

EFFECTOR COMPLEXITY ENHANCES TRANSFER LEARNING

Anonymous authors

Paper under double-blind review

ABSTRACT

Both biological and artificial embodied systems rely on effectors to interact with the world. How does this embodiment impact the way they learn? The role of embodiment in shaping learning dynamics is not well understood from a neuroscience or a machine learning perspective. In this study, we treat embodiment as a variable in artificial agents and study how changes in effector complexity reshape the dynamics of learning. Our hypothesis is that more complex effectors provide constraints that yield better transfer learning on new tasks, despite simultaneously posing a more complex control problem. We evaluated this hypothesis on area under the performance curve, and use time to sustained performance plateau as a parameter for task difficulty. Our results show that while a simpler effector excels when trained from scratch, a more complex effector yields superior performance after pre-training on another task. We further demonstrate that the improvement gained from transfer learning is greater for the complex effector. Our findings suggest that embodiment plays an important role in enabling efficient transfer, offering insights into the differences in learning dynamics between disembodied artificial systems and their biological counterparts.

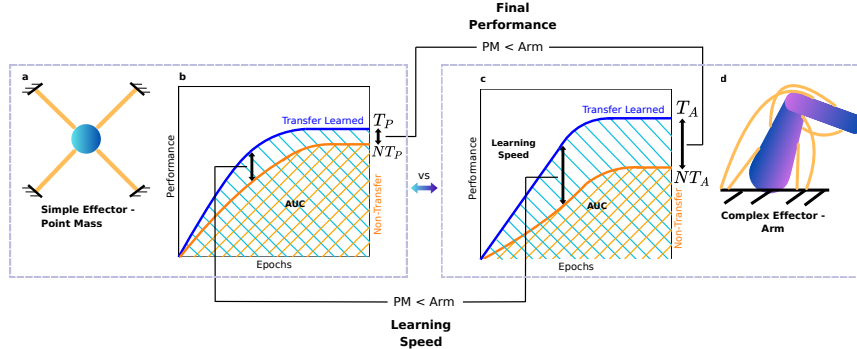


Figure 1: Summary of the Approach and Main Result: Simple effectors learn faster from scratch, but complex effectors achieve greater benefits from transfer.

(a) Simple Effector: A Point Mass (blue) controlled by four diagonal muscles (yellow) moving in the 2D plane.

(b) Point Mass performance: no-transfer (orange) vs transfer (blue). Modest improvement with transfer performance ($T_P > NT_P$) and learning speed, modest AUC.

(c) Arm performance: no-transfer (orange) vs transfer (blue), stronger transfer performance benefits ($T_A \gg NT_A$), and learning speed, indicating better knowledge reuse. Overall, larger AUC after transfer.

(d) Complex Effector: A biomechanical Arm with two links (purple) and six muscles (yellow), introducing higher embodiment complexity.

1 INTRODUCTION

Deep learning based approaches have seen significant success in many domains ranging from image to speech. However, these models usually require large amounts of high-quality data and long training durations to reach competitive performance (Talaei Khoei et al., 2023). Deep Reinforcement Learning (DRL), which combines reinforcement learning with deep neural networks, has shown promise in learning effective function approximations for tackling complex problems (Arulkumaran et al., 2017) but struggle with cross-task adaptation. In contrast, humans can quickly adapt to new tasks by transferring meta-skills acquired from diverse experiences throughout their lives (Lake et al., 2015). This is in sharp contrast to current deep reinforcement learning-based navigation methods, where policy networks are trained from scratch.

Embodiment is the principle that intelligence is not solely computational or abstract, but tightly integrated with the agent’s body and physical capabilities (Zhao et al., 2024). It allows a system to have meaningful interactions with the environment it is placed in (Roy et al., 2021). The integration of a physical body and sensory-motor feedback that allows the ability to transfer and reuse knowledge across tasks, should help by reducing the need for extensive training samples (Li et al., 2020). In robotics, perception-based feedback has traditionally dominated, with agents relying on visual inputs to estimate their state and take necessary action (Sünderhauf et al., 2018; Levine et al., 2018; 2016). In this work, we propose a shift toward actuator-based feedback, leveraging the embodiment of effectors, i.e., getting feedback from proprioception instead of depending on perception. We study transfer learning in embodied neural networks, where a model trained on a source task (referred to as the Base Model), is subsequently trained on a target task (referred to as the Transfer Model) using policy reuse. The hypothesis we would like to test is if embodiment improves learning speed and final performance. We assess this hypothesis across a diverse set of tasks and effectors and evaluate transfer quality using four metrics: peak performance, index at peak, area under the performance curve and time to sustained plateau. All biological learning systems are embodied, motivating the hypothesis that embodiment impacts a wide range of functions of the brain including decision making and learning (Cisek & Pastor-Bernier, 2014). However, precisely how embodiment shapes the dynamics of learning in the brain is not well understood. The body is often treated as a fixed substrate instead of one with tunable parameters, and the implications of different body plans remain largely unexplored. By treating embodiment as a variable in artificial agents, we examine how effector complexity impacts the dynamics of learning a new task, both from arbitrary initial conditions and in models with prior training on other tasks.

1.1 RELATED WORK

The role of embodiment in learning has received previous attention in several different domains. For example, Kulkarni & Nair (2023) examined the problem of transfer learning in the context of embodied systems with robots and neural network controllers. They propose a new method for transferring knowledge between different tasks using an evolutionary algorithm to train neural networks instead of reinforcement learning. They also demonstrate that ‘hot neurons’, neurons with significant positive contribution to learning, are particularly important for transfer learning. However, the authors point out the lack of methods to measure such transfers among different tasks. Our study seeks to overcome this problem by introducing metrics to evaluate transfer learning performance systematically.

Embodiment can also impact the performance of language models. For example, Driess et al. (2023) show the capabilities of an embodied language model, PaLM-E by injecting multi-modal information into the embedding space. PaLM-E directly incorporates continuous inputs from sensor modalities into a large language model, enabling more grounded inferences for sequential decision making in the real world. While their focus is on vision guided multi-modal grounding, our work emphasizes actuator feedback as the primary signal providing an alternative approach to embodiment that is not dependent on perception. The authors also demonstrate that when PaLM-E was evaluated on robotic manipulation tasks, multi-task training improves performance compared to training models on individual tasks and also improves data efficiency. As tasks become more complex, the performance of a neural network controller decreases, especially when learning happens through an embodied system (Kulkarni & Nair, 2023).

Recent work suggests that architectural complexity and embodiment can jointly influence how well agents can generalize across tasks. Studies such as those from Reed et al. (2022) and Das et al. (2018) also report performance improvement for multi-modal tasks. In addition to LLM-based models, the work in Lin et al. (2024) also defines the Clip on Wheels method (CoW) (Gadre et al., 2023) which uses zero-shot visual models like CLIP (Contrastive Language–Image Pre-training) (Radford et al., 2021) for embodied AI navigation and relies on a GRU-Actor + GRU-Critic formation for learning. Overall, prior work provides promising evidence that embodiment influences learning dynamics in multi-modal settings, motivating a more thorough investigation into its role in transfer learning.

2 METHODS

2.1 NETWORK ARCHITECTURE

The gated recurrent unit (GRU) is particularly suited for sequential data applications, such as time-series or trajectory modeling, where adaptive hidden states and robust initialization improve performance. Our architecture includes a custom GRU module that extends PyTorch’s *nn.GRU* with a learnable hidden state. The hidden state $hidden_0$ is parameterized and constrained through a *tanh* activation to remain bounded and numerically stable. This design allows the network to adapt its internal dynamics to the structure of the task from the outset of training. Our work builds on the Motornet library (Codol et al., 2024). For training structure, please see Appendix A.2.

2.2 TRAINING

We define baseline learning performance by training each network architecture and effector on every task independently from random initial conditions, without prior experience from other tasks. This ‘no-transfer’ model serves as a comparison to quantify the benefit of transfer learning.

1. Base model: trained on the source task.
2. Transfer model: initialized with the base model, then trained on a new task.
3. No-Transfer model: counterpart of the transfer model that was trained from scratch.

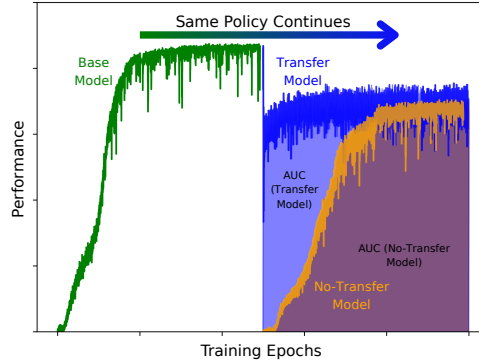


Figure 2: Overall Training: First a Base Model (green) is trained on a task (e.g. RTR). The same model is then trained on a transfer task (e.g. COR), and hence known as a Transfer Model (blue). The Transfer Model is then compared to a No-Transfer Model that is trained on the same transfer-task (i.e. COR) but from scratch.

2.3 TASKS

To test our hypothesis and model, we created six different reaching tasks, which are summarized in Table 1. All tasks are restricted to a 2-D plane, where goal locations are in (X,Y) co-ordinates.

Table 1: Task descriptions

Task Name	Description
Random Target Reach (RTR)	Reach randomly generated targets from randomly generated joint states
Center Out Reach (COR)	Reach randomly generated targets from Cartesian center joint states
Center Out and Back, Fixed Goal Reach (COBFGR)	Reach randomly generated targets from Cartesian center joint states, and return to center. 'Visual' goal remains the same but the 'rewarded' goal changes. This gives the effector mixed signals about returning to the center.
Center Out and Back, Jumping Goal Reach (COBJGR)	Reach randomly generated targets from Cartesian center joint states, and return to center. Both 'visual' and 'rewarded' goal shift giving the effector a strong signal to return to the center.
Delayed Center Out Reach (DCOR)	Start from the Cartesian center joint states and wait before moving to a randomly generated target.
Jumping Target Reach (JTR)	Reach a new randomly generated target every time a target is hit.

The start configuration of the effector is characterized in joint angles in the case of Arm, or end effector location (X,Y) in the case of Point Mass. Both, goal locations and starting configuration of the effector can be randomized depending on the task.

For all tasks except for JTR, performance was evaluated as the fraction of time the model spends in the target zone on each trial. Thus, because it takes a finite amount of time to move from the start location to the target, optimal performance is a fraction less than 1.0. For example, an average reach takes 15% of a trial in RTR, and COBJGR duplicates this reach, making optimal performance in these tasks approximately 85% and 70%, respectively. For JTR, time spent in the target region is negligible even for expert-level performance, because the target moves upon being reached. Therefore, we quantified performance in JTR as the number of new targets reached on a trial normalized by the number of targets presented during the trial.

We evaluated performance across this family of reaching tasks that have varying initial conditions, goal dynamics, and temporal constraints. These tasks are designed to vary in complexity to probe different aspects of control and adaptability. While some tasks were more difficult than others, these tasks are not organized along a predefined difficulty scale. Instead, we quantify difficulty using Time to Plateau (TTP). By training models on each task from scratch (without transfer) and measuring how long they take to reach stable performance, we obtain an empirical proxy for difficulty. This design allows us to ground our comparisons in data driven metrics rather than qualitative assumptions about complexity (see Table 1). We define TTP as the number of epochs taken for the performance of the No-Transfer models to reach and stabilize near their respective peak values. A Savitzky-Golay filter (Savitzky & Golay (1964), Schafer (2011)) is applied with a window length of 500 and a polynomial order of 3 to smooth out fluctuations in the data. Next an amplitude tolerance is set to 10% of the respective peak value, allowing for slight deviations from the peak to still be considered part of the plateau. A plateau is identified when the slope is below a small threshold ($1e^{-4}$ in this case) and when the performance difference between that point and the peak is within amplitude tolerance. At such an index, a plateau is said to have been established. We show difficulty of each task in Figure 7 (left). The time to sustained plateau (T_{plateau}) is calculated as:

$$T_{\text{plateau}} = \min \left\{ t : |y_t - y_{\text{peak}}| < \epsilon \text{ and } \frac{\Delta y_t}{\Delta t} < \delta \right\}, \quad (1)$$

where:

- y_t : Performance value at time t .
- y_{peak} : Peak performance value.
- ϵ : Amplitude tolerance threshold for deviation from y_{peak} .
- $\frac{\Delta y_t}{\Delta t}$: Slope of the performance curve.

- δ : Threshold for slope, defining a stabilized plateau.

2.4 EFFECTORS

We create the 'embodiment' for the neural network to control through two effectors. These effectors are based on the Motornet library for python Codol et al. (2024) which is a toolbox for building and controlling differentiable biomechanical effectors.

2.4.1 POINT MASS

The first effector is a simple point mass object with a central mass connected to four muscles diagonally creating an "X" formation Figure 3(a). One end of each muscle (yellow) is fixed to the ground/reference frame and the other end is attached to the point mass itself (blue). Each muscle is modeled as a rectified linear muscle whose force output is a linear function of its maximum isometric force and activation value which itself is bound between 0 and 1. The end effector itself is modeled as a bony structure and can be considered as a one-bone system with no joints. However, note that the number of degrees of freedom is not zero like the number of joints, but is two, corresponding to the freedom to move in the X-Y plane.

2.4.2 ARM

The Arm is the more complex effector shown in Figure 3(b). It is a two-link, six-muscle model with the 'end effector' being the end point of the second link. It also integrates biomechanical details and approximations for controlling a robotic arm. This effector class is an implementation of a 6-lumped muscle model from Nijhof & Kouwenhoven (2000) since lumped muscles can be considered a functional approximation of biological reality. The skeleton consists of two parts. The first bone (Humerus) between the shoulder and elbow and the second bone (Radius and Ulna, modeled as a single link) between the elbow and the end effector. This system is made to perform as a planar arm with two degrees of freedom. The six muscles represent functional approximations of biological actuators. The skeleton and each muscle has individual physical parameters such as weight, length, tendon length, inertia etc. For further numerical specifications refer to the appendix. The effector uses a third degree polynomial function to approximate the moment arms, musculotendon lengths, and musculotendon velocities during movement. A moment arm describes how changes in joint position affect tendon force transmission while musculotendon refers to the overall muscle + tendons structure. Euler integration is used as the default numerical method and control inputs directly actuate the muscles, influencing joint torques through moment arms.

2.5 PERFORMANCE METRICS

Area Under the Curve (AUC) is calculated to quantify the overall performance of the task and baseline models over a defined interval. In this case, AUC is calculated using the trapezoidal rule for numerical integration. It is extracted starting from the end of the baseline and extending till the last epoch (n) of the transfer task. AUC is calculated as:

$$\text{AUC} = \sum_{i=1}^n \frac{y_i + y_{i+1}}{2} (x_{i+1} - x_i), \quad (2)$$

where:

- y_i and y_{i+1} are the performance values at indices i and $i + 1$, respectively.
- x_i and x_{i+1} are the corresponding indices or time steps.

3 RESULTS

3.1 SIMPLER EFFECTORS LEARN FASTER FROM RANDOM INITIAL CONDITIONS

We quantified the dynamics of learning using the area under the curve (AUC) of the performance over training trials. This metric quantifies performance over the entirety of learning without additional parameters. To compare performance of 'No-transfer' models across tasks, we compute

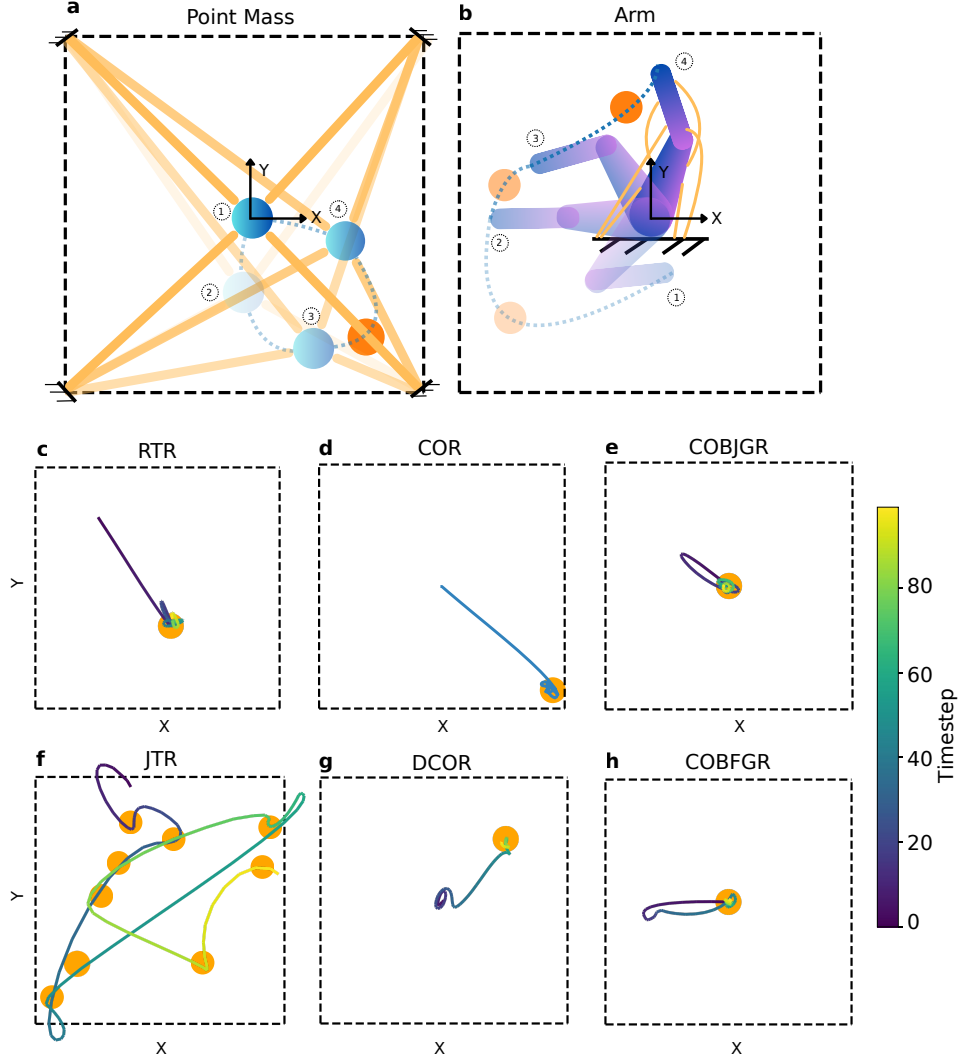


Figure 3: (a) The Point Mass is a central mass (blue), connected by four diagonal muscles. In this case the end effector is the mass itself. An example sequence of the Point Mass performing COBFGR is shown from points 1 through 4, where the dotted line (blue) shows the trajectory progression, the orange circle shows the target. The four muscles move in sequence to pass the Point Mass through the target. Once the target is hit, the Point Mass moves back to the center.

(b) The Arm is built from two links (magenta) and six muscles (yellow). Links move when the muscles contract in sequence. An example sequence of the Arm performing JTR is shown from points 1 through 4, where the dotted line (blue) shows the trajectory progression, the orange circles show the target. The tip of the arm (end effector) must pass through the target to get the reward.

(c) - (f) Different tasks (as described in Table 1) being completed in the X-Y plane, with the end-effector trajectory and target plotted. Each trial lasts 100 time steps.

$AUC_{Arm} - AUC_{PointMass}$. Figure 4 shows examples of this learning trajectory for both effector models. In this No-Transfer Scenario where models start from random initial conditions, the Point Mass model outperformed the Arm model on all tasks as shown in Figure 4 (g). Intuitively, this follows from the fact that simpler effectors are easier to control, leading to faster learning from random

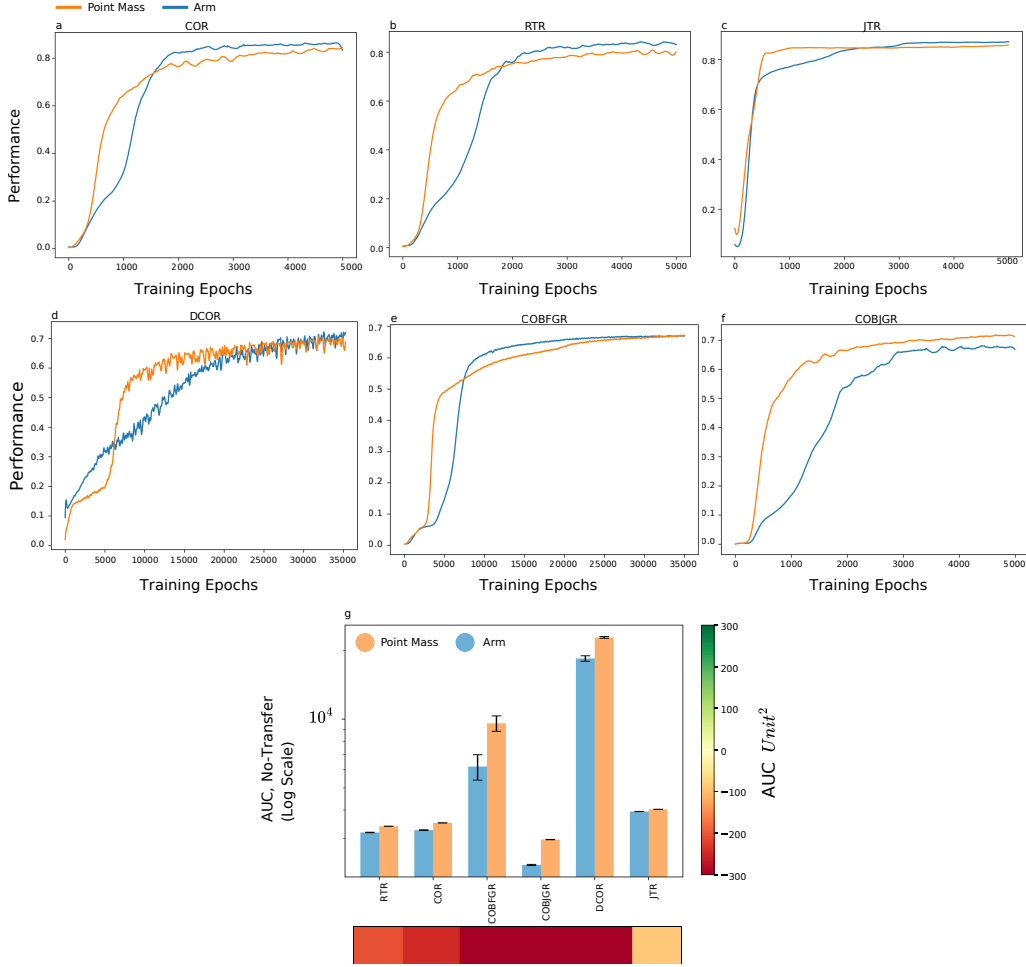


Figure 4: No-transfer models performance comparison for different tasks
(a-f) Training performance of different tasks, in no particular order. Note that the Point Mass (yellow) learns faster, and matches or sometimes, exceeds the Arm’s (blue) final performance.
(g) A bar plot with Log scale Y axis. Shows the relative difference in the AUCs covered, and a heatmap to visualize the difference ($AUC_{Arm:No-transfer} - AUC_{PointMass:No-transfer}$).

initial conditions to expert level performance in each task. This finding directly parallels prior work from Bazzi et al. (2024), which shows that complex manipulation is often substituted by simpler approximations.

3.2 INCREASED EFFECTOR COMPLEXITY IMPROVES TRANSFER PERFORMANCE

As described above, we examined transfer learning by taking models previously trained on one task and training them on one of the new tasks (Methods). Intuitively, previous training on one task should confer some benefit to the speed of acquisition of a new task, if the two tasks share some of the same features.

To quantify the improvement in learning due to transfer, we defined a metric $AUC_{Transfer} - AUC_{No-Transfer}$ separately for the arm and point mass. Across all six tasks, transfer models using the Arm showed positive improvements in AUC relative to their respective baselines (Figure 6 (right) and Figure 7 (right)).

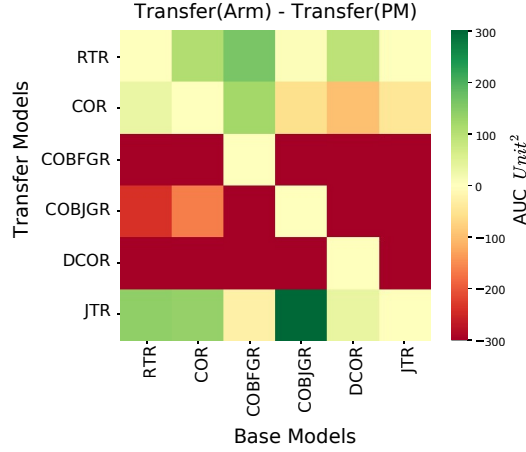


Figure 5: Transfer (Arm) - Transfer (PM)

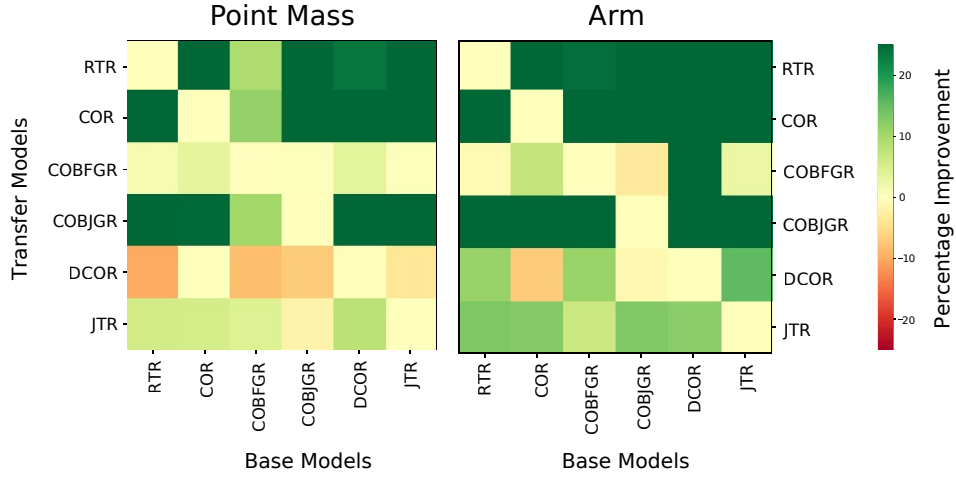


Figure 6: Percentage improvement in AUC of transfer models for Point Mass (left) and Arm (right). Each cell shows a combination of base and transfer tasks, whose performance is then compared to the No-transfer counterpart. The X-axis shows Base Models onto which the transfer models are trained, shown on the Y-axis. (For e.g. the top right cell denotes the percentage improvement of transfer task [RTR on a network previously trained on JTR] from a No-Transfer task [RTR trained from scratch])

Refer to Figure 2 for definition of base, transfer and no-transfer models. Warmer colors (red/yellow) indicate limited or negative transfer, while cooler colors (green) reflect stronger positive transfer. The Arm exhibits more consistent and higher transfer gains across all tasks.

Improvement for the arm model is higher than for the point mass for all task combinations. This trend is present in all task combinations, but it is most pronounced in the intermediate difficulty tasks (COBJGR) where the Arm’s AUC improves by 30-35% compared to only 15-20% with the Point Mass (Figure 7 (right)). For the harder tasks (COBFGR and DCOR) the Point Mass shows negligible and even negative transfer, highlighting its limited ability to generalize across temporally constrained tasks. In contrast, the Arm model consistently maintains positive transfer even on more challenging tasks. These gains are all statistically significant, with the exception being COBFGR,

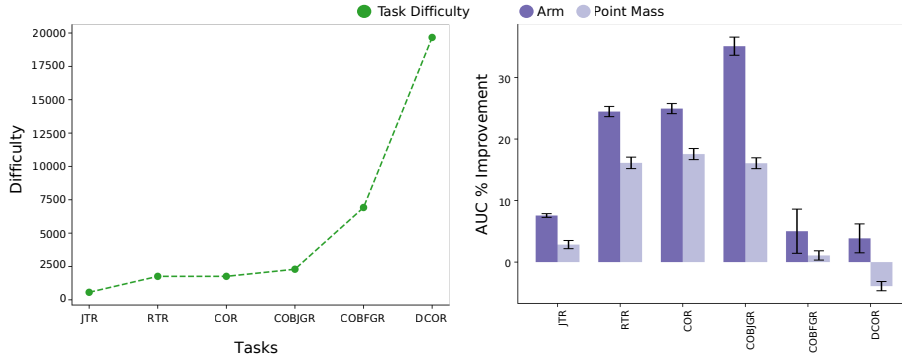


Figure 7: Left: Task difficulty measured in time taken (epochs) to reach stable performance. Right: Comparison between Arm and Point Mass

where performance is statistically indistinguishable (ns) (Figure 7 (right) and Figure 8 Appendix). Thus, averaged across all tasks, the Arm model consistently achieved larger improvements in performance due to transfer learning (Figure 6). Put simply, a complex effector shows greater improvement with transfer learning.

In principle, greater improvement of the arm model due to transfer learning could simply be a consequence of lower No-Transfer performance for the arm model. To address this, we compared the performance of both effector models after transfer learning directly. We quantified the relative performance of the models using the metric $AUC_{Arm:Transfer} - AUC_{PointMass:Transfer}$, that is, the difference in transfer performance across effectors. As shown in Figure 5, this metric is predominantly positive for half of the tasks we examined, in stark contrast to the difference in No-transfer performance, which is exclusively negative (Figure 4). Counterintuitively, this demonstrates that complex effectors not only improve more than simpler effectors through transfer learning, but can also outperform simple effectors after transfer learning for some tasks (Figure 5). Strikingly, this suggests that an increase in effector complexity can actually lead to an improvement in transfer learning efficiency.

Collectively, these observations not only support the hypothesis that embodiment impacts learning, they support the stronger hypothesis that a more complex body can actually facilitate learning.

4 CONCLUSION

The role of embodiment in transfer learning is unexplored. Our work advances our understanding by comparing embodied agents with different effectors across a family of two-dimensional reaching tasks. We find that while simpler effectors with orthogonal muscles and no joints are more effective in learning from scratch, more complex bio-mechanically inspired effectors benefit more from transfer learning. These findings suggest that embodiment is not merely a constraint but a variable that can shape an agent’s ability to transfer information across diverse tasks.

Future work may extend this framework to non-reaching tasks such as object manipulation or train a policy on more than one transfer task (e.g. having the same policy learn all tasks, combining this study with work from Lazzari & Saxena (2025)). It would also be valuable to explore the role of inductive biases and structured dynamics that could form the theoretical framework to explain our observed results. By re-framing properties of the body as influential components of a learning system, our findings open new directions for investigating embodiment as a resource for skill acquisition and transfer.

5 ETHICS STATEMENT

This work does not involve human subjects, personally identifiable information, or sensitive data. The simulation environments (Motornet) and software packages used are publicly available/ open-source and no privacy or security risks are posed. The research is intended for advancing understanding of transfer learning in embodied agents with potential applications to robotics, machine learning and neuroscience communities, with no foreseeable harmful applications. AI/LLMs were only used for minor edits to the writing and structure. The authors acknowledge that we have read and adhered to the ICLR Code of Ethics throughout this study.

6 REPRODUCIBILITY STATEMENT

We have taken several steps to ensure reproducibility. Training details are provided in Section 2.1 of the main text. We also include code snippets in the Appendix that thoroughly illustrate the structure of our tasks, effectors, and training. While these snippets are not sufficient to run the full project end-to-end, they document the critical implementation details. Our work additionally builds on an open-source package, Motornet (<https://www.motornet.org>) referenced in Section 2.4), which allows readers to reproduce the core components. Upon acceptance, we plan to release a GitHub repository containing the full implementation to further support reproducibility.

REFERENCES

- Kai Arulkumaran, Marc Peter Deisenroth, Miles Brundage, and Anil Anthony Bharath. Deep reinforcement learning: A brief survey. *IEEE Signal Processing Magazine*, 34(6):26–38, 2017.
- Salah Bazzi, Stephan Stansfield, Neville Hogan, and Dagmar Sternad. Simplified internal models in human control of complex objects. *PLOS Computational Biology*, 20(11):e1012599, 2024.
- Paul Cisek and Alexandre Pastor-Bernier. On the challenges and mechanisms of embodied decisions. *Philos. Trans. R. Soc. Lond. B Biol. Sci.*, 369(1655):20130479, November 2014.
- Olivier Codol, Jonathan A Michaels, Mehrdad Kashefi, J Andrew Pruszynski, and Paul L Gribble. Motornet, a python toolbox for controlling differentiable biomechanical effectors with artificial neural networks. *Elife*, 12:RP88591, 2024.
- Abhishek Das, Georgia Gkioxari, Stefan Lee, Devi Parikh, and Dhruv Batra. Neural modular control for embodied question answering. In *Conference on Robot Learning*, pp. 53–62. PMLR, 2018.
- Danny Driess, Fei Xia, Mehdi SM Sajjadi, Corey Lynch, Aakanksha Chowdhery, Brian Ichter, Ayzaan Wahid, Jonathan Tompson, Quan Vuong, Tianhe Yu, et al. Palm-e: An embodied multi-modal language model. *arXiv preprint arXiv:2303.03378*, 2023.
- Samir Yitzhak Gadre, Mitchell Wortsman, Gabriel Ilharco, Ludwig Schmidt, and Shuran Song. Cows on pasture: Baselines and benchmarks for language-driven zero-shot object navigation. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pp. 23171–23181, 2023.
- Divya D. Kulkarni and Shivashankar B. Nair. Transfer learning for embodied neuroevolution. In *Proceedings of the Companion Conference on Genetic and Evolutionary Computation, GECCO ’23 Companion*, pp. 2128–2135, New York, NY, USA, 2023. Association for Computing Machinery. ISBN 9798400701207. doi: 10.1145/3583133.3596400. URL <https://doi.org/10.1145/3583133.3596400>.
- Brenden M Lake, Ruslan Salakhutdinov, and Joshua B Tenenbaum. Human-level concept learning through probabilistic program induction. *Science*, 350(6266):1332–1338, 2015.
- John Lazzari and Shreya Saxena. Multitasking recurrent networks utilize compositional strategies for control of movement. *bioRxiv*, 2025. doi: 10.1101/2025.09.10.675375. URL <https://www.biorxiv.org/content/early/2025/09/16/2025.09.10.675375>.

- Sergey Levine, Chelsea Finn, Trevor Darrell, and Pieter Abbeel. End-to-end training of deep visuomotor policies. *Journal of Machine Learning Research*, 17(39):1–40, 2016.
- Sergey Levine, Peter Pastor, Alex Krizhevsky, Julian Ibarz, and Deirdre Quillen. Learning hand-eye coordination for robotic grasping with deep learning and large-scale data collection. *The International journal of robotics research*, 37(4-5):421–436, 2018.
- Juncheng Li, Xin Wang, Siliang Tang, Haizhou Shi, Fei Wu, Yueting Zhuang, and William Yang Wang. Unsupervised reinforcement learning of transferable meta-skills for embodied navigation. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pp. 12123–12132, 2020.
- Jinzhou Lin, Han Gao, Xuxiang Feng, Rongtao Xu, Changwei Wang, Man Zhang, Li Guo, and Shibiao Xu. Advances in embodied navigation using large language models: A survey, 2024. URL <https://arxiv.org/abs/2311.00530>.
- Evert-Jan Nijhof and Erik Kouwenhoven. Simulation of multijoint arm movements. In *Biomechanics and neural control of posture and movement*, pp. 363–372. Springer, 2000.
- Alec Radford, Jong Wook Kim, Chris Hallacy, Aditya Ramesh, Gabriel Goh, Sandhini Agarwal, Girish Sastry, Amanda Askell, Pamela Mishkin, Jack Clark, et al. Learning transferable visual models from natural language supervision. In *International conference on machine learning*, pp. 8748–8763. PmLR, 2021.
- Scott Reed, Konrad Zolna, Emilio Parisotto, Sergio Gomez Colmenarejo, Alexander Novikov, Gabriel Barth-Maron, Mai Gimenez, Yury Sulsky, Jackie Kay, Jost Tobias Springenberg, et al. A generalist agent. *arXiv preprint arXiv:2205.06175*, 2022.
- Nicholas Roy, Ingmar Posner, Tim Barfoot, Philippe Beaudoin, Yoshua Bengio, Jeannette Bohg, Oliver Brock, Isabelle Depatie, Dieter Fox, Dan Koditschek, Tomas Lozano-Perez, Vikash Mansinghka, Christopher Pal, Blake Richards, Dorsa Sadigh, Stefan Schaal, Gaurav Sukhatme, Denis Therien, Marc Toussaint, and Michiel Van de Panne. From machine learning to robotics: Challenges and opportunities for embodied intelligence, 2021. URL <https://arxiv.org/abs/2110.15245>.
- Abraham Savitzky and Marcel JE Golay. Smoothing and differentiation of data by simplified least squares procedures. *Analytical chemistry*, 36(8):1627–1639, 1964.
- Ronald W. Schafer. What is a savitzky-golay filter? [lecture notes]. *IEEE Signal Processing Magazine*, 28(4):111–117, 2011. doi: 10.1109/MSP.2011.941097.
- Niko Sünderhauf, Oliver Brock, Walter Scheirer, Raia Hadsell, Dieter Fox, Jürgen Leitner, Ben Upcroft, Pieter Abbeel, Wolfram Burgard, Michael Milford, et al. The limits and potentials of deep learning for robotics. *The International journal of robotics research*, 37(4-5):405–420, 2018.
- Tala Talaie Khoei, Hadjar Ould Slimane, and Naima Kaabouch. Deep learning: Systematic review, models, challenges, and research directions. *Neural Computing and Applications*, 35(31):23103–23124, 2023.
- Zikai Zhao, Qiuxuan Wu, Jian Wang, Botao Zhang, Chaoliang Zhong, and Anton A Zhilenkov. Exploring embodied intelligence in soft robotics: a review. *Biomimetics*, 9(4):248, 2024.

A APPENDIX

A.1 DETAILED FIGURES

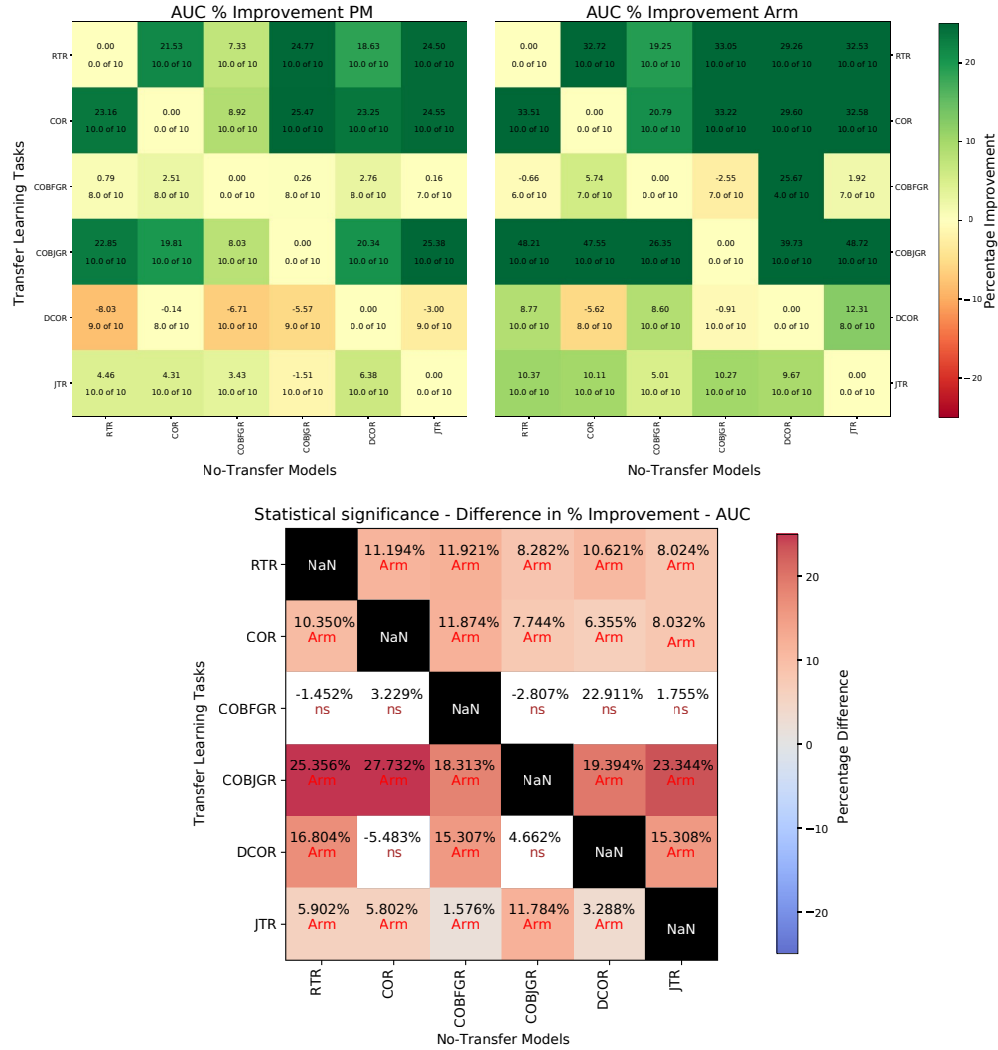


Figure 8: AUC Results

A.2 CODE FOR TRAINING

Documentation for Motornet: <http://motornet.org> Documentation for Arm and Point Mass: <https://www.motornet.org/tutorials/build-effector.html>

Training is a two-stage process involving both from-scratch learning and transfer learning. Models are instantiated as gated recurrent unit (GRU)-based policies with learnable initial hidden states, enabling adaptation to sequential dynamics. Each policy interacts with a differentiable biomechanical effector (either the Point Mass or Arm) within a Motornet environment, where proprioceptive and action signals form the state representation. The reward structure is task-specific, with both dense (distance-based) and structured feedback (e.g., delayed or return-to-center goals) designed to probe temporal and control complexity. Training runs are managed through scripted pipelines that define optimizer settings, learning schedules, and checkpointing, ensuring comparability across task-effector combinations.

Although the training environment provides trial-based rewards, our models are not trained with reinforcement learning. Instead, the entire agent-effector system is differentiable, allowing us to treat the reward signal as a supervised objective function. At each timestep, task-specific performance measures (e.g., distance to target, time in goal zone) are converted into a differentiable loss, and network parameters are updated via standard gradient descent. Thus, training proceeds in a supervised fashion: the policy is optimized directly against continuous task objectives without stochastic policy gradients, value functions, or exploration strategies typically associated with RL. This framing distinguishes our approach from reinforcement learning and situates it more accurately within differentiable control and supervised sequence learning.

```

670 1 import motornet as mn
671 2 class Environment(mn.environment.Environment):
672 3     """
673 4     This class modifies the version in Motornet by changing some aspects
674 5     of how the observations,
675 6     noise, and rewards are processed.
676 7     """
677 8     def __init__(
678 9         self,
679 10         effector: ConstrainedGoalReluPointMass |
680 11         ConstrainedGoalRigidTendonArm,
681 12         efferent_noise: float = 0.,
682 13         target_radius: float = 0.,
683 14         noise_type: Literal['additive', 'multiplicative'] | None =
684 15         None,
685 16         action_transform: Callable[[torch.Tensor], torch.Tensor] |
686 17         None = None,
687 18         action_low: float = 0.,
688 19         action_high: float = 1.,
689 20         reward_type: Literal['distance', 'disc'] | None = None,
690 21         **kwargs
691 22     ):
692 23         self.target_radius = target_radius
693 24         self._apply_noise: Callable[[torch.Tensor, torch.Tensor], torch.
694 25         Tensor] = (lambda x, noise: x * (1 + noise)) if noise_type ==
695 26         'multiplicative' else (lambda x, noise: x + noise)
696 27         self.action_transform = action_transform
697 28         # self.reward_type = 'disc' if reward_type is None else
698 29         reward_type
699 30         self.reward_type = 'distance'
700 31         super().__init__(effector, **kwargs)
701 32         self.efferent_noise = [efferent_noise]
702 33         self.effector = effector
703 34         assert isinstance(self.action_space, gym.spaces.Box), f'\
704 35         action_space is not appropriately initialized.'
705 36         self.action_space = gym.spaces.Box(low=action_low, high=
706 37         action_high, shape=self.action_space.shape, dtype=self.
707 38         action_space.dtype.type)

```

```

702 30
703 31     def get_vision(self) -> torch.Tensor:
704 32         """
705 33         Provides visual feedback to the neural network about the current
706 34         fingertip position using a normalized coordinate system.
707 35         Called by the inherited update_obs_buffer
708 36         """
709 37         return 2 * (self.states["fingertip"] - self.effector.xy_goal_min)
710 38         / (self.effector.xy_goal_max - self.effector.xy_goal_min) -
711 39         1
712 40
713 41     # def _get_vision(self):
714 42     #     self.dot_width = 0.05
715 43     #     self.num_pixels = 28
716 44     #     self.pixel_size = 2 / 28
717 45     #     xy = 2 * (self.states["fingertip"] - self.effector.xy_goal_min)
718 46     #     / (self.effector.xy_goal_max - self.effector.xy_goal_min) - 1
719 47     #     x = y = torch.linspace(-1, 1, self.num_pixels)
720 48     #     x = torch.tile(x, [self.num_pixels, 1])
721 49     #     y = torch.tile(y[:, None], [1, self.num_pixels])
722 50     #     locs = torch.stack([x, y], dim=-1)
723 51     #     return ((locs - xy) / (2 * self.dot_width))**2).sum(dim=-1).
724 52     exp().flatten() # modolo batches
725 53
726 54     def get_proprioception(self) -> torch.Tensor:
727 55     if not hasattr(self, '_mlen_idx'):
728 56         self._mlen_idx = self.muscle.state_name.index('muscle_length')
729 57     mlen = self.states["muscle"][:, self._mlen_idx:self._mlen_idx +
730 58         1] / self.muscle.l0_ce
731 59     if not hasattr(self, '_mvel_idx'):
732 60         self._mvel_idx = self.muscle.state_name.index('muscle_
733 61         velocity')
734 62     mvel = self.states["muscle"][:, self._mvel_idx:self._mvel_idx +
735 63         1] / self.muscle.vmax
736 64     return torch.cat([mlen, mvel], dim=-1).squeeze(dim=1)
737 65
738 66     def _build_spaces(self):
739 67     self.goal_delay = self.vision_delay
740 68     super()._build_spaces()
741 69
742 70     # The following properties are useful because the TorchRL libraries
743 71     don't all allow kwargs passing in reset() and step()
744 72     @property
745 73     def step_kwargs(self) -> dict[str, Any]:
746 74     if not hasattr(self, '_step_kwargs'):
747 75         self._step_kwargs = {}
748 76     return self._step_kwargs
749 77
750 78     @step_kwargs.setter
751 79     def step_kwargs(self, kwargs: dict[str, Any]) -> None:
752 80     self._step_kwargs = kwargs
753 81
754 82     @property
755 83     def reset_kwargs(self) -> dict[str, Any]:
756 84     if not hasattr(self, '_reset_kwargs'):
757 85         self._reset_kwargs = {}
758 86     return self._reset_kwargs
759 87
760 88     @reset_kwargs.setter
761 89     def reset_kwargs(self, kwargs: dict[str, Any]) -> None:
762 90     self._reset_kwargs = kwargs
763 91
764 92     @property
765 93     def th_random(self):

```

```

756 85         return self.effector.th_random
757 86
758 87     # The following properties are a hack that allow for noise sampling
759 88     with pytorch
760 89     @property
761 90     def obs_noise(self) -> torch.Tensor:
762 91         return self._obs_noise
763 92
764 93     @obs_noise.setter
765 94     def obs_noise(self, value: Sequence[float]) -> None:
766 95         self._obs_noise = torch.tensor(value, dtype=torch.float32, device
767 96         =self.device)
768 97
769 98     @property
770 99     def action_noise(self) -> torch.Tensor:
771 100         return self._action_noise
772 101
773 102     @action_noise.setter
774 103     def action_noise(self, value: Sequence[float]) -> None:
775 104         self._action_noise = torch.tensor(value, dtype=torch.float32,
776 105         device=self.device)
777 106
778 107     @property
779 108     def vision_noise(self) -> torch.Tensor:
780 109         return self._vision_noise
781 110
782 111     @vision_noise.setter
783 112     def vision_noise(self, value: Sequence[float]) -> None:
784 113         self._vision_noise = torch.tensor(value, dtype=torch.float32,
785 114         device=self.device)
786 115
787 116     @property
788 117     def efferent_noise(self) -> torch.Tensor:
789 118         return self._efferent_noise
790 119
791 120     @efferent_noise.setter
792 121     def efferent_noise(self, value: Sequence[float]) -> None:
793 122         self._efferent_noise = torch.tensor(value, dtype=torch.float32,
794 123         device=self.device)
795 124
796 125     @property
797 126     def proprioception_noise(self) -> torch.Tensor:
798 127         return self._proprioception_noise
799 128
800 129     @proprioception_noise.setter
801 130     def proprioception_noise(self, value: Sequence[float]) -> None:
802 131         self._proprioception_noise = torch.tensor(value, dtype=torch.
803 132         float32, device=self.device)
804 133
805 134     def apply_noise(self, loc, scale: torch.Tensor | Sequence[float] |
806 135         float):
807 136         scale = scale if isinstance(scale, torch.Tensor) else torch.
808 137         tensor(scale, dtype=torch.float32, device=self.device)
809 138         white_noise = torch.randn(size=loc.shape, generator=self.
810 139         th_random, dtype=torch.float32, device=self.device) * scale
811 140         return self._apply_noise(loc, white_noise)
812 141
813 142     def get_obs(self, action=None, deterministic: bool = False) -> torch.
814 143         Tensor | np.ndarray:
815 144         """
816 145         The parent version of this method applies noise twice, once when
817 146         saved to 'obs_buffer' and once
818 147         when retrieved from 'obs_buffer'. This version only applies noise
819 148         during retrieval and ignores
820 149         the 'obs_noise' parameter

```

```

810 138 """
811 139 self.update_obs_buffer(action=action)
812 140 no_op: Callable[[torch.Tensor, torch.Tensor]] = lambda x, y: x
813 141 f = no_op if deterministic else self.apply_noise
814 142 obs = (
815 143     [
816 144         self.obs_buffer["goal"][0] if deterministic else self.
            apply_noise(self.obs_buffer["goal"][0], self.
            vision_noise),
817         self.obs_buffer["vision"][0] if deterministic else self.
            apply_noise(self.obs_buffer["vision"][0], self.
            vision_noise),
818 145         self.obs_buffer["proprioception"][0] if deterministic
819         else self.apply_noise(self.obs_buffer["proprioception"]
820         "[0], self.proprioception_noise),
821         ] +
822         [
823 147         action if deterministic else self.apply_noise(action,
824 148         self.efferent_noise) for action in self.obs_buffer["
825 149         action"]
826         ]
827 150     )
828 151 obs = torch.cat(obs, dim=-1)
829 152 return obs if self.differentiable else self.detach(obs)
830 153
831 154 def step( # type: ignore
832 155     self,
833 156     action: torch.Tensor | np.ndarray | Sequence[float],
834 157     **kwargs
835 158 ):
836 159     kwargs = self.step_kwargs | kwargs
837 160     deterministic = kwargs.pop("deterministic", False)
838 161     self.elapsed += self.dt
839 162
840 163     action = action if isinstance(action, torch.Tensor) else torch.
            tensor(action, dtype=torch.float32, device=self.device)
841 164     if self.action_transform is not None:
842 165         action = self.action_transform(action)
843 166     noisy_action = action if deterministic else self.apply_noise(
            action, self.action_noise)
844 167
845 168     self.effector.step(noisy_action, **kwargs)
846 169
847 170     reward = self._get_reward(self.effector.batch_size)
848 171     obs = self.get_obs(action=action, deterministic=deterministic)
849 172
850 173     truncated = torch.zeros(reward.shape, dtype=torch.bool, device=
            self.device) if self.differentiable else np.zeros(reward.
            shape, dtype=bool)
851 174     terminated = ~truncated & (self.elapsed > self.max_ep_duration or
852 175     bool(np.isclose(self.elapsed, self.max_ep_duration)))
853 176
854 177     info = {
855 178         "states": self._maybe_detach_states(),
856 179         "action": action if self.differentiable else self.detach(
            action),
857 180         "noisy_action": noisy_action if self.differentiable else self.
            detach(noisy_action),
858         "goal": self.goal if self.differentiable else self.detach(
            self.goal),
859 181     }
860 182
861 183     return obs, reward, terminated, truncated, info
862 184
863 185

```

```

864 186 def _get_reward(self, batch_size: int) -> np.ndarray | torch.Tensor:
865     # type: ignore
866 187     error = (self.goal - self.states['fingertip']).norm(dim=-1,
867     keepdim=True) / self.target_radius
868 188     self._at_goal = error < 1
869 189     if self.reward_type == 'distance':
870 190         reward = torch.where(self._at_goal, 1 - error**2, 2 - 2 *
871         error) * self.dt
872 191     else:
873 192         reward = torch.where(self._at_goal, self.dt, 0.)
874 193     return reward if self.differentiable else self.detach(reward)
875 194
876 195 def reset(self, **kwargs) -> tuple[Any, dict[str, Any]]:
877 196     return super().reset(**(self.reset_kwargs | kwargs))

```

Listing 1: Training Environment

```

878
879 1 class RandomTargetReach(Environment): # % time on target
880 2     def _get_goal(self, batch_size: int) -> torch.Tensor:
881 3         return self.joint2cartesian(self.effector.draw_random_uniform_states(
882         batch_size)).chunk(2, dim=-1)[0]

```

Listing 2: Random Target Reach

```

883
884
885 1 class CenterOutReach(RandomTargetReach): # % time on target
886 2     def __init__(self, effector: ConstrainedGoalReluPointMass |
887     ConstrainedGoalRigidTendonArm, *args, **kwargs):
888 3         super().__init__(effector, *args, q_init=effector.center, **
889         kwargs)
890 4
891 5     @property
892 6     def cartesian_center(self):
893 7         if not hasattr(self, '_cartesian_center'):
894 8             joint_state = torch.cat((self.effector.center, torch.
895             zeros_like(self.effector.center)), dim=-1)
896 9             self._cartesian_center = self.joint2cartesian(joint_state).
897             chunk(2, dim=-1)[0]
898 10        return self._cartesian_center

```

Listing 3: Center Out Reach

```

899
900 1 class CenterOutAndBackWithJumpingGoalReach(CenterOutReach):
901 2     def _get_reward(self, batch_size: int):
902 3         reward = super()._get_reward(batch_size)
903 4         self.goal = torch.where(self._at_goal, self.cartesian_center,
904         self.goal)
905 5         return reward

```

Listing 4: COBJGR

```

906
907 1 class CenterOutAndBackWithFixedGoalReach(CenterOutReach):
908 2     def __init__(self, effector: ConstrainedGoalReluPointMass |
909     ConstrainedGoalRigidTendonArm, *args, **kwargs):
910 3         super().__init__(effector, *args, **kwargs)
911 4         self.reward_offset = 0
912 5         self.rewarded_goal = self.goal
913 6         self.previous_fingertip = None # Store the previous fingertip
914 7         self.goal_hold_counter = None # Counter for how long the goal
915         has been held, originally 0
916 8     def _get_reward(self, batch_size: int):
917 9         visual_goal = self.goal # Store the current goal temporarily
918 10        self.goal = self.rewarded_goal # Use the rewarded goal for
919        reward calculation

```

```

918 11         current_reward = super()._get_reward(batch_size)
919 12         at_goal = self._at_goal
920 13
921 14         reward = current_reward + self.reward_offset
922 15
923 16         # Initialize hold counter
924 17         if self.goal_hold_counter is None:
925 18             self.goal_hold_counter = torch.zeros((batch_size,1),device=
926 19                 self.device)
927 20
928 21         # Hold at the goal before switching to the return phase
929 22         self.goal_hold_counter = torch.where(at_goal, self.
930 23             goal_hold_counter+1,torch.zeros_like(self.goal_hold_counter))
931 24         # Increment the counter if at goal, reset otherwise
932 25         hold_steps = 5 # Define how many steps to hold at the goal before
933 26             switching to the return phase
934 27         self.switch_return = self.goal_hold_counter >= hold_steps #
935 28             Switch to return phase if hold steps are reached
936 29
937 30         # Adjust the rewarded goal if the effector reaches the current
938 31             goal & waits enough steps
939 32         self.goal = self.rewarded_goal = torch.where(self.switch_return,
940 33             self.cartesian_center, self.rewarded_goal)
941 34         temp_reward = super()._get_reward(batch_size) # Calculate reward
942 35             for the new goal
943 36         self.reward_offset = torch.where(self.switch_return,
944 37             current_reward - temp_reward, 0) + self.reward_offset
945 38         self.goal = visual_goal # Restore the original current goal
946 39         return reward
947 40
948 41     def step(self, *args, **kwargs):
949 42         obs, reward, terminated, truncated, info = super().step(*args, **
950 43             kwargs)
951 44         info['goal'] = self.rewarded_goal if self.differentiable else
952 45             self.detach(self.rewarded_goal)
953 46
954 47         # Calculate feedback indicating if the current goal matches the
955 48             rewarded goal
956 49         feedback = (self.rewarded_goal == self.goal).to(dtype=torch.float
957 50             )
958 51
959 52         # Taking mean
960 53         feedback = feedback.mean(dim=-1, keepdim=True) # Convert [
961 54             batch_size, 2] to [batch_size, 1]
962 55
963 56         # Concatenate feedback with observations
964 57         obs = torch.cat((obs, feedback), dim=1)
965 58
966 59         return obs, reward, terminated, truncated, info
967 60
968 61     def reset(self, *args, **kwargs):
969 62         obs, info = super().reset(*args, **kwargs)
970 63         self.rewarded_goal = self.goal
971 64         self.previous_fingertip = None # Reset previous fingertip
972 65             position
973 66         self.goal_hold_counter = None
974 67
975 68         # Concatenate initial feedback (set to 1) with observations
976 69         obs = torch.cat((obs, torch.ones_like(obs[:, :1])), dim=1)
977 70
978 71         return obs, info

```

Listing 5: COBFGR

```

972
973 1 class DelayedCenterOutReach(CenterOutReach):
974 2     def __init__(self, *args, delay: float = 0., **kwargs):
975 3         self.delay = delay
976 4         super().__init__(*args, **kwargs)
977 5     @property
978 6     def waiting(self):
979 7         return self.elapsed < self.delay or bool(np.isclose(self.elapsed,
980 8             self.delay))
981 9     def _get_reward(self, batch_size: int):
982 10         goal = self.goal
983 11         if self.waiting:
984 12             self.goal = self.cartesian_center.expand_as(goal)
985 13             reward = super()._get_reward(batch_size)
986 14             self.goal = goal
987 15             return reward
988 16     def step(self, *args, **kwargs):
989 17         obs, reward, terminated, truncated, info = super().step(*args, **
990 18             kwargs)
991 19         goal = self.cartesian_center.expand_as(self.goal) if self.waiting
992 20         else self.goal
993 21         info['goal'] = goal if self.differentiable else self.detach(goal)
994 22         return obs, reward, terminated, truncated, info
995 23     def reset(self, *args, **kwargs):
996 24         obs, info = super().reset(*args, **kwargs)
997 25         goal = self.cartesian_center.expand_as(self.goal) if self.waiting
998 26         else self.goal
999 27         info['goal'] = goal if self.differentiable else self.detach(goal)
1000 28         return obs, info

```

Listing 6: DCOR

```

999 1 class JumpingTargetReach(RandomTargetReach): #num targets hit
1000 2     def _get_reward(self, batch_size: int):
1001 3         current_reward = super()._get_reward(batch_size) # negative
1002 4         distance to the current goal
1003 5         reward = current_reward + self.reward_offset
1004 6         # save reward data at every timestep
1005 7         num_jumps = int(self._at_goal.sum())
1006 8         self.num_goals_reached += self._at_goal
1007 9         if num_jumps > 0:
1008 10             new_goals = self._get_goal(num_jumps)
1009 11             idx = self._at_goal.nonzero()[0, :].expand_as(new_goals)
1010 12             self.goal = self.goal.scatter(0, idx, new_goals)
1011 13             temp_reward = super()._get_reward(batch_size) # negative
1012 14             distance to the new goals
1013 15             self.reward_offset = current_reward - temp_reward + self.
1014 16             reward_offset
1015 17         return reward

```

Listing 7: JTR

```

1015
1016 1 import subprocess
1017 2 from time import sleep
1018 3 import os
1019 4 import time
1020 5 import torch
1021 6
1022 7 def get_checkpoint_info(model_path):
1023 8     """
1024 9     Extracts the checkpoint information from the model file.
1025 10     """
1026 11     try:
1027 12         checkpoint = torch.load(model_path, map_location='cpu')
1028 13         epochs_completed = checkpoint.get('epochs_run', 0)

```

```

1026     return epochs_completed
1027 except Exception as e:
1028     print(f"Error_reading_checkpoint_{model_path}:{e}")
1029     return 0
1030
1031 def launch_training(model_filename, effector_type, task, name, lr, epochs
1032 , batch_size, hidden_sizes, muscle_weights, hidden_act_weight,
1033 hidden_act_deriv_weight, debug=False, log_dir="logs11"):
1034     outstr = f"#!/bin/sh
1035
1036     #_Training_with_model_{model_filename}_on_task_{task}
1037     #
1038     #SBATCH--account=X
1039     #SBATCH--job-name={model_filename}_{task}#_The_job_name.
1040     #SBATCH--c_l#_The_number_of_cpu_cores_to_use
1041     #SBATCH--t_0-12:00#_Runtime_in_D-HH:MM
1042     #SBATCH--mem=32gb
1043     #SBATCH--o_{log_dir}/{task}_{name}.out#_Create_output_files_for_each_
1044     training_run
1045     #SBATCH--e_{log_dir}/{task}_{name}.err#_Create_error_files_for_each_
1046     training_run
1047     echo_Job_ID:_$SLURM_JOBID
1048
1049     #_module_load_anaconda
1050     source_activate_embodied_metalearning
1051     export_PYTHONUNBUFFERED=TRUE
1052     echo_"SLURM job running on: $(hostname)"
1053     echo_"Python path: $(which python)"
1054     echo_"Conda env: $CONDA_DEFAULT_ENV"
1055     python--version
1056
1057     python_retrain_models1.py\\
1058     --effector-type_{effector_type}\\
1059     --task_{task}\\
1060     --file_{model_filename}\\
1061     --epochs_{epochs}\\
1062     --batch-size_{batch_size}\\
1063     --hidden-sizes_{hidden_sizes}\\
1064     --lr_{lr}\\
1065     --load-file_{model_filename}
1066     """
1067
1068     script_filename = "tmp.sh"
1069     with open(script_filename, "w") as text_file:
1070         text_file.write(outstr)
1071
1072     subprocess.run(["sbatch", script_filename])
1073
1074     print(f"Script_saved_as_{script_filename}")
1075
1076 # Base model training epochs
1077 base_model_epochs = {
1078     'RandomTargetReach': 5000,
1079     'CenterOutReach': 5000,
1080     'CenterOutAndBackWithJumpingGoalReach': 5000,
1081     'CenterOutAndBackWithFixedGoalReach': 35000,
1082     'DelayedCenterOutReach': 35000,
1083     'JumpingTargetReach': 5000,
1084 }
1085
1086 task_parameters = {
1087     'RandomTargetReach': {
1088         'lr': 1e-3,
1089         'epochs': 5000,
1090         'batch_size': 128,

```

```

1080         'hidden_sizes': "32_16",
1081         'muscle-weight': 0,
1082         'hidden-act-weight': 100,
1083         'hidden-act-deriv-weight': 0.005
1084     },
1085     'CenterOutReach': {
1086         'lr': 1e-3,
1087         'epochs': 5000,
1088         'batch_size': 128,
1089         'hidden_sizes': "32_16",
1090         'muscle-weight': 0,
1091         'hidden-act-weight': 100,
1092         'hidden-act-deriv-weight': 0.005
1093     },
1094     'CenterOutAndBackWithJumpingGoalReach': {
1095         'lr': 1e-3,
1096         'epochs': 5000,
1097         'batch_size': 128,
1098         'hidden_sizes': "32_16",
1099         'muscle-weight': 0,
1100         'hidden-act-weight': 100,
1101         'hidden-act-deriv-weight': 0.005
1102     },
1103     'CenterOutAndBackWithFixedGoalReach': {
1104         'lr': 1e-4,
1105         'epochs': 35000,
1106         'batch_size': 512,
1107         'hidden_sizes': "128_16",
1108         'muscle-weight': 0,
1109         'hidden-act-weight': 100,
1110         'hidden-act-deriv-weight': 0.0005
1111     },
1112     'DelayedCenterOutReach': {
1113         'lr': 1e-3,
1114         'epochs': 35000,
1115         'batch_size': 128,
1116         'hidden_sizes': "32_16",
1117         'muscle-weight': 0,
1118         'hidden-act-weight': 100,
1119         'hidden-act-deriv-weight': 0.005
1120     },
1121     'JumpingTargetReach': {
1122         'lr': 1e-3,
1123         'epochs': 5000,
1124         'batch_size': 128,
1125         'hidden_sizes': "32_16",
1126         'muscle-weight': 0,
1127         'hidden-act-weight': 100,
1128         'hidden-act-deriv-weight': 0.005
1129     },
1130 },
1131
1132 def main():
1133     models = 'leftmodels' # Directory containing model files
1134     # tasks = [
1135     #     "RandomTargetReach",
1136     #     "CenterOutReach",
1137     #     "CenterOutAndBackWithJumpingGoalReach",
1138     #     "CenterOutAndBackWithFixedGoalReach",
1139     #     "DelayedCenterOutReach",
1140     #     "JumpingTargetReach"
1141     # ]
1142     log_dir = "logs11"

```

```

1134140 os.makedirs(log_dir, exist_ok=True)
1135141
1136142 # List of models
1137143 model_filenames = [f for f in os.listdir(models) if f.endswith('.pt')
1138144 ]
1139145 print(f"Found_{len(model_filenames)}_models.")
1140146
1141147 # Iterate through models and tasks
1142148 for model_filename in model_filenames:
1143149     model_path = os.path.join(models, model_filename)
1144150
1145151     #get the number of epochs completed from the model file
1146152     #epochs_completed = get_checkpoint_info(model_path)
1147153
1148154     effector_type = "arm" if "arm" in model_filename else "point-mass"
1149155
1150156     parts = model_filename.split('_over_')
1151157     target_task = parts[0].split('_')[1] # the task to be trained
1152158     name = parts[1].split('.pt')[0] # the name of the base model file
1153159     trained_task = name.split('_')[1] # base task
1154160
1155161     # Get the number of epochs the base model was trained on
1156162     base_model_trained_epochs = base_model_epochs.get(trained_task,
1157163                                                         0)
1158164
1159165     print(f"Continue_training_model_{model_filename}_on_task_{
1160166           target_task}_with_effector_type_{effector_type}")
1161167
1162168     params = task_parameters[target_task]
1163169
1164170     # Calculate total epochs (base model epochs + new task epochs)
1165171     total_epochs = base_model_trained_epochs + params["epochs"]
1166172
1167173     #left_epochs = total_epochs - epochs_completed
1168174
1169175     launch_training(
1170176         model_filename=model_filename,
1171177         effector_type=effector_type,
1172178         task=target_task,
1173179         name=name,
1174180         lr=params["lr"],
1175181         epochs= total_epochs, # Pass total epochs
1176182         batch_size=params["batch_size"],
1177183         hidden_sizes=params["hidden_sizes"],
1178184         muscle_weights=params["muscle-weight"],
1179185         hidden_act_weight=params["hidden-act-weight"],
1180186         hidden_act_deriv_weight=params['hidden-act-deriv-weight']
1181187     )
1182188
1183189     print(f"Training_command_sent.")
1184190
1185191 if __name__ == "__main__":
1186192     main()

```

Listing 8: Transfer Learning Hyperparameters

```

1182 1 def build_env(seed):
1183 2     env = make_motornet_env(differentiable=True, effector_type=
1184     effector_type, device=device, task=task, target_radius_frac=
1185     target_radius_frac, efferent_noise=noise_level, vision_noise=
1186     noise_level, proprioception_noise=noise_level, action_noise=
1187 3     noise_level)
1188 4     env._set_generator(seed)
1189     return env

```

```

1188 5
1189 6 def build_actor(obs_dim, action_dim):
1190 7     actor = GRUActor(obs_dim, action_dim, hidden_sizes, hidden0_sigma=
1191     h0_scale, positive_func=positive_func, separate_scale=
1192     separate_scale, scale_hidden_sizes=hidden_sizes[1:], manual_BG=
1193     rl_for_BG, device=device, dtype=dtype)
1194 8     if distributed:
1195 9         actor = DDP(actor)
1196 10    return actor
1197 11
1198 12 def build_critic(obs_dim):
1199 13     critic = GRUCritic(obs_dim, hidden_sizes, hidden0_sigma=h0_scale,
1200 14     device=device, dtype=dtype)
1201 15     if distributed:
1202 16         critic = DDP(critic)
1203 17    return critic
1204 18
1205 19 # %% [markdown]
1206 20 # Helper functions for saving and loading checkpoints. Model will be
1207 21 # loaded if 'load_file' exists. (Note: if command line option '--load-
1208 22 # file' is not used then 'load_file' will equal 'file'.)
1209 23
1210 24 # %%
1211 25 def _load(env, actor, critic, optimizer, scheduler):
1212 26     defaults = 0, {'losses': [], 'actor_losses': [], 'critic_losses': [],
1213     'hidden_unit_losses': [],
1214     'muscle_activation_losses': [], '
1215     avg_hidden_activity_derivative': [], 'tar_frac':
1216     []}
1217 27     if load_file is None:
1218 28         return defaults
1219 29     try:
1220 30         snapshot = th.load(load_file)
1221 31     except FileNotFoundError:
1222 32         return defaults
1223 33     print(f'Restoring_model_from_checkpoint_in_{load_file}.', file = sys.
1224 34     stdout)
1225 35
1226 36     env.load_state_dict(snapshot['env_dict'])
1227 37     (actor.module if distributed else actor).load_state_dict(snapshot['
1228 38     actor_dict'])
1229 39     (critic.module if distributed else critic).load_state_dict(snapshot['
1230 40     critic_dict'])
1231 41
1232 42     optimizer.load_state_dict(snapshot['optimizer_dict'])
1233 43     scheduler.load_state_dict(snapshot['scheduler_dict'])
1234 44
1235 45     if len(snapshot['torch_states']) != world_size:
1236 46         warn(f'Cannot_restore_the_random_state_because_number_saved_
1237 47         states_{len(snapshot["torch_states"])}_and_world_size_{
1238 48         world_size}.', category=RuntimeWarning)
1239 49     else:
1240 50         torch.set_rng_state(snapshot['torch_states'][rank])
1241 51         env.np_random.bit_generator.state = snapshot['numpy_env_states'][
1242 52         rank]
1243 53     if rank == 0:
1244 54         return snapshot['epochs_run'], snapshot['stats']
1245 55     else:
1246 56         return snapshot['epochs_run'], defaults[1]
1247 57
1248 58 def load_actor_weights(actor, snapshot):
1249 59     actor_dict = snapshot['actor_dict']
1250 60     model_dict = actor.state_dict()

```

```

1242 55 # Filter out mismatched layers
1243 56 filtered_actor_dict = {k: v for k, v in actor_dict.items() if k in
1244 model_dict and v.size() == model_dict[k].size()}
1245 57
1246 58 # Load the filtered weights
1247 59 (actor.module if distributed else actor).load_state_dict(
1248 filtered_actor_dict, strict=False)
1249 60 print(f"Actor_weights_partially_loaded_{len(filtered_actor_dict)}_
1250 layers_were_loaded_successfully.")
1251 61
1252 62 def load_critic_weights(critic, snapshot):
1253 63 critic_dict = snapshot['critic_dict']
1254 64 model_dict = critic.state_dict()
1255 65
1256 66 # Filter out mismatched layers
1257 67 filtered_critic_dict = {k: v for k, v in critic_dict.items() if k in
1258 model_dict and v.size() == model_dict[k].size()}
1259 68
1260 69 # Load the filtered weights
1261 70 (critic.module if distributed else critic).load_state_dict(
1262 filtered_critic_dict, strict=False)
1263 71 print(f"Critic_weights_partially_loaded_{len(filtered_critic_dict)}_
1264 layers_were_loaded_successfully.")
1265 72
1266 73 def _load_weights(env, actor, critic, optimizer, scheduler):
1267 74 print('I_am_inside_loading_weights', file=sys.stdout)
1268 75 defaults = 0, {'losses': [], 'actor_losses': [], 'tar_frac': []}
1269 76 if load_file is None:
1270 77     print('Loading_defaults', file=sys.stdout)
1271 78     return defaults
1272 79
1273 80 load_file_path = os.path.join("leftmodels", load_file)
1274 81 print(f'Loading_file_from_path:{load_file_path}', file=sys.stdout)
1275 82
1276 83 if not os.path.exists(load_file_path):
1277 84     print(f'File_{load_file_path}_does_not_exist.', file=sys.stdout)
1278 85     return defaults
1279 86
1280 87 try:
1281 88     snapshot = torch.load(load_file_path)
1282 89     print('Restoring_model_weights_from_checkpoint', file=sys.stdout)
1283 90 except Exception as e:
1284 91     print(f'Error_loading_file:{e}', file=sys.stdout)
1285 92     return defaults
1286 93
1287 94 # Load actor weights (with filtering)
1288 95 try:
1289 96     load_actor_weights(actor, snapshot)
1290 97 except Exception as e:
1291 98     print(f'Error_loading_Actor_weights_completely:{e}', file=sys.
1292 stdout)
1293 99
1294 100 # Load critic weights (with filtering)
1295 101 try:
1296 102     load_critic_weights(critic, snapshot)
1297 103 except Exception as e:
1298 104     print(f'Error_loading_Critic_weights_completely:{e}', file=sys.
1299 stdout)
1300 105
1301 106 return snapshot['epochs_run'], snapshot['stats']
1302 107
1303 108
1304 109 def _save(env, actor, critic, optimizer, scheduler, stats, run,
1305 current_epoch=None): #Command to save .pt file
1306 foldername = 'continue_models3'

```

```

1296 111 #if folder does not exist, create it
1297 112 if not os.path.exists(foldername):
1298 113     os.makedirs(foldername)
1299 114
1300 115 #filename = os.path.join(foldername, f'{run}_{task}_over_{file}')
1301 116 filename = os.path.join(foldername, f'{file}')
1302 117 torch_state = torch.get_rng_state()
1303 118 numpy_env_state = env.np_random.bit_generator.state
1304 119 if distributed:
1305 120     torch_states = [torch.zeros_like(torch_state) for _ in range(
1306 121         world_size)]
1307 122     torch.distributed.all_gather(torch_states, torch_state)
1308 123     numpy_env_states = [None] * world_size
1309 124     torch.distributed.all_gather_object(numpy_env_states,
1310 125         numpy_env_state)
1311 126 else:
1312 127     torch_states = [torch_state]
1313 128     numpy_env_states = [numpy_env_state]
1314 129 if rank == 0:
1315 130     snapshot = dict(
1316 131         env_dict=env.state_dict(),
1317 132         actor_dict=(actor.module if distributed else actor).
1318 133             state_dict(),
1319 134         critic_dict=(critic.module if distributed else critic).
1320 135             state_dict(),
1321 136         optimizer_dict=optimizer.state_dict(),
1322 137         scheduler_dict=scheduler.state_dict(),
1323 138         #epochs_run=len(stats['losses']),
1324 139         epochs_run=current_epoch + 1 if current_epoch is not None
1325 140         else len(stats['losses']),
1326 141         stats=stats,
1327 142         torch_states=torch_states,
1328 143         numpy_env_states=numpy_env_states,
1329 144     )
1330 145     th.save(snapshot, filename)
1331 146
1332 147 # %% [markdown]
1333 148 # This function updates the basal ganglia weights using the advantages
1334 149 when the command line option '--rl-for-BG' is used. If the command
1335 150 line option '--check-grad' is also used, then automatic
1336 151 differentiation is also used to make sure the gradient calculations
1337 152 are correct. Gradients are clipped to be less than 'max_norm' when
1338 153 that argument is not 'None'.
1339 154
1340 155 # %%
1341 156 def update_BG_weights(adv, ac, loc, scale, ctx, str, actor, lr=1e-3,
1342 157 max_norm=None, detach=False):
1343 158     for x in (ac, loc, scale, ctx, str):
1344 159         assert adv.shape == x[... :1].shape
1345 160     if detach:
1346 161         adv = adv.detach()
1347 162         ac = ac.detach()
1348 163         loc = loc.detach()
1349 164         scale = scale.detach()
1350 165         ctx = ctx.detach()
1351 166         str = str.detach()
1352 167
1353 168     a = ac
1354 169     c = ctx
1355 170     s = str
1356 171     d = torch.where(s > 0, 1., 0.)
1357 172     = loc
1358 173     = scale
1359 174
1360 175 actor = actor.module if distributed else actor

```

```

1350 cs = actor.bg.corticostratial[0]
1351 s = actor.bg.loc
1352 if not separate_scale:
1353     s = actor.bg.scale[0]
1354
1355 def apply_adv(w, two_d=False):
1356     tmp = ((adv[...], None] if two_d else adv) * w).mean(dim=(0, 1))
1357     if distributed:
1358         dist.all_reduce(tmp)
1359     return tmp / world_size
1360
1361 z = (a - ) /
1362 grads = {'cs': {}, 's ': {}}
1363 tmp_ = z /
1364 grads['s ']['bias'] = apply_adv(tmp_ )
1365 grads['s ']['weight'] = apply_adv(tmp_ .unsqueeze(-1) * s.unsqueeze
1366     (-2), two_d=True)
1367 tmp_s = tmp_ @ s .weight
1368 if not separate_scale:
1369     grads['s '] = {}
1370     d _over_ = 1 if positive_func == 'exp' else -(- ).expml() /
1371     tmp_ = -d _over_ * (1 - z**2)
1372     grads['s ']['bias'] = apply_adv(tmp_ )
1373     grads['s ']['weight'] = apply_adv(tmp_ .unsqueeze(-1) * s.
1374         unsqueeze(-2), two_d=True)
1375     tmp_s = tmp_s + tmp_ @ s .weight
1376 tmp_s = d * tmp_s
1377 grads['cs']['bias'] = apply_adv(tmp_s)
1378 grads['cs']['weight'] = apply_adv(tmp_s.unsqueeze(-1) * c.unsqueeze
1379     (-2), two_d=True)
1380
1381 if check_grad:
1382     assert not distributed, 'Gradient_checking_currently_only_works_
1383         for_single_process_training.'
1384     _params = {}
1385     sync = lambda source, target: target.load_state_dict(source.
1386         state_dict())
1387
1388     _cs = nn.Linear(cs.in_features, cs.out_features, device=device,
1389         dtype=dtype)
1390     _s = nn.Linear(s .in_features, s .out_features, device=device
1391         , dtype=dtype)
1392     sync(cs, _cs)
1393     sync(s , _s )
1394     if distributed:
1395         _cs, _s = DDP(_cs), DDP(_s )
1396     _params['cs'] = dict(weight=( _cs.module if distributed else _cs).
1397         weight, bias=( _cs.module if distributed else _cs).bias)
1398     _params['s '] = dict(weight=( _s .module if distributed else
1399         _s ).weight, bias=( _s .module if distributed else _s ).
1400         bias)
1401
1402 if not separate_scale:
1403     _s = nn.Linear(s .in_features, s .out_features, device=
1404         device, dtype=dtype)
1405     sync(s , _s )
1406     if distributed:
1407         _s = DDP(_s )
1408     _params['s '] = dict(weight=( _s .module if distributed else
1409         _s ).weight, bias=( _s .module if distributed else _s
1410         ).bias)
1411
1412 _c = c.detach()
1413 _a = a.detach()

```

```

1404216         _ = actor.bg.corticostriatal[1]
1405217         _s = _ (_cs(_c))
1406218         _ = _s (_s)
1407219         if not separate_scale:
1408220             _ = actor.bg.scale[1]
1409221             _ = _ (_s (_s))
1410222         else:
1411223             _ = _ .detach()
1412224
1413225         _pi = th.distributions.Normal(_ , _ )
1414226         _pi = th.distributions.Independent(_pi, 1)
1415227         _lp = _pi.log_prob(_a)
1416228         _adv = adv.detach()
1417229         (_adv[... , 0] * _lp).mean().backward()
1418230
1419231         if rank == 0:
1420232             for layer in grads:
1421233                 for type in ('weight', 'bias'):
1422234                     allclose = torch.allclose(_params[layer][type].grad,
1423235                         grads[layer][type], rtol=1e-4, atol=1e-7)
1424236                     if not allclose:
1425237                         print(f'Gradients_mismatch_for_{layer}.{type}._
1426238                             Max_error:_{(((_params[layer][type].grad-_
1427239                                 grads[layer][type])).abs()_/_ (1e-3+_grads[
1428240                                     layer][type].abs()))}.max():.3e}', file = sys.
1429241                                     stdout)
1430242
1431243         _norms = []
1432244         for layer in grads:
1433245             for type in ('weight', 'bias'):
1434246                 _norms.append(grads[layer][type].norm())
1435247         norm = torch.stack(_norms).norm()
1436248         info_str = ('Clipping_' if max_norm is not None and max_norm < norm
1437249             else '') + f'BG_grad_{norm:.0e},_'
1438250         if max_norm is not None and max_norm < norm:
1439251             for layer in grads:
1440252                 for type in ('weight', 'bias'):
1441253                     grads[layer][type] = grads[layer][type] * max_norm / norm
1442254
1443255         actor.bg.corticostriatal[0].update(lr * grads['cs']['weight'], lr *
1444256             grads['cs']['bias'])
1445257         actor.bg.loc.update(lr * grads['s ']['weight'], lr * grads['s ']['
1446258             bias'])
1447259         if not separate_scale:
1448260             actor.bg.scale[0].update(lr * grads['s ']['weight'], lr * grads[
1449261                 's ']['bias'])
1450262
1451263         return info_str + f"BG_lr:_{lr:.0e},_"
1452264
1453265 # %% [markdown]
1454266 # The function computes the generalized advantage estimate. In standard
1455267 RL, this uses 'rewards' from an environment rollout and 'values'
1456268 computed by the critic. This function alternatively accepts 'deltas',
1457269 which are analogous to the TD0 values in standard RL. In our case,
1458270 we will try to learn 'deltas' directly rather than compute them from
1459271 rewards and values.
1460272
1461273 # %%
1462274 def gae(adv_mat, values=None, rewards=None, deltas=None):
1463275     if deltas is None:
1464276         deltas = rewards + gamma * values[:, 1:] - values[:, :-1]
1465277     adv = (adv_mat * deltas).sum(-2)[..., None]
1466278     vtar = adv + values[:, :-1] if values is not None else None
1467279     return adv, vtar

```

```

1458 # %% [markdown]
1459 # The find_unique_goals function find the number of goals that appeared
1460 for each batch, where performance is measured as the fraction of the
1461 number of goals reached over the number of goals that appeared. This
1462 function is needed when training on the task "JumpingTargetReach".
1463 For this task, tar_frac does not accurately assess performance, as
1464 the effector could spend a very small fraction of time at each goal,
1465 but have perfect performance.
1466
1467 # %%
1468 def find_unique_goals(tg):
1469     # Initialize an empty tensor to store the counts
1470     num_unique_pairs_per_batch = torch.zeros((tg.size(0), 1), dtype=torch
1471 .int32)
1472
1473     # Iterate over each batch
1474     for i in range(tg.size(0)):
1475         reshaped_batch = tg[i].view(-1, tg.size(-1))
1476
1477         # Find and count unique goals
1478         unique_rows = torch.unique(reshaped_batch, dim=0)
1479         num_unique_pairs = unique_rows.size(0)
1480         num_unique_pairs_per_batch[i] = num_unique_pairs
1481
1482     return num_unique_pairs_per_batch
1483
1484 # %% [markdown]
1485 # This function generates 'batch_size' parallel rollouts of the
1486 environment using the current policy of 'actor'.
1487
1488 # %%
1489 def rollout(env, actor, batch_size):
1490     o, info = env.reset(options={"batch_size": batch_size})
1491     done = False
1492     h = None
1493     is_init = th.ones(batch_size, 1, dtype=bool, device=device)
1494
1495     ob = [o]
1496     xy = [info["states"]["fingertip"]]
1497     joint_states = [info["states"]["joint"]]
1498     tg = [info["goal"]]
1499     noisy_action = []
1500     ac = []
1501     lp = []
1502     rw = []
1503     loc = []
1504     scale = []
1505     ctx = []
1506     str = []
1507
1508     while not done: # will run until 'max_ep_duration' is reached
1509         pi, , , striatum, cortex, h = actor((o.to(dtype), is_init, h)
1510 )
1511         a = pi.rsample()
1512         # logp = pi.log_prob(a)
1513
1514         # compute forces acting on fingertip position
1515         # cartesian_fingertip_pos = env.effector.joint2cartesian(info["
1516 states"]["fingertip"])
1517         # force_vector = force_field_selector(cartesian_fingertip_pos,
1518 force_field_type, force_field_strength)
1519
1520         o, r, terminated, truncated, info = env.step(action=a.to(torch.
1521 float32), endpoint_load=th.tensor(force_field).unsqueeze(0))
1522         done = terminated.any() | truncated.any()

```

```

1512 ob.append(o)
1513 xy.append(info["states"]["fingertip"]) # trajectories
1514 joint_states.append(info["states"]["joint"]) # joint states
1515 tg.append(info["goal"]) # targets
1516 noisy_action.append(info["noisy_action"])
1517 ac.append(a)
1518 # lp.append(logp)
1519 rw.append(r)
1520 loc.append( )
1521 scale.append( )
1522 ctx.append(cortex)
1523 str.append(striatum)
1524
1525 ob = th.stack(ob, dim=1)
1526 xy = th.stack(xy, dim=1)
1527 joint_states = th.stack(joint_states, dim=1)
1528 tg = th.stack(tg, dim=1)
1529 noisy_action = th.stack(noisy_action, dim=1)
1530 ac = th.stack(ac, dim=1)
1531 # lp = th.stack(lp, dim=1)
1532 rw = th.stack(rw, dim=1)
1533 loc = th.stack(loc, dim=1)
1534 scale = th.stack(scale, dim=1)
1535 ctx = th.stack(ctx, dim=1)
1536 str = th.stack(str, dim=1)
1537 return ob, xy, joint_states, tg, noisy_action, ac, lp, rw, loc, scale
1538 , ctx, str
1539
1540 # %% [markdown]
1541 # This function runs and trains the model. Model training can take
1542 # multiple forms according to the command line options. The default is
1543 # supervised learning with no RL.
1544
1545 # %%
1546 def train_model():
1547     th.manual_seed(torch_seed + rank)
1548     bs = batch_size // world_size
1549     if world_size > 1:
1550         print(f'Worker_{rank}_with_batch_size_{bs}', file = sys.stdout)
1551
1552     env = build_env(np_seed + rank)
1553     actor = build_actor(env.observation_space.shape[0], env.action_space.
1554                        shape[0])
1555     critic = build_critic(env.observation_space.shape[0])
1556     optimizer = th.optim.AdamW(*actor.parameters(), *critic.parameters()
1557                               ), lr=lr)
1558
1559     scheduler = th.optim.lr_scheduler.CosineAnnealingLR(optimizer, epochs
1560                , eta_min=lr if lr_min is None else lr_min)
1561     if exp_state == 'train':
1562         epochs_run, stats = _load(env, actor, critic, optimizer,
1563                                scheduler)
1564     else:
1565         epochs_run, stats = _load_weights(env, actor, critic, optimizer,
1566                                scheduler)
1567         print('epochs_run_is_', epochs_run, file = sys.stdout)
1568         print('Epoch_length_is_', len(stats['losses']), file = sys.stdout)
1569
1570     # if exp_state == 'train':
1571     #     epochs_run, stats = _load(env, actor, critic, None, None)
1572     # else:
1573     #     epochs_run, stats = _load_weights(env, actor, critic, None, None)
1574 )

```

```

1566 # print('epochs run is ', epochs_run, file = sys.stdout)
1567
1568 # remaining_epochs = epochs - epochs_run
1569 # scheduler = th.optim.lr_scheduler.CosineAnnealingLR(optimizer,
1570 # remaining_epochs, eta_min=lr if lr_min is None else lr_min)
1571
1572 # # If we have a checkpoint, load the scheduler state
1573 # if exp_state == 'train' and epochs_run > 0:
1574 #     try:
1575 #         snapshot = th.load(load_file)
1576 #         scheduler.load_state_dict(snapshot['scheduler_dict'])
1577 #     except:
1578 #         print("Could not load scheduler state, using new scheduler")
1579
1580 adv_mat = th.tensor(
1581     circulant((gamma * lambda)**np.arange(int(round(env.
1582         max_ep_duration / env.dt)))),
1583     dtype=dtype,
1584     device=device
1585 ).tril()
1586 print('before_gt-trl', file = sys.stdout)
1587 get_rl_lr = lambda batch: rl_lr if rl_lr_min is None else (rl_lr -
1588     rl_lr_min) / 2 * (1 + np.cos(batch * np.pi / epochs)) + rl_lr_min
1589 print('About_to_enter_loop', file = sys.stdout)
1590 for batch in range(epochs_run-1, epochs):
1591
1592     optimizer.zero_grad()
1593     info_str = ''
1594
1595     # this is pretty ugly code...
1596     ob, xy, joint_states, tg, noisy_action, ac, lp, rw, loc, scale,
1597     ctx, str = rollout(env, actor, bs)
1598
1599     is_init = th.ones(bs, 1, 1, dtype=bool, device=device)
1600     vs, _ = critic((ob.to(dtype), is_init, None))
1601     adv, vtar = gae(adv_mat, values=vs, rewards=rw)
1602
1603     if rl_for_BG:
1604         if not auto_grad_through_rl:
1605             info_str += update_BG_weights(adv, ac, loc, scale, ctx,
1606                 str, actor, lr=get_rl_lr(batch), max_norm=
1607                 max_grad_norm, detach=True)
1608         else:
1609             info_str += update_BG_weights(adv, ac, loc, scale, ctx,
1610                 str, actor, lr=get_rl_lr(batch), max_norm=
1611                 max_grad_norm)
1612             ob, xy, tg, _, _, rw, *_ = rollout(env, actor, bs)
1613             actor.detach()
1614             vs, _ = critic((ob.to(dtype), is_init, None))
1615             _, vtar = gae(adv_mat, values=vs, rewards=rw)
1616
1617     # Loss function computation with hidden unit + muscle activity
1618     regularization
1619     actor_loss = -rw.sum(dim=1).mean()
1620     # critic_loss = smooth_l1_loss(vs[:, :-1], vtar.detach())
1621
1622     # hidden unit terms
1623     avg_hidden_activity = ((ctx**2).sum(dim=-1) / ctx.shape[-1]).
1624     mean()
1625     hidden_unit_loss = avg_hidden_activity
1626     aug_ctx = torch.cat((actor.gru.hidden0.expand(bs, 1, -1), ctx),
1627         dim=1)
1628     avg_hidden_activity_derivative = (((aug_ctx[:, 1:] - aug_ctx[:,
1629         :-1]) / env.dt)**2)/ctx.shape[-1]).sum(dim=-1)

```

```

1620428 # hidden_unit_loss = hidden_act_deriv_weight*
1621431     avg_hidden_activity_derivative).sum(dim=1).mean()
1622429 avg_hidden_activity_derivative_loss = (
1623430     avg_hidden_activity_derivative).sum(dim=1).mean()
1624431
1625432 # muscle activation terms
1626432 normed_max_iso_force = (env.effector.muscle.max_iso_force / (
1627433     torch.sqrt(torch.sum((env.effector.muscle.max_iso_force**2)))
1628433     )).squeeze(1)
1629434 muscle_activation_loss = (torch.matmul(noisy_action,
1630435     normed_max_iso_force.T)**2).squeeze(-1).sum(dim=1).mean()
1631435
1632436 # MotorNet Version - L2-norm is squared? (now denominator is an
1633436     inner product)
1634437 normed_max_iso_force = (env.effector.muscle.max_iso_force / (
1635438     torch.sum((env.effector.muscle.max_iso_force**2)))
1636439     ).squeeze(1)
1637439 # hidden_act_weight*hidden_act_deriv_weight
1638440
1639441 # NEW LOSS
1640442 loss = actor_loss + hidden_act_weight*hidden_unit_loss +
1641443     hidden_act_deriv_weight*avg_hidden_activity_derivative_loss +
1642444     muscle_weight*muscle_activation_loss
1643444
1644445 # backward pass & update weights
1645446 loss.backward()
1646447 norm = nn.utils.clip_grad_norm_((actor.parameters(), *critic.
1647448     parameters()), max_norm=max_grad_norm) # important!
1648449 info_str += ('Clipping_g' if max_grad_norm < norm else 'G') + f'
1649450     rad_{norm:.0e},_ '
1650451 optimizer.step()
1651452 scheduler.step()
1652452
1653453 if rank == 0:
1654454     if task == "JumpingTargetReach":
1655455         avg_num_unique_goals = find_unique_goals(tg).to(dtype).
1656456             mean().item()
1657457         num_goals_reached = env.num_goals_reached.to(dtype).mean
1658458             ().item()
1659459         tar_frac = num_goals_reached/avg_num_unique_goals
1660460     else:
1661461         tar_frac = ((xy - tg).norm(dim=-1) < env.target_radius).
1662462             to(dtype).mean().item()
1663463     stats['losses'].append(loss.item())
1664464     stats['actor_losses'].append(actor_loss.item())
1665465     # stats['critic_losses'].append(critic_loss.item())
1666466     stats['hidden_unit_losses'].append(hidden_unit_loss.item()) #
1667467         add in later
1668468     stats['muscle_activation_losses'].append(
1669469         muscle_activation_loss.item())
1670470     stats['avg_hidden_activity_derivative'].append(
1671471         avg_hidden_activity_derivative_loss.item())
1672472     stats['tar_frac'].append(tar_frac)
1673473
1674473 # Save checkpoint every 'save_every' epochs
1675474 if batch % SAVE_EVERY_N_EPOCHS == 0:
1676475     _save(env, actor, critic, optimizer, scheduler, stats,
1677476         run)
1678477     print(f'Saved_model_at_batch_{batch}', file=sys.stdout)
1679478
1680479 if batch % print_every == 0:
1681480     print('I_am_in_the_batch_loop_epochs_run_is_', epochs_run
1682481         , file = sys.stdout)
1683482     print('but_epochs_is_', epochs, file = sys.stdout)

```

```

1674 info_str = (
1675     f"Batch_{batch:5d}/{epochs:5d}, " +
1676     f"tot_loss:_{stats['losses'][-1]:.1e}, " +
1677     f"a_loss:_{stats['actor_losses'][-1]:.1e}, " +
1678     # f"v_loss: {stats['critic_losses'][-1]:.1e}, " +
1679     f"h_u_loss:_{stats['hidden_unit_losses'][-1]:.1e}, "
1680     +
1681     f"ma_loss:_{stats['muscle_activation_losses'][-1]:.1e}
1682     }, " +
1683     f"h_u_derivative:_{stats['
1684     avg_hidden_activity_derivative'][-1]:.1e}, " +
1685     f"goal_time:_{int(round(stats['tar_frac'][-1]*100))
1686     :2d}%, " +
1687     info_str +
1688     f"lr:_{optimizer.param_groups[0]['lr']:.0e}"
1689 )
1690 if 'cur' in locals():
1691     info_str += f', batch_time:_{(time()-cur)/_
1692     print_every:.2f}_s'
1693 cur = time()
1694 print(info_str, file = sys.stdout)
1695 # if file is not None and batch % save_every == 0:
1696     # _save(env, actor, critic, optimizer, scheduler, stats)
1697 return env, actor, critic, optimizer, scheduler, stats
1698
1699 # %% [markdown]
1700 # ## Main program
1701 #
1702 # Setup distributed computing if using.
1703
1704 # %%
1705 distributed = "LOCAL_RANK" in os.environ and "WORLD_SIZE" in os.environ
1706 world_size = int(os.environ["WORLD_SIZE"]) if distributed else 1
1707 rank = int(os.environ["LOCAL_RANK"]) if distributed else 0
1708
1709 # %% [markdown]
1710 # Extract the command line arguments and run the model training loop.
1711
1712 # %%
1713 print('Parser_begins', file = sys.stdout)
1714 parser = argparse.ArgumentParser(prog='random-reach-trainer', description
1715     ='Trains either a point mass with relu muscles or a 2d arm with hill
1716     muscles (specially \'RigidTendonHillMuscleThelen\' muscles) on the
1717     reach_to_random_target task')
1718
1719 parser.add_argument('--task', type=str, default='RandomTargetReach', help
1720     ='task_name_(default:_(default)s)')
1721 parser.add_argument('--exp_state', type=str, default='test', help='state_
1722     of_experiment:_training_or_testing_(default:_(default)s)')
1723 parser.add_argument('--epochs', type=int, default=3000, help='number_of_
1724     training_epochs_(default:_(default)s)')
1725 parser.add_argument('--batch-size', type=int, default=32, help='
1726     enviroment_rollouts_per_epoch_(default:_(default)s)')
1727 parser.add_argument('--effector-type', type=str, default='arm', choices=[
1728     'arm', 'point-mass'], help='type_of_effector_(default:_(default)s;_
1729     options:_(choices)s)')
1730 parser.add_argument('--target-radius-frac', type=float, default=0.05,
1731     help='tolerance_allowed_around_target_location_as_fraction_of_world_
1732     size_(default:_(default)s)')
1733 parser.add_argument('--noise-level', type=float, default=0.02, help='
1734     noise_value_that_is_used_for_efferent,_visual,_proprioceptive,_and_
1735     action_noise_(default:_(default)s)')
1736 parser.add_argument('--force-field', type=float, nargs=2, default=[0, 0],
1737     help='constants_that_determine_the_strength_of_the_force_applied_in_
1738     SingleConstantForceFieldHold task_(default:_(default)s)')

```

```

1728 parser.add_argument('--hidden-sizes', type=int, nargs=2, default=[32,
1729 16], help='size_of_the "cortex" and "striatum" in the model_BG (
1730 default:_%(default)s)')
1731 parser.add_argument('--h0-scale', type=float, default=0., help='scale_
1732 between_0 (inclusive) and_1 (exclusive) of the initial recurrent_
1733 activity_of the actor_and critic GRUs at the beginning_of training (
1734 default:_%(default)s)')
1735 parser.add_argument('--positive-func', type=str, default='exp', choices=[
1736 'exp', 'softplus'], help='function_used_to_ensure_positivity_of_the_
1737 scale_of_the_policy (default:_%(default)s; options:_%(choices)s)')
1738 parser.add_argument('--critic-weight', type=float, default=1., help='
1739 weight_of_critic_loss_in_the_total_loss (default:_%(default)s)')
1740 parser.add_argument('--muscle-weight', type=float, default=0.0, help='
1741 weight_of_muscle_activation_loss_in_the_total_loss (default:_%(
1742 default)s)')
1743 parser.add_argument('--hidden-act-weight', type=float, default=1000, help
1744 = 'weight_of_hidden_unit_activity_loss_in_the_total_loss (default:_%(
1745 default)s)')
1746 parser.add_argument('--hidden-act-deriv-weight', type=float, default=0.0,
1747 help='weight_of_derivative_of_hidden_unit_activity_in_the_hidden_
1748 unit_activity_loss (default:_%(default)s)')
1749 parser.add_argument('--lr', type=float, default=1e-3, help='initial_
1750 learning_rate (default:_%(default)s)')
1751 parser.add_argument('--lr-min', type=float, help='minimum_learning_rate_
1752 using_cosine_annealing (default:_no_annealing)')
1753 parser.add_argument('--max-grad-norm', type=float, default=1., help='
1754 Clipping_will_be_applied_if_gradient_norm_exceeds_this_value (default
1755 :_%(default)s)')
1756 parser.add_argument('--rl-for-BG', action='store_true', help='train_basal
1757 ganglia_weights_with_RL')
1758 parser.add_argument('--gamma', type=float, default=0.99, help='reward_
1759 discount_rate (default:_%(default)s)')
1760 parser.add_argument('--lambda', type=float, default=0.95, help='
1761 Generalized_advantage_estimator_parameter (default:_%(default)s)')
1762 parser.add_argument('--separate-scale', action='store_true', help='when_
1763 training_with_RL, exclude_the_pathway_which_computes_the_scale_of_the
1764 policy')
1765 parser.add_argument('--check-grad', action='store_true', help="check_the_
1766 manually_computed_policy_gradient_for_the_basal_ganglia_weights_with_
1767 pytorch's_auto_grad")
1768 parser.add_argument('--auto-grad-through-rl', action='store_true', help='
1769 differentiate_through_RL_dynamics (this_results_in_two_rollouts_per_
1770 epoch)')
1771 parser.add_argument('--rl-lr', type=float, default=1e-3, help='learning_
1772 rate_for_RL (default:_%(default)s)')
1773 parser.add_argument('--rl-lr-min', type=float, help='minimum_RL_learning_
1774 rate_using_cosine_annealing (default:_no_annealing)')
1775 parser.add_argument('--print-every', type=int, default=100, help='epochs_
1776 between_information_string_prints (default:_%(default)s)')
1777 parser.add_argument('--save-every', type=int, default=100, help='epochs_
1778 between_model_checkpoint_saves (default:_%(default)s)')
1779 parser.add_argument('--file', type=str, help='filename_for_loading_and_
1780 saving_checkpoints')
1781 parser.add_argument('--load-file', type=str, help='separate_file_for_
loading_checkpoints (default:_same_as_file)')
1782 parser.add_argument('--seed', type=float, default=0, help='random_seed (
default:_%(default)s)')
1783 parser.add_argument('--dtype', type=str, default='float32', choices=['
float32', 'float64'], help='floating_point_type (WARNING:_Motornet_
converts_to_float32_so_this_only_applies_to_actor_and_critic_networks
; default:_%(default)s; options:_%(choices)s)')
1784 parser.add_argument('--device', type=str, default='cpu', choices=['cpu',
'mps', 'cuda'], help='currently_tested_only_on "cpu", _try "mps" and_
"cuda" _at_your_own_risk (default:_%(default)s; options:_%(choices)s)')
1785 ;

```

```

1782 print('parser_ends', file = sys.stdout)
1783
1784 # %% [markdown]
1785 # The default behavior is to load and save from the filename given by the
1786 # '--file' option. The '--load-file' option
1787 # overrides this behavior such that loading and saving use different
1788 # files.
1789 #
1790 # If you load from a file and the number of epochs requested via the '--
1791 # epochs' option is less than or equal to the number in the file,
1792 # then the model from the file is loaded but no training occurs. This is
1793 # useful for plotting after training at the
1794 # command line.
1795 #
1796 # If you load from a file and the number of epochs requested is greater
1797 # than the number in the file, training will
1798 # resume from the loaded checkpoint. In this case you will need to make
1799 # sure that all the other options are the
1800 # same as during the prior training.
1801 #
1802 # 'args1', 'args2', 'args3', and 'args4' below correspond to the
1803 # following examples.
1804 #
1805 # Example 1: train from scratch with supervised learning and cosine
1806 # annealing.
1807 #
1808 # Example 2: load saved model which was trained for 3000 epochs with
1809 # supervised learning.
1810 #
1811 # Example 3: continue training from checkpoint after training was aborted
1812 # partway through. Note that after training completes, 'completed.pt'
1813 # will be the same as 'supervised.pt'. You can check this in the
1814 # terminal with the following command: 'python check-model-equivalence.
1815 # py supervised.pt completed.pt'.
1816 #
1817 # Example 4: load saved model which was trained for 3000 epochs with RL
1818 # for the basal ganglia. The file 'rl-for-BG.pt' was originally created
1819 # from the terminal using a distributed 8 CPU core training scheme: `
1820 $CONDA_PREFIX/bin/torchrun --standalone --nnodes=1 --nproc-per-node=8
1821 random-reach-task.py --epochs 3000 --batch-size 1024 --file rl-for-
1822 BG.pt --lr 1e-3 --lr-min 1e-5 --rl-for-BG --rl-lr .1 --rl-lr-min
1823 .001`.
1824 # %%
1825 if is_notebook():
1826     args1 = [
1827         '--task', 'RandomTargetReach',
1828         # '--file', 'RandomTargetReach_012_pointmass.pt',
1829         '--file', 'RandomTargetReach_arm.pt',
1830         '--epochs', '3000',
1831         '--batch-size', '128',
1832         # '--effector-type', 'point-mass',
1833         '--effector-type', 'arm',
1834         '--hidden-sizes', '32', '16',
1835         '--lr', '1e-3',
1836         '--muscle-weight', '0.0',
1837         '--hidden-act-weight', '1000.0',
1838         '--hidden-act-deriv-weight', '0.0',
1839         '--load-file', 'none'
1840     ]
1841
1842     args2 = [
1843         '--task', 'CenterOutReach',

```

```

1836      # '--file', 'CenterOutReach_pointmass.pt',
1837      '--file', 'CenterOutReach_arm.pt',
1838      '--epochs', '3000',
1839      '--batch-size', '128',
1840      # '--effector-type', 'point-mass',
1841      '--effector-type', 'arm',
1842      '--hidden-sizes', '32', '16',
1843      '--lr', '1e-3',
1844      '--muscle-weight', '0.0',
1845      '--hidden-act-weight', '1000.0',
1846      '--hidden-act-deriv-weight', '0.0',
1847      '--load-file', 'none',
1848  ]
1849
1850  args3 = [
1851      '--task', 'CenterOutAndBackWithJumpingGoalReach',
1852      # '--file', 'CenterOutAndBackWithJumpingGoalReach_pointmass.pt',
1853      '--file', 'CenterOutAndBackWithJumpingGoalReach_arm.pt',
1854      '--epochs', '3000',
1855      '--batch-size', '128',
1856      # '--effector-type', 'point-mass',
1857      '--effector-type', 'arm',
1858      '--hidden-sizes', '32', '16',
1859      '--lr', '1e-3',
1860      '--muscle-weight', '0.0',
1861      '--hidden-act-weight', '1000.0',
1862      '--hidden-act-deriv-weight', '0.0',
1863      '--load-file', 'none'
1864  ]
1865
1866  args4 = [
1867      '--task', 'CenterOutAndBackWithFixedGoalReach',
1868      # '--file', 'CenterOutAndBackWithFixedGoalReach_pointmass.pt',
1869      '--file', 'CenterOutAndBackWithFixedGoalReach_arm.pt',
1870      '--epochs', '3000',
1871      '--batch-size', '128',
1872      # '--effector-type', 'point-mass',
1873      '--effector-type', 'arm',
1874      '--hidden-sizes', '32', '16',
1875      '--lr', '1e-3',
1876      '--muscle-weight', '0.0',
1877      '--hidden-act-weight', '1000.0',
1878      '--hidden-act-deriv-weight', '0.0',
1879      '--load-file', 'none',
1880  ]
1881
1882  args5 = [
1883      '--task', 'DelayedCenterOutReach',
1884      # '--file', 'DelayedCenterOutReach_pointmass.pt',
1885      '--file', 'DelayedCenterOutReach_arm.pt',
1886      '--epochs', '3000',
1887      '--batch-size', '128',
1888      # '--effector-type', 'point-mass',
1889      '--effector-type', 'arm',
1890      '--hidden-sizes', '32', '16',
1891      '--lr', '1e-3',
1892      '--muscle-weight', '0.0',
1893      '--hidden-act-weight', '1000.0',
1894      '--hidden-act-deriv-weight', '0.0',
1895      '--load-file', 'none',
1896  ]
1897
1898  args6 = [
1899      '--task', 'JumpingTargetReach',
1900      # '--file', 'JumpingTargetReach_pointmass.pt',

```

```

1890         '--file', 'JumpingTargetReach_arm.pt',
1891         '--epochs', '3000',
1892         '--batch-size', '128',
1893         # '--effector-type', 'point-mass',
1894         '--effector-type', 'arm',
1895         '--hidden-sizes', '32', '16',
1896         '--lr', '1e-3',
1897         '--muscle-weight', '0.0',
1898         '--hidden-act-weight', '1000.0',
1899         '--hidden-act-deriv-weight', '0.0',
1900         '--load-file', 'none',
1901     ]
1902     args7 = [
1903         '--task', 'CenterHold',
1904         '--file', 'CenterHold_pointmass.pt',
1905         # '--file', 'CenterHold_arm.pt',
1906         '--epochs', '3000',
1907         '--batch-size', '128',
1908         '--effector-type', 'point-mass',
1909         # '--effector-type', 'arm',
1910         '--hidden-sizes', '32', '16',
1911         '--lr', '1e-3',
1912         '--load-file', 'none'
1913     ]
1914     arguments = [args1, args2, args3, args4, args5, args6]
1915     # args = parser.parse_args(args1)
1916     print('Running_in_IPython_mode.', file = sys.stdout)
1917     sys.stdout.flush()
1918 else:
1919     args = parser.parse_args()
1920     if rank == 0:
1921         print('Running_in_script_mode.', file = sys.stdout)
1922
1923     # %% [markdown]
1924     # Extract arguments and train model.
1925
1926     # %%%
1927     # arguments = [ # this has to come from SLURM script
1928     #     '--effector-type', 'point-mass',
1929     #     '--task', 'RandomTargetReach',
1930     #     '--file', 'RandomTargetReach_pointmass.pt'
1931     # ]
1932     # args = parser.parse_args(arguments)
1933     task = args.task
1934     print("The_task_inside_the_main_script_is_now", task, file = sys.stdout)
1935     exp_state = args.exp_state
1936     print('Exp_state_is_set_to', exp_state, file = sys.stdout)
1937     epochs = args.epochs
1938     print("The_epochs_inside_the_main_script_now_are", epochs, file = sys.
1939           stdout)
1940     batch_size = args.batch_size
1941     effector_type = args.effector_type
1942     print("The_effector_type_in_the_main_script_is_", effector_type, file = sys.
1943           stdout)
1944     target_radius_frac = args.target_radius_frac
1945     noise_level = args.noise_level
1946     force_field = args.force_field
1947     hidden_sizes = args.hidden_sizes
1948     h0_scale = args.h0_scale
1949     positive_func = args.positive_func
1950     critic_weight = args.critic_weight
1951     muscle_weight = args.muscle_weight
1952     hidden_act_weight = args.hidden_act_weight

```

```

1944717 print("The_hidden_act_weight_in_the_main_script_is", hidden_act_weight,
1945718       file = sys.stdout)
1946719 hidden_act_deriv_weight = args.hidden_act_deriv_weight
1947720 lr = args.lr
1948721 lr_min = args.lr_min
1949722 max_grad_norm = args.max_grad_norm
1950723 rl_for_BG = args.rl_for_BG
1951724 gamma = args.gamma
1952725 lmbda = getattr(args, 'lambda')
1953726 separate_scale = args.separate_scale
1954727 check_grad = args.check_grad
1955728 auto_grad_through_rl = args.auto_grad_through_rl
1956729 rl_lr = args.rl_lr
1957730 rl_lr_min = args.rl_lr_min
1958731 print_every = args.print_every
1959732 save_every = args.save_every
1960733 file = args.file
1961734 load_file = args.load_file if args.load_file is not None else file
1962735 torch_seed = args.seed
1963736 dtype = torch.float64 if args.dtype == 'float64' else torch.float32
1964737 device = torch.device(args.device)
1965738
1966739 np_seed = torch_seed + 42
1967740
1968741 for run in range(1):
1969742     print(f'Starting_{task}_for_{effector_type}_at_{run}_of_1', file =
1970743           sys.stdout)
1971744     env, actor, critic, optimizer, scheduler, stats = train_model()
1972745     _save(env, actor, critic, optimizer, scheduler, stats, run)
1973746
1974747 print('done', file = sys.stdout)

```

Listing 9: Transfer Learning Main Script

```

19731 class ConstrainedGoalReluPointMass(TorchRandomNumberGenerator,
19742                                     ReluPointMass24):
19753     """
19764     Bounds_the_range_of_allowable_goals_and_start_states.
19775     """
19786     def __init__(self, *args, xy_goal_min=None, xy_goal_max=None, **
19797                  kwargs):
19808         super().__init__(*args, **kwargs)
19819         self.xy_goal_min = torch.tensor(xy_goal_min, dtype=torch.float32,
198210                                         device=self.device)
198311         self.xy_goal_max = torch.tensor(xy_goal_max, dtype=torch.float32,
198412                                         device=self.device)
198513
198614     @property
198715     def center(self) -> torch.Tensor:
198816         if not hasattr(self, '_center'):
198917             self._center = 0.5 * (self.xy_goal_max - self.xy_goal_min) +
199018                                   self.xy_goal_min
199119         return self._center
199220
199321     def draw_random_uniform_states(self, batch_size):
199422         """Draws_joint_states_according_to_a_random_uniform_distribution,
199523         bounded_by_the_position_and_velocity_boundary
199624         attributes_defined_at_initialization.
199725
199826         Args:
199927         batch_size: Integer, the_desired_batch_size.
200028
200129         Returns:
200230         A_tensor_containing_batch_size_joint_states.

```

```

1998 25 """
1999 26 """
2000 27 sz = (batch_size, self.dof)
2001 28 rnd = torch.rand(size=sz, generator=self.th_random, dtype=torch.
2002 29 float32, device=self.device)
2003 30 pos = (self.xy_goal_max - self.xy_goal_min) * rnd + self.
2004 31 xy_goal_min
2005 32 vel = torch.zeros_like(pos)
2006 33 return torch.cat([pos, vel], dim=1)

```

Listing 10: Effector Class: Point Mass

```

2007 1 class ConstrainedGoalRigidTendonArm(TorchRandomNumberGenerator,
2008 2 RigidTendonArm26):
2009 3 """
2010 4 """Bounds_the_range_of_allowable_goals_and_start_states._This_requires_
2011 5 us_to_solve_an_inverse_kinematics_problem_(i.e.,
2012 6 going_from_cartesian_to_joint_space)_which_we_do_numerically.
2013 7 """
2014 8 def __init__(self, *args, xy_goal_min=None, xy_goal_max=None, c2j_n:
2015 9 None | int = None, **kwargs):
2016 10 super().__init__(*args, **kwargs)
2017 11 self.xy_goal_min = torch.tensor(xy_goal_min, dtype=torch.float32,
2018 12 device=self.device)
2019 13 self.xy_goal_max = torch.tensor(xy_goal_max, dtype=torch.float32,
2020 14 device=self.device)
2021 15 if c2j_n is not None:
2022 16 self._c2j_n = c2j_n
2023 17 self._c2j_xs = torch.linspace(self.xy_goal_min[0].item(),
2024 18 self.xy_goal_max[0].item(), c2j_n)
2025 19 self._c2j_ys = torch.linspace(self.xy_goal_min[1].item(),
2026 20 self.xy_goal_max[1].item(), c2j_n)
2027 21 self._c2j_dx = self._c2j_xs[1] - self._c2j_xs[0]
2028 22 self._c2j_dy = self._c2j_ys[1] - self._c2j_ys[0]
2029 23 xs_grid, ys_grid = torch.meshgrid(self._c2j_xs, self._c2j_ys)
2030 24 xy = torch.stack((xs_grid.flatten(), ys_grid.flatten())).T
2031 25 self._c2j_thetas = self.cartesian2joint(xy)
2032 26
2033 27 @property
2034 28 def center(self) -> torch.Tensor:
2035 29 if not hasattr(self, '_center'):
2036 30 self._center = self.cartesian2joint(0.5 * (self.xy_goal_max -
2037 31 self.xy_goal_min) + self.xy_goal_min[None, :])[0]
2038 32 return self._center
2039 33
2040 34 def draw_random_uniform_states(self, batch_size):
2041 35 """Draws_joint_states_according_to_a_random_uniform_distribution,
2042 36 bounded_by_the_position_and_velocity_boundary
2043 37 attributes_defined_at_initialization.
2044 38
2045 39 Args:
2046 40 batch_size: `Integer`, the_desired_batch_size.
2047 41
2048 42 Returns:
2049 43 A `tensor` containing `batch_size` joint_states.
2050 44 """
2051 45 max_sample_attempts = 1000
2052 46
2053 47 sz = (batch_size, self.dof)
2054 48
2055 49 # Attempt to sample valid joint positions
2056 50 # If the sampling fails, we will retry up to max_sample_attempts
2057 51 times
2058 52 # If we still cannot sample valid joint positions, we will raise
2059 53 an error
2060 54 for attempt in range(max_sample_attempts):

```

```

2052 44         rnd = torch.rand(size=sz, generator=self.th_random, dtype=
2053         torch.float32, device=self.device)
2054 45         xy_pos = (self.xy_goal_max - self.xy_goal_min) * rnd + self.
2055         xy_goal_min
2056 46         try:
2057 47             pos = self.cartesian2joint(xy_pos)
2058 48             vel = torch.zeros_like(pos)
2059 49             return torch.cat([pos, vel], dim=1)
2059 50         except AssertionError:
2060 51             if attempt == max_sample_attempts - 1:
2061 52                 raise RuntimeError(f'Failed_to_sample_valid_joint_
2062                 positions_after_{max_sample_attempts}_attempts._
2063                 Please_check_the_goal_bounds_and_the_
2064                 cartesian2joint()_implementation.')
2064 53             continue
2065 54
2066 55
2067 56 def cartesian2joint(self, pos: torch.Tensor) -> torch.Tensor:
2068 57     if self.dof != 2:
2069 58         raise NotImplementedError('We_require_2D')
2070 60
2071 61     if pos.requires_grad:
2072 62         raise NotImplementedError('We_currently_cannot_take_gradients_
2073         through_cartesian2joint()')
2074 63
2074 64     if hasattr(self, '_c2j_thetas'):
2075 65         i_x = (pos[:, :1] < self._c2j_xs).to(torch.int).argmax(dim
2076         =-1)
2077 66         i_y = (pos[:, 1:] < self._c2j_ys).to(torch.int).argmax(dim
2078         =-1)
2078 67         w_x = (self._c2j_xs[i_x] - pos[:, 0]) / self._c2j_dx
2079 68         w_y = (self._c2j_ys[i_y] - pos[:, 1]) / self._c2j_dy
2080 69         tmp0 = w_x[:, None] * self._c2j_thetas[(i_x - 1) * self.
2081         _c2j_n + i_y - 1] + (1 - w_x)[:, None] * self._c2j_thetas
2082         [i_x * self._c2j_n + i_y - 1]
2083 70         tmp1 = w_x[:, None] * self._c2j_thetas[(i_x - 1) * self.
2084         _c2j_n + i_y] + (1 - w_x)[:, None] * self._c2j_thetas[i_x
2085         * self._c2j_n + i_y]
2085 71         thetas = w_y[:, None] * tmp0 + (1 - w_y)[:, None] * tmp1
2086 72     else:
2087 73         cs: torch.Tensor = (self.skeleton.L2**2 - self.skeleton.L1**2
2088         - (pos**2).sum(dim=-1)) / (-2 * self.skeleton.L1)
2089 74
2089 75         theta_param = torch.zeros_like(pos[:, 0]).requires_grad_(True
2090         )
2091 76         lbfgs = torch.optim.LBFGS([theta_param], max_iter=100,
2092         line_search_fn='strong_wolfe')
2093 77         def closure():
2094 78             lbfgs.zero_grad()
2095 79             loss = ((pos[:, 0] * theta_param.cos() + pos[:, 1] *
2096             theta_param.sin() - cs)**2).sum()
2097 80             loss.backward()
2098 81             return loss
2099 82         lbfgs.step(closure=closure) # type: ignore
2100 83         theta0 = theta_param.detach()
2101 84
2101 85         theta1 = torch.arctan2(pos[:, 1] - self.skeleton.L1 * theta0.
2102         sin(), pos[:, 0] - self.skeleton.L1 * theta0.cos()) -
2103         theta0
2104 86         theta1 = (theta1 - self.skeleton.pos_lower_bound[1]) % (2 *
2105         torch.pi) + self.skeleton.pos_lower_bound[1]
2106 87
2106 88         thetas = torch.stack((theta0, theta1), dim=-1)

```

```
210689         assert (thetas >= self.skeleton.pos_lower_bound).all() and (  
2107             thetas <= self.skeleton.pos_upper_bound).all(), '  
2108             cartesian2joint()_optimization_failed._Sampled_joint_position  
2109             _is_out_of_the_allowed_range'  
211090         return thetas
```

Listing 11: Effector Class: Arm

2111
2112
2113
2114
2115
2116
2117
2118
2119
2120
2121
2122
2123
2124
2125
2126
2127
2128
2129
2130
2131
2132
2133
2134
2135
2136
2137
2138
2139
2140
2141
2142
2143
2144
2145
2146
2147
2148
2149
2150
2151
2152
2153
2154
2155
2156
2157
2158
2159