

HiPolicy: Hierarchical Multi-Frequency Action Chunking for Policy Learning

Jiyao Zhang^{1,2}, Zimu Han³, Junhan Wang¹, Xionghao Wu⁴, Shihong Lin⁵,
Jinzhou Li¹, Hongwei Fan^{1,2}, Ruihai Wu¹, Dongjiang Li⁶, and Hao Dong^{1,2}

¹ CFCS, School of CS, PKU, China

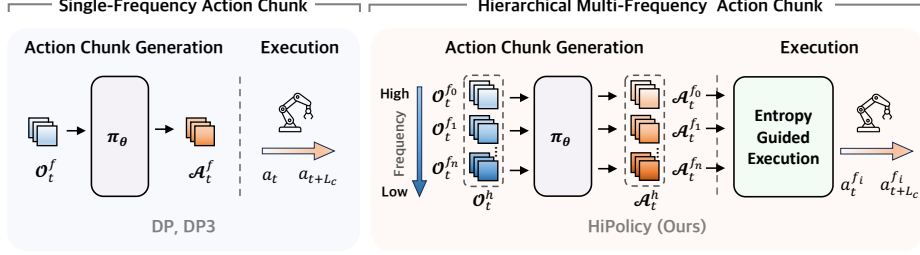
² National Key Laboratory for Multimedia Information Processing, School of CS, PKU, China

³ XJTU, China ⁴ THU, China ⁵ BUAA, China

⁶ Jingdong Technology Information Technology Co., Ltd, China
jiyaozhang@stu.pku.edu.cn

<https://hipolicy.github.io>

Abstract. Robotic imitation learning faces a fundamental trade-off between modeling long-horizon dependencies and enabling fine-grained closed-loop control. Existing fixed-frequency action chunking approaches struggle to achieve both. Building on this insight, we propose **HiPolicy**, a hierarchical multi-frequency action chunking framework that jointly predicts action



predicts action sequences at multiple frequencies from hierarchically aligned observation histories, fusing their representations to model long-horizon and fine-control behavior concurrently. To resolve the execution trade-off between speed and precision, we propose an action-entropy-guided adaptive execution strategy: Low-entropy frames indicate concentrated action distributions and stable predictions, prompting execution of high-frequency actions to enable precise closed-loop refinement. High-entropy frames reflect dispersed distributions and greater uncertainty, corresponding to broader phase-level decisions, where executing low-frequency actions aligns with high-level intent while increasing operational speed [10]. By dynamically selecting execution frequency based on policy uncertainty, our method speeds up the execution while balancing fine-grained control with long-horizon consistency.

Our contributions are threefold:

- We introduce HiPolicy, a novel hierarchical multi-frequency action chunking framework for policy learning that addresses the trade-off between long-horizon dependency modeling and high-frequency reactive control.
- We propose an action-entropy-guided adaptive execution mechanism that dynamically selects action frequency based on policy uncertainty, enhancing both robustness and execution efficiency.
- We validate our approach on extensive simulation benchmarks and real-world manipulation tasks, demonstrating consistent improvements in both performance and efficiency across 2D and 3D generative policies.

2 Related Work

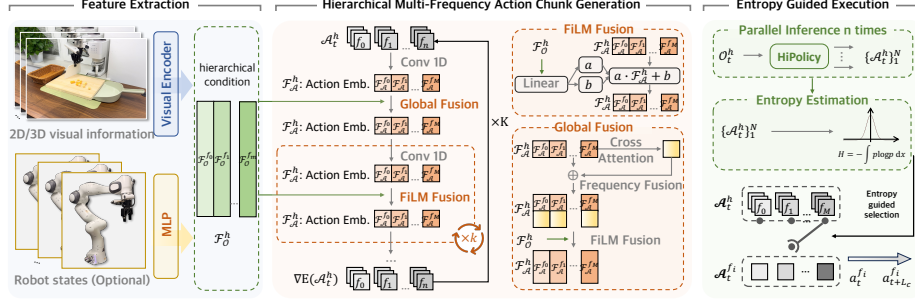
2.1 Imitation Learning for Manipulation

Imitation learning (IL) has become a dominant learning paradigm for robots to acquire manipulation skills from human demonstrations [

2.2 Hierarchical Manipulation Policy

Hierarchical structures in the area of computer vision can effectively perceive and integrate environmental information at different semantic levels, thereby enabling the model to take into account both global information and local details, and improving model performance.

In the robot learning community, recent research has sought to integrate hierarchical structures into imitation learning framework. [21, 23] H³DP [22] proposes a visuomotor learning framework that explicitly strengthens the coupling between visual features and action generation by hierarchically conditioning on multi-scale visual representations. Reactive Diffusion Policy [39] introduces a hierarchical slow-fast visual-tactile imitation learning algorithm that uses a slow



within the chunk, enabling coarse high-level planning and fine-grained reactive control to communicate and synergize.

HiPolicy Overview. As shown in Figure 3, our HiPolicy is designed as a diffusion-based model, predicting action chunks with hierarchical time resolutions simultaneously. First, we extract hierarchical observation features from visual and proprioception input. Then we obtain the overall action chunk by concatenating noisy action chunks at different frequencies and predict the noise through a 1D-CNN-based U-Net [28]. Hierarchical FiLM fusion [26] and global feature fusion are two key fusion modules designed to promote fusing observation features with action and injecting global information across different frequencies. After predicting action chunks at hierarchical temporal resolutions, we perform the inference in parallel to estimate the action entropy and

Algorithm 1: Entropy-Guided Execution

Input : Observation $\mathcal{O}_{t,L_h}^h$
Output: Action for Execution $\mathcal{A}_{t,L_c}^{\text{exec}}$
 $\mathcal{A}_{t,L_c}^h \leftarrow \emptyset$
for $i \leftarrow 1$ **to** N **do**
 $\mathcal{A}_{t,L_c}^i \leftarrow \pi_{\text{HiPolicy}}(\mathcal{O}_{t,L_h}^h)$ ▷ Independently sample for N

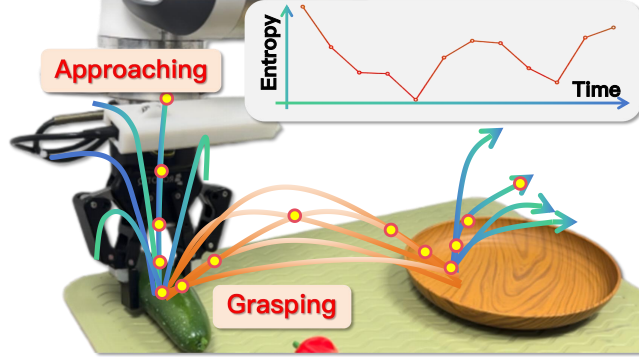


Fig. 4: Entropy-Guided Execution.</

Table 1: Detailed evaluation results on RoboTwin 1.0 and RoboTwin 2.0 Benchmark. Results (success rate) are based on 100 evaluation episodes, distinguishing between precise and other tasks. Bold red highlights the relative percentage improvement over baseline methods, while color shades are used to categorize tasks by their horizon (short, long, and super long). Tasks are sorted according to demo steps.

Task name	Task Source	Precise	Demo steps	DP [4]	DP+HiPolicy	DP3 [40]	DP3+HiPolicy
<i>Blocks stack hard</i>							

To highlight HiPolicy’s ability to jointly model long-horizon dependencies and high-precision closed-loop controls shown in Table 1, tasks are further categorized, and we present various temporal dependency profiles and levels of control granularity (precise tasks or non-precise tasks). Precise tasks require tight end-effector pose tolerances to succeed (e.g., cube stacking, QR-code scanning), while non-precise tasks have more lenient success criteria.

All demo data was collected according to the methods provided by the official. Details of the selected tasks, demo collection, and preprocessing are provided in Appendix A.

Baselines. We compare our method against two representative generative imitation learning policies, selected

Table 2: Ablation Experiment. In the RoboTwin 1.0 platform, we compare our complete HiPolicy with that removing the hierarchical condition, including all conditions only with the highest frequency (High) or the lowest (Low) frequency, the fusion module, and the hierarchical frequency structure (Hier.), respectively.

Task name	Condition			w/o Fusion w/o Hier.	
	Hier.	Low	High		

Table 3: Acceleration Experiment. We measure the success rate (SR) and the execution steps for DP baseline, DP+HiPolicy without entropy-guided (EG) execution, and DP+HiPolicy across 8 tasks from the RoboTwin 2.0 platform. We use red to highlight our improvement against DP in the success rate and execution speed.

Task name	DP [4]		DP+HiPolicy w/o EG		DP+HiPolicy	
	SR					

Table 5: Runtime and memory comparison between DP and DP with HiPolicy on an RTX 4090 GPU.

Method	VRAM↓	Infer.↓	Steps↓	Wall-clock↓	SR↑
DP	6560MB	100ms	133	10.57s	24%
DP + HiPolicy	7457MB	108ms	100 (25%↓)	8.07s (24%↓)	41% (+

Table 6: Detailed evaluation results on 8 real-world tasks. Our results are reported based on an average of 10 evaluation trials per task, each conducted with randomized object initializations to ensure statistical robustness. And the bold red color is used to highlight the relative percentage improvement over baseline methods. In *close microwave oven*, we respectively report the execution length of closing the door to nearly-closed position and closing the door tightly.

Task name	Success Rate↑		Avg. Execution Steps↓	
	DP [4]	DP+HiPolicy		

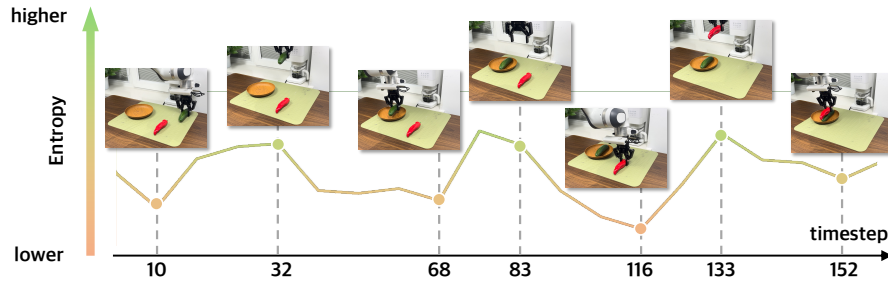


Fig. 6: Action entropy curve

References

1. Andrychowicz, M., Raichuk, A., Stanek, P., Shipley, T., Debiak, H., Chociej, S., Durkaya, M., Sherborne, R., Trzciński, B., Tabaka, M., et al.: Learning dexterous in-hand manipulation. arXiv preprint arXiv:1808.00177 (2018)
2. Black, K., Brown, N., Driess, D., Esmail, A., Equi, M., Finn, C., Fusai, N., Groom, L., Hausman, K., Ichter, B., Jakubczak, S., Jones, T., Ke

- Black, K., Chi, C., Hatch, K.B., Lin, S., Lu, J., Mercat, J., Rehman, A., Sanketi, P.R., Sharma, A., Simpson, C., Vuong, Q., Walke, H.R., Wulfe, B., Xiao, T., Yang, J.H., Yavary, A., Zhao, T.Z., Agia, C., Baijal, R., Castro, M.G., Chen, D., Chen, Q., Chung, T., Drake, J., Foster, E.P., Gao, J., Guizilini, V., Herrera, D.A., Heo, M., Hsu, K., Hu, J., Irshad, M.Z., Jackson, D., Le, C., Li, Y., Lin, K., Lin, R., Ma, Z., Maddukuri, A., Mirchandani, S., Morton, D., Nguyen, T., O'Neill, A., Scalise, R., Seale, D., Son, V

29. Ross, S., Gordon, G., Bagnell, D.: A reduction of imitation learning and structured prediction to no-regret online learning. In: Proceedings of the fourteenth international conference on artificial intelligence and statistics. pp. 627–635. JMLR Workshop and Conference Proceedings (2011)
30. Shannon, C.E.: A mathematical theory of communication. The Bell System Technical Journal **27**(3), 379–423 (Jul 1948). <https://doi.org/10.1002/j.1538-7305.1948.tb01338.x>, part I of two parts.
- 3