

Safe Imitation Learning via Fast Bayesian Reward Inference from Preferences

Daniel S. Brown¹ Russell Coleman^{1,2} Ravi Srinivasan² Scott Niekum¹

Abstract

Bayesian reward learning from demonstrations enables rigorous safety and uncertainty analysis when performing imitation learning. However, Bayesian reward learning methods are typically computationally intractable for complex control problems. We propose Bayesian Reward Extrapolation (Bayesian REX), a highly efficient Bayesian reward learning algorithm that scales to high-dimensional imitation learning problems by pre-training a low-dimensional feature encoding via self-supervised tasks and then leveraging preferences over demonstrations to perform fast Bayesian inference. Bayesian REX can learn to play Atari games from demonstrations, without access to the game score and can generate 100,000 samples from the posterior over reward functions in only 5 minutes on a personal laptop. Bayesian REX also results in imitation learning performance that is competitive with or better than state-of-the-art methods that only learn point estimates of the reward function. Finally, Bayesian REX enables efficient high-confidence policy evaluation without having access to samples of the reward function. These high-confidence performance bounds can be used to rank the performance and risk of a variety of evaluation policies and provide a way to detect reward hacking behaviors.

demonstrations, it is important for an agent to be able to provide high-confidence bounds on its performance with respect to the demonstrator; however, while there exists much work on high-confidence off-policy evaluation in the reinforcement learning (RL) setting, there has been much less work on high-confidence policy evaluation in the imitation learning setting, where the reward samples are unavailable.

Prior work on high-confidence policy evaluation for imitation learning has used Bayesian inverse reinforcement learning (IRL) (Ramachandran & Amir, 2007) to allow an agent to reason about reward uncertainty and policy generalization error (Brown et al., 2018). However, Bayesian IRL is typically intractable for complex problems due to the need to repeatedly solve an MDP in the inner loop, resulting in high computational cost as well as high sample cost if a model is not available. This precludes robust safety and uncertainty analysis for imitation learning in high-dimensional problems or in problems in which a model of the MDP is unavailable. We seek to remedy this problem by proposing and evaluating a method for safe and efficient Bayesian reward learning via preferences over demonstrations. Preferences over trajectories are intuitive for humans to provide (Akrouf et al., 2011; Wilson et al., 2012; Sadigh et al., 2017; Christiano et al., 2017; Palan et al., 2019) and enable better-than-demonstrator performance (Brown et al., 2019b;a). To the best of our knowledge, we are the first to show that preferences over demonstrations enable both fast Bayesian reward learning in high-dimensional, visual control tasks as well as efficient high-confidence performance bounds.

1. Introduction

It is important that robots and other autonomous agents can safely learn from and adapt to a variety of human preferences and goals. One common way to learn preferences and goals is via imitation learning, in which an autonomous agent learns how to perform a task by observing demonstrations of the task (Argall et al., 2009). When learning from

We first formalize the problem of high-confidence policy evaluation (Thomas et al., 2015) for imitation learning. We then propose a novel algorithm, Bayesian Reward Extrapolation (Bayesian REX), that uses a pairwise ranking likelihood to significantly increase the efficiency of generating samples from the posterior distribution over reward functions. We demonstrate that Bayesian REX can leverage neural network function approximation to learn useful reward features via self-supervised learning in order to efficiently perform deep Bayesian reward inference from visual demonstrations. Finally, we demonstrate that samples obtained from Bayesian REX can be used to solve the high-confidence policy evaluation problem for imitation learning. We evaluate our method on imitation learning for Atari games and demonstrate that we can efficiently compute high-confidence bounds on pol-

¹Computer Science Department, The University of Texas at Austin. ²Applied Research Laboratories, The University of Texas at Austin. Correspondence to: Daniel Brown <ds-brown@cs.utexas.edu>.

icy performance, without access to samples of the reward function. We use these high-confidence performance bounds to rank different evaluation policies according to their risk and expected return under the posterior distribution over the unknown ground-truth reward function. Finally, we provide evidence that bounds on uncertainty and risk provide a useful tool for detecting reward hacking/gaming (Amodei et al., 2016), a common problem in reward inference from demonstrations (Ibarz et al., 2018) as well as reinforcement learning (Ng et al., 1999; Leike et al., 2017).

2. Related work

2.1. Imitation Learning

Imitation learning is the problem of learning a policy from demonstrations and can roughly be divided into techniques that use behavioral cloning and techniques that use inverse reinforcement learning. Behavioral cloning methods (Pomerleau, 1991; Torabi et al., 2018) seek to solve the imitation learning problem via supervised learning, in which the goal is to learn a mapping from states to actions that mimics the demonstrator. While computationally efficient, these methods suffer from compounding errors (Ross et al., 2011). Methods such as DAgger (Ross et al., 2011) and DART (Laskey et al., 2017) avoid this problem by repeatedly collecting additional state-action pairs from an expert.

Inverse reinforcement learning (IRL) methods seek to solve the imitation learning problem by estimating the reward function that the demonstrator is optimizing (Ng & Russell, 2000). Classical approaches repeatedly alternate between a reward estimation step and a full policy optimization step (Abbeel & Ng, 2004; Ziebart et al., 2008; Ramachandran & Amir, 2007). Bayesian IRL (Ramachandran & Amir, 2007) samples from the posterior distribution over reward functions, whereas other methods seek a single reward function that induces the demonstrator’s feature expectations (Abbeel & Ng, 2004), often while also seeking to maximize the entropy of the resulting policy (Ziebart et al., 2008).

Most deep learning approaches for IRL use maximum entropy policy optimization and divergence minimization between marginal state-action distributions (Ho & Ermon, 2016; Fu et al., 2017; Ghasemipour et al., 2019) and are related to Generative Adversarial Networks (Finn et al., 2016a). These methods scale to complex control problems by iterating between reward learning and policy learning steps. Alternatively, Brown et al. (2019b) use ranked demonstrations to learn a reward function via supervised learning without requiring an MDP solver or any inference time data collection. The learned reward function can then be used to optimize a potentially better-than-demonstrator policy. Brown et al. (2019a) automatically generate preferences over demonstrations via noise injection, allowing better-

than-demonstrator performance even in the absence of explicit preference labels. However, despite their successes, deep learning approaches to IRL typically only return a point estimate of the reward function, precluding uncertainty and robustness analysis.

2.2. Safe Imitation Learning

While there has been much interest in imitation learning, less attention has been given to problems related to safety. SafeDAgger (Zhang & Cho, 2017) and EnsembleDAgger (Menda et al., 2019) are extensions of DAgger that give control to the demonstrator in states where the imitation learning policy is predicted to have a large action difference from the demonstrator. Other approaches to safe imitation learning seek to match the tail risk of the expert as well as find a policy that is indistinguishable from the demonstrations (Majumdar et al., 2017; Lacotte et al., 2019).

Brown & Niekum (2018) propose a Bayesian sampling approach to provide explicit high-confidence performance bounds in the imitation learning setting, but require an MDP solver in the inner-loop. Their method uses samples from the posterior distribution $P(R|D)$ to compute sample efficient probabilistic upper bounds on the policy loss of any evaluation policy. Other work considers robust policy optimization over a distribution of reward functions conditioned on demonstrations or a partially specified reward function, but these methods require an MDP solver in the inner loop, limiting their scalability (Hadfield-Menell et al., 2017; Brown et al., 2018; Huang et al., 2018). We extend and generalize the work of Brown & Niekum (2018) by demonstrating, for the first time, that high-confidence performance bounds can be efficiently obtained when performing imitation learning from high-dimensional visual demonstrations without requiring an MDP solver or model during reward inference.

2.3. Value Alignment and Active Preference Learning

Safe imitation learning is closely related to the problem of value alignment, which seeks to design methods that prevent AI systems from acting in ways that violate human values (Hadfield-Menell et al., 2016; Fisac et al., 2020). Research has shown that difficulties arise when an agent seeks to align its value with a human who is not perfectly rational (Milli et al., 2017) and there are fundamental impossibility results regarding value alignment unless the objective is represented as a set of partially ordered preferences (Eckersley, 2018).

Prior work has used active queries to perform Bayesian reward inference on low-dimensional, hand-crafted reward features (Sadigh et al., 2017; Brown et al., 2018; Bıyık et al., 2019). Christiano et al. (2017) and Ibarz et al. (2018) use deep networks to scale active preference learning to high-dimensional tasks, but require large numbers of active queries during policy optimization and do not perform

Bayesian reward inference. Our work complements and extends prior work by: (1) removing the requirement for active queries during reward inference or policy optimization, (2) showing that preferences over demonstrations enable efficient Bayesian reward inference in high-dimensional visual control tasks, and (3) providing an efficient method for computing high-confidence bounds on the performance of any evaluation policy in the imitation learning setting.

2.4. Safe Reinforcement Learning

Research on safe reinforcement learning (RL) usually focuses on safe exploration strategies or optimization objectives other than expected return (Garcia & Fernández, 2015). Recently, objectives based on measures of risk such as value at risk (VaR) and conditional VaR have been shown to provide tractable and useful risk-sensitive measures of performance for MDPs (Tamar et al., 2015; Chow et al., 2015). Other work focuses on finding robust solutions to MDPs (Ghavamzadeh et al., 2016; Petrik & Russell, 2019), using model-based RL to safely improve upon suboptimal demonstrations (Thananjeyan et al., 2019), and obtaining high-confidence off-policy bounds on the performance of an evaluation policy (Thomas et al., 2015; Hanna et al., 2019). Our work provides an efficient solution to the problem of high-confidence policy evaluation in the imitation learning setting, in which samples of rewards are *not observed* and the demonstrator’s policy is unknown.

2.5. Bayesian Neural Networks

Bayesian neural networks typically either perform Markov Chain Monte Carlo (MCMC) sampling (MacKay, 1992), variational inference (Sun et al., 2019; Khan et al., 2018), or use hybrid methods such as particle-based inference (Liu & Wang, 2016) to approximate the posterior distribution over neural network weights. Alternative approaches such as ensembles (Lakshminarayanan et al., 2017) or approximations such as Bayesian dropout (Gal & Ghahramani, 2016; Kendall & Gal, 2017) have also been used to obtain a distribution on the outputs of a neural network in order to provide uncertainty quantification (Maddox et al., 2019). We are not only interested in the uncertainty of the output of the reward function, but also in the uncertainty over the performance of a policy when evaluated under an uncertain reward function. Thus, we face the difficult problem of measuring the uncertainty in the evaluation of a policy, which depends on the stochasticity of the policy and the environment, as well as the uncertainty over the unobserved reward function.

3. Preliminaries

We model the environment as a Markov Decision Process (MDP) consisting of states $\mathcal{S}$, actions $\mathcal{A}$, transition dynamics $T : \mathcal{S} \times \mathcal{A} \times \mathcal{S} \rightarrow [0, 1]$, reward function $R : \mathcal{S} \rightarrow \mathbb{R}$,

initial state distribution S_0 , and discount factor γ . Our approach extends naturally to rewards defined as $R(s, a)$ or $R(s, a, s')$; however, state-based rewards have some advantages. Fu et al. (2017) prove that a state-only reward function is a necessary and sufficient condition for a reward function that is disentangled from dynamics. Learning a state-based reward also allows the learned reward to be used as a potential function for reward shaping (Ng et al., 1999), if a sparse ground-truth reward function is available.

A policy π is a mapping from states to a probability distribution over actions. We denote the value of a policy π under reward function R as $V_R^\pi = \mathbb{E}_\pi[\sum_{t=0}^\infty \gamma^t R(s_t) | s_0 \sim S_0]$ and denote the value of executing policy π starting at state $s \in \mathcal{S}$ as $V_R^\pi(s) = \mathbb{E}_\pi[\sum_{t=0}^\infty \gamma^t R(s_t) | s_0 = s]$. Given a reward function R , the Q-value of a state-action pair (s, a) is $Q_R^\pi(s, a) = \mathbb{E}_\pi[\sum_{t=0}^\infty \gamma^t R(s_t) | s_0 = s, a_0 = a]$. We also denote $V_R^* = \max_\pi V_R^\pi$ and $Q_R^*(s, a) = \max_\pi Q_R^\pi(s, a)$.

Bayesian inverse reinforcement learning (IRL) (Ramachandran & Amir, 2007) models the environment as an MDP in which the reward function is unavailable. Bayesian IRL seeks to infer the latent reward function of a Boltzmann-rational demonstrator that executes the following policy

$$\pi_R^\beta(a|s) = \frac{e^{\beta Q_R^*(s,a)}}{\sum_{b \in \mathcal{A}} e^{\beta Q_R^*(s,b)}}, \quad (1)$$

in which R is the true reward function of the demonstrator, and $\beta \in [0, \infty)$ represents the confidence that the demonstrator is acting optimally. Under the assumption of Boltzmann rationality, the likelihood of a set of demonstrated state-action pairs, $D = \{(s, a) : (s, a) \sim \pi_D\}$, given a specific reward function hypothesis R , can be written as

$$P(D|R) = \prod_{(s,a) \in D} \pi_R^\beta(a|s) = \prod_{(s,a) \in D} \frac{e^{\beta Q_R^*(s,a)}}{\sum_{b \in \mathcal{A}} e^{\beta Q_R^*(s,b)}}. \quad (2)$$

Bayesian IRL generates samples from the posterior distribution $P(R|D) \sim P(D|R)P(R)$ via Markov Chain Monte Carlo (MCMC) sampling, but this requires solving for $Q_{R'}^*$ to compute the likelihood of each new proposal R' . Thus, Bayesian IRL methods are only used for low-dimensional problems with reward functions that are often linear combinations of a small number of hand-crafted features (Bobu et al., 2018; Bıyık et al., 2019). One of our contributions is an efficient Bayesian reward inference algorithm that leverages preferences over demonstrations in order to significantly improve the efficiency of Bayesian reward inference.

4. High Confidence Policy Evaluation for Imitation Learning

Before detailing our approach, we first formalize the problem of high-confidence policy evaluation for imitation learn-

ing. We assume access to an MDP $\mathcal{M}$, an evaluation policy π_{eval} , a set of demonstrations, $D = \{\tau_1, \dots, \tau_m\}$, in which τ_i is either a complete or partial trajectory comprised of states or state-action pairs, a confidence level δ , and performance statistic $g : \Pi \times \mathcal{R} \rightarrow \mathbb{R}$, in which $\mathcal{R}$ denotes the space of reward functions and Π is the space of all policies.

The *High-Confidence Policy Evaluation problem for Imitation Learning* (HCPE-IL) is to find a high-confidence lower bound $\hat{g} : \Pi \times \mathcal{D} \rightarrow \mathbb{R}$ such that

$$\Pr(g(\pi_{\text{eval}}, R^*) \geq \hat{g}(\pi_{\text{eval}}, D)) \geq 1 - \delta, \quad (3)$$

in which R^* denotes the demonstrator’s true reward function and $\mathcal{D}$ denotes the space of all possible demonstration sets. HCPE-IL takes as input an evaluation policy π_{eval} , a set of demonstrations D , and a performance statistic, g , which evaluates a policy under a reward function. The goal of HCPE-IL is to return a high-confidence lower bound $\hat{g}$ on the performance statistic $g(\pi_{\text{eval}}, R^*)$.

5. Deep Bayesian Reward Extrapolation

We now describe our main contribution: a method for scaling Bayesian reward inference to high-dimensional visual control tasks as a way to efficiently solve the HCPE-IL problem for complex imitation learning tasks. Our first insight is that the main bottleneck for standard Bayesian IRL (Ramachandran & Amir, 2007) is computing the likelihood function in Equation (2) which requires optimal Q-values. Thus, to make Bayesian reward inference scale to high-dimensional visual domains, it is necessary to either efficiently approximate optimal Q-values or to formulate a new likelihood. Value-based reinforcement learning focuses on efficiently learning optimal Q-values; however, for complex visual control tasks, RL algorithms can take several hours or even days to train (Mnih et al., 2015; Hessel et al., 2018). This makes MCMC, which requires evaluating large numbers of likelihood ratios, infeasible given the current state-of-the-art in value-based RL. Methods such as transfer learning have great potential to reduce the time needed to calculate Q_R^* for a new proposed reward function R ; however, transfer learning is not guaranteed to speed up reinforcement learning (Taylor & Stone, 2009). Thus, we choose to focus on reformulating the likelihood function as a way to speed up Bayesian reward inference.

An ideal likelihood function requires little computation and minimal interaction with the environment. To accomplish this, we leverage recent work on learning control policies from preferences (Christiano et al., 2017; Palan et al., 2019; Bıyık et al., 2019). Given ranked demonstrations, Brown et al. (2019b) propose Trajectory-ranked Reward Extrapolation (T-REX): an efficient reward inference algorithm that transforms reward function learning into classification problem via a pairwise ranking loss. T-REX removes the

need to repeatedly sample from or partially solve an MDP in the inner loop, allowing it to scale to visual imitation learning domains such as Atari and to extrapolate beyond the performance of the best demonstration. However, T-REX only solves for a point estimate of the reward function. We now discuss how a similar approach based on a pairwise preference likelihood allows for efficient sampling from the posterior distribution over reward functions.

We assume access to a sequence of m trajectories, $D = \{\tau_1, \dots, \tau_m\}$, along with a set of pairwise preferences over trajectories $\mathcal{P} = \{(i, j) : \tau_i \prec \tau_j\}$. Note that we do not require a total-ordering over trajectories. These preferences may come from a human demonstrator or could be automatically generated by watching a learner improve at a task (Jacq et al., 2019; Brown et al., 2019b) or via noise injection (Brown et al., 2019a). Given trajectory preferences, we can formulate a pair-wise ranking likelihood to compute the likelihood of a set of preferences over demonstrations $\mathcal{P}$, given a parameterized reward function hypothesis R_θ . We use the standard Bradley-Terry model (Bradley & Terry, 1952) to obtain the following pairwise ranking likelihood function, commonly used in learning to rank applications such collaborative filtering (Volkovs & Zemel, 2014):

$$P(D, \mathcal{P} \mid R_\theta) = \prod_{(i,j) \in \mathcal{P}} \frac{e^{\beta R_\theta(\tau_j)}}{e^{\beta R_\theta(\tau_i)} + e^{\beta R_\theta(\tau_j)}}, \quad (4)$$

in which $R_\theta(\tau) = \sum_{s \in \tau} R_\theta(s)$ is the predicted return of trajectory τ under the reward function R_θ , and β is the inverse temperature parameter that models the confidence in the preference labels. We can then perform Bayesian inference via MCMC to obtain samples from $P(R_\theta \mid D, \mathcal{P}) \propto P(D, \mathcal{P} \mid R_\theta)P(R_\theta)$. We call this approach Bayesian Reward Extrapolation or Bayesian REX.

Note that using the likelihood function defined in Equation (4) does not require solving an MDP. In fact, it does not require any rollouts or access to the MDP. All that is required is that we first calculate the return of each trajectory under R_θ and compare the relative predicted returns to the preference labels to determine the likelihood of the demonstrations under the reward hypothesis R_θ . Thus, given preferences over demonstrations, Bayesian REX is significantly more efficient than standard Bayesian IRL. In the following section, we discuss further optimizations that improve the efficiency of Bayesian REX and make it more amenable to our end goal of high-confidence policy evaluation bounds.

5.1. Optimizations

In order to learn rich, complex reward functions, it is desirable to use a deep network to represent the reward function R_θ . While MCMC remains the gold-standard for Bayesian Neural Networks, it is often challenging to scale to deep networks. To make Bayesian REX more efficient and practical,

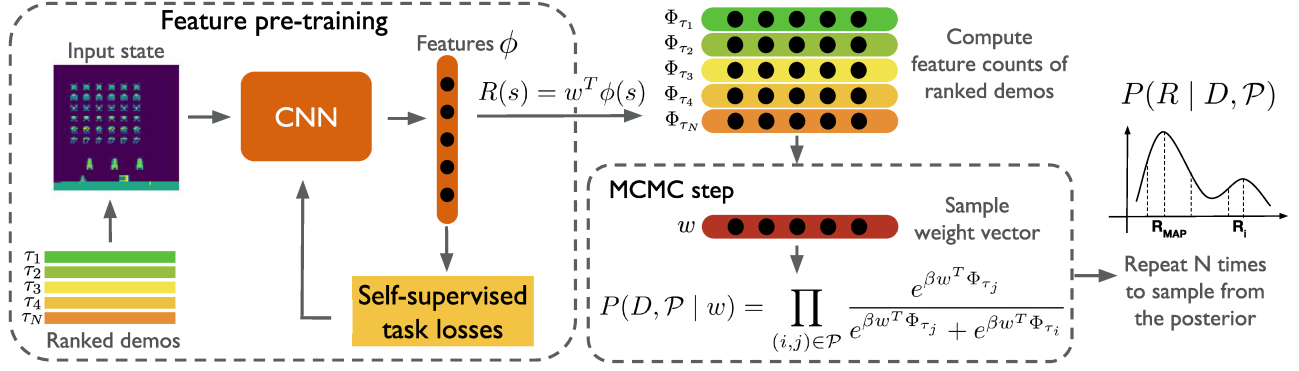


Figure 1. Bayesian Reward Extrapolation uses ranked demonstrations to pre-train a low-dimensional state feature embedding $\phi(s)$ via self-supervised losses. After pre-training, the latent embedding function $\phi(s)$ is frozen and the reward function is represented as a linear combination of the learned features: $R(s) = w^T \phi(s)$. MCMC proposal evaluations use a pairwise ranking likelihood that gives the likelihood of the preferences $\mathcal{P}$ over demonstrations D , given a proposal w . By pre-computing the embeddings of the ranked demonstrations, Φ_{τ_i} , MCMC sampling is highly efficient—it does not require access to an MDP solver or data collection during inference.

we propose to limit the proposal to only change the last layer of weights in R_θ when generating MCMC proposals—we will discuss pre-training the bottom layers of R_θ in the next section. After pre-training, we freeze all but the last layer of weights and use the activations of the penultimate layer as the latent reward features $\phi(s) \in \mathbb{R}^k$. This allows the reward at a state to be represented as a linear combination of k features: $R_\theta(s) = w^T \phi(s)$. Similar to work by Pradier et al. (2018), operating in a lower-dimensional latent space makes full Bayesian inference tractable.

A second advantage of using a learned linear reward function is that it allows us to efficiently compute likelihood ratios when performing MCMC. Consider the likelihood function in Equation (4). If we do not represent R_θ as a linear combination of pretrained features, and instead let any parameter in R_θ change during each proposal, then for m demonstrations of length T , computing $P(D, \mathcal{P} | R_\theta)$ for a new proposal R_θ requires $O(mT)$ forward passes through the entire network to compute $R_\theta(\tau_i)$. Thus, the complexity of generating N samples from the posterior results is $O(mTN|R_\theta|)$, where $|R_\theta|$ is the number of computations required for a full forward pass through the entire network R_θ . Given that we would like to use a deep network to parameterize R_θ and generate thousands of samples from the posterior distribution over R_θ , this many computations will significantly slow down MCMC proposal evaluation.

If we represent R_θ as a linear combination of pre-trained features, we can reduce this computational cost because

$$R_\theta(\tau) = \sum_{s \in \tau} w^T \phi(s) = w^T \sum_{s \in \tau} \phi(s) = w^T \Phi_\tau. \quad (5)$$

Thus, we can precompute and cache $\Phi_{\tau_i} = \sum_{s \in \tau_i} \phi(s)$ for

$i = 1, \dots, m$ and rewrite the likelihood as

$$P(D, \mathcal{P} | R_\theta) = \prod_{(i,j) \in \mathcal{P}} \frac{e^{\beta w^T \Phi_{\tau_j}}}{e^{\beta w^T \Phi_{\tau_j}} + e^{\beta w^T \Phi_{\tau_i}}}. \quad (6)$$

Note that demonstrations only need to be passed through the reward network once to compute Φ_{τ_i} since the pre-trained embedding remains constant during MCMC proposal generation. This results in an initial $O(mT)$ passes through all but the last layer of R_θ to obtain Φ_{τ_i} , for $i = 1, \dots, m$, and then only $O(mk)$ multiplications per proposal evaluation thereafter—each proposal requires that we compute $w^T \Phi_{\tau_i}$ for $i = 1, \dots, m$ and $\Phi_{\tau_i} \in \mathbb{R}^k$. Thus, when using feature pre-training, the total complexity is only $O(mT|R_\theta| + mkN)$ to generate N samples via MCMC. This reduction in the complexity of MCMC from $O(mTN|R_\theta|)$ to $O(mT|R_\theta| + mkN)$ results in significant and practical computational savings because (1) we want to make N and R_θ large and (2) the number of demonstrations, m , and the size of the latent embedding, k , are typically several orders of magnitude smaller than N and $|R_\theta|$.

A third, and critical advantage of using a learned linear reward function is that it makes solving the HCPE-IL problem discussed in Section 4 tractable. Performing a single policy evaluation is a non-trivial task (Sutton et al., 2000) and even in tabular settings has complexity $O(|S|^3)$ in which $|S|$ is the size of the state-space (Littman et al., 1995). Because we are in an imitation learning setting, we would like to be able to efficiently evaluate any given policy across the posterior distribution over reward functions found via Bayesian REX. Given a posterior distribution over N reward function hypotheses we would need to solve N policy evaluations. However, note that given $R(s) = w^T \phi(s)$, the

value function of a policy can be written as

$$V_R^\pi = \mathbb{E}_\pi \left[\sum_{t=0}^T R(s_t) \right] = w^T \mathbb{E}_\pi \left[\sum_{t=0}^T \phi(s_t) \right] = w^T \Phi_\pi, \quad (7)$$

in which we assume a finite horizon MDP with horizon T and in which Φ_π are the expected feature counts (Abbeel & Ng, 2004; Barreto et al., 2017) of π . Thus, given any evaluation policy π_{eval} , we only need to solve one policy evaluation problem to compute Φ_{eval} . We can then compute the expected value of π_{eval} over the entire posterior distribution of reward functions via a single matrix vector multiplication $W\Phi_{\text{eval}}$, where W is an N -by- k matrix with each row corresponding to a single reward function weight hypothesis w^T . This significantly reduces the complexity of policy evaluation over the reward function posterior distribution from $O(N|S|^3)$ to $O(|S|^3 + Nk)$.

When we refer to Bayesian REX we will refer to the optimized version described in this section (see the Appendix for full implementation details and pseudo-code)¹. Running MCMC with 66 preference labels to generate 100,000 reward hypothesis for Atari imitation learning tasks takes approximately 5 minutes on a Dell Inspiron 5577 personal laptop with an Intel i7-7700 processor without using the GPU. In comparison, using standard Bayesian IRL to generate *one sample* from the posterior takes 10+ hours of training for a parallelized PPO reinforcement learning agent (Dhariwal et al., 2017) on an NVIDIA TITAN V GPU.

5.2. Pre-training the Reward Function Network

The previous section presupposed access to a pretrained latent embedding function $\phi : S \rightarrow \mathbb{R}^k$. We now discuss our pre-training process. Because we are interested in imitation learning problems, we need to be able to train $\phi(s)$ from the demonstrations without access to the ground-truth reward function. One potential method is to train R_θ using the pairwise ranking likelihood function in Equation (4) and then freeze all but the last layer of weights; however, the learned embedding may overfit to the limited number of preferences over demonstrations and fail to capture features relevant to the ground-truth reward function. Thus, we supplement the pairwise ranking objective with auxiliary objectives that can be optimized in a self-supervised fashion using data from the demonstrations.

We use the following self-supervised tasks to pre-train R_θ : (1) Learn an inverse dynamics model that uses embeddings $\phi(s_t)$ and $\phi(s_{t+1})$ to predict the corresponding action a_t (Torabi et al., 2018; Hanna & Stone, 2017), (2) Learn a forward dynamics model that predicts s_{t+1} from $\phi(s_t)$ and a_t (Oh et al., 2015; Thananjeyan et al., 2019), (3) Learn an

Table 1. Self-supervised learning objectives used to pre-train $\phi(s)$.

Inverse Dynamics	$f_{\text{ID}}(\phi(s_t), \phi(s_{t+1})) \rightarrow a_t$
Forward Dynamics	$f_{\text{FD}}(\phi(s_t), a_t) \rightarrow s_{t+1}$
Temporal Distance	$f_{\text{TD}}(\phi(s_t), \phi(s_{t+x})) \rightarrow x$
Variational Autoencoder	$f_A(\phi(s_t)) \rightarrow s_t$

embedding $\phi(s)$ that predicts the temporal distance between two randomly chosen states from the same demonstration (Aytar et al., 2018), and (4) Train a variational pixel-to-pixel autoencoder in which $\phi(s)$ is the learned latent encoding (Makhzani & Frey, 2017; Doersch, 2016). Table 1 summarizes the self-supervised tasks used to train $\phi(s)$.

There are many possibilities for pre-training $\phi(s)$. We used the objectives described above to encourage the embedding to encode different features. For example, an accurate inverse dynamics model can be learned by only attending to the movement of the agent. Learning forward dynamics supplements this by requiring $\phi(s)$ to encode information about short-term changes to the environment. Learning to predict the temporal distance between states in a trajectory forces $\phi(s)$ to encode long-term progress. Finally, the autoencoder loss acts as a regularizer to the other losses as it seeks to embed all aspects of the state (see the Appendix for details). The Bayesian REX pipeline for sampling from the reward function posterior is shown in Figure 1.

5.3. HCPE-IL via Bayesian REX

We now discuss how to use Bayesian REX to find an efficient solution to the high-confidence policy evaluation for imitation learning (HCPE-IL) problem (see Section 4). Given samples from the distribution $P(w \mid D, \mathcal{P})$, where $R(s) = w^T \phi(s)$, we compute the posterior distribution over any performance statistic $g(\pi_{\text{eval}}, R^*)$ as follows. For each sampled weight vector w produced by Bayesian REX, we compute $g(\pi_{\text{eval}}, w)$. This results in a sample from the posterior distribution $P(g(\pi_{\text{eval}}, R) \mid D, \mathcal{P})$, i.e., the posterior distribution over performance statistic g . We then compute a $(1 - \delta)$ confidence lower bound, $\hat{g}(\pi_{\text{eval}}, D)$, by finding the δ -quantile of $g(\pi_{\text{eval}}, w)$ for $w \sim P(w \mid D, \mathcal{P})$.

While there are many potential performance statistics g , we chose to focus on bounding the expected value of the evaluation policy, i.e., $g(\pi_{\text{eval}}, R^*) = V_{R^*}^{\pi_{\text{eval}}} = w^{*T} \Phi_{\pi_{\text{eval}}}$. To compute a $1 - \delta$ confidence bound on $V_{R^*}^{\pi_{\text{eval}}}$, we take advantage of the learned linear reward representation to efficiently calculate the posterior distribution over policy returns given preferences and demonstrations. This distribution over returns is calculated via a matrix vector product, $W\Phi_{\pi_{\text{eval}}}$, in which each row of W is a sample, w , from the MCMC chain and π_{eval} is the evaluation policy. We then sort the resulting vector and select the δ -quantile lowest value. This results in

¹Project page, code, and demonstration data are available at <https://sites.google.com/view/bayesianrex/>

a $1 - \delta$ confidence lower bound on $V_{R^*}^{\pi_{\text{eval}}}$ and corresponds to the δ -Value at Risk (VaR) over $V_R^{\pi_{\text{eval}}} \sim P(R \mid D, \mathcal{P})$ (Jorion, 1997; Brown & Niekum, 2018).

6. Experimental Results

6.1. Bayesian IRL vs. Bayesian REX

As noted previously, Bayesian IRL does not scale to high-dimensional tasks due to the requirement of repeatedly solving for an MDP in the inner loop. However, for low-dimensional problems it is still interesting to compare Bayesian IRL with Bayesian REX. We performed a large number of experiments on a variety of randomly generated gridworlds with low-dimensional reward features. We summarize our results here for three different ablations and give full results and implementation details in the appendix.

Ranked Suboptimal vs. Optimal Demos: Given a sufficient number of suboptimal ranked demonstrations (> 5), Bayesian REX performs on par and occasionally better than Bayesian IRL when given the same number of optimal demonstrations.

Only Ranked Suboptimal Demos B-REX always significantly outperforms Bayesian IRL when both algorithms receive suboptimal ranked demonstrations. For fairer comparison, we used a Bayesian IRL algorithm designed to learn from both good and bad demonstrations (Cui & Niekum, 2018). We labeled the top $X\%$ ranked demonstrations as good and bottom $X\%$ ranked as bad. This improved results for Bayesian IRL, but Bayesian REX still performed significantly better across all X .

Only Optimal Demos: Given a sufficient number of optimal demonstrations (> 5), Bayesian IRL significantly outperforms Bayesian REX. To use Bayesian REX with only optimal demonstrations, we followed prior work (Brown et al., 2019a) and auto-generated pairwise preferences using uniform random rollouts that were labeled as less preferred than the demonstrations. In general, this performed much worse than Bayesian IRL, but for small numbers of demonstrations (≤ 5) Bayesian REX leverages self-supervised rankings to perform nearly as well as full Bayesian IRL.

These results demonstrate that if a very small number of unlabeled near-optimal demonstrations are available, then classical Bayesian IRL is the natural choice for performing reward inference. However, if any of these assumptions are not true, then Bayesian REX is a competitive and often superior alternative for performing Bayesian reward inference even in low-dimensional problems where an MDP solver is tractable. If a highly efficient MDP solver is not available, then Bayesian IRL is infeasible and Bayesian REX is the natural choice for Bayesian reward inference.

6.2. Visual Imitation Learning via Bayesian REX

We next tested the imitation learning performance of Bayesian REX for high-dimensional problems where classical Bayesian reward inference is infeasible. We pre-trained a 64 dimensional latent state embedding $\phi(s)$ using the self-supervised losses shown in Table 1 and the T-REX pairwise preference loss. We found via ablation studies that combining the T-REX loss with the self-supervised losses resulted in better performance than training only with the T-REX loss or only with the self-supervised losses (see Appendix for details). We then used Bayesian REX to generate 200,000 samples from the posterior $P(R \mid D, \mathcal{P})$. To optimize a control policy, we used Proximal Policy Optimization (PPO) (Schulman et al., 2017) with the MAP and mean reward functions from the posterior (see Appendix for details).

To test whether Bayesian REX scales to complex imitation learning tasks we selected five Atari games from the Arcade Learning Environment (Bellemare et al., 2013). We do not give the RL agent access to the ground-truth reward signal and mask the game scores and number of lives in the demonstrations. Table 2 shows the imitation learning performance of Bayesian REX. We also compare against the results reported by (Brown et al., 2019b) for T-REX, and GAIL (Ho & Ermon, 2016) and use the same 12 suboptimal demonstrations used by Brown et al. (2019b) to train Bayesian REX (see Appendix for details).

Table 2 shows that Bayesian REX is able to utilize preferences over demonstrations to infer an accurate reward function that enables better-than-demonstrator performance. The average ground-truth return for Bayesian REX surpasses the performance of the best demonstration across all 5 games. In comparison, GAIL seeks to match the demonstrator’s state-action distributions which makes imitation learning difficult when demonstrations are suboptimal and noisy. In addition to providing uncertainty information, Bayesian REX remains competitive with T-REX (which only finds a maximum likelihood estimate of the reward function) and achieves better performance on 3 out of 5 games.

6.3. High-Confidence Policy Performance Bounds

Next, we ran an experiment to validate whether the posterior distribution generated by Bayesian REX can be used to solve the HCPE-IL problem described in Section 4. We evaluated four different evaluation policies, $A \prec B \prec C \prec D$, created by partially training a PPO agent on the ground-truth reward function and checkpointing the policy at various stages of learning. We ran Bayesian REX to generate 200,000 samples from $P(R \mid D, \mathcal{P})$. To address some of the ill-posedness of IRL, we normalize the weights w such that $\|w\|_2 = 1$. Given a fixed scale for the reward weights, we can compare the relative performance of the different evaluation policies when evaluated over the posterior.

Table 2. Ground-truth average scores when optimizing the mean and MAP rewards found using Bayesian REX. We also compare against the performance of T-REX (Brown et al., 2019b) and GAIL (Ho & Ermon, 2016). Bayesian REX and T-REX are each given 12 demonstrations with ground-truth pairwise preferences. GAIL cannot learn from preferences so it is given 10 demonstrations comparable to the best demonstration given to the other algorithms. The average performance for each IRL algorithm is the average over 30 rollouts.

Game	Ranked Demonstrations		Bayesian REX Mean	Bayesian REX MAP	T-REX	GAIL
	Best	Avg	Avg (Std)	Avg (Std)	Avg	Avg
Beam Rider	1332	686.0	5,504.7 (2121.2)	5,870.3 (1905.1)	3,335.7	355.5
Breakout	32	14.5	390.7 (48.8)	393.1 (63.7)	221.3	0.28
Enduro	84	39.8	487.7 (89.4)	135.0 (24.8)	586.8	0.28
Seaquest	600	373.3	734.7 (41.9)	606.0 (37.6)	747.3	0.0
Space Invaders	600	332.9	1,118.8 (483.1)	961.3 (392.3)	1,032.5	370.2

Table 3. Beam Rider policy evaluation bounds compared with ground-truth game scores. Policies A-D correspond to evaluation policies of varying quality obtained by checkpointing an RL agent during training. The No-Op policy seeks to hack the learned reward by always playing the no-op action, resulting in very long trajectories with high mean predicted performance but a very negative 95%-confidence (0.05-VaR) lower bound on expected return.

Policy	Predicted		Ground Truth Avg.	
	Mean	0.05-VaR	Score	Length
A	17.1	7.9	480.6	1372.6
B	22.7	11.9	703.4	1,412.8
C	45.5	24.9	1828.5	2,389.9
D	57.6	31.5	2586.7	2,965.0
No-Op	102.5	-1557.1	0.0	99,994.0

The results for Beam Rider are shown in Table 3. We show results for partially trained RL policies A–D. We found that the ground-truth returns for the checkpoints were highly correlated with the mean and 0.05-VaR (5th percentile policy return) returns under the posterior. However, we also noticed that the trajectory length was also highly correlated with the ground-truth reward. If the reward function learned via IRL gives a small positive reward at every time step, then long policies that do the wrong thing may look good under the posterior. To test this hypothesis we used a No-Op policy that seeks to exploit the learned reward function by not taking any actions. This allows the agent to live until the Atari emulator times out after 99,994 steps.

Table 3 shows that while the No-Op policy has a high expected return over the chain, looking at the 0.05-VaR shows that the No-Op policy has high risk under the distribution, much lower than evaluation policy A. Our results demonstrate that reasoning about probabilistic worst-case performance may be one potential way to detect policies that exhibit so-called reward hacking (Amodei et al., 2016) or that have overfit to certain features in the demonstrations that are correlated with the intent of the demonstrations,

Table 4. Breakout policy evaluation bounds compared with ground-truth game scores. Top Half: No-Op never releases the ball, resulting in high mean predicted performance but a low 95%-confidence bound (0.05-VaR). The MAP policy has even higher risk but also high expected return. Bottom Half: After rerunning MCMC with a ranked trajectory from both the MAP and No-Op policies, the posterior distribution matches the true preferences.

Risk profiles given initial preferences				
Policy	Predicted		Ground Truth Avg.	
	Mean	0.05-VaR	Score	Length
A	1.5	0.5	1.9	202.7
B	6.3	3.7	15.8	608.4
C	10.6	5.8	27.7	849.3
D	13.9	6.2	41.2	1020.8
MAP	98.2	-370.2	401.0	8780.0
No-Op	41.2	1.0	0.0	7000.0

Risk profiles after rankings w.r.t. MAP and No-Op				
A	0.7	0.3	1.9	202.7
B	8.7	5.5	15.8	608.4
C	18.3	12.1	27.7	849.3
D	26.3	17.1	41.2	1020.8
MAP	606.8	289.1	401.0	8780.0
No-Op	-5.0	-13.5	0.0	7000.0

but do not lead to desired behavior, a common problem in imitation learning (Ibarz et al., 2018; de Haan et al., 2019).

Table 4 contains policy evaluation results for the game Breakout. The top half of the table shows the mean return and 95%-confidence lower bound on the expected return under the reward function posterior for four evaluation policies as well as the MAP policy found via Bayesian IRL and a No-Op policy that never chooses to release the ball. Both the MAP and No-Op policies have high expected returns under the reward function posterior, but also have high risk (low 0.05-VaR). The MAP policy has much higher risk than the

Table 5. Beam Rider human demonstrations.

Policy	Predicted		Ground Truth	
	Mean	0.05-VaR	Avg.	Length
good	12.4	5.8	1092	1000.0
bad	10.7	4.5	396	1000.0
pessimist	6.6	0.8	0	1000.0
adversarial	8.4	2.4	176	1000.0

No-Op policy, despite good true performance. One likely reason is that, as shown in Table 2, the best demonstrations given to Bayesian REX only achieved a game score of 32. Thus, the MAP policy represents an out of distribution sample and thus has potentially high risk, since Bayesian REX was not trained on policies that hit any of the top layers of bricks. The ranked demonstrations do not give enough evidence to eliminate the possibility that only lower layers of bricks should be hit.

To test whether active learning can help, we incorporated two active queries: a single rollout from the MAP policy and a single rollout from the No-Op policy and ranked them as better and worse, respectively, than the original set of 12 suboptimal demonstrations. As the bottom of Table 4 shows, adding two more ranked demonstrations and re-running Bayesian inference, results in a significant change in the risk profiles of the MAP and No-Op policy—the No-Op policy is now correctly predicted to have high risk and low expected returns and the MAP policy now has a much higher 95%-confidence lower bound on performance.

6.4. Human Demonstrations

To investigate whether Bayesian REX is able to correctly rank human demonstrations, we used Bayesian REX to calculate high-confidence performance bounds for a variety of human demonstrations (see the Appendix for full details and additional results).

We generated four human demonstrations for Beam Rider: (1) *good*, a good demonstration that plays the game well, (2) *bad*, a bad demonstration that seeks to play the game but does a poor job, (3) *pessimist*, a demonstration that does not shoot enemies and seeks enemy bullets, and (4) *adversarial* a demonstration that pretends to play the game by moving and shooting but tries to avoid actually shooting enemies. The resulting high-confidence policy evaluations are shown in Table 5. The high-confidence bounds and average performance over the posterior correctly rank the behaviors. This provides evidence that the learned linear reward correctly rewards actually destroying aliens and avoiding getting shot, rather than just flying around and shooting.

Next we demonstrated four different behaviors when play-

Table 6. Enduro evaluation of a variety of human demonstrations.

Policy	Predicted		Ground Truth	
	Mean	0.05-VaR	Avg.	Length
good	246.7	-113.2	177	3325.0
periodic	230.0	-130.4	44	3325.0
neutral	190.8	-160.6	0	3325.0
ram	148.4	-214.3	0	3325.0

ing Enduro: (1) *good* a demonstration that seeks to play the game well, (2) *periodic* a demonstration that alternates between speeding up and passing cars and then slowing down and being passed, (3) *neutral* a demonstration that stays right next to the last car in the race and doesn’t try to pass or get passed, and (4) *ram* a demonstration that tries to ram into as many cars while going fast. Table 6 shows that Bayesian REX is able to accurately predict the performance and risk of each of these demonstrations and gives the highest (lowest 0.05-VaR) risk to the *ram* demonstration and the least risk to the *good* demonstration.

7. Conclusion

Bayesian reasoning is a powerful tool when dealing with uncertainty and risk; however, existing Bayesian reward learning algorithms often require solving an MDP in the inner loop, rendering them intractable for complex problems in which solving an MDP may take several hours or even days. In this paper we propose a novel deep learning algorithm, Bayesian Reward Extrapolation (Bayesian REX), that leverages preference labels over demonstrations to make Bayesian reward inference tractable for high-dimensional visual imitation learning tasks. Bayesian REX can sample tens of thousands of reward functions from the posterior in a matter of minutes using a consumer laptop. We tested our approach on five Atari imitation learning tasks and showed that Bayesian REX achieves state-of-the-art performance in 3 out of 5 games. Furthermore, Bayesian REX enables efficient high-confidence performance bounds for arbitrary evaluation policies. We demonstrated that these high-confidence bounds allow an agent to accurately rank different evaluation policies and provide a potential way to detect reward hacking and value misalignment.

We note that our proposed safety bounds are only safe with respect to the assumptions that we make: good feature pre-training, rapid MCMC mixing, and accurate preferences over demonstrations. Future work includes using exploratory trajectories for better pre-training of the latent feature embeddings, developing methods to determine when a relevant feature is missing from the learned latent space, and using high-confidence performance bounds to perform safe policy optimization in the imitation learning setting.

Acknowledgements

This work has taken place in the Personal Autonomous Robotics Lab (PeARL) at The University of Texas at Austin. PeARL research is supported in part by the NSF (IIS-1724157, IIS-1638107, IIS-1617639, IIS-1749204) and ONR (N00014-18-2243, N00014-17-1-2143). This research was also sponsored by the Army Research Office and was accomplished under Cooperative Agreement Number W911NF-19-2-0333. The views and conclusions contained in this document are those of the authors and should not be interpreted as representing the official policies, either expressed or implied, of the Army Research Office or the U.S. Government. The U.S. Government is authorized to reproduce and distribute reprints for Government purposes notwithstanding any copyright notation herein.

References

- Abbeel, P. and Ng, A. Y. Apprenticeship learning via inverse reinforcement learning. In *Proceedings of the 21st International Conference on Machine Learning*, 2004.
- Akrou, R., Schoenauer, M., and Sebag, M. Preference-based policy learning. In *Joint European Conference on Machine Learning and Knowledge Discovery in Databases*, pp. 12–27. Springer, 2011.
- Amodei, D., Olah, C., Steinhardt, J., Christiano, P., Schulman, J., and Mané, D. Concrete problems in ai safety. *arXiv preprint arXiv:1606.06565*, 2016.
- Argall, B. D., Chernova, S., Veloso, M., and Browning, B. A survey of robot learning from demonstration. *Robotics and autonomous systems*, 57(5):469–483, 2009.
- Aytar, Y., Pfaff, T., Budden, D., Paine, T. L., Wang, Z., and de Freitas, N. Playing hard exploration games by watching youtube. *arXiv preprint arXiv:1805.11592*, 2018.
- Barreto, A., Dabney, W., Munos, R., Hunt, J. J., Schaul, T., van Hasselt, H. P., and Silver, D. Successor features for transfer in reinforcement learning. In *Advances in neural information processing systems*, pp. 4055–4065, 2017.
- Bellemare, M. G., Naddaf, Y., Veness, J., and Bowling, M. The arcade learning environment: An evaluation platform for general agents. *Journal of Artificial Intelligence Research*, 47:253–279, 2013.
- Bıyık, E., Palan, M., Landolfi, N. C., Losey, D. P., and Sadigh, D. Asking easy questions: A user-friendly approach to active reward learning. In *Conference on Robot Learning (CoRL)*, 2019.
- Bobu, A., Bajcsy, A., Fisac, J. F., and Dragan, A. D. Learning under misspecified objective spaces. *arXiv preprint arXiv:1810.05157*, 2018.
- Bradley, R. A. and Terry, M. E. Rank analysis of incomplete block designs: I. the method of paired comparisons. *Biometrika*, 39(3/4):324–345, 1952.
- Brown, D. S. and Niekum, S. Efficient Probabilistic Performance Bounds for Inverse Reinforcement Learning. In *AAAI Conference on Artificial Intelligence*, 2018.
- Brown, D. S., Cui, Y., and Niekum, S. Risk-aware active inverse reinforcement learning. In *Conference on Robot Learning (CoRL)*, 2018.
- Brown, D. S., Goo, W., and Niekum, S. Better-than-demonstrator imitation learning via automatically-ranked demonstrations. In *Conference on Robot Learning (CoRL)*, 2019a.
- Brown, D. S., Goo, W., Prabhat, N., and Niekum, S. Extrapolating beyond suboptimal demonstrations via inverse reinforcement learning from observations. In *Proceedings of the 36th International Conference on Machine Learning, ICML 2019*, 2019b.
- Chow, Y., Tamar, A., Mannor, S., and Pavone, M. Risk-sensitive and robust decision-making: a cvar optimization approach. In *Advances in Neural Information Processing Systems*, 2015.
- Christiano, P. F., Leike, J., Brown, T., Martic, M., Legg, S., and Amodei, D. Deep reinforcement learning from human preferences. In *Advances in Neural Information Processing Systems*, pp. 4299–4307, 2017.
- Cui, Y. and Niekum, S. Active reward learning from critiques. In *IEEE International Conference on Robotics and Automation (ICRA)*, 2018.
- de Haan, P., Jayaraman, D., and Levine, S. Causal confusion in imitation learning. In *Advances in Neural Information Processing Systems*, pp. 11693–11704, 2019.
- Dhariwal, P., Hesse, C., Klimov, O., Nichol, A., Plappert, M., Radford, A., Schulman, J., Sidor, S., Wu, Y., and Zhokhov, P. Openai baselines. <https://github.com/openai/baselines>, 2017.
- Doersch, C. Tutorial on variational autoencoders. *arXiv preprint arXiv:1606.05908*, 2016.
- Eckersley, P. Impossibility and uncertainty theorems in ai value alignment (or why your agi should not have a utility function). *arXiv preprint arXiv:1901.00064*, 2018.
- Finn, C., Christiano, P., Abbeel, P., and Levine, S. A connection between generative adversarial networks, inverse reinforcement learning, and energy-based models. *arXiv preprint arXiv:1611.03852*, 2016a.

- Finn, C., Levine, S., and Abbeel, P. Guided cost learning: Deep inverse optimal control via policy optimization. In *International Conference on Machine Learning*, 2016b.
- Fisac, J. F., Gates, M. A., Hamrick, J. B., Liu, C., Hadfield-Menell, D., Palaniappan, M., Malik, D., Sastry, S. S., Griffiths, T. L., and Dragan, A. D. Pragmatic-pedagogic value alignment. In *Robotics Research*, pp. 49–57. Springer, 2020.
- Fu, J., Luo, K., and Levine, S. Learning robust rewards with adversarial inverse reinforcement learning. *arXiv preprint arXiv:1710.11248*, 2017.
- Gal, Y. and Ghahramani, Z. Dropout as a bayesian approximation: Representing model uncertainty in deep learning. In *international conference on machine learning*, pp. 1050–1059, 2016.
- Garcia, J. and Fernández, F. A comprehensive survey on safe reinforcement learning. *Journal of Machine Learning Research*, 16(1):1437–1480, 2015.
- Ghasemipour, S. K. S., Zemel, R., and Gu, S. A divergence minimization perspective on imitation learning methods. *arXiv preprint arXiv:1911.02256*, 2019.
- Ghavamzadeh, M., Petrik, M., and Chow, Y. Safe policy improvement by minimizing robust baseline regret. In *Advances in Neural Information Processing Systems*, pp. 2298–2306, 2016.
- Hadfield-Menell, D., Russell, S. J., Abbeel, P., and Dragan, A. Cooperative inverse reinforcement learning. In *Advances in neural information processing systems*, pp. 3909–3917, 2016.
- Hadfield-Menell, D., Milli, S., Abbeel, P., Russell, S. J., and Dragan, A. Inverse reward design. In *Advances in neural information processing systems*, pp. 6765–6774, 2017.
- Hanna, J., Niekum, S., and Stone, P. Importance sampling policy evaluation with an estimated behavior policy. In *Proceedings of the 36th International Conference on Machine Learning (ICML)*, June 2019.
- Hanna, J. P. and Stone, P. Grounded action transformation for robot learning in simulation. In *Thirty-First AAAI Conference on Artificial Intelligence*, 2017.
- Hessel, M., Modayil, J., Van Hasselt, H., Schaul, T., Ostrovski, G., Dabney, W., Horgan, D., Piot, B., Azar, M., and Silver, D. Rainbow: Combining improvements in deep reinforcement learning. In *Thirty-Second AAAI Conference on Artificial Intelligence*, 2018.
- Ho, J. and Ermon, S. Generative adversarial imitation learning. In *Advances in Neural Information Processing Systems*, pp. 4565–4573, 2016.
- Hron, J., Matthews, D. G., Ghahramani, Z., et al. Variational bayesian dropout: Pitfalls and fixes. In *35th International Conference on Machine Learning, ICML 2018*, volume 5, pp. 3199–3219, 2018.
- Huang, J., Wu, F., Precup, D., and Cai, Y. Learning safe policies with expert guidance. In *Advances in Neural Information Processing Systems*, pp. 9105–9114, 2018.
- Ibarz, B., Leike, J., Pohlen, T., Irving, G., Legg, S., and Amodei, D. Reward learning from human preferences and demonstrations in atari. In *Advances in Neural Information Processing Systems*, 2018.
- Jacq, A., Geist, M., Paiva, A., and Pietquin, O. Learning from a learner. In *International Conference on Machine Learning*, pp. 2990–2999, 2019.
- Jorion, P. *Value at risk*. McGraw-Hill, New York, 1997.
- Kendall, A. and Gal, Y. What uncertainties do we need in bayesian deep learning for computer vision? In *Advances in neural information processing systems*, pp. 5574–5584, 2017.
- Khan, M. E., Nielsen, D., Tangkaratt, V., Lin, W., Gal, Y., and Srivastava, A. Fast and scalable bayesian deep learning by weight-perturbation in adam. *arXiv preprint arXiv:1806.04854*, 2018.
- Kingma, D. P. and Ba, J. Adam: A method for stochastic optimization. *arXiv preprint arXiv:1412.6980*, 2014.
- Lacotte, J., Ghavamzadeh, M., Chow, Y., and Pavone, M. Risk-sensitive generative adversarial imitation learning. In *The 22nd International Conference on Artificial Intelligence and Statistics*, pp. 2154–2163, 2019.
- Lakshminarayanan, B., Pritzel, A., and Blundell, C. Simple and scalable predictive uncertainty estimation using deep ensembles. In *Advances in neural information processing systems*, pp. 6402–6413, 2017.
- Laskey, M., Lee, J., Fox, R., Dragan, A., and Goldberg, K. Dart: Noise injection for robust imitation learning. *Conference on Robot Learning (CoRL)*, 2017.
- Leike, J., Martic, M., Krakovna, V., Ortega, P. A., Everitt, T., Lefrancq, A., Orseau, L., and Legg, S. Ai safety gridworlds. *arXiv preprint arXiv:1711.09883*, 2017.
- Littman, M. L., Dean, T. L., and Kaelbling, L. P. On the complexity of solving markov decision problems. *Proceedings of the Eleventh Conference on Uncertainty in Artificial Intelligence*, 1995.
- Liu, Q. and Wang, D. Stein variational gradient descent: A general purpose bayesian inference algorithm. In *Advances in neural information processing systems*, pp. 2378–2386, 2016.

- MacKay, D. J. A practical bayesian framework for backpropagation networks. *Neural computation*, 4(3):448–472, 1992.
- Maddox, W. J., Izmailov, P., Garipov, T., Vetrov, D. P., and Wilson, A. G. A simple baseline for bayesian uncertainty in deep learning. In *Advances in Neural Information Processing Systems*, pp. 13132–13143, 2019.
- Majumdar, A., Singh, S., Mandlekar, A., and Pavone, M. Risk-sensitive inverse reinforcement learning via coherent risk models. In *Robotics: Science and Systems*, 2017.
- Makhzani, A. and Frey, B. J. Pixelgan autoencoders. In *Advances in Neural Information Processing Systems*, pp. 1975–1985, 2017.
- Menda, K., Driggs-Campbell, K., and Kochenderfer, M. J. Ensembledagger: A bayesian approach to safe imitation learning. In *2019 IEEE/RSJ International Conference on Intelligent Robots and Systems (IROS)*, pp. 5041–5048. IEEE, 2019.
- Milli, S., Hadfield-Menell, D., Dragan, A., and Russell, S. Should robots be obedient? In *Proceedings of the 26th International Joint Conference on Artificial Intelligence*, pp. 4754–4760, 2017.
- Mnih, V., Kavukcuoglu, K., Silver, D., Rusu, A. A., Veness, J., Bellemare, M. G., Graves, A., Riedmiller, M., Fidjeland, A. K., Ostrovski, G., et al. Human-level control through deep reinforcement learning. *Nature*, 518(7540): 529, 2015.
- Ng, A. Y. and Russell, S. J. Algorithms for inverse reinforcement learning. In *Proceedings of the International Conference on Machine Learning*, pp. 663–670, 2000.
- Ng, A. Y., Harada, D., and Russell, S. Policy invariance under reward transformations: Theory and application to reward shaping. In *ICML*, volume 99, pp. 278–287, 1999.
- Oh, J., Guo, X., Lee, H., Lewis, R. L., and Singh, S. Action-conditional video prediction using deep networks in atari games. In *Advances in neural information processing systems*, pp. 2863–2871, 2015.
- Osband, I. Risk versus uncertainty in deep learning: Bayes, bootstrap and the dangers of dropout. In *NIPS Workshop on Bayesian Deep Learning*, 2016.
- Palan, M., Landolfi, N. C., Shevchuk, G., and Sadigh, D. Learning reward functions by integrating human demonstrations and preferences. In *Proceedings of Robotics: Science and Systems (RSS)*, June 2019.
- Petrik, M. and Russell, R. H. Beyond confidence regions: Tight bayesian ambiguity sets for robust mdps. *arXiv preprint arXiv:1902.07605*, 2019.
- Pomerleau, D. A. Efficient training of artificial neural networks for autonomous navigation. *Neural Computation*, 3(1):88–97, 1991.
- Pradier, M. F., Pan, W., Yao, J., Ghosh, S., and Doshi-Velez, F. Projected bnns: Avoiding weight-space pathologies by learning latent representations of neural network weights. *arXiv preprint arXiv:1811.07006*, 2018.
- Ramachandran, D. and Amir, E. Bayesian inverse reinforcement learning. In *Proceedings of the 20th International Joint Conference on Artificial intelligence*, pp. 2586–2591, 2007.
- Ross, S., Gordon, G., and Bagnell, D. A reduction of imitation learning and structured prediction to no-regret online learning. In *Proceedings of the fourteenth international conference on artificial intelligence and statistics*, pp. 627–635, 2011.
- Sadigh, D., Dragan, A. D., Sastry, S. S., and Seshia, S. A. Active preference-based learning of reward functions. In *Proceedings of Robotics: Science and Systems (RSS)*, 2017.
- Schulman, J., Wolski, F., Dhariwal, P., Radford, A., and Klimov, O. Proximal policy optimization algorithms. *arXiv preprint arXiv:1707.06347*, 2017.
- Sun, S., Zhang, G., Shi, J., and Grosse, R. Functional variational bayesian neural networks. *arXiv preprint arXiv:1903.05779*, 2019.
- Sutton, R. S., McAllester, D. A., Singh, S. P., and Mansour, Y. Policy gradient methods for reinforcement learning with function approximation. In *Advances in neural information processing systems*, pp. 1057–1063, 2000.
- Tamar, A., Glassner, Y., and Mannor, S. Optimizing the cvar via sampling. In *Proceedings of the Twenty-Ninth AAAI Conference on Artificial Intelligence*, pp. 2993–2999, 2015.
- Taylor, M. E. and Stone, P. Transfer learning for reinforcement learning domains: A survey. *Journal of Machine Learning Research*, 10(Jul):1633–1685, 2009.
- Thananjeyan, B., Balakrishna, A., Rosolia, U., Li, F., McAllister, R., Gonzalez, J. E., Levine, S., Borrelli, F., and Goldberg, K. Safety augmented value estimation from demonstrations (saved): Safe deep model-based rl for sparse cost robotic tasks. *arXiv preprint arXiv:1905.13402*, 2019.
- Thomas, P. S., Theodorou, G., and Ghavamzadeh, M. High-confidence off-policy evaluation. In *Proceedings of the AAAI Conference on Artificial Intelligence*, pp. 3000–3006, 2015.

- Torabi, F., Warnell, G., and Stone, P. Behavioral cloning from observation. In *Proceedings of the 27th International Joint Conference on Artificial Intelligence (IJCAI)*, July 2018.
- Volkovs, M. N. and Zemel, R. S. New learning methods for supervised and unsupervised preference aggregation. *The Journal of Machine Learning Research*, 15(1):1135–1176, 2014.
- Wilson, A., Fern, A., and Tadepalli, P. A bayesian approach for policy learning from trajectory preference queries. In *Advances in neural information processing systems*, pp. 1133–1141, 2012.
- Zhang, J. and Cho, K. Query-efficient imitation learning for end-to-end simulated driving. In *Thirty-First AAAI Conference on Artificial Intelligence*, 2017.
- Ziebart, B. D., Maas, A. L., Bagnell, J. A., and Dey, A. K. Maximum entropy inverse reinforcement learning. In *Proceedings of the 23rd AAAI Conference on Artificial Intelligence*, pp. 1433–1438, 2008.

A. Source Code and Videos

See the project webpage <https://sites.google.com/view/bayesianrex/>.

B. MCMC Details

We represent R_θ as a linear combination of pre-trained features:

$$R_\theta(\tau) = \sum_{s \in \tau} w^T \phi(s) = w^T \sum_{s \in \tau} \phi(s) = w^T \Phi_\tau. \quad (8)$$

We pre-compute and cache $\Phi_{\tau_i} = \sum_{s \in \tau_i} \phi(s)$ for $i = 1, \dots, m$ and the likelihood becomes

$$P(\mathcal{P}, D \mid R_\theta) = \prod_{(i,j) \in \mathcal{P}} \frac{e^{\beta w^T \Phi_{\tau_j}}}{e^{\beta w^T \Phi_{\tau_j}} + e^{\beta w^T \Phi_{\tau_i}}}. \quad (9)$$

We use $\beta = 1$ and enforce constraints on the weight vectors by normalizing the output of the weight vector proposal such that $\|w\|_2 = 1$ and use a Gaussian proposal function centered on w with standard deviation σ . Thus, given the current sample w_t , the proposal is defined as $w_{t+1} = \text{normalize}(\mathcal{N}(w_t, \sigma))$, in which normalize divides by the L2 norm of the sample to project back to the surface of the L2-unit ball.

For all experiments, except Seaquest, we used a default step size of 0.005. For Seaquest increased the step size to 0.05. We run 200,000 steps of MCMC and use a burn-in of 5000

Algorithm 1 Bayesian REX: Bayesian Reward Extrapolation

```

1: Input: demonstrations  $D$ , pairwise preferences  $\mathcal{P}$ ,
   MCMC proposal width  $\sigma$ , number of proposals to gen-
   erate  $N$ , deep network architecture  $R_\theta$ , and prior  $P(w)$ .
2: pre-train  $R_\theta$  using auxiliary tasks (see Section 5.2).
3: Freeze all but last layer,  $w$ , of  $R_\theta$ .
4:  $\phi(s) :=$  activations of the penultimate layer of  $R_\theta$ .
5: Precompute and cache  $\Phi_\tau = \sum_{s \in \tau} \phi(s)$  for all  $\tau \in D$ .
6: Initialize  $w$  randomly.
7: Chain[0]  $\leftarrow w$ 
8: Compute  $P(\mathcal{P}, D|w)P(w)$  using Equation (6)
9: for  $i \leftarrow 1$  to  $N$  do
10:  $\tilde{w} \leftarrow \text{normalize}(\mathcal{N}(w_t, \sigma))$ 
11: Compute  $P(\mathcal{P}, D|\tilde{w})P(\tilde{w})$  using Equation (6)
12:  $u \leftarrow \text{Uniform}(0, 1)$ 
13: if  $u < \frac{P(\mathcal{P}, D|\tilde{w})P(\tilde{w})}{P(\mathcal{P}, D|w)P(w)}$  then
14:   Chain[i]  $\leftarrow \tilde{w}$ 
15:    $w \leftarrow \tilde{w}$ 
16: else
17:   Chain[i]  $\leftarrow w$ 
18: end if
19: end for
20: return Chain
    
```

and skip every 20th sample to reduce auto-correlation. We initialize the MCMC chain with a randomly chosen vector on the L2-unit ball. Because the inverse reinforcement learning is ill-posed there are an infinite number of reward functions that could match any set of demonstrations. Prior work by Finn et al. (2016b) demonstrates that strong regularization is needed when learning cost functions via deep neural networks. To ensure that the rewards learned allow good policy optimization when fed into an RL algorithm we used a non-negative return prior on the return of the lowest ranked demonstration. The prior takes the following form:

$$\log P(w) = \begin{cases} 0 & \text{if } e^{\beta w^T \Phi_{\tau_1}} < 0 \\ -\infty & \text{otherwise} \end{cases} \quad (10)$$

This forces MCMC to not only find reward function weights that match the rankings, but to also find weights such that the return of the worse demonstration is non-negative. If the return of the worse demonstration was negative during proposal generation, then we assigned it a prior probability of $-\infty$. Because the ranking likelihood is invariant to affine transformations of the rewards, this prior simply shifts the range of learned returns and does not affect the log likelihood ratios.

C. Bayesian IRL vs. Bayesian REX

Bayesian IRL does not scale to high-dimensional tasks due to the requirement of repeatedly solving for an MDP in the inner loop. In this section we focus on low-dimensional problems where it is tractable to repeatedly solve an MDP. We compare the performance of Bayesian IRL with Bayesian REX when performing reward inference. Because both algorithms make very different assumptions, we compare their performance across three different tasks. The first task attempts to give both algorithms the demonstrations they were designed for. The second evaluation focuses on the case where all demonstrations are optimal and is designed to put Bayesian IRL at a disadvantage. The third evaluation focuses on the case where all demonstrations are optimal and is designed to put Bayesian REX at a disadvantage. Note that we focus on sample efficiency rather than computational efficiency as Bayesian IRL is significantly slower than Bayesian REX as it requires repeatedly solving an MDP, whereas Bayesian REX requires no access to an MDP during reward inference.

All experiments were performed using 6x6 gridworlds with 4 binary features placed randomly in the environment. The ground-truth reward functions are sampled uniformly from the L1-ball (Brown & Niekum, 2018). The agent can move in the four cardinal directions and stays in place if it attempts to move off the grid. Transitions are deterministic, $\gamma = 0.9$, and there are no terminal states. We perform evaluations over 100 random gridworlds for varying numbers of demonstrations. Each demonstration is truncated to a horizon of 20. We use $\beta = 50$ for both Bayesian IRL and Bayesian REX and we remove duplicates from demonstrations. After performing MCMC we used a 10% burn-in period for both algorithms and only used every 5th sample after the burn-in when computing the mean reward under the posterior. We then optimized a policy under the mean reward from the Bayesian IRL posterior and under the mean reward from the Bayesian REX posterior. We then compare the average policy loss for each algorithm when compared with optimal performance under the ground-truth reward function.

C.1. Ranked Suboptimal vs. Optimal Demonstrations

We first compare Bayesian IRL when it is given optimal demonstrations with Bayesian REX when it receives suboptimal demonstrations. We give each algorithm the demonstrations best suited for its assumptions while keeping the number of demonstrations equal and using the same starting states for each algorithm. To generate suboptimal demonstrations we simply use random rollouts and then rank them according to the ground-truth reward function.

Table 7 shows that, given a sufficient number of suboptimal ranked demonstrations (> 5), Bayesian REX performs on par or slightly better than Bayesian IRL when given the same

Table 7. Ranked Suboptimal vs. Optimal Demos: Average policy loss over 100 random 6x6 grid worlds with 4 binary features.

	2	5	10	20	30
B-IRL	0.044	0.033	0.020	0.009	0.006
B-REX	1.779	0.421	0.019	0.006	0.006

Table 8. Ranked Suboptimal Demos: Average policy loss for Bayesian IRL versus Bayesian REX over 100 random 6x6 grid worlds with 4 binary features.

	2	5	10	20	30
B-IRL	3.512	3.319	2.791	3.078	3.158
B-REX	1.796	0.393	0.026	0.006	0.006

number of optimal demonstrations starting from the same states as the suboptimal demonstrations. This result shows that not only is Bayesian REX much more computationally efficient, but it also has sample efficiency comparable to Bayesian IRL as long as there are a sufficient number of ranked demonstrations. Note that 2 ranked demonstrations induces only a single constraint on the reward function so it is not surprising that it performs much worse than running full Bayesian IRL with all the counterfactuals afforded by running an MDP solver in the inner-loop.

C.2. Only Ranked Suboptimal Demonstrations

For the next experiment we consider what happens when Bayesian IRL receives suboptimal ranked demonstrations. Table 8 shows that B-REX always significantly outperforms Bayesian IRL when both algorithms receive suboptimal ranked demonstrations. To achieve a fairer comparison, we also compared Bayesian REX with a Bayesian IRL algorithm designed to learn from both good and bad demonstrations (Cui & Niekum, 2018). We labeled the top $x\%$ ranked demonstrations as good and bottom $x\%$ ranked as bad. Table 9 shows that leveraging the ranking significantly improves the performance of Bayesian IRL, but Bayesian REX still performed significantly better across all x .

C.3. Only Optimal Demonstrations

Finally, we compared Bayesian REX with Bayesian IRL when both algorithms are given optimal demonstrations. As an attempt to use Bayesian REX with only optimal demonstrations, we followed prior work (Brown et al., 2019a) and auto-generated pairwise preferences using uniform random rollouts that are labeled as less preferred than the demonstrations. Table 10 shows that Bayesian IRL outperforms Bayesian REX. This demonstrates the value of giving a variety of ranked trajectories to Bayesian REX. For small numbers of optimal demonstrations (≤ 5) we found that

Table 9. Ranked Suboptimal Demos: Average policy loss for Bayesian REX and Bayesian IRL using the method proposed by (Cui & Niekum, 2018)* which makes use of good and bad demonstrations. We used the top $x\%$ of the ranked demos as good and bottom $x\%$ as bad. Results are averaged over 100 random 6x6 grid worlds with 4 binary features.

	Top/bottom percent of 20 ranked demos			
	x=5%	x=10%	x=25%	x=50%
B-IRL(x)*	1.283	0.956	1.065	2.096
B-REX	0.006			

Table 10. Ranked Suboptimal Demos: Average policy loss for Bayesian IRL versus Bayesian REX over 100 random 6x6 grid worlds with 4 binary features.

	2	5	10	20	30
B-IRL	0.045	0.034	0.018	0.009	0.006
B-REX	0.051	0.045	0.040	0.034	0.034

Bayesian REX leveraged the self-supervised rankings to only perform slightly worse than full Bayesian IRL. This result is encouraging since it is possible that a more sophisticated method for auto-generating suboptimal demonstrations and rankings could be used to further improve the performance of Bayesian REX even when demonstrations are not ranked (Brown et al., 2019a).

C.4. Summary

The results above demonstrate that if a very small number of unlabeled near-optimal demonstrations are available, then classical Bayesian IRL is the natural choice for performing reward inference. However, if any of these assumptions are not true, then Bayesian REX is a competitive and often superior alternative for performing Bayesian reward inference. Also implicit in the above results is the assumption that a highly tractable MDP solver is available for performing Bayesian IRL. If this is not the case, then Bayesian IRL is infeasible and Bayesian REX is the natural choice for Bayesian reward inference.

D. Pre-training Latent Reward Features

We experimented with several pretraining methods. One method is to train R_θ using the pairwise ranking likelihood function in Equation (6) and then freeze all but the last layer of weights; however, the learned embedding may overfit to the limited number of preferences over demonstrations and fail to capture features relevant to the ground-truth reward function. Thus, we supplement the pairwise ranking objective with auxiliary objectives that can be optimized in a self-supervised fashion using data from the demonstrations.

Table 11. Self-supervised learning objectives used to pre-train $\phi(s)$.

Inverse Dynamics	$f_{ID}(\phi(s_t), \phi(s_{t+1})) \rightarrow a_t$
Forward Dynamics	$f_{FD}(\phi(s_t), a_t) \rightarrow s_{t+1}$
Temporal Distance	$f_{TD}(\phi(s_t), \phi(s_{t+x})) \rightarrow x$
Variational Autoencoder	$f_A(\phi(s_t)) \rightarrow s_t$

We use the following self-supervised tasks to pre-train R_θ : (1) Learn an inverse dynamics model that uses embeddings $\phi(s_t)$ and $\phi(s_{t+1})$ to predict the corresponding action a_t (Torabi et al., 2018; Hanna & Stone, 2017), (2) Learn a forward dynamics model that predicts s_{t+1} from $\phi(s_t)$ and a_t (Oh et al., 2015; Thananjeyan et al., 2019), (3) Learn an embedding $\phi(s)$ that predicts the temporal distance between two randomly chosen states from the same demonstration (Aytar et al., 2018), and (4) Train a variational pixel-to-pixel autoencoder in which $\phi(s)$ is the learned latent encoding (Makhzani & Frey, 2017; Doersch, 2016). Table 1 summarizes the auxiliary tasks used to train $\phi(s)$.

There are many possibilities for pre-training $\phi(s)$; however, we found that each objective described above encourages the embedding to encode different features. For example, an accurate inverse dynamics model can be learned by only attending to the movement of the agent. Learning forward dynamics supplements this by requiring $\phi(s)$ to encode information about short-term changes to the environment. Learning to predict the temporal distance between states in a trajectory forces $\phi(s)$ to encode long-term progress. Finally, the autoencoder loss acts as a regularizer to the other losses as it seeks to embed all aspects of the state.

In the Atari domain, input to the network is given visually as grayscale frames resized to 84×84 . To provide temporal information, four sequential frames are stacked one on top of another to create a *framestack* which provides a brief snapshot of activity. The network architecture takes a framestack, applies four convolutional layers following a similar architecture to Christiano et al. (2017) and Brown et al. (2019b), with leaky ReLU units as non-linearities following each convolution layer. The convolutions follow the following structure:

#	Filter size	Image size	Stride
Input	-	$84 \times 84 \times 4$	-
1	7×7	$26 \times 26 \times 16$	3
2	5×5	$11 \times 11 \times 32$	2
3	5×5	$9 \times 9 \times 32$	1
4	3×3	$7 \times 7 \times 16$	1

The convolved image is then flattened. Two sequential fully connected layers, with leaky ReLU applied to the first layer, transform the flattened image into the encoding, $\phi(s)$ where s is the initial framestack. The width of these layers

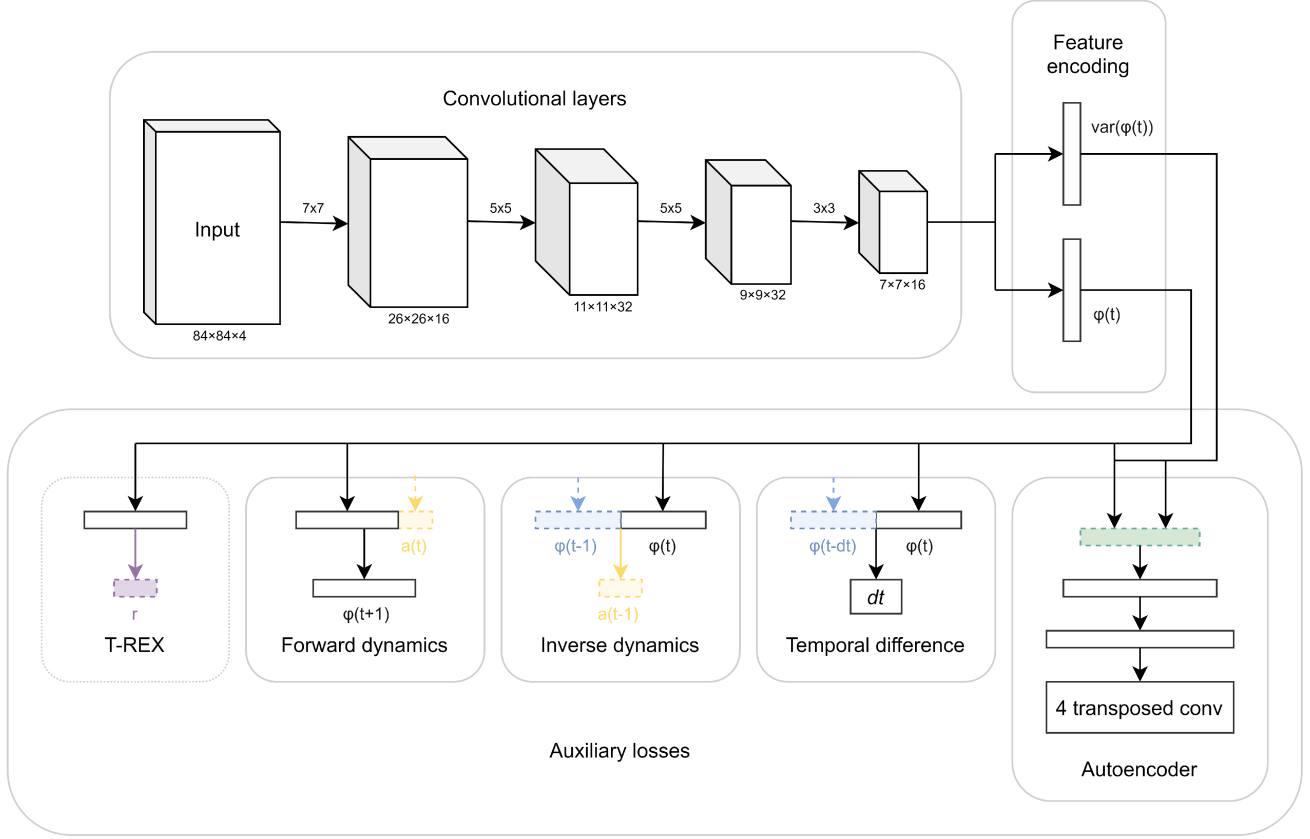


Figure 2. Diagram of the network architecture used when training feature encoding $\phi(s)$ with self-supervised and T-REX losses. Yellow denotes actions, blue denotes feature encodings sampled from elsewhere in a demonstration trajectory, and green denotes random samples for the variational autoencoder.

depends on the size of the feature encoding chosen. In our experiments with a latent dimension of 64, the first layer transforms from size 784 to 128 and the second from 128 to 64.

See Figure 2 for a complete diagram of this process.

Architectural information for each auxiliary task is given below.

1. The variational autoencoder (VAE) tries to reconstruct the original framestack from the feature encoding using transposed convolutions. Mirroring the structure of the initial convolutions, two fully connected layers precede four transposed convolution layers. These first two layers transform the 64-dimensional feature encoding from 64 to 128, and from 128 to 1568. The following four layers' structures are summarized below:

#	Filter size	Image size	Stride
Input	-	$28 \times 28 \times 2$	-
1	3×3	$30 \times 30 \times 4$	1
2	6×6	$35 \times 35 \times 16$	1
3	7×7	$75 \times 75 \times 16$	2
4	10×10	$84 \times 84 \times 4$	1

A cross-entropy loss is applied between the reconstructed image and the original, as well as a term added to penalize the KL divergence of the distribution from the unit normal.

2. A temporal difference estimator, which takes two random feature encodings from the same demonstration and predicts the number of timesteps in between. It is a single fully-connected layer, transforming the concatenated feature encodings into a scalar time difference. A mean-squared error loss is applied between the real difference and predicted.
3. An inverse dynamics model, which takes two sequential feature encodings and predicts the action taken in between. It is again a single fully-connected layer,

trained as a classification problem with a binary cross-entropy loss over the discrete action set.

4. A forward dynamics model, which takes a concatenated feature encoding and action and predicts the next feature encoding with a single fully-connected layer. This is repeated 5 times, which increases the difference between the initial and final encoding. It is trained using a mean-squared error between the predicted and real feature encoding.
5. A T-REX loss, which samples feature encodings from two different demonstrations and tries to predict which one of them has preference over the other. This is done with a single fully-connected layer that transforms an encoding into scalar reward, and is then trained as a classification problem with a binary cross-entropy loss. A 1 is assigned to the demonstration sample with higher preference and a 0 to the demonstration sample with lower preference.

In order to encourage a feature encoding that has information easily interpretable via linear combinations, the temporal difference, T-REX, inverse dynamics, and forward dynamics tasks consist of only a single layer atop the feature encoding space rather than multiple layers.

To compute the final loss on which to do the backwards pass, all of the losses described above are summed with weights determined empirically to balance out their values.

D.1. Training specifics

We used an NVIDIA TITAN V GPU for training the embedding. We used the same 12 demonstrations used for MCMC to train the self-supervised and ranking losses described above. We sample 60,000 trajectory snippets pairs from the demonstration pool, where each snippet is between 50 and 100 timesteps long. We use a learning rate of 0.001 and a weight decay of 0.001. We make a single pass through all of the training data using batch size of 1 resulting in 60,000 updates using the Adam (Kingma & Ba, 2014) optimizer. For Enduro prior work (Brown et al., 2019b) showed that full trajectories resulted in better performance than subsampling trajectories. Thus, for Enduro we subsample 10,000 pairs of entire trajectories by randomly selecting a starting time between 0 and 5 steps after the initial state and then skipping every t frames where t is chosen uniformly from the range $[3, 7)$ and train with two passes through the training data. When performing subsampling for either snippets or full trajectories, we subsample pairs of trajectories such that one is from a worse ranked demonstration and one is from a better ranked demonstration following the procedure outlined in (Brown et al., 2019b).

E. Visualizations of Learned Features

Viewable [here](#)² is a video containing an Enduro demonstration trajectory, its decoding with respect to the pre-trained autoencoder, and a plot of the dimensions in the latent encoding over time. Observe how changes in the demonstration, such as turning right or left or a shift, correspond to changes in the plots of the feature embedding. We noticed that certain features increase when the agent passes other cars while other features decrease when the agent gets passed by other cars. This is evidence that the pretraining has learned features that are relevant to the ground truth reward which gives +1 every time the agent passes a car and -1 every time the agent gets passed.

Viewable [here](#)³ is a similar visualization of the latent space for Space Invaders. Notice how it tends to focus on the movement of enemy ships, useful for game progress in things such as the temporal difference loss, but seems to ignore the player’s ship despite its utility in inverse dynamics loss. Likely the information exists in the encoding but is not included in the output of the autoencoder.

Viewable [here](#)⁴ is visualization of the latent space for Breakout. Observe that breaking a brick often results in a small spike in the latent encoding. Many dimensions, like the dark green curve which begins at the lowest value, seem to invert as game progress continues on, thus acting as a measure of how much time has passed.

F. Imitation Learning Ablations for Pre-training $\phi(s)$

Table 12 shows the results of pre-training reward features only using different losses. We experimented with using only the T-REX Ranking loss (Brown et al., 2019b), only the self-supervised losses shown in Table 1 of the main paper, and using both the T-REX ranking loss plus the self-supervised loss function. We found that performance varied over the different pre-training schemes, but that using Ranking + Self-Supervised achieved high performance across all games, clearly outperforming only using self-supervised losses and achieving superior performance to only using the ranking loss on 3 out of 5 games.

G. Suboptimal Demonstration Details

We used the same suboptimal demonstrations used by Brown et al. (2019b) for comparison. These demonstrations

²<https://www.youtube.com/watch?v=DMf8kNH9nVg>

³<https://www.youtube.com/watch?v=2uN5uD17H6M>

⁴<https://www.youtube.com/watch?v=8zgbD1fZOH8>

Table 12. Comparison of different reward feature pre-training schemes. Ground-truth average returns for several Atari games when optimizing the mean and MAP rewards found using Bayesian REX. Each algorithm is given the same 12 demonstrations with ground-truth pairwise preferences. The average performance for each IRL algorithm is the average over 30 rollouts.

Game	Ranking Loss		Self-Supervised		Ranking + Self-Supervised	
	Mean	MAP	Mean	MAP	Mean	MAP
Beam Rider	3816.7	4275.7	180.4	143.7	5870.3	5504.7
Breakout	389.9	409.5	360.1	367.4	393.1	390.7
Enduro	472.7	479.3	0.0	0.0	135.0	487.7
Seaquest	675.3	670.7	674.0	683.3	606.0	734.7
Space Invaders	1482.0	1395.5	391.2	396.2	961.3	1118.8

were obtained by running PPO on the ground truth reward and checkpointing every 50 updates using OpenAI Baselines (Dhariwal et al., 2017). Brown et al. (2019b) make the checkpoint files available, so to generate the demonstration data we used their saved checkpoints and followed the instructions in their released code to generate the data for our algorithm⁵. We gave Bayesian REX these demonstrations as well as ground-truth rankings using the game score; however, other than the rankings, Bayesian REX has no access to the true reward samples. Following the recommendations of Brown et al. (2019b), we mask the game score and other parts of the game that are directly indicative of the game score such as the number of enemy ships left, the number of lives left, the level number, etc. See (Brown et al., 2019b) for full details.

H. Reinforcement Learning Details

We used the OpenAI Baselines implementation of Proximal Policy Optimization (PPO) (Schulman et al., 2017; Dhariwal et al., 2017). We used the default hyperparameters for all games and all experiments. We run RL for 50 million frames and then take the final checkpoint to perform evaluations. We adapted the OpenAI Baselines code so even though the RL agent receives a standard preprocessed observation, it only receives samples of the reward learned via Bayesian REX, rather than the ground-truth reward. T-REX (Brown et al., 2019b) uses a sigmoid to normalize rewards before passing them to the RL algorithm; however, we obtained better performance for Bayesian REX by feeding the unnormalized predicted reward $R_\theta(s)$ into PPO for policy optimization. We follow the OpenAI baselines default preprocessing for the framestacks that are fed into the RL algorithm as observations. We also apply the default OpenAI baselines wrappers the environments. We run PPO with 9 workers on an NVIDIA TITAN V GPU.

⁵Code from (Brown et al., 2019b) is available here <https://github.com/hiwonjoon/ICML2019-TREX>

I. High-Confidence Policy Performance Bounds

In this section we describe the details of the policy performance bounds.

I.1. Policy Evaluation Details

We estimated $\Phi_{\pi_{\text{eval}}}$ using C Monte Carlo rollouts for each evaluation policy. Thus, after generating C rollouts, $\tau_1, \dots, \tau_C$ from π_{eval} the feature expectations are computed as

$$\Phi_{\pi_{\text{eval}}} = \frac{1}{C} \left[\sum_{i=1}^C \sum_{s \in \tau_i} \phi(s) \right]. \quad (11)$$

We used $C = 100$ for all experiments.

I.2. Evaluation Policies

We evaluated several different evaluation policies. To see if the learned reward function posterior can interpolate and extrapolate we created four different evaluation policies: A, B, C, and D. These policies were created by running RL via PPO on the ground truth reward for the different Atari games. We then checkpointed the policy and selected checkpoints that would result in different levels of performance. For all games except for Enduro these checkpoints correspond to 25, 325, 800, and 1450 update steps using OpenAI baselines. For Enduro, PPO performance was stuck at 0 return until much later in learning. To ensure diversity in the evaluation policies, we chose to use evaluation policies corresponding to 3125, 3425, 3900, and 4875 steps. We also evaluated each game with a No-Op policy. These policies are often adversarial for some games, such as Seaquest, Breakout, and Beam Rider, since they allow the agent to live for a very long time without actually playing the game—a potential way to hack the learned reward since most learned rewards for Atari will incentivize longer gameplay.

The results for Beam Rider and Breakout are shown in the main paper. For completeness, we have included the high-confidence policy evaluation results for the other games here

Table 13. Policy evaluation statistics for Enduro over the return distribution from the learned posterior $P(R|D, \mathcal{P})$ compared with the ground truth returns using game scores. Policies A-D correspond to checkpoints of an RL policy partially trained on the ground-truth reward function and correspond to 25, 325, 800, and 1450 training updates to PPO. No-Op that always plays the no-op action, resulting in high mean predicted performance but low 95%-confidence return (0.05-VaR).

Policy	Predicted		Ground Truth	
	Mean	0.05-VaR	Avg.	Length
A	324.7	48.2	7.3	3322.4
B	328.9	52.0	26.0	3322.4
C	424.5	135.8	145.0	3389.0
D	526.2	192.9	199.8	3888.2
Mean	1206.9	547.5	496.7	7249.4
MAP	395.2	113.3	133.6	3355.7
No-Op	245.9	-31.7	0.0	3322.0

in the Appendix. Table 13 shows the high-confidence policy evaluation results for Enduro. Both the average returns over the posterior as well as the the high-confidence performance bounds ($\delta = 0.05$) demonstrate accurate predictions relative to the ground-truth performance. The No-Op policy results in the racecar slowly moving along the track and losing the race. This policy is accurately predicted as being much worse than the other evaluation policies. We also evaluated the Mean and MAP policies found by optimizing the Mean reward and MAP reward from the posterior obtained using Bayesian REX. We found that the learned posterior is able to capture that the MAP policy is more than twice as good as the evaluation policy D and that the Mean policy has performance somewhere between the performance of policies B and C. These results show that Bayesian REX has the potential to predict better-than-demonstrator performance (Brown et al., 2019a).

Table 14 shows the results for high-confidence policy evaluation for Seaquest. The results show that high-confidence performance bounds are able to accurately predict that evaluation policies A and B are worse than C and D. The ground truth performance of policies C and D are too close and the mean performance over the posterior and 0.05-VaR bound on the posterior are not able to find any statistical difference between them. Interestingly the no-op policy has very high mean and 95%-confidence lower bound, despite not scoring any points. However, as shown in the bottom half of Table 14, adding one more ranked demonstration from a 3000 length segment of a no-op policy solves this problem. These results motivate a natural human-in-the-loop approach for safe imitation learning.

Finally, Table 15 shows the results for high-confidence policy evaluation for Space Invaders. The results show that us-

Table 14. Policy evaluation statistics for Seaquest over the return distribution from the learned posterior $P(R|D, \mathcal{P})$ compared with the ground truth returns using game scores. Policies A-D correspond to checkpoints of an RL policy partially trained on the ground-truth reward function and correspond to 25, 325, 800, and 1450 training updates to PPO. No-Op always plays the no-op action, resulting in high mean predicted performance but low 0.05-quantile return (0.05-VaR). Results predict that No-Op is much better than it really is. However, simply adding a single ranked rollout from the No-Op policy and rerunning MCMC results in correct relative rankings with respect to the No-Op policy

Policy	Predicted		Ground Truth	
	Mean	0.05-VaR	Avg.	Length
A	24.3	10.8	338.6	1077.8
B	53.6	24.1	827.2	2214.1
C	56.0	25.4	872.2	2248.5
D	55.8	25.3	887.6	2264.5
No-Op	2471.6	842.5	0.0	99994.0
Results after adding one ranked demo from No-Op				
A	0.5	-0.5	338.6	1077.8
B	3.7	2.0	827.2	2214.1
C	3.8	2.1	872.2	2248.5
D	3.2	1.5	887.6	2264.5
No-Op	-321.7	-578.2	0.0	99994.0

ing both the mean performance and 95%-confidence lower bound are good indicators of ground truth performance for the evaluation policies. The No-Op policy for Space Invaders results in the agent getting hit by alien lasers early in the game. The learned reward function posterior correctly assigns low average performance and indicates high risk with a low 95%-confidence lower bound on the expected return of the evaluation policy.

J. Different Evaluation Policies

To test Bayesian REX on different learned policies we took a policy trained with RL on the ground truth reward function for Beam Rider, the MAP policy learned via Bayesian REX for Beam Rider, and a policy trained with an earlier version of Bayesian REX (trained without all of the auxiliary losses) that learned a novel reward hack where the policy repeatedly presses left and then right, enabling the agent’s ship to stay in between two of the firing lanes of the enemies. The imitation learning reward hack allows the agent to live for a very long time. We took a 2000 step prefix of each policy and evaluated the expected and 5th percentile worst-case predicted returns for each policy. We found that Bayesian REX is able to accurately predict that the reward hacking policy is worse than both the RL policy and the policy optimizing the Bayesian REX reward. However, we found

Table 15. Policy evaluation statistics for Space Invaders over the return distribution from the learned posterior $P(R|D, \mathcal{P})$ compared with the ground truth returns using game scores. Policies A-D correspond to checkpoints of an RL policy partially trained on the ground-truth reward function and correspond to 25, 325, 800, and 1450 training updates to PPO. The mean and MAP policies are the results of PPO using the mean and MAP rewards, respectively. No-Op that always plays the no-op action, resulting in high mean predicted performance but low 0.05-quantile return (0.05-VaR).

Policy	Predicted		Ground Truth	
	Mean	0.05-VaR	Avg.	Length
A	45.1	20.6	195.3	550.1
B	108.9	48.7	436.0	725.7
C	148.7	63.6	575.2	870.6
D	150.5	63.8	598.2	848.2
Mean	417.4	171.7	1143.7	1885.7
MAP	360.2	145.0	928.0	1629.5
NoOp	18.8	3.8	0.0	504.0

Table 16. Beamrider policy evaluation for an RL policy trained on ground truth reward, an imitation learning policy, and a reward hacking policy that exploits a game hack to live for a long time by moving quickly back and forth.

Policy	Predicted		Ground Truth	
	Mean	0.05-VaR	Avg.	Length
RL	36.7	19.5	2135.2	2000
B-REX	68.1	38.1	649.4	2000
Hacking	28.8	10.2	2.2	2000

that the Bayesian REX policy, while not performing as well as the RL policy, was given higher expected return and a higher lower bound on performance than the RL policy. Results are shown in Table 16.

K. Human Demonstrations

To investigate whether Bayesian REX is able to correctly rank human demonstrations, one of the authors provided demonstrations of a variety of different behaviors and then we took the latent embeddings of the demonstrations and used the posterior distribution to find high-confidence performance bounds for these different rollouts.

K.1. Beamrider

We generated four human demonstrations: (1) *good*, a good demonstration that plays the game well, (2) *bad*, a bad demonstration that seeks to play the game but does a poor job, (3) *pessimal*, a demonstration that does not shoot enemies and seeks enemy bullets, and (4) *adversarial* a demonstration that pretends to play the game by moving and shoot-

Table 17. Beam Rider evaluation of a variety of human demonstrations.

Policy	Predicted		Ground Truth	
	Mean	0.05-VaR	Avg.	Length
good	12.4	5.8	1092	1000.0
bad	10.7	4.5	396	1000.0
pessimal	6.6	0.8	0	1000.0
adversarial	8.4	2.4	176	1000.0

Table 18. Space Invaders evaluation of a variety of human demonstrations.

Policy	Predicted		Ground Truth	
	Mean	0.05-VaR	Avg.	Length
good	198.3	89.2	515	1225.0
every other	56.2	25.9	315	728.0
hold 'n fire	44.3	18.6	210	638.0
shoot shelters	47.0	20.6	80	712.0
flee	45.1	19.8	0	722.0
hide	83.0	39.0	0	938.0
miss	66.0	29.9	0	867.0
pessimal	0.5	-13.2	0	266.0

ing as much as possibly but tries to avoid actually shooting enemies. The results of high-confidence policy evaluation are shown in Table 17. The high-confidence bounds and average performance over the posterior correctly rank the behaviors. This provides evidence that the learned linear reward correctly rewards actually destroying aliens and avoiding getting shot, rather than just flying around and shooting.

K.2. Space Invaders

For Space Invaders we demonstrated an even wider variety of behaviors to see how Bayesian REX would rank their relative performance. We evaluated the following policies: (1) *good*, a demonstration that attempts to play the game as well as possible, (2) *every other*, a demonstration that only shoots aliens in the 2nd and 4th columns, (3) *flee*, a demonstration that did not shoot aliens, but tried to always be moving while avoiding enemy lasers, (4) *hide*, a demonstration that does not shoot and hides behind one of the barriers to avoid enemy bullets, (5) *pessimal*, a policy that seeks enemy bullets while not shooting, (6) *shoot shelters*, a demonstration that tries to destroy its own shelters by shooting at them, (7) *hold 'n fire*, a demonstration where the player rapidly fires but does not move to avoid enemy lasers, and (8) *miss*, a demonstration where the demonstrator tries to fire but not hit any aliens while avoiding enemy lasers.

Table 18 shows the results of evaluating the different demonstrations. The good demonstration is clearly the best per-

Table 19. Space Invaders evaluation of a variety of human demonstrations when considering only the first 6000 steps.

Policy	Predicted		Ground Truth	
	Mean	0.05-VaR	Avg.	Length
good	47.8	22.8	515	600.0
every other	34.6	15.0	315	600.0
hold 'n fire	40.9	17.1	210	600.0
shoot shelters	33.0	13.3	80	600.0
flee	31.3	11.9	0	600.0
hide	32.4	13.8	0	600.0
miss	30.0	11.3	0	600.0

forming demonstration in terms of mean performance and 95%-confidence lower bound on performance and the pessimistic policy is correctly given the lowest performance lower bound. However, we found that the length of the demonstration appears to have a strong effect on the predicted performance for Space Invaders. Demonstrations such as hide and miss are able to live for a longer time than policies that actually hit aliens. This results in them having higher 0.05-quantile worst-case predicted performance and higher mean performance.

To study this further we looked at only the first 600 timesteps of each policy, to remove any confounding by the length of the trajectory. The results are shown in Table 19. With a fixed length demonstration, Bayesian REX is able to correctly rank *good*, *every other*, and *hold 'n fire* as the best demonstrations, despite evaluation policies that are deceptive.

K.3. Enduro

For Enduro we tested four different human demonstrations: (1) *good* a demonstration that seeks to play the game well, (2) *periodic* a demonstration that alternates between speeding up and passing cars and then slowing down and being passed, (3) *neutral* a demonstration that stays right next to the last car in the race and doesn't try to pass or get passed, and (4) *ram* a demonstration that tries to ram into as many cars while going fast. Table 20 shows that Bayesian REX is able to accurately predict the performance and risk of each of these demonstrations and gives the highest (lowest 0.05-VaR) risk to the *ram* demonstration and the least risk to the *good* demonstration.

L. Comparison with Other Methods for Uncertainty Quantification

Bayesian REX is only one possible method for measure uncertainty. Other popular methods for measuring epistemic uncertainty include using bootstrapping to create an ensemble

Table 20. Enduro evaluation of a variety of human demonstrations.

Policy	Predicted		Ground Truth	
	Mean	0.05-VaR	Avg.	Length
good	246.7	-113.2	177	3325.0
periodic	230.0	-130.4	44	3325.0
neutral	190.8	-160.6	0	3325.0
ram	148.4	-214.3	0	3325.0

of neural networks (Lakshminarayanan et al., 2017) and using dropout as an approximation of MCMC sampling (Gal & Ghahramani, 2016). In this section we compare our fully Bayesian approach with these two approximations.

L.1. T-REX Ensemble

We used the same implementation used by Brown et al. (2019b)⁶, but trained an ensemble of five T-REX networks using the same training demonstrations but with randomized seeds so each network is initialized differently and has a different training set of subsampled snippets from the full length ranked trajectories. To estimate the uncertainty over the return of a trajectory or policy we run the trajectory through each network to get a return estimate or run multiple rollouts of the policy through each member of the ensemble to get a distribution over returns. We used 100 rollouts for the evaluation policies.

L.2. MC Dropout

For the MC Dropout baseline we used the same base architecture as T-REX and Bayesian REX, except that we did not add additional auxiliary losses, but simply trained the base network to predict returns using dropout during training. For each training pair of trajectories we randomly sample a dropout mask on the last layer of weights. Because MC dropout is supposed to approximate a large ensemble, we kept the dropout mask consistent across each sampled preference pair such that the same portions of the network are dropped out for each frame of each trajectory and for both the more preferred and less preferred trajectories. Thus, for each training sample, consisting of a more and less preferred trajectory, we sample a random dropout mask and then apply this same mask across all states in both trajectories. To keep things as similar to Bayesian REX as possible, we used full trajectories with the same pairwise preferences used by Bayesian REX.

To estimate the posterior distribution over returns for a trajectory or policy we simply sampled 50 random masks on the last layer of weights. Thus, this method corresponds to the MC dropout equivalent of Bayesian REX where the latent

⁶<https://github.com/hiwonjoon/icml2019-trex>

state encoding is trained end-to-end via dropout rather than pre-trained and where the posterior distribution is estimated via randomly dropping out weights on the corresponding linear reward function. We applied these 50 random dropouts to each of 100 rollouts for each evaluation policy. We used a dropout probability of 0.5.

Table 21 shows the results for running RL on the learned reward functions. The results show that Bayesian REX is superior or competitive with T-REX Ensemble and MC Dropout across all games except Beam Rider, where MC Dropout performs much better.

L.3. T-REX Ensemble High-Confidence Bounds

Tables 22–26 show the results for evaluating different evaluation policies via high-confidence performance bounds. Table 22 shows that the ensemble has accurate expected returns, but that the 95% confidence lower bounds are not informative and do not represent risk as accurately as Bayesian REX since policy D is viewed as much worse than policy A. Note that we normalized the predicted scores by calculating the average predicted return of each ensemble member for rollouts from policy A and then using this as a baseline for all other predictions of each ensemble member by subtracting off the average predicted return for policy A from return predictions of other policies. Tables 23 and 25 show that the T-REX Ensemble can sometimes fail to produce meaningful predictions for the expectation or the 95% worst-case bounds. Table 24 and 26 show good predictions.

L.4. MC Dropout Results

Tables 27–31 show the results for high-confidence bounds for MC Dropout. Tables 27 and 31 show that MC Dropout is able to accurately predict high risk for the Beam Rider and Space Invaders No-Op policies. However, table 28 29, and 30 show that MC Dropout often fails to predict that the No-Op policy has high risk. Recent work has shown that MC Dropout is not a principled Bayesian approximation since the distribution obtained from MC Dropout does not concentrate in the limit as the number of data samples goes to infinity and thus does not necessarily measure the kind of epistemic risk we are interested in (Osband, 2016). Thus, while MC Dropout does not perform full Bayesian inference like Bayesian REX, it appears to work sometimes in practice. Future work should examine more sophisticated applications of dropout to uncertainty estimation that seek to solve the theoretical and practical problems with vanilla MC Dropout (Hron et al., 2018).

Table 21. Comparison of policy performance when using a reward function learned by Bayesian REX, a T-REX ensemble, and a dropout version of Bayesian REX. The results show averages (standard deviations over 30 rollouts. Bayesian REX results in comparable or better performance across all games except Beam Rider.

Game	Bayesian REX		T-REX Ensemble	MC Dropout
	Mean	MAP		
Beam Rider	5870.3 (1905.1)	5504.7 (2121.2)	5925.0 (2097.9)	7243.1 (2543.6)
Breakout	393.1 (63.7)	390.7 (48.8)	253.7 (136.2)	52.6 (10.1)
Enduro	135.0 (24.8)	487.7 (89.4)	281.5 (95.2)	111.8 (17.9)
Seaquest	606.0 (37.6)	734.7 (41.9)	0.0 (0.0)	0.0 (0.0)
Space Invaders	961.3 (392.3)	1118.8 (483.1)	1164.8 (546.3)	387.5 (166.3)

Table 22. Beam Rider T-REX ensemble.

Policy	Predicted		Ground Truth Avg.	
	Mean	0.05-VaR	Score	Length
A	0.0	-119.5	454.4	1372.6
B	76.1	-63.7	774.8	1412.8
C	201.3	-173.8	1791.8	2389.9
D	282.0	-304.0	2664.5	2965.0
Mean	956.9	-2294.8	5736.8	9495.6
MAP	1095.7	-2743.4	5283.0	11033.4
No-Op	-1000.0	-5643.7	0.0	99,994.0

Table 25. Seaquest T-REX ensemble.

Policy	Predicted		Ground Truth Avg.	
	Mean	0.05-VaR	Score	Length
A	0.0	-320.9	321.0	1077.8
B	-425.2	-889.3	826.6	2214.1
C	-336.7	-784.3	863.4	2248.5
D	-386.3	-837.3	884.4	2264.5
Mean	-1013.1	-2621.8	721.8	2221.7
MAP	-636.8	-1820.1	607.4	2247.2
No-Op	-19817.9	-28209.7	0.0	99,994.0

Table 23. Breakout T-REX ensemble.

Policy	Predicted		Ground Truth Avg.	
	Mean	0.05-VaR	Score	Length
A	0.0	-506.3	1.8	202.7
B	-305.8	-1509.1	16.1	608.4
C	-241.1	-1780.7	24.8	849.3
D	-57.9	-2140.0	42.7	1020.8
Mean	-5389.5	-867.4	388.9	13762.1
MAP	-3168.5	-1066.5	401.0	8780.0
No-Op	-39338.4	-95987.2	0.0	99,994.0

Table 26. Space Invaders T-REX ensemble.

Policy	Predicted		Ground Truth Avg.	
	Mean	0.05-VaR	Score	Length
A	0.0	-136.8	159.4	550.1
B	257.3	-47.9	425.0	725.7
C	446.5	-6.0	553.1	870.6
D	443.3	9.0	591.5	848.2
Mean	1105.6	-392.4	1143.7	1885.7
MAP	989.0	-387.2	928.0	1629.5
No-Op	-211.9	-311.9	0.0	504.0

Table 24. Enduro T-REX ensemble.

Policy	Predicted		Ground Truth Avg.	
	Mean	0.05-VaR	Score	Length
A	-0.0	-137.7	9.4	3322.4
B	23.6	-157.4	23.2	3322.4
C	485.4	232.2	145.6	2289.0
D	1081.1	270.3	214.2	3888.2
Mean	3408.9	908.0	496.7	7249.4
MAP	23.2	-1854.1	133.6	3355.7
No-Op	-1618.1	-3875.9	0.0	3322.0

Table 27. Beam Rider MC Dropout.

Policy	Predicted		Ground Truth Avg.	
	Mean	0.05-VaR	Score	Length
A	20.9	-1.5	454.4	1372.6
B	27.9	2.3	774.8	1412.8
C	48.7	8.3	1791.8	2389.9
D	63.5	11.0	2664.5	2965.0
Mean	218.2	-89.2	5736.8	1380
MAP	211.2	-148.7	5283.0	708
No-Op	171.2	-3385.7	0.0	99,994.0

Table 28. Breakout MC Dropout.

Policy	Predicted		Ground Truth Score	Avg. Length
	Mean	0.05-VaR		
A	10.8	5.2	1.8	202.7
B	33.1	17.7	16.1	608.4
C	43.5	24.1	24.8	849.3
D	56.0	28.5	42.7	1020.8
Mean	822.9	77.3	388.9	13762.1
MAP	519.7	73.8	401.0	8780.0
No-Op	6050.7	3912.4	0.0	99,994.0

Table 29. Enduro MC Dropout.

Policy	Predicted		Ground Truth Score	Avg. Length
	Mean	0.05-VaR		
A	541.7	398.0	7.3	3322.4
B	543.6	401.0	26.4	3322.4
C	556.7	409.3	142.5	3389.0
D	663.3	422.3	200.3	3888.2
Mean	2473.0	1701.7	496.7	7249.4
MAP	1097.3	799.5	133.6	3355.7
No-Op	1084.1	849.8	0.0	3322.0

Table 30. Seaquest MC Dropout.

Policy	Predicted		Ground Truth Score	Avg. Length
	Mean	0.05-VaR		
A	98.9	49.5	321.0	1077.8
B	258.8	194.8	826.6	2214.1
C	277.7	213.2	863.4	2248.5
D	279.6	214.2	884.4	2264.5
Mean	375.6	272.8	721.8	2221.7
MAP	426.3	319.8	607.4	2247.2
No-Op	16211.1	10478.5	0.0	99,994.0

Table 31. Space Invaders MC Dropout.

Policy	Predicted		Ground Truth Score	Avg. Length
	Mean	0.05-VaR		
A	10.6	0.8	195.3	550.1
B	22.3	8.8	434.9	725.7
C	26.7	9.8	535.3	870.6
D	28.9	15.6	620.9	848.2
Mean	125.9	54.4	1143.7	848.2
MAP	110.6	52.5	928.0	1885.7
No-Op	8.4	-8.6	0.0	504.0