanisms to monitor how a system learns from feedback over time, improving the efficiency and accessibility of ML).

11. Safeguards

Question: Does the paper describe safeguards that have been put in place for responsible release of data or models that have a high risk for misuse (e.g., pre-trained language models, image generators, or scraped datasets)?

Answer: [N/A]

Justification: Our IBRL agent does not run a risk of outputting harmful or sensitive data or artifacts.

Guidelines:

- The answer [N/A] means that the paper poses no such risks.
- Released models that have a high risk for misuse or dual-use should be released with necessary safeguards to allow for controlled use of the model, for example by requiring that users adhere to usage guidelines or restrictions to access the model or implementing safety filters.
- Datasets that have been scraped from the Internet could pose safety risks. The authors should describe how they avoided releasing unsafe images.
- We recognize that providing effective safeguards is challenging, and many papers do not require this, but we encourage authors to take this into account and make a best faith effort.

12. Licenses for existing assets

Question: Are the creators or original owners of assets (e.g., code, data, models), used in the paper, properly credited and are the license and terms of use explicitly mentioned and properly respected?

Answer: [Yes]

Justification: All code authors are properly credited in the repository, and no extra assets have been used.

Guidelines:

- The answer [N/A] means that the paper does not use existing assets.
- The authors should cite the original paper that produced the code package or dataset.
- The authors should state which version of the asset is used and, if possible, include a URL.
- The name of the license (e.g., CC-BY 4.0) should be included for each asset.

- For scraped data from a particular source (e.g., website), the copyright and terms of service of that source should be provided.
- If assets are released, the license, copyright information, and terms of use in the package should be provided. For popular datasets, paperswithcode.com/datasets has curated licenses for some datasets. Their licensing guide can help determine the license of a dataset.
- For existing datasets that are re-packaged, both the original license and the license of the derived asset (if it has changed) should be provided.
- If this information is not available online, the authors are encouraged to reach out to the asset's creators.

13. New assets

Question: Are new assets introduced in the paper well documented and is the documentation provided alongside the assets?

Answer: [\[Yes\]](#)

Justification: All details of any new assets we have generated can be readily found in the repository.

Guidelines:

- The answer [\[N/A\]](#) means that the paper does not release new assets.
- Researchers should communicate the details of the dataset/code/model as part of their submissions via structured templates. This includes details about training, license, limitations, etc.
- The paper should discuss whether and how consent was obtained from people whose asset is used.
- At submission time, remember to anonymize your assets (if applicable). You can either create an anonymized URL or include an anonymized zip file.

14. Crowdsourcing and research with human subjects

Question: For crowdsourcing experiments and research with human subjects, does the paper include the full text of instructions given to participants and screenshots, if applicable, as well as details about compensation (if any)?

Answer: [\[N/A\]](#)

Justification: Our research did not involve crowdsourcing or research with human subjects.

Guidelines:

- The answer [\[N/A\]](#) means that the paper does not involve crowdsourcing nor research with human subjects.
- Including this information in the supplemental material is fine, but if the main contribution of the paper involves human subjects, then as much detail as possible should be included in the main paper.
- According to the NeurIPS Code of Ethics, workers involved in data collection, curation, or other labor should be paid at least the minimum wage in the country of the data collector.

15. Institutional review board (IRB) approvals or equivalent for research with human subjects

Question: Does the paper describe potential risks incurred by study participants, whether such risks were disclosed to the subjects, and whether Institutional Review Board (IRB) approvals (or an equivalent approval/review based on the requirements of your country or institution) were obtained?

Answer: [\[N/A\]](#)

Justification: Our research did not involve crowdsourcing or research with human subjects.

Guidelines:

- The answer [\[N/A\]](#) means that the paper does not involve crowdsourcing nor research with human subjects.

- 789
- Depending on the country in which research is conducted, IRB approval (or equivalent) may be required for any human subjects research. If you obtained IRB approval, you should clearly state this in the paper.
 - 790
 - 791
 - 792
 - We recognize that the procedures for this may vary significantly between institutions and locations, and we expect authors to adhere to the NeurIPS Code of Ethics and the guidelines for their institution.
 - 793
 - 794
 - For initial submissions, do not include any information that would break anonymity (if applicable), such as the institution conducting the review.
 - 795
 - 796

797 **16. Declaration of LLM usage**

798 Question: Does the paper describe the usage of LLMs if it is an important, original, or
799 non-standard component of the core methods in this research? Note that if the LLM is used
800 only for writing, editing, or formatting purposes and does *not* impact the core methodology,
801 scientific rigor, or originality of the research, declaration is not required.

802 Answer: [N/A]

803 Justification: LLMs played no role in the core methodology, scientific rigor, or originality
804 of our research.

805 Guidelines:

- 806 • The answer [N/A] means that the core method development in this research does not
807 involve LLMs as any important, original, or non-standard components.
- 808 • Please refer to our LLM policy in the NeurIPS handbook for what should or should
809 not be described.