
Blindfolded Experts Generalize Better: Insights from Robotic Manipulation and Videogames

Ev Zisselman*, Mirco Mutti, Shelly Francis-Meretzki, Elisei Shafer, Aviv Tamar
Technion – Israel Institute of Technology

Abstract

Behavioral cloning is a simple yet effective technique for learning sequential decision-making from demonstrations. Recently, it has gained prominence as the core of foundation models for the physical world, where achieving generalization requires countless demonstrations of a multitude of tasks. Typically, a human expert with full information on the task demonstrates a (nearly) optimal behavior. In this paper, we propose to hide some of the task’s information from the demonstrator. This “blindfolded” expert is compelled to employ non-trivial *exploration* to solve the task. We show that cloning the blindfolded expert generalizes better to unseen tasks than its fully-informed counterpart. We conduct experiments of real-world robot peg insertion tasks with (limited) human demonstrations, alongside videogames from the Procgen benchmark. Additionally, we support our findings with theoretical analysis, which confirms that the generalization error scales with $\sqrt{I/m}$, where I measures the amount of task information available to the demonstrator, and m is the number of demonstrated tasks. Both theory and practice indicate that cloning blindfolded experts generalizes better with fewer demonstrated tasks. Project page with videos and code: <https://sites.google.com/view/blindfoldedexperts/home>.

1 Introduction

Behavioral cloning (BC) is a simple yet effective method for training policies in sequential decision-making problems [32, 6]. In BC, an expert demonstrates how to perform a task, and the sequence of observation-action data is input to a supervised learning algorithm for training a policy.

A key question in BC is generalization—how many demonstrations are required to train an effective policy. For single tasks, a well-investigated challenge is compounding errors—small mistakes in the trained policy may lead to visit states that the expert did not visit, further increasing the prediction errors [43]. However, recent results show that using appropriate neural-network architectures, BC can learn to solve complex tasks even with a modest number of demonstrations [6, 59], and these results are reinforced by recent theory [13]. For multiple tasks (or significant variations of a single task), on the other hand, BC still requires abundant data, and recent methods for mitigating data requirements include augmentations [26], simulation [52], and fine tuning foundation models trained on large scale demonstration data [30, 49, 21]. In this work, we hence focus on generalization to task variations.

While various works study how to improve generalization via the BC algorithm [43], the policy representation [6], and the data diversity [23], one aspect that remains unexplored is *the experts themselves*. Many tasks can be solved in various ways—can some behaviors generalize better than others? Recently, in the context of zero-shot reinforcement learning, Zisselman et al. [61] showed that certain exploratory behaviors generalize better than goal-oriented, reward-maximizing behavior. Intuitively, since exploratory behavior is less goal-oriented, it is less dependent on any particular task

*Correspondence E-mail: ev_zis@campus.technion.ac.il

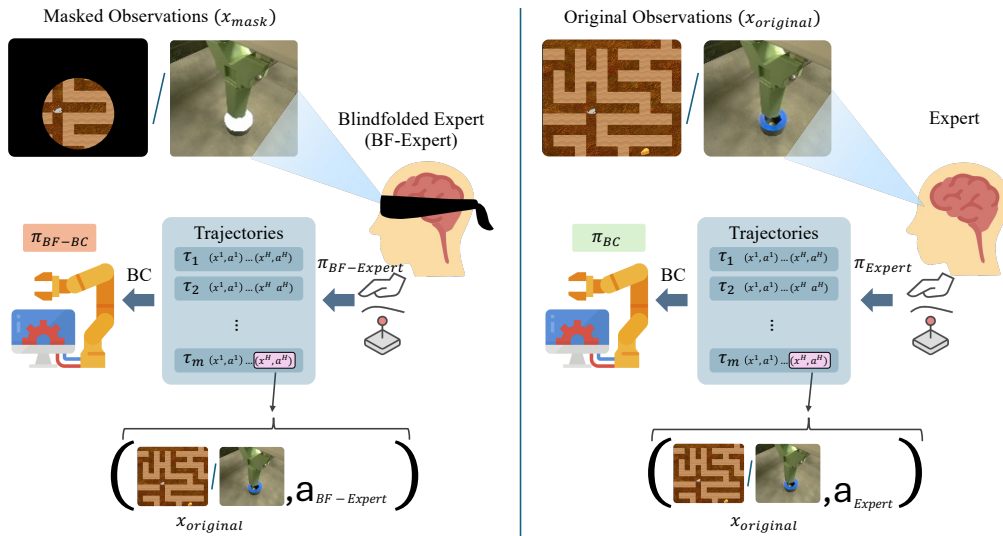


Figure 1: Illustration of the learning process. Note that the mask on observations only applies to the blindfolded expert, while the observations in the logged trajectories are unmasked in both cases.

instance and therefore, more likely to generalize to novel tasks. In this work we ask—can a similar principle be useful also for imitation learning?

Our main idea (depicted in Figure 1) is that by introducing a *blindfold*—an information bottleneck on the expert’s observation that makes it harder to identify the particular task, we can induce the expert to express a more exploratory and less task-dependent behavior, which we conjecture will generalize better. Importantly, *our method does not change the observations used for training the policy* but only the expert’s behavior, and is therefore compliant with any BC algorithm, and complementary to the methods for improving generalization mentioned above. Since the expert’s exploratory behavior is typically history dependent, our method makes use of policy architectures that can process a sequence of observations, such as recurrent neural networks or transformers [7, 53].

Theoretically, we prove an upper bound on the generalization that scales with $\mathcal{E}_{gen} + \sqrt{I/m}$, where I measures the amount of task information available to the demonstrator, m is the number of demonstrated tasks, and $\mathcal{E}_{gen}$ is a cost associated with the expert not taking optimal actions. For domains where the exploratory behavior can still solve the task, $\mathcal{E}_{gen}$ is zero, and thus by lowering I using a “blindfold” we reduce the generalization error without paying any price. To our knowledge, this result is the first of its kind in relating non-trivial properties of the expert’s behavior to multi-task generalization of the resulting BC policy.

Empirically, we demonstrate our approach on simulated games from the Procgen suite [8]—a standard benchmark for task generalization, and on a real-robot peg insertion task, based on the FMB challenge [25], where the shapes of the peg and the hole define the task. For the Procgen games, our blindfold hides the observation and reveals only the agent’s immediate surroundings to the expert. For the peg insertion domain, we let the expert teleoperate the robot by observing images from robot-mounted cameras, and mask out the hole shape from the image. In both domains, we find that the blindfold induces more exploratory behavior from the expert, which in turn yields significantly better generalization to different tasks.

Our results pave the way to a new and principled approach for collecting demonstrations, both for specific problems, and also for more general foundation-model scale endeavors.

2 Problem formulation

Throughout the paper, we will focus on a *multi-task* imitation learning setup in which we aim to clone an expert’s behavior from demonstrations of a (small) selection of tasks with the goal to generalize the behavior to (many) unseen tasks. As we shall see both theoretically and empirically, the generalization is not only affected by the number of available demonstrations, but also by the information available to the expert when performing demonstrations. First, let us introduce the setting formally.

Setting. We consider a set of tasks $\Theta := \{\theta_i\}_{i=1}^M$ and a task distribution $P_0 \in \Delta(\Theta)$, where $\Delta(S)$ denotes the probability simplex over a set S . Each task $\theta \in \Theta$ is defined through a Markov Decision Process (MDP [33]) $\mathcal{M}_\theta := (\mathcal{X}, \mathcal{A}, p_\theta, r_\theta, H)$, where $\mathcal{X}$, $\mathcal{A}$, H respectively denote the observation space, the action space, and the horizon of an episode, which we assume common across all the tasks in Θ .² Instead, each task may have their own transition model $p_\theta : \mathcal{X} \times \mathcal{A} \rightarrow \Delta(\mathcal{X})$ and reward function $r_\theta : \mathcal{X} \times \mathcal{A} \rightarrow [0, 1]$. A history-based randomized policy is a sequence of functions $\pi := \{\pi_h : \mathcal{T}_h \rightarrow \Delta(\mathcal{A})\}_{h=0}^{H-1}$ where $\mathcal{T}_h$ is the set of h -steps trajectories $\tau^h = (x^0, a^0, r^0, \dots, x^h)$ and $\mathcal{T} = \cup_{h=0}^{H-1} \mathcal{T}_h$. A policy π on the MDP $\mathcal{M}_\theta$ induces a distribution $\mathbb{P}_\theta^\pi$ over trajectories with the following process. An initial observation is sampled $x^0 \sim p_\theta(\cdot)$. Then, for every step $h \geq 0$, an action is sampled from the policy $a^h \sim \pi(\tau^h)$, the reward $r^h = r_\theta(x^h, a^h)$ is collected, and the MDP emits the next observation $x^{h+1} \sim p_\theta(x^h, a^h)$. The process goes on until the step H is reached.³ The Reinforcement Learning (RL [47]) objective for an MDP $\mathcal{M}_\theta$ is the cumulative sum of rewards $J(\pi) := \mathbb{E}_{\mathbb{P}_\theta^\pi} [\sum_{h=0}^{H-1} r^h]$, where the sequence $(r^0, \dots, r^{H-1})$ is taken on expectation over trajectories $\tau \sim \mathbb{P}_\theta^\pi$. An optimal policy for $\mathcal{M}_\theta$ is denoted as $\pi^* \in \arg \max J(\pi)$. For some $R \in \mathbb{N}$, we assume $J(\pi^*) \leq R$, where typically $R = 1$ when rewards are sparse, as large as H when rewards are dense.

Behavioral cloning. In the setting described above, we assume to have access to a dataset of expert demonstrations $E = \{\theta_i \sim P_0, (\tau_{i1}, \dots, \tau_{in}) \sim \mathbb{P}_{\theta_i}^{\pi^E}\}_{i=1}^m$ where $\tau_{ij} = (x_{ij}^0, a_{ij}^0, r_{ij}^0, \dots, x_{ij}^H, a_{ij}^H, r_{ij}^H)$ is a H -steps trajectory sampled independently from a policy π^E in the MDP θ_i . Thus, the total number of trajectories is $|E| = mn$ and the total number of transitions is mnH . With the available data, we aim to clone the expert's behavior π^E , a problem that is known as *behavioral cloning* [41]. The idea is to train a policy $\hat{\pi}$ to mimic the expert's policy π^E by minimizing a supervised learning loss on the demonstrations. While several choice of loss functions could be made [55], here we opt for the negative log likelihood as in [13]. The behavioral cloning problem is then

$$\hat{\pi} \in \arg \min_{\pi \in \Pi} \mathcal{L}(\pi) := \sum_{i=1}^m \sum_{j=1}^n \sum_{h=0}^{H-1} \log \left(\frac{1}{\pi(a_{ij}^h | \tau_{ij}^h)} \right) \quad (1)$$

where Π is a policy space of our choice and τ_{ij}^h is h -steps chunk of the trajectories τ_{ij} in the dataset of demonstrations E , a_{ij}^h is the action taken at step h in τ_{ij} . While a sufficiently expressive policy space Π may allow for a cloned policy $\hat{\pi}$ that closely approximates the expert on the training data $\mathcal{L}(\hat{\pi}) \approx 0$, we typically aim for a policy $\hat{\pi}$ that can mimic the expert's behavior on unseen data as well. Differently from the common setting [41, 42, 56, 38, 36, 37, 13], here we are not only concerned with *generalization* across unseen observations in $\mathcal{X}$, but also across unseen tasks in Θ . Before proceeding with the study of generalization in the next section, we introduce additional notation for later use.

Additional notation. In our behavioral cloning problem (1), a single data point is given by the triplet $(\theta_i, \tau_{ij}^h, a_{ij}^h)$, which we intend as realizations from the random variables (T, X, A) distributed as $T \sim P_0$ and $(X, A) \sim \mathbb{P}_T^{\pi^E}$ respectively. We will turn to one or the other notation when convenient. For a random variable A taking values $a_1, a_2, \dots$ with probabilities $p(a_1), p(a_2), \dots$, we denote its *entropy* $H(A) = -\sum_i p(a_i) \log p(a_i)$. For two random variables A, B , we denote their *mutual information* $I_{A;B} = H(A) - H(A|B) = H(B) - H(B|A)$, where $H(A|B)$ is the conditional entropy. Finally, we will use the symbol $\lesssim$ to hide constant and lower order terms from inequalities.

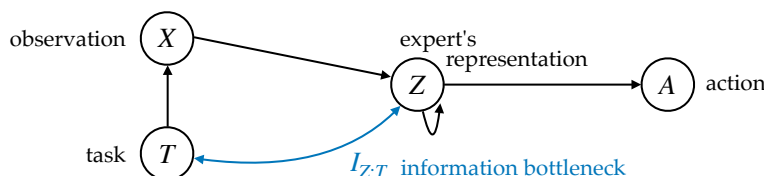
3 Generalization analysis

In the previous section, we detailed how an expert's policy can be "cloned" from data by solving the optimization problem (1). Obviously, fully cloning the expert's behavior is a far fetched objective when limited demonstrations are available: When training data spans only a small portion of the observation space $\mathcal{X}$ and the set of tasks Θ , how can we extract information on what would the expert do in unseen observations and tasks? Nonetheless, we aim for our cloned policy $\hat{\pi}$ to *transfer* at least part of the expert's behavior beyond the demonstrated observations and tasks. In this section, we

²Note that this does not hinder generality, as we can always take $\mathcal{X} = \cup_{\theta \in \Theta} \mathcal{X}_\theta$ when observation spaces vary across tasks (ditto for the action space) and $H = \max_{\theta \in \Theta} H_\theta$ when the episode horizons vary.

³Oftentimes, the episode horizon is an upper bound to the episode length, while secondary termination conditions may end the episode early, as it will be the case in our experimental setting. For the ease of presentation, we ignore early termination in our setup and consider episodes of length H .

provide a formal study of the *generalization* guarantees of the cloned policy $\hat{\pi}$, showing an original dependence with the *information* available to the expert when collecting demonstrations. To specify what do we mean by “information” in this setting, let us consider the figure below.



The latter is a graphical illustration of the expert’s behavior.⁴ At each step, the expert takes as input an observation (X) which depends on the task they are going to demonstrate (T), processing them into an internal representation Z . The demonstrated action (A) is then conditioned on Z . Note that Z is recurrent and retains any information that is relevant to select A , including the task information that may be available in the observation X . The mutual information $I_{Z;T}$ measures the task-related information that goes into Z and, consequently, how much the strategy to select A relies on it.

We typically expect an expert with full information on the task—large $I_{Z;T}$ —to demonstrate an optimal policy specific to the task, without any exploratory actions. Instead, whenever there is an information bottleneck between the task information and the expert—small $I_{Z;T}$ —we expect them to take exploratory actions to first “understand” the task in order to solve it. We still call them *experts* as we assume their behavior to be optimal *with the given information*, a concept that has been formalized with Bayes-optimal policies [16]. We conjecture that the latter behavior may generalize better to new tasks, as the process of understanding the task is more general than just solving it. In the following, we provide a formal result based on this conjecture, in which we analyze the generalization gap of the cloned policy as a function of the information $I_{Z;T}$ available to the demonstrator.

Analogously to previous works [e.g., 41, 13], we are interested in deriving an upper bound on the performance gap between an optimal policy for each task (π^*) and the cloned policy ($\hat{\pi}$). Since we are considering a multi-task setting, we average the gap across the task distribution P_0 .⁵

$$\mathbb{E}_{T \sim P_0} [J(\pi_T^*) - J(\hat{\pi})] = \mathbb{E}_{T \sim P_0} \left[\mathbb{E}_{\mathbb{P}_T^{\pi^*}} \left[\sum_{h=0}^{H-1} r^h \right] - \mathbb{E}_{\mathbb{P}_T^{\hat{\pi}}} \left[\sum_{h=0}^{H-1} r^h \right] \right]. \quad (2)$$

Before stating the result, we introduce a few technical assumptions. First, we define the *generalization error* of a policy π as

$$\mathcal{E}_{gen}(\pi) := \mathbb{E}_{T \sim P_0} \mathbb{E}_{X \sim \mathbb{P}_T^{\pi^*}} [\mathbf{1}(\pi(X) \neq A)] \quad (3)$$

for a single point indicator loss. We make the following assumptions on the expert’s policy.

Assumption 1. *The expert’s policy π^E is deterministic.*

Assumption 2. *The generalization error of the expert’s policy is given by $\mathcal{E}_{gen}(\pi^E)$.*

Note that we allow the expert’s policy to depend on the history, for which assuming determinism is reasonable even when an information bottleneck is applied to the expert. Whereas it is standard in the literature to assume the expert is optimal, i.e., $\mathcal{E}_{gen}(\pi^E) = 0$, in the presence of an information bottleneck we do not take for granted that the expert is optimal *in all the tasks*. Nonetheless, if the expert’s behavior is Bayes-optimal, non-trivial worst-case bounds on $\mathcal{E}_{gen}(\pi^E)$ hold [5], for which the generalization error only scales with $\log(H)$ under our assumptions. Moreover, in settings where the reward is sparse denoting task success, i.e., $R = 1$, the Bayes-optimal policy may still have $\mathcal{E}_{gen}(\pi^E) = 0$ w.r.t. some optimal policy, albeit inefficient in the number of steps.

Then, regarding the behavioral cloning problem (1), we make a pair of assumptions as follows.

Assumption 3. *The expert’s policy is realizable in the policy space Π , i.e., $\pi^E \in \Pi$.*

Assumption 4. *We have access to an optimization oracle that solves problem (1) with bounded error*

$$\mathcal{E}_{opt}(\hat{\pi}) := \frac{1}{mnH} \sum_{i=1}^m \sum_{j=1}^n \sum_{h=0}^{H-1} \mathbf{1}(\hat{\pi}(\tau_{ij}^h) \neq a_{ij}^h).$$

⁴Note that this is not related to the architecture of the cloned policy, which will be discussed later on.

⁵Note that the optimal policy π^* depends on the task T whereas $\hat{\pi}$ is a single policy cloned from data.

In principle, one can fulfill Asm. 3 by cloning the demonstrations into a rich enough policy space Π . However, a more expressive policy space may lead to a harder optimization problem, especially when the policies are represented through large neural networks, for which (1) is non-convex.

We now have all the ingredients to state our main result.

Theorem 3.1. *For a confidence $\delta \in (0, 1)$, it holds with probability at least $1 - 2\delta$*

$$\mathbb{E}_{T \sim P_0} [J(\pi_T^*) - J(\hat{\pi})] \lesssim RH \left(\mathcal{E}_{gen}(\pi^E) + \mathcal{E}_{opt}(\hat{\pi}) + \sqrt{\frac{I_{T;Z} |\mathcal{A}| \log(|\mathcal{A}|/\delta)}{m}} + \frac{\log(|\Pi|m/\delta)}{n} \right).$$

All of the derivations and the hidden constants can be found in Appendix A. Instead, here we unpack the bound and discuss the meaning of each term. First, the RH factor accounts for the cost of a “mistake” of the cloned policy, i.e., choosing an action different from π^* . The terms $\mathcal{E}_{gen}(\pi^E)$ and $\mathcal{E}_{opt}(\hat{\pi})$ depends on the quality of the expert’s policy and the solver for (1), hence they cannot be reduced with additional data. The last term, scaling with the number of trajectories in each demonstrated task n^{-1} , comes from a typical behavioral cloning analysis of generalization within the training task [e.g., 13]. The more demonstrations we have from a task, the better we can clone the expert’s policy in that task. The third term, scaling with the number of demonstrated tasks m^{-1} , controls the generalization across tasks and comes from the analysis of generalization induced by an information bottleneck [20], which is expressed by $I_{Z;T}$. The most important finding of our result lies in this term: To improve generalization of the cloned policy, we can either increase the number of tasks m or apply an information bottleneck—a “blindfold”—to the demonstrator to reduce $I_{Z;T}$ without paying any meaningful price, as it is typically $\mathcal{E}_{gen}(\pi^E) = 0$ for tasks with sparse rewards and $\mathcal{E}_{gen}(\pi^E) \lesssim \log(H)$ for dense rewards, while other terms remain the same.

In the next sections, we provide an extensive empirical evaluation showing that this result is far from being a theoretical fluke, but translates to practical scenarios as well.

Overcoming assumptions. The result presented in this section holds for deterministic expert’s policies (Asm. 1), finite action space (due to $|\mathcal{A}|$ dependency), and finite policy class (due to $|\Pi|$ dependency). Deterministic expert’s policy and finite policy class can be easily overcome by extending the in-task generalization result to stochastic policies and infinite policy classes, as done in [13]. Instead, the dependency on $|\mathcal{A}|$ comes from reducing the cloning problem to classification, in order to invoke information bottleneck generalization results [20]. In principle, extending the analysis to continuous action requires an analogous generalization bound for the regression problem [29] and to circumvent established negative results for this setting [46].

4 Experiments

In this section, we report an experimental campaign to validate the results of previous sections and to demonstrate the importance of the expert’s behavior for multi-task BC generalization. To this end, we train two policies, with the same BC algorithm, on human demonstrations collected by either a traditional expert or a blindfolded expert⁶. We refer to the resulting policies as π_{BC} and π_{BC-BF} respectively. We compare them in the success rate achieved on both the demonstrated tasks and unseen test tasks. We repeat the experiment twice, first in simulation on the Procgen maze and heist (Section 4.1), then on a real robot peg insertion task (Section 4.2). For both domains, we describe how the information bottleneck for the blindfolded expert is obtained in practice.

4.1 Procgen maze and heist

Procgen [8] is a popular benchmark for measuring sample efficiency and generalization [35, 34, 9, 22, 61]. It consists of 16 different procedurally generated environments in which new levels are randomly generated for every episode, forcing agents to handle changing layouts, colors, and textures. Here we focus on the maze and heist games in the “easy” setting, in which each level is a 2D maze-like layout that the agent navigates. These tasks are the most challenging tasks for generalization in the Procgen suite, and until the results of [61] have seen only minor improvements over random walk. We test how training on a set of demonstrated maze and heist tasks generalizes to unseen tasks for the different experts.

⁶Our human demonstration dataset for Procgen maze and heist is available on the project website.

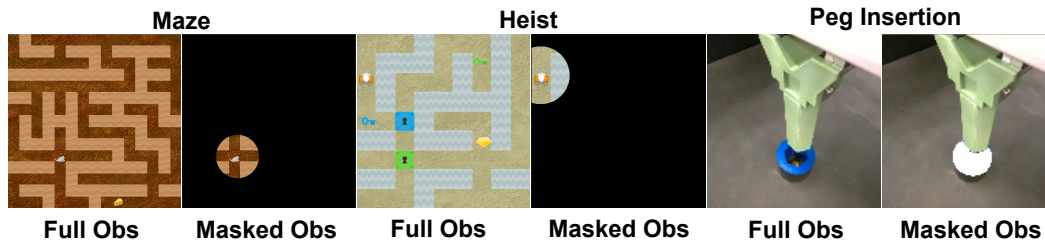


Figure 2: Demonstration of actual experts' observations. For the Procgen maze (left-pair), Procgen heist (middle-pair), and robotic peg insertion (right-pair). We show the full observation of the Expert and the masked observation of the Blindfolded-Expert.

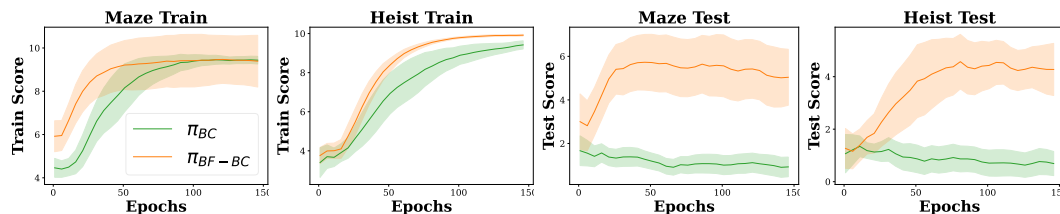


Figure 3: Game score as a function of training epochs in the Procgen maze and heist. Left-pair: performance on 200 training levels. Right-pair: performance on unseen test levels. The mean and standard deviation are computed over 10 seeds.

Maze. In Procgen maze, the agent (mouse) must find and reach a single goal (piece of cheese) in a 2D maze. Procedural level generation produces variation in maze layout (sprawling corridors), goal positions, and the color and texture of navigable spaces.

Heist. The Procgen heist is based on a maze-like level, in which the agent has to additionally solve a sequence of sub-goals in order to solve the task (reaching the gem). At each level, the goal (gem) is hidden behind a series of color-coded locks, accessible through keys of corresponding colors, that are scattered throughout the level (see Figure 2). This task requires multi-stage planning in the form of lock-waypoints and to overcome variations in sub-task ordering, making the task more complex and challenging than a regular maze.

Setup. For training, we consider 200 procedurally generated levels. There are 4 discrete **actions**, move up, down, left, and right in the maze game, and additional 4 actions in heist: left-down, left-up, right-down, and right-up. The **observations** are 64×64 RGB images of the current state as a top-down view of the maze.

Data collection. We collect a total of $4K$ demonstrated trajectories from 200 different training levels (20 trajectories per level). The environment allows for at most 500 input steps for maze and 1000 for heist, and we only retain trajectories of successful expert demonstrations. The **standard experts** are humans who see the top-view of the entire game, and therefore are likely to follow the shortest path to the goal. The **blindfolded experts** are humans playing the game with occluded observations of its layout, such that only the immediate proximity of the agent is visible and the rest of the level is concealed (see Figure 2, middle and left). As a result, human experts cannot directly plan a path to the goal location, and exploration of the level is needed. Note that while the observations are masked to the expert, the stored data contains the original (unmasked) observations for training the cloning algorithm. The **trajectories** are tuples of observation, action, reward, and done flag $(o_t, a_t, r_t, \text{done})$.

Policy architecture and training. Training is conducted from scratch on the demonstrated trajectories by minimizing the negative log likelihood 1. We use the architecture from [27] for both the Expert (π_{BC}) and BF-Expert (π_{BF-BC})—a ResNet [17] to encode the observations, which are then processed by two fully-connected layers. To capture the exploratory behavior, we add a single GRU [7] before the Softmax policy layer (further details are in Appendix C).

Table 1: Average number of steps $\pm$ std in the trajectories demonstrated by standard Experts and BF-Experts. The latter takes more steps due to the information bottleneck, which forces exploration.

Mode	⬢	★	■	+	◄	Maze
Experts	61.2 ± 1.4	69.4 ± 3.3	64.8 ± 2.1	96.6 ± 4.9	68.8 ± 2.2	28.6 ± 1.2
BF-Experts	74.3 ± 5.9	165.6 ± 42.9	163.6 ± 43.6	148.6 ± 33.1	194.4 ± 55.0	51.6 ± 1.8

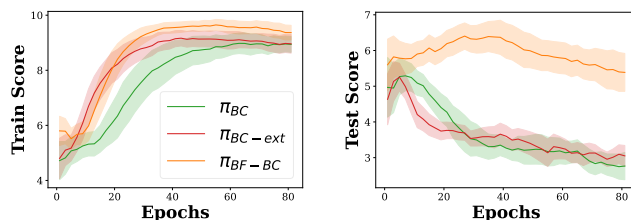


Figure 4: Procgen maze score as a function of training epochs. Left: performance on 100 training levels. Right: performance on unseen test levels. Policies π_{BC} and π_{BF-BC} are trained on $2K$ trajectories, the former clearly overfits. We also train π_{BC-ext} on $4.6K$ non-blindfolded experts’ trajectories, matching the total number of steps of π_{BF-BC} . The resulting π_{BC-ext} overfits despite access to twice the amount of data. The mean and standard deviation are computed over 10 seeds.

Results. In Figure 3, we compare the game score over the training epochs achieved by policy cloning of the standard Expert (π_{BC}) and the BF-Expert (π_{BC-BF}). As we can see from Figure 3 (left-pair), the performance of the two policies is similar on the training levels, approaching the maximum task score. However, as evident from Figure 3 (right-pair), while the π_{BC-BF} gracefully improves the test score, the π_{BC} policy overfits to the training-set, and its test score slightly degrades with further training. This is a testament to the inherent generalization capabilities granted by the BF-Expert w.r.t. a standard Expert.

Number of steps. Table 1 (left-most column) details the average number of steps taken by the experts (Experts) and the blindfolded experts (BF-Experts) across 100 demonstrated levels, accruing $2K$ maze trajectories. Trajectories taken by the Experts are shorter on average than the BF-Experts, supporting the assumption that the BF-Experts exhibit a more exploratory behavior, whereas the Experts take the shortest path to the goal.

To further highlight the contribution of exploration toward generalization and to show that it isn’t merely the result of additional training steps, we train a policy π_{BC-ext} on double the amount of non-blindfolded experts’ trajectories—to match the total number of steps produced by the blindfolded experts. Figure 4 shows that even with an equal number of total steps from both experts, π_{BC-ext} overfits. Further details on the number of trajectories and steps are provided in Appendix C.

4.2 Robotic peg insertion

Peg insertion is a standard problem in robotic manipulation. Here, we consider the insertion task in the Functional Manipulation Benchmark [FMB, 25], which focuses on inserting variously shaped pegs into tightly matching holes. Different from [25], however, we investigate *generalization*: How training on a fixed set of shapes generalizes to inserting previously unseen shapes. We simplify less relevant technical aspects of the benchmark by fixing the peg to the robot gripper, and 3D-printing individual holes, so that discerning the target hole becomes trivial (see Figure 5-left). In addition, we added several new shapes to the benchmark to increase its variations.

Setup. The task comprises ten pegs of various shapes with corresponding slots. In our setting, the robot initiates with the peg already in its grip and needs to insert it into a single-slotted board anchored to the surface. The **robot setup** is shown in Figure 5 in full, alongside a few examples of peg shapes. We use a Franka Emika Panda robot arm and teleoperate the robot using a SpaceMouse. For operating the robot and the SpaceMouse, we use the SERL open-sourced package [24] with the same settings as the authors. The **actions** are input as SpaceMouse commands: 6-DoF end-effector move and twist (location and Euler angles) at 10Hz, tracked by a low-level impedance controller running at 1KHz. The **observations** are obtained as RGB-only images from two Intel RealSense

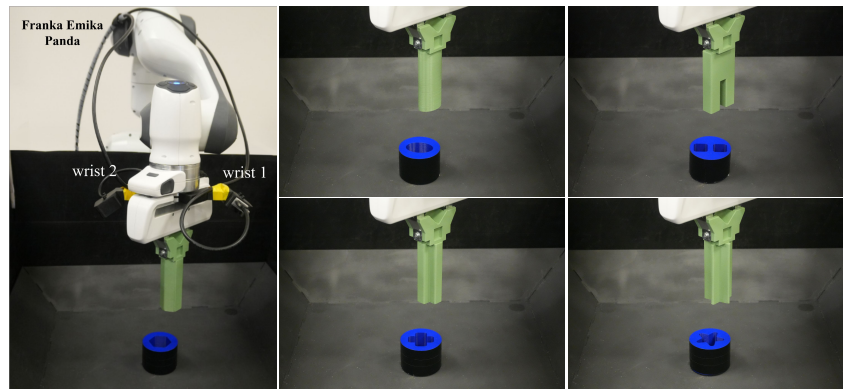


Figure 5: The robotic arm configured for peg insertion (left). Close-up view of various peg shapes and their boards (right).

D405 cameras, mounted on the robot end-effector, which simultaneously capture images⁷. To avoid background distractions in image observations, we place each shaped board inside a black bin. In addition to the 6-DoF pose of the robot end-effector, the framework also includes the force, torque, and velocity information provided by the Franka Panda robot.

Data collection. We split the ten shapes into training and test shapes. When collecting demonstrations, only the training shapes are used, while the test shapes serve as a withheld subset for evaluation. We collect 400 trajectories from human demonstrations of each training shape, first with a standard Experts and then with BF-Experts. Each trajectory begins at a random initial pose. When the full information from the wrist cameras is provided to the human experts, they easily succeed in inserting all training shapes. The **blindfolded experts** are human experts exposed to redacted observations from the wrist cameras, which occlude the articulation of the slot as they attempt to insert the peg (see Figure 2). We use SegmentAnything2 (SAM2) [39] modified to segment a live video stream and prompt it to mask out the shape of the target hole. In addition to the masked-out images from the wrist cameras, the initial robot pose varies with each insertion attempt, thus preventing the expert from memorizing or inferring the articulation of the target hole. Note, however, that the recorded trajectories collected by both the Expert and BF-Expert contain the full unmasked observation. The distinction is then the behavior of the two experts, with blindfolded experts taking exploratory actions to cope with masked-out images in an attempt to complete the task. The **trajectories** are collected on the training shapes only. Table 1 details the training shapes and the average number of steps taken by the standard Expert and BF-Expert until successful peg insertion. Clearly, the trajectories demonstrated by the BF-Expert are longer, indicating more exploratory behavior, whereby the expert must rely on masked-out images until resolving the correct articulation for inserting the peg. In the next sections, we show that the resulting exploratory behavior is useful for generalization.

Policy architecture and training. Training was conducted for $k \in \{2, 3, 4, 5\}$ peg shapes, with the remaining shapes serving as a withheld test set. We use the same network architecture for cloning both the expert π_{BC} and the blindfolded expert π_{BF-BC} . Specifically, we use a weight-shared frozen ResNet-10 encoder [17] pretrained on the ImageNet dataset [10] for encoding the incoming images from both wrist cameras. The resulting embeddings are concatenated with the MLP-embedding of the proprioceptive information before entering a single GRU [7] that outputs a Gaussian policy. The use of a memory-based architecture is crucial to fully capture the non-Markovian exploratory behavior of the blindfolded expert [28]. For more detailed specifications regarding our experimental setup and hyperparameters, please refer to Appendix B.

Results. Figure 6 shows the success rate for $k = 2, 5$ training shapes (results for $k = 3, 4$ are in Appendix B). The success rate is the average of 24 insertion attempts per peg shape by the robotic arm. The results demonstrate that cloning the BF-Expert with π_{BF-BC} achieves better generalization compared with the standard Expert, cloned with π_{BC} , across all peg shapes and over all test subsets. Importantly, the advantage of the proposed blindfolding approach is more significant when fewer

⁷Following the conclusions of FMB [25], we omit the depth data, as they showed it has a marginal benefit.

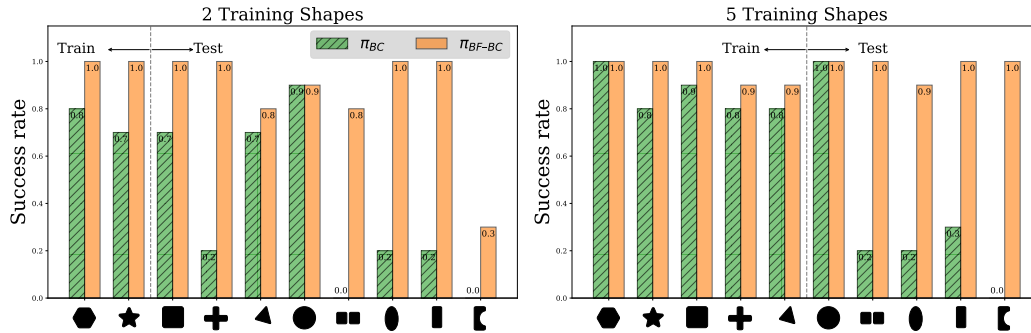


Figure 6: Success rate of robotic peg insertion for 10 peg shapes (horizontal axis). We train on a subset of shapes $k = \{2, 5\}$, with the remaining shapes withheld as test-set. Both for $k = 2$ (left) and $k = 5$ (right), cloning blindfolded experts generalizes better than cloning standard experts.

shapes are used to train the model, i.e., when a larger portion of shapes are withheld during training. Interestingly, the figure shows that even for shapes encountered in training, the π_{BF-BC} policy generalizes better than cloning the standard expert π_{BC} . This is a result of the limited ability of expert demonstrations to account for all real-world variations, e.g., lighting and control-loop errors. In comparison, the π_{BF-BC} is more robust to these kinds of errors, since it imitates blindfolded experts who cannot rely on visual cues, and must compensate in order to solve the task. Additionally, it is worth mentioning that certain shapes are more challenging than others. We observe that shapes with greater radial symmetry and convexity (such as the circle, triangle, hexagon) are simpler to learn than shapes with non-radial symmetry (such as the ellipse, rectangle) or non-convex shapes (such as the cut-out-rectangle), as they require a more specific orientation for insertion.

The sensitivity to the choice of a specific blindfold. To demonstrate the generality of our approach, we conduct an experiment of robotic peg insertion, where we apply random noise to the experts' observations instead of masking the shape socket. We collected a total of 1600 expert trajectories with random noise $p \sim U[0, P]$ added to all pixel values, from varying maximum noise level $P \in [0, 100, 170, 200]$ and 2 training shapes (square and star). We then train a network on each noise level P (i.e., 400 trajectories per level). Table 2 details the Success Rate (SR) [%] and State Entropy (SE) of the demonstrated trajectories, alongside the networks' performance on train (square and star) and test shapes (plus and ellipse). The table shows several interesting results:

1. Introducing more noise increases the State Entropy (SE), but a severe noise corruption ($P = 200$) impairs the expert, lowering the Success Rate (SR) of the demonstrations, and degrading the robot's performance.
2. Adding the "right" amount of noise ($P = 100$) helps generalize to unseen shapes. Although less effective than masking, it still significantly outperforms the traditional approach, i.e., without any blindfolding (0 noise).
3. This experiment demonstrates that even a general form of blindfold (non-task specific) may still be superior to the conventional imitation learning approach.

Table 2: Success Rate (SR) [%] and State Entropy (SE) for the various blindfold types: range of noise levels or SAM2 mask. The column "Demos" details the SR and SE of experts' demonstration subject to the blindfold. The left-most columns are their corresponding evaluation performance [%] on Train (square, star) and Test (plus, ellipse).

Max-Noise Level	Demos		Train		Test	
	SR	SE	square	star	plus	ellipse
0	100%	3.17	71%	67%	20%	17%
100	96%	3.26	100%	88%	96%	79%
170	95%	3.39	92%	71%	58%	62%
200	58%	3.46	46%	38%	21%	54%
SAM2 Mask	100%	3.52	96%	96%	100%	96%

5 Related works

Our work closely relates to imitation learning, information bottleneck, and robotic manipulation.

Imitation learning theory. The imitation learning literature counts a plethora of contributions, for which we refer to recent surveys [19, 60]. Here we are concerned with the generalization of behavioral cloning – a dominant imitation learning technique. Previous works [41, 42, 56, 38, 36, 37, 50, 13] have studied the theoretical limits of behavior cloning and settled the generalization gap on the training task as $J(\pi^*) - J(\hat{\pi}) \lesssim R \log(|\Pi|/\delta)/n$ [13]. These results are mostly limited to cloning a Markovian policy in a single MDP. Other works have considered more general settings, including cloning history-based policies [4] and imitation learning under partial observability [48]. However, the generalization of behavior cloning in a meta learning setup, which is popular in empirical works [11, 12], is understudied. The only generalization analysis that strikes close to this setting is, to the best of our knowledge, the one in [40]. Differently from ours, they assume online access to the set of tasks to further fine-tune the cloned policy and they do not study how the information available to the demonstrator affects generalization, which is our main theoretical contribution.

Imitation learning in robotics. Imitation learning has been fundamental to various robotic domains, including autonomous driving [32], locomotion [31], flight [1], and manipulation [12]. A recent survey on manipulation, our case of interest here, is provided in [2]. Several works showed that imitation is useful for learning complex visuo-motor policies for robotic manipulation, where key ideas include predicting a sequence of future actions, and using a diffusion generative model to learn a distribution over the action sequence [6, 59, 14]. However, it is known that imitation learning requires a large number of demonstrations in order to generalize to variations in the task. Previously studied mitigations include 3-dimensional priors in the representation [58], automatic data augmentation [26, 15, 57], interactive data collection [18], and using simulations [52]. Another approach is leveraging large-scale data, either by collecting diverse task variations [23], or by fine tuning robotics foundations models [30, 49, 21]. Differently from the approaches above, we postulate that the way a demonstrator performs a task affects generalization, showing that by blindfolding the demonstrator we obtain an exploratory behavior, which induces better generalization when cloned. The generalization of exploratory behavior has been demonstrated in zero-shot reinforcement learning [61]. In comparison, we apply this idea to imitation learning, which requires a different approach, and also develop a theoretical explanation for the improved generalization. The idea that exploration at test time helps generalization has also been explored in the sim-to-real context in [54].

Information bottleneck. When learning a $X \rightarrow Y$ relationship between random variables, the information bottleneck [51] prescribes to “squeeze” X into a representation that only retains information to predict Y . This principle is believed to be a factor beyond the generalization capabilities of deep learning [44] and formal generalization bounds through the information bottleneck have been derived [45, 29, 20]. In imitation learning, the information bottleneck has been used to analyze generalization in [3]. However, they consider generalization on the training task only and the information bottleneck is applied to the representation of the cloned policy. Our work advocates for applying an information bottleneck to the demonstrator to improve generalization of the cloned policy.

6 Conclusion

We showed that cloning the behavior of blindfolded experts leads to better generalization to unseen tasks. We supported this with theoretical analysis and conducted empirical tests that, for the first time, explored the concept of blindfolding experts in the context of a real-world robotic task, as well as a maze videogame. We observed that in both peg insertion and maze-solving tasks, blindfolding the experts encouraged them to enact a more exploratory behavior, cloned to produce policies that better generalize. Importantly, our approach achieves better generalization while accommodating any imitation learning algorithm.

Finally, we point out a limitation of the proposed approach: each domain may require a different kind of blindfold (e.g., concealing the field of view in the maze videogame, or masking out the shape of the hole in the peg insertion task). Too little obstruction does not elicit exploration, while redacting too aggressively would impede any informative exploratory behavior (may resort to near random walk). An interesting question is how to find the optimal balance, which we reserve for future research.

Acknowledgments

This research was Funded by the European Union (ERC, Bayes-RL, 101041250). Views and opinions expressed are however those of the author(s) only and do not necessarily reflect those of the European Union or the European Research Council Executive Agency (ERCEA). Neither the European Union nor the granting authority can be held responsible for them.

References

- [1] Pieter Abbeel, Adam Coates, and Andrew Y Ng. Autonomous helicopter aerobatics through apprenticeship learning. *The International Journal of Robotics Research*, 29(13):1608–1639, 2010.
- [2] Shan An, Ziyu Meng, Chao Tang, Yuning Zhou, Tengyu Liu, Fangqiang Ding, Shufang Zhang, Yao Mu, Ran Song, Wei Zhang, et al. Dexterous manipulation through imitation learning: A survey. *arXiv preprint arXiv:2504.03515*, 2025.
- [3] Shuanghai Bai, Wanqi Zhou, Pengxiang Ding, Wei Zhao, Donglin Wang, and Badong Chen. Rethinking latent representations in behavior cloning: An information bottleneck approach for robot manipulation. *arXiv preprint arXiv:2502.02853*, 2025.
- [4] Adam Block, Daniel Pfrommer, and Max Simchowitz. On the imitation of non-markovian demonstrations: From low-level stability to high-level planning. In *ICML Workshop on New Frontiers in Learning, Control, and Dynamical Systems*, 2023.
- [5] Xiaoyu Chen, Jiachen Hu, Chi Jin, Lihong Li, and Liwei Wang. Understanding domain randomization for sim-to-real transfer. In *International Conference on Learning Representations*, 2022.
- [6] Cheng Chi, Zhenjia Xu, Siyuan Feng, Eric Cousineau, Yilun Du, Benjamin Burchfiel, Russ Tedrake, and Shuran Song. Diffusion policy: Visuomotor policy learning via action diffusion. *The International Journal of Robotics Research*, page 02783649241273668, 2023.
- [7] Kyunghyun Cho, Bart Van Merriënboer, Dzmitry Bahdanau, and Yoshua Bengio. On the properties of neural machine translation: Encoder-decoder approaches. *arXiv preprint arXiv:1409.1259*, 2014.
- [8] Karl Cobbe, Chris Hesse, Jacob Hilton, and John Schulman. Leveraging procedural generation to benchmark reinforcement learning. In *International conference on machine learning*, pages 2048–2056. PMLR, 2020.
- [9] Karl W Cobbe, Jacob Hilton, Oleg Klimov, and John Schulman. Phasic policy gradient. In *International Conference on Machine Learning*, pages 2020–2027. PMLR, 2021.
- [10] Jia Deng, Wei Dong, Richard Socher, Li-Jia Li, Kai Li, and Li Fei-Fei. Imagenet: A large-scale hierarchical image database. In *2009 IEEE conference on computer vision and pattern recognition*, pages 248–255. Ieee, 2009.
- [11] Yan Duan, Marcin Andrychowicz, Bradly Stadie, Jonathan Ho, Jonas Schneider, Ilya Sutskever, Pieter Abbeel, and Wojciech Zaremba. One-shot imitation learning. *Advances in neural information processing systems*, 2017.
- [12] Chelsea Finn, Tianhe Yu, Tianhao Zhang, Pieter Abbeel, and Sergey Levine. One-shot visual imitation learning via meta-learning. In *Conference on robot learning*, pages 357–368. PMLR, 2017.
- [13] Dylan J Foster, Adam Block, and Dipendra Misra. Is behavior cloning all you need? understanding horizon in imitation learning. *Advances in Neural Information Processing Systems*, 2024.
- [14] Zipeng Fu, Tony Z Zhao, and Chelsea Finn. Mobile aloha: Learning bimanual mobile manipulation with low-cost whole-body teleoperation. *arXiv preprint arXiv:2401.02117*, 2024.

- [15] Caelan Garrett, Ajay Mandlekar, Bowen Wen, and Dieter Fox. Skillmimicgen: Automated demonstration generation for efficient skill learning and deployment. *arXiv preprint arXiv:2410.18907*, 2024.
- [16] Mohammad Ghavamzadeh, Shie Mannor, Joelle Pineau, Aviv Tamar, et al. Bayesian reinforcement learning: A survey. *Foundations and Trends® in Machine Learning*, 8(5-6):359–483, 2015.
- [17] Kaiming He, Xiangyu Zhang, Shaoqing Ren, and Jian Sun. Deep residual learning for image recognition. In *Proceedings of the IEEE conference on computer vision and pattern recognition*, pages 770–778, 2016.
- [18] Ryan Hoque, Ajay Mandlekar, Caelan Garrett, Ken Goldberg, and Dieter Fox. Intervengen: Interventional data generation for robust and data-efficient robot imitation learning. In *2024 IEEE/RSJ International Conference on Intelligent Robots and Systems (IROS)*, pages 2840–2846. IEEE, 2024.
- [19] Ahmed Hussein, Mohamed Medhat Gaber, Eyad Elyan, and Chrisina Jayne. Imitation learning: A survey of learning methods. *ACM Computing Surveys*, 50(2):1–35, 2017.
- [20] Kenji Kawaguchi, Zhun Deng, Xu Ji, and Jiaoyang Huang. How does information bottleneck help deep learning? In *International Conference on Machine Learning*, 2023.
- [21] Moo Jin Kim, Karl Pertsch, Siddharth Karamcheti, Ted Xiao, Ashwin Balakrishna, Suraj Nair, Rafael Rafailov, Ethan Foster, Grace Lam, Pannag Sanketi, et al. Openvla: An open-source vision-language-action model. *arXiv preprint arXiv:2406.09246*, 2024.
- [22] Robert Kirk, Amy Zhang, Edward Grefenstette, and Tim Rocktäschel. A survey of generalisation in deep reinforcement learning. *arXiv preprint arXiv:2111.09794*, 1:16, 2021.
- [23] Fanqi Lin, Yingdong Hu, Pingyue Sheng, Chuan Wen, Jiacheng You, and Yang Gao. Data scaling laws in imitation learning for robotic manipulation. *arXiv preprint arXiv:2410.18647*, 2024.
- [24] Jianlan Luo, Zheyuan Hu, Charles Xu, You Liang Tan, Jacob Berg, Archit Sharma, Stefan Schaal, Chelsea Finn, Abhishek Gupta, and Sergey Levine. Serl: A software suite for sample-efficient robotic reinforcement learning. In *2024 IEEE International Conference on Robotics and Automation (ICRA)*, pages 16961–16969. IEEE, 2024.
- [25] Jianlan Luo, Charles Xu, Fangchen Liu, Liam Tan, Zipeng Lin, Jeffrey Wu, Pieter Abbeel, and Sergey Levine. FMB: a functional manipulation benchmark for generalizable robotic learning. *The International Journal of Robotics Research*, page 02783649241276017, 2023.
- [26] Ajay Mandlekar, Soroush Nasiriany, Bowen Wen, Iretiayo Akinola, Yashraj Narang, Linxi Fan, Yuke Zhu, and Dieter Fox. Mimicgen: A data generation system for scalable robot learning using human demonstrations. *arXiv preprint arXiv:2310.17596*, 2023.
- [27] Ishita Mediratta, Qingfei You, Minqi Jiang, and Roberta Raileanu. The generalization gap in offline reinforcement learning. *arXiv preprint arXiv:2312.05742*, 2023.
- [28] Mirco Mutti, Riccardo De Santi, and Marcello Restelli. The importance of non-markovianity in maximum state entropy exploration. In *International Conference on Machine Learning*, pages 16223–16239. PMLR, 2022.
- [29] Vudtiwat Ngampruetikorn and David J Schwab. Information bottleneck theory of high-dimensional regression: relevancy, efficiency and optimality. *Advances in Neural Information Processing Systems*, 2022.
- [30] Abby O’Neill, Abdul Rehman, Abhiram Maddukuri, Abhishek Gupta, Abhishek Padalkar, Abraham Lee, Acorn Pooley, Agrim Gupta, Ajay Mandlekar, Ajinkya Jain, et al. Open x-embodiment: Robotic learning datasets and rt-x models: Open x-embodiment collaboration. In *2024 IEEE International Conference on Robotics and Automation (ICRA)*, pages 6892–6903. IEEE, 2024.

- [31] Xue Bin Peng, Pieter Abbeel, Sergey Levine, and Michiel Van de Panne. Deepmimic: Example-guided deep reinforcement learning of physics-based character skills. *ACM Transactions On Graphics (TOG)*, 37(4):1–14, 2018.
- [32] Dean A Pomerleau. Alvin: An autonomous land vehicle in a neural network. *Advances in neural information processing systems*, 1, 1988.
- [33] Martin L Puterman. *Markov decision processes: discrete stochastic dynamic programming*. John Wiley & Sons, 2014.
- [34] Roberta Raileanu and Rob Fergus. Decoupling value and policy for generalization in reinforcement learning. In *International Conference on Machine Learning*, pages 8787–8798. PMLR, 2021.
- [35] Roberta Raileanu, Maxwell Goldstein, Denis Yarats, Ilya Kostrikov, and Rob Fergus. Automatic data augmentation for generalization in reinforcement learning. *Advances in Neural Information Processing Systems*, 34:5402–5415, 2021.
- [36] Nived Rajaraman, Yanjun Han, Lin Yang, Jingbo Liu, Jiantao Jiao, and Kannan Ramchandran. On the value of interaction and function approximation in imitation learning. *Advances in Neural Information Processing Systems*, 2021.
- [37] Nived Rajaraman, Yanjun Han, Lin F Yang, Kannan Ramchandran, and Jiantao Jiao. Provably breaking the quadratic error compounding barrier in imitation learning, optimally. *arXiv preprint arXiv:2102.12948*, 2021.
- [38] Nived Rajaraman, Lin Yang, Jiantao Jiao, and Kannan Ramchandran. Toward the fundamental limits of imitation learning. *Advances in Neural Information Processing Systems*, 2020.
- [39] Nikhila Ravi, Valentin Gabeur, Yuan-Ting Hu, Ronghang Hu, Chaitanya Ryali, Tengyu Ma, Haitham Khedr, Roman Rädle, Chloe Rolland, Laura Gustafson, et al. Sam 2: Segment anything in images and videos. *arXiv preprint arXiv:2408.00714*, 2024.
- [40] Allen Ren, Sushant Veer, and Anirudha Majumdar. Generalization guarantees for imitation learning. In *Conference on Robot Learning*, 2021.
- [41] Stéphane Ross and Drew Bagnell. Efficient reductions for imitation learning. In *International Conference on Artificial Intelligence and Statistics*, 2010.
- [42] Stéphane Ross and J Andrew Bagnell. Reinforcement and imitation learning via interactive no-regret learning. *arXiv preprint arXiv:1406.5979*, 2014.
- [43] Stéphane Ross, Geoffrey Gordon, and Drew Bagnell. 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*, pages 627–635. JMLR Workshop and Conference Proceedings, 2011.
- [44] Ohad Shamir, Sivan Sabato, and Naftali Tishby. Learning and generalization with the information bottleneck. *Theoretical Computer Science*, 411(29-30):2696–2711, 2010.
- [45] Ravid Shwartz-Ziv, Amichai Painsky, and Naftali Tishby. Representation compression and generalization in deep neural networks. *OpenReview*, 2018.
- [46] Max Simchowitz, Daniel Pfrommer, and Ali Jadbabaie. The pitfalls of imitation learning when actions are continuous. *arXiv preprint arXiv:2503.09722*, 2025.
- [47] Richard S Sutton, Andrew G Barto, et al. *Reinforcement learning: An introduction*. MIT press Cambridge, 1998.
- [48] Gokul Swamy, Sanjiban Choudhury, J Bagnell, and Steven Z Wu. Sequence model imitation learning with unobserved contexts. *Advances in Neural Information Processing Systems*, 2022.
- [49] Octo Model Team, Dibya Ghosh, Homer Walke, Karl Pertsch, Kevin Black, Oier Mees, Sudeep Dasari, Joey Hejna, Tobias Kreiman, Charles Xu, et al. Octo: An open-source generalist robot policy. *arXiv preprint arXiv:2405.12213*, 2024.

- [50] Daniil Tiapkin, Denis Belomestny, Daniele Calandriello, Eric Moulines, Alexey Naumov, Pierre Perrault, Michal Valko, and Pierre Menard. Demonstration-regularized rl. In *International Conference on Learning Representations*, 2024.
- [51] Naftali Tishby, Fernando C Pereira, and William Bialek. The information bottleneck method. *arXiv preprint physics/0004057*, 1999.
- [52] Marcel Torne, Anthony Simeonov, Zechu Li, April Chan, Tao Chen, Abhishek Gupta, and Pulkit Agrawal. Reconciling reality through simulation: A real-to-sim-to-real approach for robust manipulation. *arXiv preprint arXiv:2403.03949*, 2024.
- [53] Ashish Vaswani, Noam Shazeer, Niki Parmar, Jakob Uszkoreit, Llion Jones, Aidan N Gomez, Łukasz Kaiser, and Illia Polosukhin. Attention is all you need. *Advances in neural information processing systems*, 30, 2017.
- [54] Andrew Wagenmaker, Kevin Huang, Liyiming Ke, Kevin Jamieson, and Abhishek Gupta. Overcoming the sim-to-real gap: Leveraging simulation to learn to explore for real-world rl. *Advances in Neural Information Processing Systems*, 37:78715–78765, 2024.
- [55] Qi Wang, Yue Ma, Kun Zhao, and Yingjie Tian. A comprehensive survey of loss functions in machine learning. *Annals of Data Science*, 9(2):187–212, 2022.
- [56] Tian Xu, Ziniu Li, and Yang Yu. Error bounds of imitating policies and environments. *Advances in Neural Information Processing Systems*, 2020.
- [57] Zhengrong Xue, Shuying Deng, Zhenyang Chen, Yixuan Wang, Zhecheng Yuan, and Huazhe Xu. Demogen: Synthetic demonstration generation for data-efficient visuomotor policy learning. *arXiv preprint arXiv:2502.16932*, 2025.
- [58] Yanjie Ze, Gu Zhang, Kangning Zhang, Chenyuan Hu, Muhan Wang, and Huazhe Xu. 3d diffusion policy: Generalizable visuomotor policy learning via simple 3d representations. *arXiv preprint arXiv:2403.03954*, 2024.
- [59] Tony Z Zhao, Vikash Kumar, Sergey Levine, and Chelsea Finn. Learning fine-grained bimanual manipulation with low-cost hardware. In *Robots: Science and Systems (RSS)*, 2023.
- [60] Boyuan Zheng, Sunny Verma, Jianlong Zhou, Ivor W Tsang, and Fang Chen. Imitation learning: Progress, taxonomies and challenges. *IEEE Transactions on Neural Networks and Learning Systems*, 35(5):6322–6337, 2022.
- [61] Ev Zisselman, Itai Lavie, Daniel Soudry, and Aviv Tamar. Explore to generalize in zero-shot rl. *Advances in Neural Information Processing Systems*, 2023.

A Proofs

Here we provide the derivations for the generalization bound in Section 3. The proof requires a non-trivial combination of previous results in provable imitation learning, mostly from [13], and the generalization guarantees of information bottleneck [20]. These results are reported in Lemma A.3 and A.4. Before going through the proofs, we state the main theorem again for convenience.

Theorem A.1 (Theorem 3.1). *For a confidence $\delta \in (0, 1)$, it holds with probability at least $1 - 2\delta$*

$$\begin{aligned} & \mathbb{E}_{T \sim P_0} [J(\pi_T^*) - J(\hat{\pi})] \\ & \lesssim RH \left(\mathcal{E}_{gen}(\pi^E) + \mathcal{E}_{opt}(\hat{\pi}) + C(E, \pi^E, \hat{\pi}) \sqrt{\frac{I_{T;Z} |\mathcal{A}| \log(|\mathcal{A}|/\delta)}{m}} + \frac{8 \log(|\Pi|m/\delta)}{n} \right) \end{aligned}$$

where

- R is an upper bound to the cumulative reward of any policy in any task $\theta \in \Theta$ (Section 2);
- $\mathcal{E}_{gen}(\pi^E)$ is the generalization error of the expert's policy (Asm. 2);
- $\mathcal{E}_{opt}(\hat{\pi})$ is a bound to the optimization error of solving (1) (Asm. 4);
- $C(E, \pi^E, \hat{\pi})$ is a constant that depends on the training data E , the expert's policy π^E , the cloned policy $\hat{\pi}$, and other absolute constants as detailed in Lemma A.4.

Proof. We derive the result as follows

$$\begin{aligned} & \mathbb{E}_{T \sim P_0} [J(\pi_T^*) - J(\hat{\pi})] \\ & \leq RH \mathbb{E}_{T \sim P_0} \mathbb{E}_{\tau \sim \mathbb{P}_T^{\pi^*}} [\mathbf{1}(\pi_T^*(a^h|x^h) \neq \hat{\pi}(a^h|x^h))] \end{aligned} \quad (4)$$

$$\leq RH \mathcal{E}_{gen}(\hat{\pi}) \quad (5)$$

$$\leq RH \left(\mathcal{E}_{gen}(\pi^E) + \mathbb{E}_{T \sim P_0} \mathbb{E}_{X A \sim \mathbb{P}_T^{\pi^E}} [\mathbf{1}(\hat{\pi}(X) \neq A)] \right) \quad (6)$$

$$\lesssim RH \left(\mathcal{E}_{gen}(\pi^E) + \mathcal{E}_{opt}(\hat{\pi}) + \sqrt{\frac{I_{T;Z} |\mathcal{A}| \log(|\mathcal{A}|/\delta)}{m}} + \frac{8 \log(|\Pi|m/\delta)}{n} \right) \quad (7)$$

where (4) and (5) are straightforward from the definitions of the performance $J(\pi)$ and the generalization error $\mathcal{E}_{gen}(\pi)$ (see Section 2 and (3) respectively), (6) follows from Assumption 2, and (7) holds with probability at least $1 - 2\delta$ through Lemma A.2. $\square$

We provide below the lemmas we need to prove the result above.

Lemma A.2. *For a confidence $\delta \in (0, 1)$, it holds with probability at least $1 - 2\delta$*

$$\mathbb{E}_{T \sim P_0} \mathbb{E}_{X A \sim \mathbb{P}_T^{\pi^E}} [\mathbf{1}(\hat{\pi}(X) \neq A)] \lesssim \sqrt{\frac{I_{T;Z} |\mathcal{A}| \log(|\mathcal{A}|/\delta)}{m}} + \frac{8 \log(|\Pi|m/\delta)}{n} + \mathcal{E}_{opt}(\hat{\pi})$$

where $I_{T;Z}$ is the mutual information between the task T and the internal representation of the demonstrator Z .

Proof. We derive the result as follows

$$\begin{aligned} & \mathbb{E}_{T \sim P_0} \mathbb{E}_{X A \sim \mathbb{P}_T^{\pi^E}} [\mathbf{1}(\hat{\pi}(X) \neq A)] \\ & \leq \mathbb{E}_{T \sim P_0} \mathbb{E}_{X A \sim \mathbb{P}_T^{\pi^E}} [\mathbf{1}(\hat{\pi}(X) \neq A)] - \frac{1}{mnH} \sum_{i=1}^m \sum_{j=1}^n \sum_{h=0}^{H-1} \mathbf{1}(\hat{\pi}(\tau_{ij}^h) \neq a_{ij}^h) + \mathcal{E}_{opt}(\hat{\pi}) \quad (8) \\ & \leq \mathbb{E}_{T \sim P_0} \mathbb{E}_{X A \sim \mathbb{P}_T^{\pi^E}} [\mathbf{1}(\hat{\pi}(X) \neq A)] - \frac{1}{m} \sum_{i=1}^m \mathbb{E}_{X A \sim \mathbb{P}_{\theta_i}^{\pi^E}} [\mathbf{1}(\hat{\pi}(X) \neq A)] + \frac{8 \log(|\Pi|m/\delta)}{n} + \mathcal{E}_{opt}(\hat{\pi}) \quad (9) \end{aligned}$$

$$\lesssim \sqrt{\frac{I_{T;Z}|\mathcal{A}|\log(|\mathcal{A}|/\delta)}{m}} + \frac{8\log(|\Pi|m/\delta)}{n} + \mathcal{E}_{opt}(\hat{\pi}) \quad (10)$$

where (8) is a trivial consequence of Assumption 4 on the optimization error for solving (1), (9) holds with probability at least $1 - \delta$ through Lemma A.3 and a union bound on the m training tasks, and (10) holds with probability $1 - 2\delta$ from Lemma A.4 by omitting constant and lower order terms and applying a union bound. $\square$

Lemma A.3 (Sample complexity of behavioral cloning [13]). *For a confidence $\delta \in (0, 1)$, an MDP θ , and a deterministic expert's policy π^E , it holds with probability at least $1 - \delta$*

$$\mathbb{E}_{XA \sim \mathbb{P}_{\theta}^{\pi^E}} [\mathbf{1}(\hat{\pi}(X) \neq A)] \leq \frac{8\log(|\Pi|/\delta)}{n}.$$

Proof. This result can be obtained through a combination of results in [13]. First, for a policy $\hat{\pi}$ obtained by minimizing the negative log likelihood of the data, as in (1), from Proposition 2.1 [13] we have with probability at least $1 - \delta$ that

$$D_H^2(\mathbb{P}_{\theta}^{\hat{\pi}}, \mathbb{P}_{\theta}^{\pi^E}) \leq \frac{2\log(|\Pi|/\delta)}{n}$$

where $D_H^2(\mathbb{P}, \mathbb{Q}) = \int (\sqrt{d\mathbb{P}} - \sqrt{d\mathbb{Q}})^2$ is the squared Hellinger distance between the probability measures $\mathbb{P}$ and $\mathbb{Q}$. Then, through Lemma F.3 [13] we have

$$\mathbb{E}_{XA \sim \mathbb{P}_{\theta}^{\pi^E}} [\mathbf{1}(\hat{\pi}(X) \neq A)] \leq 4D_H^2(\mathbb{P}_{\theta}^{\hat{\pi}}, \mathbb{P}_{\theta}^{\pi^E})$$

which concludes the proof. $\square$

Lemma A.4 (Information bottleneck generalization gap [20]). *For a dataset $E = \{\theta_i \sim P_0\}_{i=1}^m$ of m tasks and a single-point convex loss $\ell(\hat{\pi}(X), A)$, let us define the generalization gap across the prior P_0 as*

$$\Gamma(E) := \mathbb{E}_{T \sim P_0} \mathbb{E}_{XA \sim \mathbb{P}_T^{\pi^E}} [\ell(\hat{\pi}(X), A)] - \frac{1}{m} \sum_{i=1}^m \mathbb{E}_{XA \sim \mathbb{P}_{\theta_i}^{\pi^E}} [\ell(\hat{\pi}(X), A)].$$

For a confidence $\delta \in (0, 1)$, $\Gamma(E)$ is upper bounded with probability at least $1 - \delta$ by

$$\beta \sqrt{\frac{I_{T;Z|A} \log 2 + \alpha_{\gamma} \log 2 + H(Z|T, A) + \log(2|\mathcal{A}|/\delta)}{m}} + \frac{f(\hat{\pi}) \sqrt{2\gamma|\mathcal{A}|\log(2|\mathcal{A}|/\delta)}}{m^{3/4}} + \frac{\gamma g(\hat{\pi})}{m^{1/2}}$$

where $\alpha_{\gamma}, \beta, \gamma$ are constant values, $f(\hat{\pi}) = \max_{i \in [m]} \mathbb{E}_{XA \sim \mathbb{P}_{\theta_i}^{\pi^E}} [\ell(\hat{\pi}(X), A)]$ is the maximum training loss, $g(\hat{\pi}) = \sup_{XA} \ell(\hat{\pi}(X), A)$ is the maximum generalization loss, and $H(Z|T, A)$ is the entropy of the demonstrator internal representation given TA .

Proof. This result is based on Theorem 1 in [20], in which the notation adapted to our setting of interest. All of the derivations can be found in [20]. A coarser version of the bound is given as

$$\Gamma(E) \leq \left(\beta \sqrt{\alpha_{\gamma} \log 2 + H(Z|T, A) + f(\hat{\pi}) \sqrt{2\gamma} + \gamma g(\hat{\pi})} \right) \sqrt{\frac{I_{T;Z|A}|\mathcal{A}|\log(2|\mathcal{A}|/\delta)}{m}}$$

where the first factor can be incorporated into a constant $C(E, \pi^E, \hat{\pi})$. We note that the term $H(Z|T, A) = 0$ whenever the demonstrator internal representation is deterministic, which is a fair assumption in our setting. Further, the maximum training error $f(\hat{\pi})$ is close to zero and upper bounded by the optimization error $\mathcal{E}_{opt}(\hat{\pi})$. The value of $g(\hat{\pi})$ is upper bounded by 1 for the indicator loss $\ell(\hat{\pi}(X), A) = \mathbf{1}(\hat{\pi}(X) \neq A)$. Finally, we note that $I_{T;Z} \geq I_{T;Z|A}$. With these considerations, by omitting all of the constants, we have

$$\Gamma(E) \lesssim \sqrt{\frac{I_{T;Z}|\mathcal{A}|\log(|\mathcal{A}|/\delta)}{m}}$$

as it is reported elsewhere in the paper. $\square$

B Peg insertion extended results

This section describes in detail the setting, hyperparameters, and constants used for the peg insertion task, as well as provides extended evaluation results.

B.1 Experimental setup

Figure 7 shows a close-up view of all peg shapes (10 shapes in total) and their corresponding boards, where the training shapes are in the top row and the test shapes are in the bottom row. Each peg insertion attempt starts from 10cm above the hole (Z axis) and a random reset position within a 0.5cm box in the XY plane, centered above the hole position. In this experiment, initial rotations about the X and Y axes are fixed (0 degrees), while the Z angle starts from a random rotation ranging from -60 to 60 degrees. Two Realsense cameras are mounted on the robot's wrist. For the blindfolded expert experiment, we mask out the hole to hide its orientation from the experts, such that they cannot infer the orientation of the peg and must explore the domain in order to insert the peg.

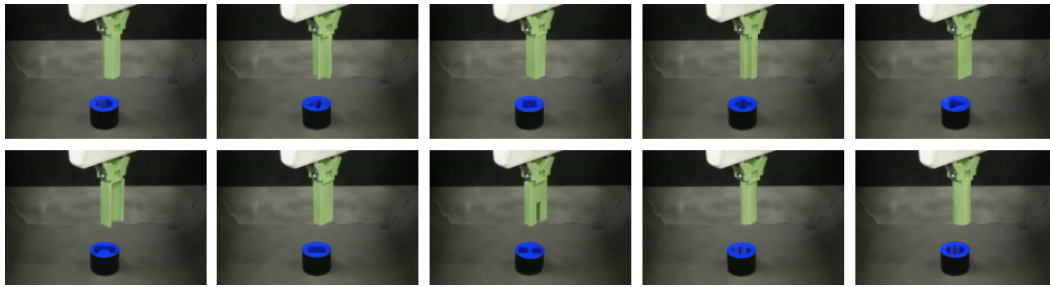


Figure 7: Close-up view of all peg insertion tasks. Top row: Training pegs. Bottom row: Test pegs.

B.2 Hyperparameters and constants

As described in section 4.2, the same architecture is used for learning both π_{BC} and π_{BF-BC} . Specifically, we use ResNet-10 [17] encoder pretrained on the ImageNet dataset [10], and a GRU of 1024, which we found to produce the best performance for both policies π_{BC} and π_{BF-BC} independently. Throughout our experiments, we train our networks using the Adam optimizer. The hyperparameters of the networks are the learning rate, learning rate decay, and the batch size. To ensure that the best performance of each approach is achieved, we perform a separate hyperparameter search for each policy, trained on each subset of shapes. The best hyperparameters are listed in Table 3. The network outputs a 6-dimensional vector for the mean and a 6-dimensional vector for the diagonal covariance matrix of a Gaussian policy (for 6-DoF action space). We train our networks using the log-likelihood loss. In all our evaluations, the action is chosen as the maximum likelihood of the distribution.

Table 3: List of hyperparameters used in the peg insertion experiment. The learning rate (lr) schedule indicates the iteration number for multiplying the lr by 0.5.

Hyperparameter	π_{BC}	π_{BF-BC}
batch size	1024	1024
hidden size	1024	1024
initial lr	0.0003	0.0003
lr schedule	$\text{lr} \times 0.5$ at $\{10, 100, 150, 200\}K$	$\text{lr} \times 0.5$ at $\{50, 100, 150, 200\}K$

B.3 Measuring exploratory behavior

We compute two additional measures for the exploratory behavior of the different experts: the map coverage score (Table 4) and the entropy of state visitation (Table 5).

Map coverage score is the ratio $C = N_v/N_{total}$ given by the number of visited states N_v divided by all accessible states N_{total} , averaged over all episodes. For the peg insertion experiment, we consider

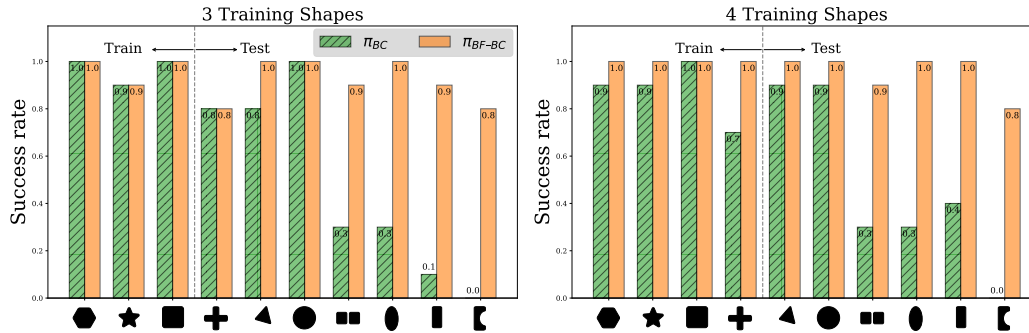


Figure 8: Success rate of robotic peg insertion for 10 peg shapes (horizontal axis). We train on a subset of shapes $k = \{3, 4\}$, with the remaining shapes withheld as test-set. Both for $k = 3$ (left) and $k = 4$ (right), cloning blindfolded experts generalizes better than cloning standard experts.

the rotation of the robotic arm around the Z-axis as the crucial component of the state space for obtaining the correct articulation for insertion. We compute the ratio of the rotation performed (in radians) divided by 2π in each trajectory, averaged over all trajectories.

The entropy of state visitation is defined by $H = -\sum_s p(s) \log p(s)$. We calculate $p(s)$ using a histogram (with 20 bins) of rotation angles along the trajectory, averaged over all trajectories.

The results confirm that blindfolded experts explore a larger portion of the state space to compensate for the redacted information in the observations.

Table 4: Map coverage score of the trajectories demonstrated by fully-informed experts (Experts) and blindfolded experts (BF-Experts).

Mode	hexagon	star	square	plus	triangle	Average
Experts	0.078	0.113	0.150	0.153	0.191	0.137
BF-Experts	0.114	0.259	0.267	0.270	0.327	0.247

Table 5: Entropy of the state visitation of the trajectories demonstrated by fully-informed experts (Experts) and blindfolded experts (BF-Experts).

Mode	hexagon	star	square	plus	triangle	Average
Experts	2.876	3.121	3.266	3.228	3.418	3.182
BF-Experts	3.100	3.416	3.546	3.577	3.631	3.454

B.4 Results for different combinations of training shapes

Figure 8 shows the success rate for $k = 3, 4$ training shapes (and the rest serve as a test set, out of a total of 10 peg shapes). The results on the varying amounts of training shapes, further support that cloning blindfolded experts generalizes better than the standard BC approach.

C Progen maze and heist extended results

This section describes in detail the hyperparameters and constants used for the Progen maze task.

C.1 Hyperparameters and constants

For a fair comparison, we conduct a separate hyperparameter search for both π_{BC} and π_{BF-BC} . We perform a hyperparameter search for the batch size $b \in \{128, 256, 512, 1024\}$, for the learning rate $\text{lr} \in \{1e^{-3}, 1e^{-4}, 1e^{-5}, 5e^{-3}, 5e^{-4}, 5e^{-5}\}$ and for hidden size $h \in \{128, 256, 512, 1024\}$. We also evaluate the performance with and without using a learning decay schedule. Our networks are trained using the Adam optimizer. The best hyperparameters are chosen based on the lowest training loss and the highest training success rate. We evaluate performance over an average of 10 different random

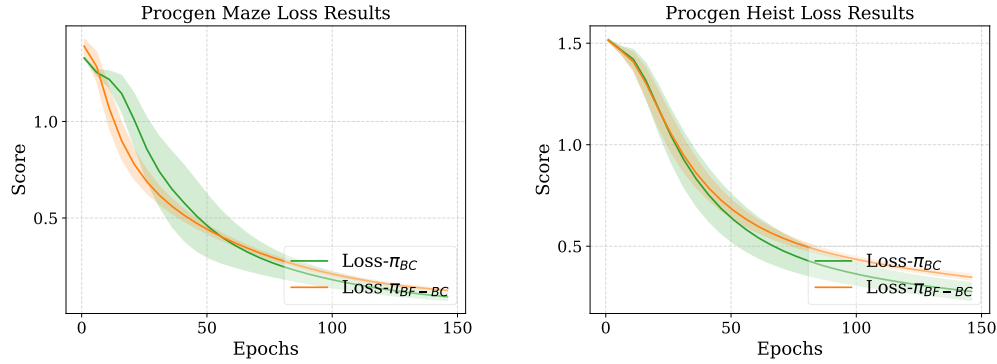


Figure 9: Loss as a function of training epochs for the Procgen maze (left) and heist (right). Mean (color line) and std (shaded region) are computed across 10 seeds.

training seeds. Figure 9 and Table 6 show the loss function and the chosen hyperparameters. Note that in most cases, the best hyperparameters for π_{BC} and π_{BF-BC} turned out to be fairly similar.

Table 6: List of hyperparameters used in the Procgen Maze experiment.

Hyperparameter	π_{BC}	π_{BF-BC}
batch size	256	256
hidden size	1024	1024
learning rate	10^{-5}	10^{-5}
learning rate schedule?	No	No

Table 7: List of hyperparameters used in the Procgen Heist experiment.

Hyperparameter	π_{BC}	π_{BF-BC}
batch size	128	128
hidden size	512	1024
learning rate	10^{-5}	10^{-5}
learning rate schedule?	No	No

C.2 Number of steps vs. number of trajectories

As described in Table 1, the blindfolded expert takes more steps on average to complete each trajectory. When comparing the different approaches, we match the number of trajectories, which leads to a greater total number of environment steps for the blindfolded expert. Table 8 shows the total number of steps available for training the different BC policies, alongside their performance. We also compare our results to a standard BC approach with twice the number of trajectories (from the same 100 seeds) to match the number of environment steps produced by the blindfolded expert. In addition, we compare our results to the results reported by [27], who train a BC policy on the Procgen maze with a dataset of $1M$ environment steps taken from a trained PPO expert, on 200 training seeds⁸.

We can see in Table 8 that π_{BF-BC} achieves better performance than all other contending policies. When compared to the results reported in [27], we can see that our results are better despite significantly less training data (an order of magnitude fewer trajectories) and half the number of training seeds.

⁸For $1M$ expert dataset, we report the results from [27] who evaluated over 5 seeds.

Table 8: Performance comparison on the Procgen maze experiment. Our results are reported at epoch 40 (early stopping) when training performance plateaus. The mean and std are computed over 10 seeds. Top performer in bold.

Parameter	1M Expert Dataset in [27]	π_{BC}	π_{BC-ext}	π_{BF-BC}
# of trajectories	15385	2000	4608	2000
# of total env steps	1000000	57166	115290	103238
# of seeds	200	100	100	100
Test performance	4.46 ± 0.16	3.35 ± 0.29	3.65 ± 0.3	6.37 ± 0.47

D Data collection

Peg insertion. For both π_{BC} and π_{BF-BC} , we use 400 trajectories for each of the training shapes. We, the authors, collected the data by operating the robot manually using a Spacemouse control. Recall that the blindfolded expert observes a masked-out view of the board (through the robot wrist cameras) such that the orientation of the peg is not directly visible and must be inferred through exploration. However, recorded observations in favor of cloning the blindfolded policy π_{BF-BC} are unmasked, i.e., only the human expert is blindfolded. In addition, we rescale the images from 480×480 to 128×128 for both π_{BC} and π_{BF-BC} to facilitate computations.

Procgen maze and heist.

To train our experts (BC) and blindfolded experts (BF-BC) policies, we collected 4000 human demonstrations on 200 levels. We conducted crowd-sourced data collection for the maze and heist videogames by recruiting 20 volunteers who played the games. Each expert played all 200 levels of maze and all 200 levels of heist twice, once with the mask and once without the mask. Their game trajectories were recorded to serve toward the imitation-learning of the experts’ policy π_{BC} and blindfolded experts’ policy π_{BF-BC} . The participants moved the mouse using the keyboard’s arrow keys and relied on the Procgen “interactive” GUI for maze and heist observations in full resolution (512×512). For the blindfolded experts’ data, their observations are modified to reveal only the agent’s immediate surroundings (a diameter of $\frac{1}{8}$ of the width for maze and $\frac{1}{6}$ for heist) with the rest of the observation masked out. Note that the state observations that are provided to the cloning networks are a lower resolution of 64×64 of the unmasked observations.

All participants were compensated with vouchers for their efforts. The experiment’s GUI environment, alongside the training code and the recorded data, are available at:
<https://github.com/EvZissel/blindfolded-experts/>

NeurIPS Paper Checklist

1. Claims

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

Answer: [Yes]

Justification: The abstract claims that the paper shows that cloning blindfolded experts leads to better generalization. We support this claim both theoretically (see Theorem 3.1 in Section 3) and empirically (see Section 4.1 and Section 4.2 where we provide the Progen and peg insertion experiments advertised in the abstract).

Guidelines:

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

2. Limitations

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

Answer: [Yes]

Justification: While the paper does not include a "Limitations" section, limitations are discussed in the paper. On the theoretical side, a discussion on the limitations given by the considered assumptions, and how they may be overcome, is provided at the end of Section 3. For the general approach, a limitation on the design of the information bottleneck for the blindfolded expert is reported in the Conclusion section.

Guidelines:

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

3. Theory assumptions and proofs

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

Answer: [Yes]

Justification: The paper has a theoretical section (Section 3) with a main theoretical result (Theorem 3.1). The assumptions to derive the result are explicitly stated in the same section, above the theorem (Asm. 1 2 3 4). The proof of the theorem, together with the necessary lemmas, is provided in the Appendix A.

Guidelines:

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

4. Experimental result reproducibility

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

Answer: [Yes] .

Justification: The experimental setup, network architecture and hyperparameters, data collection and training schemes are all detailed in Sections 4, C, B for full reproducibility. In addition, the abstract provides a link to the project home page that contains the full source code and documentation.

Guidelines:

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

- (d) We recognize that reproducibility may be tricky in some cases, in which case authors are welcome to describe the particular way they provide for reproducibility. In the case of closed-source models, it may be that access to the model is limited in some way (e.g., to registered users), but it should be possible for other researchers to have some path to reproducing or verifying the results.

5. Open access to data and code

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

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

Justification: The abstract section provides a link to the project home page that will contain the full source code, environment setup and documentation, tutorial, components CAD and network weights (complete soon).

Guidelines:

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

6. Experimental setting/details

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

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

Justification: The experiments, their setup, parameters and process are detailed at length in the Experiments section (4). Also, the abstract provides a link to the project’s page with all the necessary code, documentation, components CAD and tutorials.

Guidelines:

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

7. Experiment statistical significance

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

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

Justification: The paper contains error bars in table 1 and figure 3 in the form of mean and standard deviation and the method for their computation is detailed in the corresponding captions.

Guidelines:

- The answer NA means that the paper does not include experiments.
- The authors should answer "Yes" if the results are accompanied by error bars, confidence intervals, or statistical significance tests, at least for the experiments that support the main claims of the paper.
- The factors of variability that the error bars are capturing should be clearly stated (for example, train/test split, initialization, random drawing of some parameter, or overall run with given experimental conditions).
- The method for calculating the error bars should be explained (closed form formula, call to a library function, bootstrap, etc.)
- The assumptions made should be given (e.g., Normally distributed errors).
- It should be clear whether the error bar is the standard deviation or the standard error of the mean.
- It is OK to report 1-sigma error bars, but one should state it. The authors should preferably report a 2-sigma error bar than state that they have a 96% CI, if the hypothesis of Normality of errors is not verified.
- For asymmetric distributions, the authors should be careful not to show in tables or figures symmetric error bars that would yield results that are out of range (e.g. negative error rates).
- If error bars are reported in tables or plots, The authors should explain in the text how they were calculated and reference the corresponding figures or tables in the text.

8. Experiments compute resources

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

Answer: [Yes] .

Justification: For both experiments, the compute-resources are detailed in sections B and C.

Guidelines:

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

9. Code of ethics

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

Answer: [Yes]

Justification: The paper fully respect the NeurIPS Code of Ethics.

Guidelines:

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

10. Broader impacts

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

Answer: [NA]

Justification: The paper can be categorized as foundational research. At the current stage of development, it is not tied to any product that may be deployed in the real world and have an impact on society.

Guidelines:

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

11. Safeguards

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

Answer: [NA]

Justification: The paper does not pose such risks such as potential for dual-use or risk to society.

Guidelines:

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

12. Licenses for existing assets

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

Answer: [Yes] .

Justification: All relevant papers and works mentioned are properly cited, and supportive code implementations are recognized and accredited (Procgen [8], SERL [25])

Guidelines:

- The answer NA means that the paper does not use existing assets.
- The authors should cite the original paper that produced the code package or dataset.
- The authors should state which version of the asset is used and, if possible, include a URL.
- The name of the license (e.g., CC-BY 4.0) should be included for each asset.
- For scraped data from a particular source (e.g., website), the copyright and terms of service of that source should be provided.
- If assets are released, the license, copyright information, and terms of use in the package should be provided. For popular datasets, paperswithcode.com/datasets has curated licenses for some datasets. Their licensing guide can help determine the license of a dataset.
- For existing datasets that are re-packaged, both the original license and the license of the derived asset (if it has changed) should be provided.
- If this information is not available online, the authors are encouraged to reach out to the asset's creators.

13. **New assets**

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

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

Justification: Our dataset of human expert trajectories for the Procgen maze and peg-insertion tasks will be made available online on the project page <https://sites.google.com/view/blindfoldedexperts/home>.

Guidelines:

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

14. **Crowdsourcing and research with human subjects**

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

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

Justification: In section D we detail how data collection was conducted, including the instructions for the participants and the compensation provided.

Guidelines:

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

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

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

Answer: [Yes]

Justification: There are no potential risks incurred by study participants, and we verified our study conforms to the guidelines in our institution.

Guidelines:

- The answer NA means that the paper does not involve crowdsourcing nor research with human subjects.
- Depending on the country in which research is conducted, IRB approval (or equivalent) may be required for any human subjects research. If you obtained IRB approval, you should clearly state this in the paper.
- We recognize that the procedures for this may vary significantly between institutions and locations, and we expect authors to adhere to the NeurIPS Code of Ethics and the guidelines for their institution.
- For initial submissions, do not include any information that would break anonymity (if applicable), such as the institution conducting the review.

16. Declaration of LLM usage

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

Answer: [NA]

Justification: This paper does not involve LLMs.

Guidelines:

- The answer NA means that the core method development in this research does not involve LLMs as any important, original, or non-standard components.
- Please refer to our LLM policy (<https://neurips.cc/Conferences/2025/LLM>) for what should or should not be described.