

## Problem

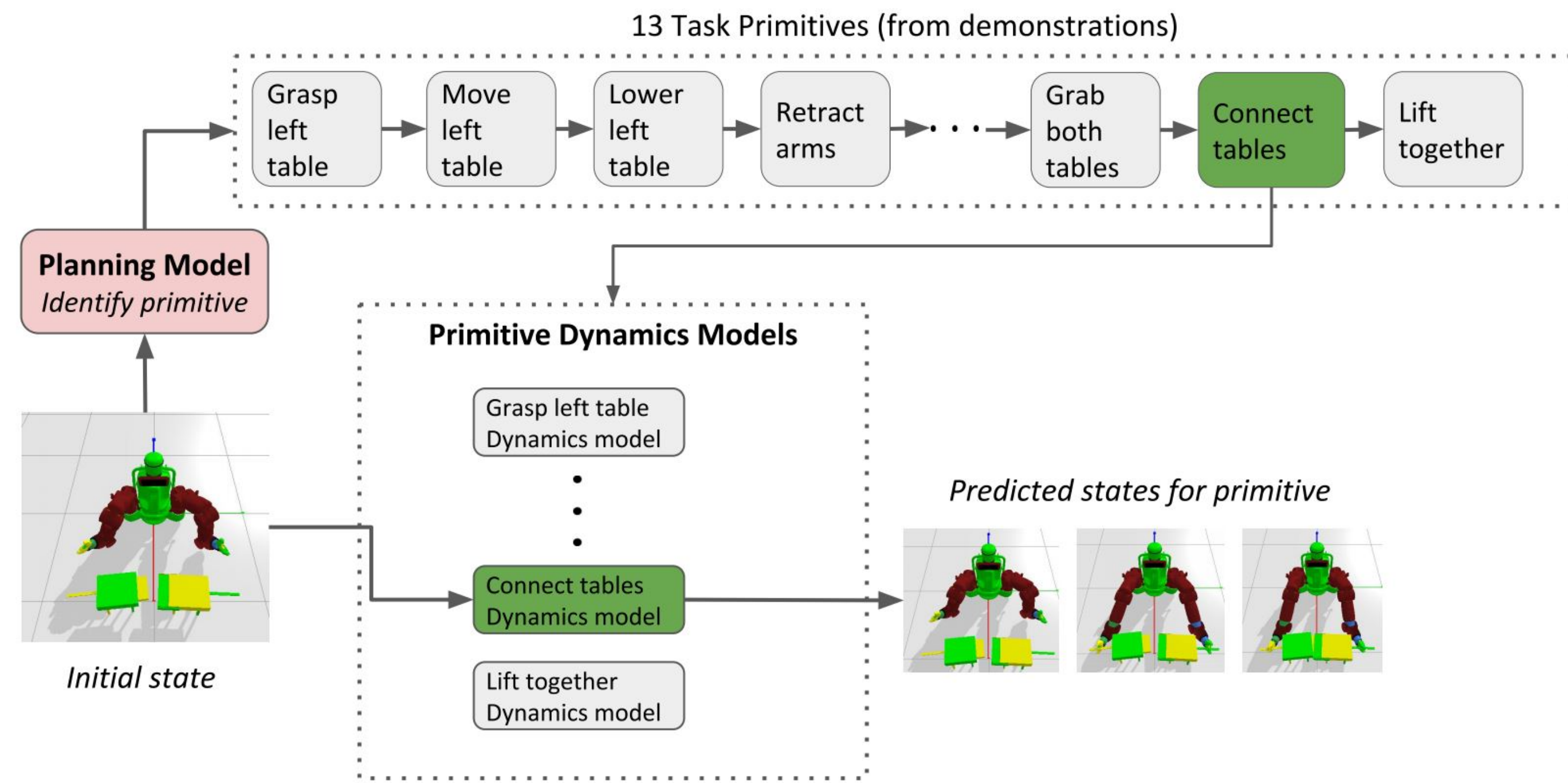
How to design deep imitation learning [2] strategies such that the learned skills can generalize to different object locations in robotic bimanual manipulation [1]?

## Our Results

We propose a hierarchical deep relational imitation learning model (HDR-IL). Our model generalizes better and achieves significantly higher success rates on two bimanual robotic table lifting experiments in simulation. To achieve this, our model:

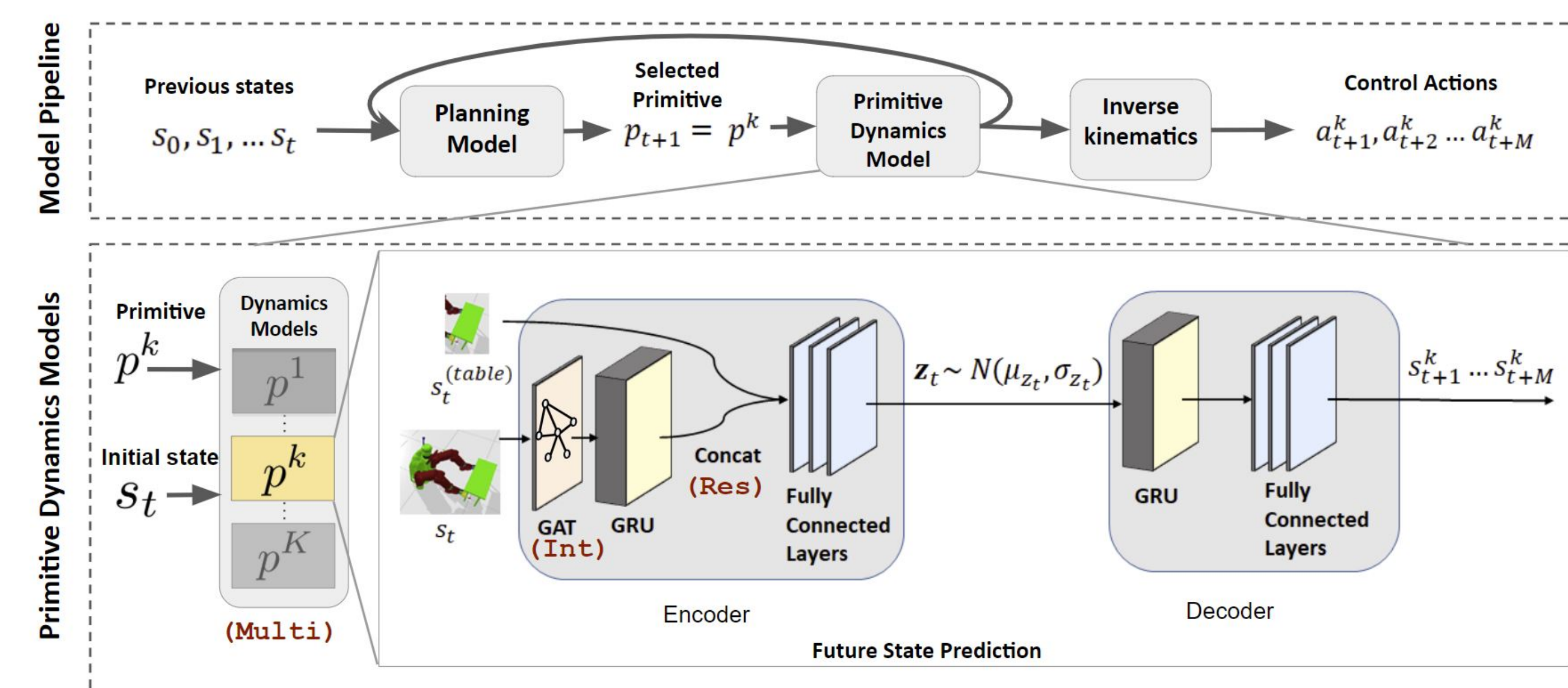
- Decomposes multi-modal dynamics into elemental movement primitives;
- Capture primitive dynamics in low-level dynamics models using graph neural networks to model interactions and residual connections to emphasize goal state;
- Integrate a high-level planner that composes low-level primitives dynamics sequentially to produce trajectories;
- We open source the code for simulation, data, and models at:  
Code: <https://github.com/Rose-STL-Lab/HDR-IL>

## Hierarchical Modular Model Architecture



The hierarchical structure of our HDR-IL model. The high-level planning model selects the next primitive in the sequence, and the corresponding dynamics model is used to predict the next trajectory.

## Model Features in Dynamics Models



- **(Multi)** Multiple dynamics models, which are selected by the high-level planning model.
- **(Int)** Graph attention layer (GAT) [3] in the graph neural network capture interactions between objects in the environment
- **(Res)** Residual connection of the object state helps direct the projection towards the target.

## Task #1: Table Lifting

**Task:** Lift a 35cm × 85cm table onto a platform. The location of the table varies between demonstrations.

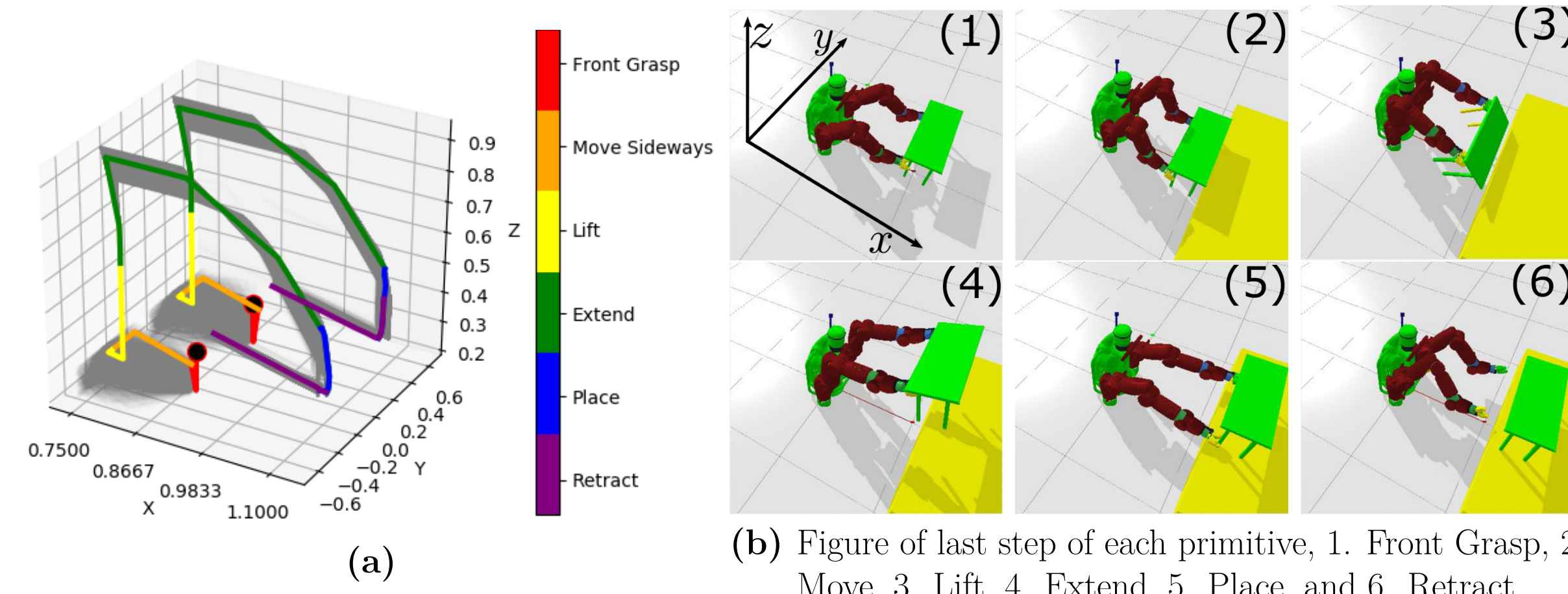
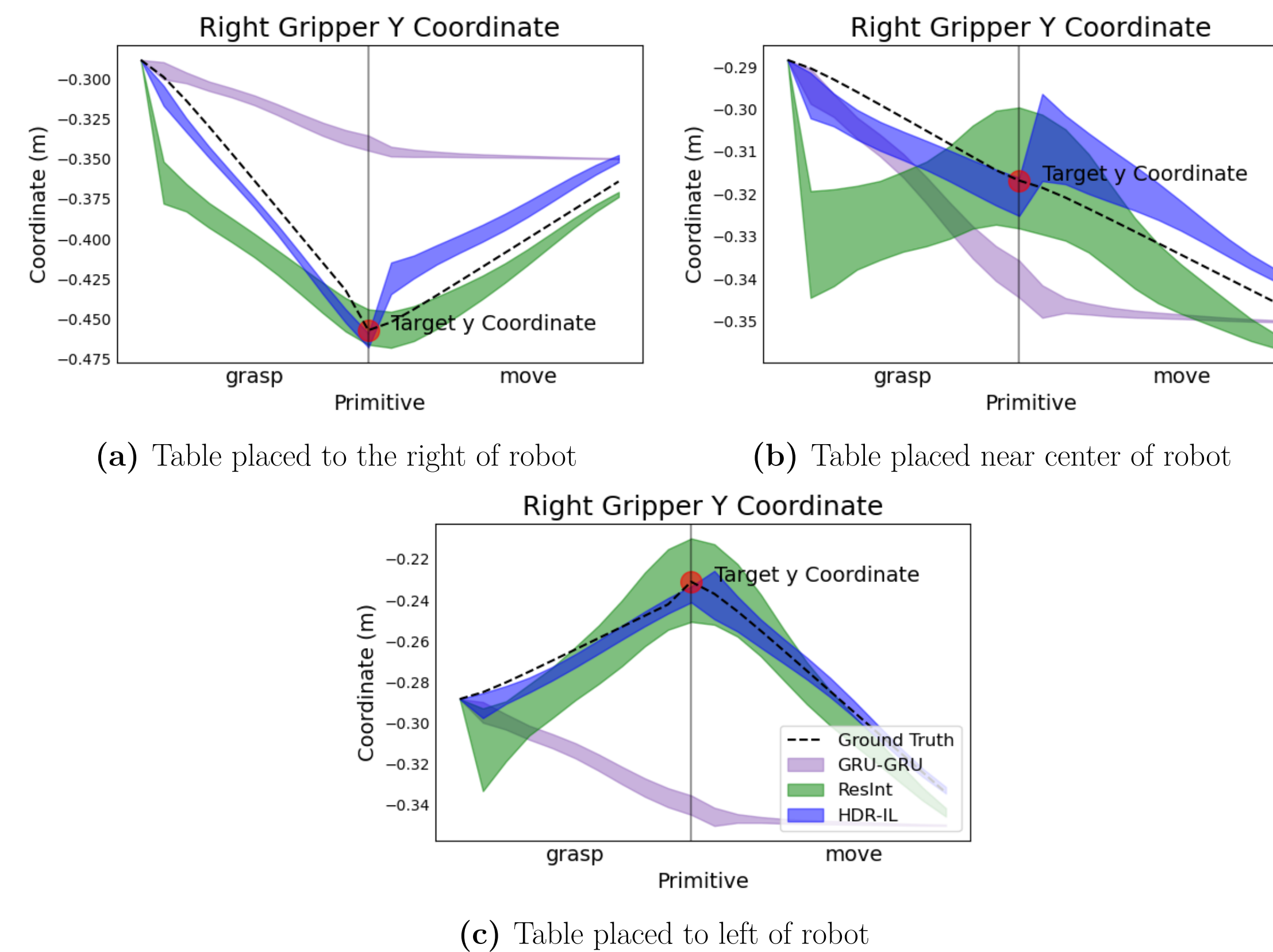


Figure (a) shows training trajectories in  $(x, y, z)$  of the left and right grippers for 2500 demonstrations. One sample trajectory is shown in color to highlight the trajectory for each primitive, while the rest are grey. The black dot is the starting location. Figure (b) shows a snapshot of each primitive.

**Prediction Performance:** Comparison of model performance by test error and by percent success in 127 test simulations. The errors are measured on the left and right grippers, and the range represents one standard deviation. The % success is calculated by running projections through simulation. The ablation study show graph and skip connections together improve success rates in single and multi-model designs.

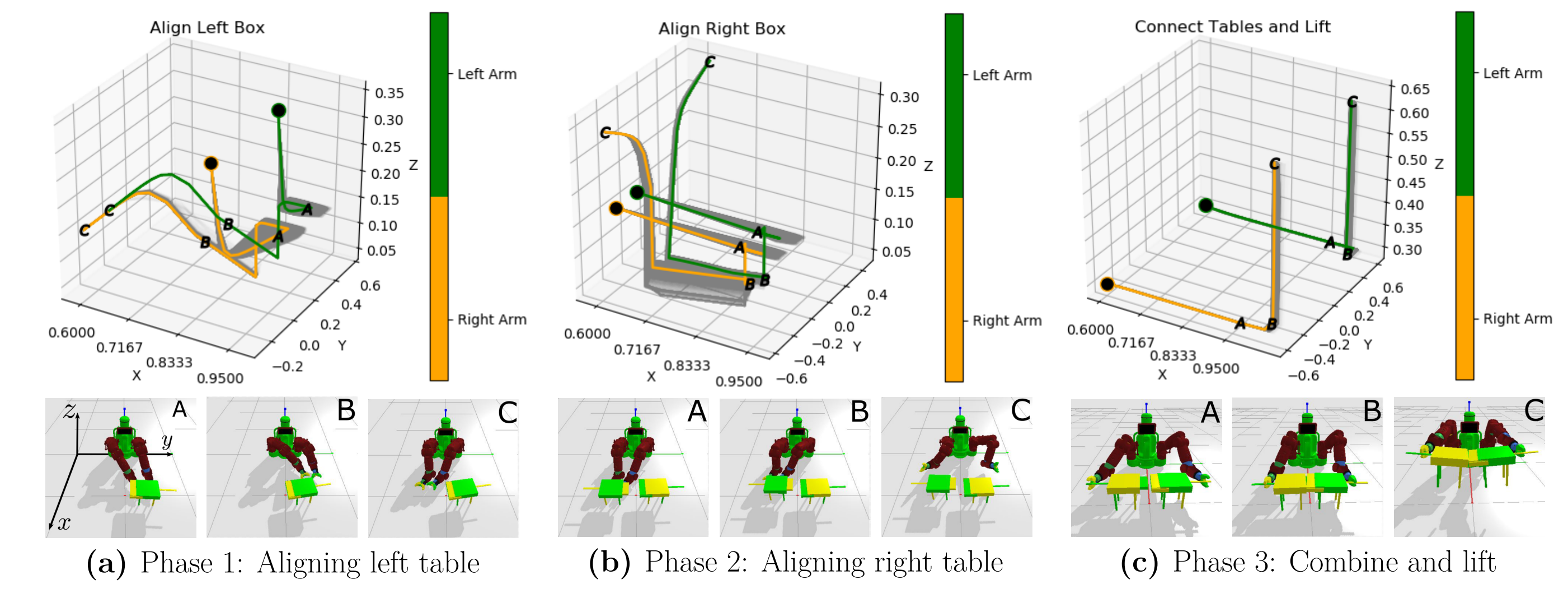
| Model         | Graph | Skip | Conn | Multi-Model | Euclidean Dist     | Angular Dist      | DTW Dist     | % Success   |
|---------------|-------|------|------|-------------|--------------------|-------------------|--------------|-------------|
| GRU-GRU       |       |      |      |             | $6.53 \pm 7.05$    | $0.139 \pm 0.182$ | 0.135        | 13%         |
| Res           |       | ✓    |      |             | $7.74 \pm 5.88$    | $0.143 \pm 0.194$ | 0.124        | 13%         |
| Int           | ✓     |      |      |             | $6.67 \pm 5.80$    | $0.145 \pm 0.177$ | 0.123        | 17%         |
| ResInt        | ✓     | ✓    |      |             | $5.64 \pm 5.17$    | $0.121 \pm 0.205$ | 0.128        | 72%         |
| GRU-GRU Multi |       |      |      | ✓           | $6.53 \pm 7.05$    | $0.139 \pm 0.182$ | 0.131        | 14%         |
| Res Multi     |       | ✓    |      | ✓           | $4.97 \pm 5.83$    | $0.123 \pm 0.191$ | 0.121        | 92%         |
| Int Multi     | ✓     |      |      | ✓           | $11.69 \pm 10.142$ | $0.246 \pm 0.269$ | 0.126        | 13%         |
| HDR-IL        | ✓     | ✓    |      | ✓           | $5.01 \pm 5.33$    | $0.112 \pm 0.208$ | <b>0.119</b> | <b>100%</b> |

**Prediction Visualization:** Sample Y coordinate predictions for 3 test samples in the table lifting task. The robot gripper reaches for the table leg, denoted by the "Target", whose location is randomized for each demonstration. The HDR-IL shows the best generalization compared to single model ResInt and the GRU-GRU baseline.



## Task #2: Peg-In-Hole Lifting

**Task:** Connect two halves of a table before lifting. The locations of the two halves varies between demonstrations.

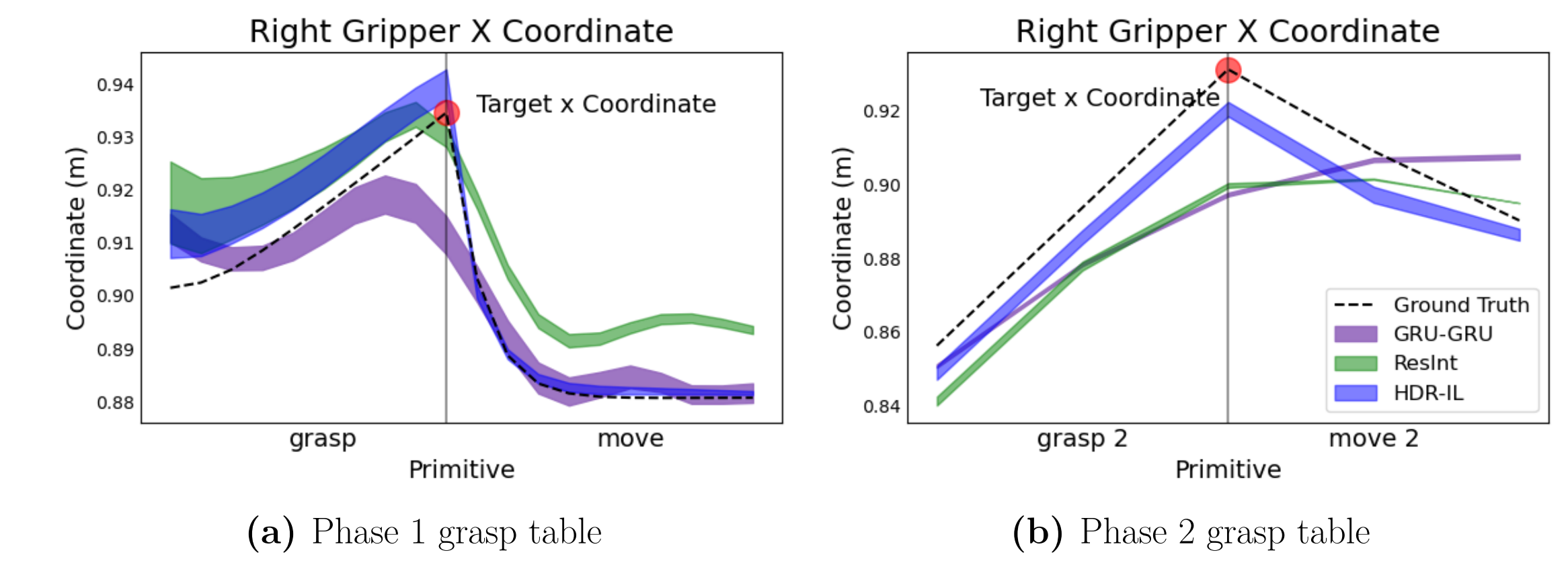


From left to right, these figures sequentially show  $(x, y, z)$  trajectories of the left and right grippers for 4700 demonstrations. The black dot is the starting location. The figure corresponds to snapshots along the demonstration.

**Prediction Performance:** Comparison of model performances in the peg-in-hole task by test error and percent success in 281 test simulations. The errors are measured as the average of the left and right grippers, and the range represents one standard deviation. The % success is calculated by running projections through simulation.

| Model   | Eculclidean Dist | Angular Dist      | DTW Dist     | % Success - Lift |
|---------|------------------|-------------------|--------------|------------------|
| GRU-GRU | $2.11 \pm 1.11$  | $0.029 \pm 0.022$ | 0.121        | 1%               |
| ResInt  | $1.59 \pm 0.83$  | $0.024 \pm 0.012$ | 0.117        | 15%              |
| HDR-IL  | $0.90 \pm 1.01$  | $0.013 \pm 0.010$ | <b>0.113</b> | <b>29%</b>       |

**Prediction Visualization:** Sample X coordinate predictions comparison for two phases: (a) The beginning of the projection. (b) 60 time steps into the projection. The robot gripper reaches for the table leg, denoted by the "Target", whose location is randomized for each demonstration. Accuracy decreases in the later generalizations, resulting in lower success rates.



## References

- [1] R. Chitnis, S. Tulsiani, S. Gupta, and A. Gupta. Efficient bimanual manipulation using learned task schemas, 2020.
- [2] T. Kipf, Y. Li, H. Dai, V. Zambaldi, A. Sanchez-Gonzalez, E. Grefenstette, P. Kohli, and P. Battaglia. Compile: Compositional imitation learning and execution, 2019.
- [3] P. Velićković, G. Cucurull, A. Casanova, A. Romero, P. Liò, and Y. Bengio. Graph attention networks. In *International Conference on Learning Representations*, 2018.

**Acknowledgments** This work was supported in part by U. S. Army Research Office under Grant W911NF-20-1-0334. Additional revenues related to this work: NSF #1850349, ONR # N68335-19-C-0310, Google Faculty Research Award, Adobe Data Science Research Awards, GPUs donated by NVIDIA, and computing allocation awarded by DOE.