

# Machine Learning for Robotic Napkin Folding

Joshua Himmens<sup>1</sup>, Sloan Sobie<sup>2</sup>, Dawson March<sup>2</sup>, Genevieve Merz<sup>3</sup>, Cameron Powell<sup>3</sup>, Jaden Legate<sup>3</sup>

<sup>1</sup>UBC Engineering Physics, <sup>2</sup>UBC Mechanical Engineering, <sup>3</sup>UBC Sauder School of Business



THE UNIVERSITY OF BRITISH COLUMBIA

Engineering Physics

Faculty of Applied Science | Department of Physics and Astronomy

## Motivation

Robotics and robotics control are becoming increasingly capable and inexpensive due to simultaneous innovations in robotic control and the cost of robotics hardware. Our team hopes to leverage these advances in robotics technology to develop commercial products which are able to reduce menial labour. Specifically, we are exploring the application of the open-source and low-cost SO-101 arms from The Robot Studio and Hugging Face to the domain of napkin folding.

Napkin folding is a technically and commercially interesting application for two primary reasons:

1. Thousands of hours are spent annually by service staff in fine dining establishments, making it a problem of significance.
2. While many similar problems are solved through reinforcement learning (RL), deformable materials like linen napkins are notoriously hard to both simulate and use policies trained in simulation (sim-to-real gap) due to their highly non-linear dynamics.

## Policy Choice

There are many choices of policies which might be suitable for this task. As discussed previously, as deformable materials are involved, our policy cannot be trained in simulation. Due to this constraint, it is initially imperative to leverage imitation learning or behaviour cloning models, as reinforcement models would take unsuitably long times to train.

As shown by Zhao et al [1], an action chunking with transformers (ACT) policy is likely suitable for this type of task. ACT is well-suited for this task as it only requires a few episodes (~30), and can train in under 24 hours.

A current trend in robotic control is the use of vision-language-action (VLA) models, which also use transformers. While VLAs are often more performant, they require significantly more data and compute. [2] VLA-type policies have also been shown to transfer between tasks and robotic systems well. For this reason, there is significant focus in the industry to further develop them; their size and complexity make them unsuitable for this task. [3]

## Data Collection Setup

In order to minimize distribution shift, a data-collection jig was fabricated from aluminum extrusions and composite wood to mount one overhead camera, lights, and all four arms. Data was collected by having a human control the leader arms, which would have their pose imitated by the follower arms in real time.

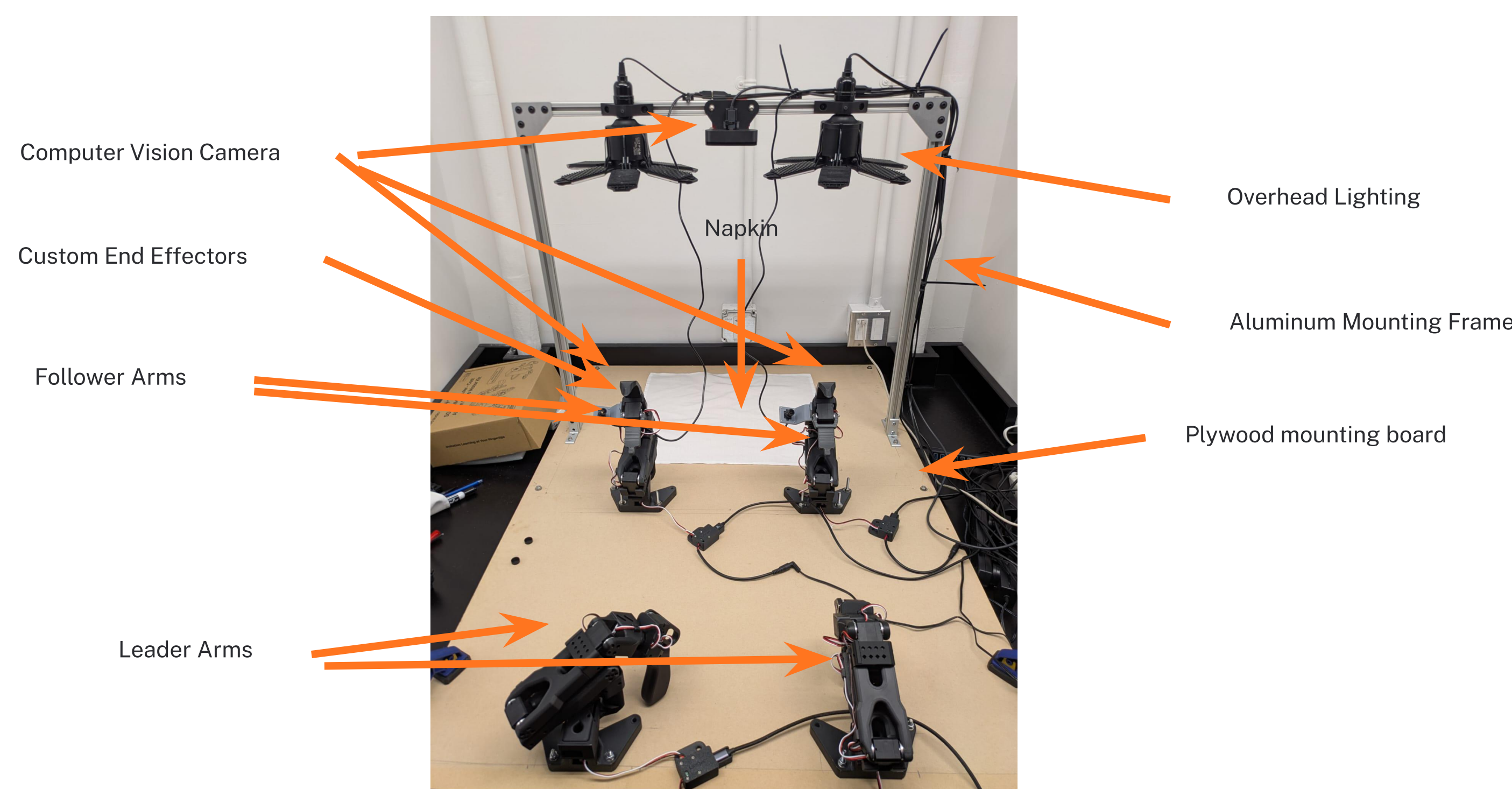
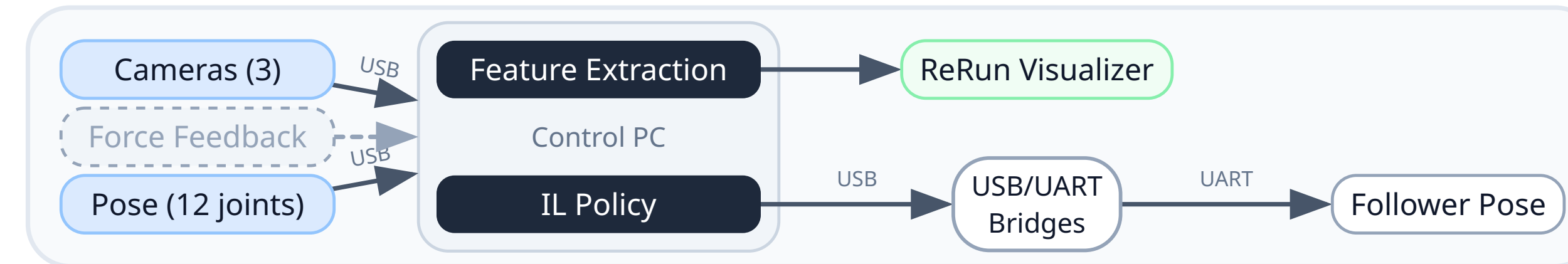


Figure 1: Data Collection Platform

Collected data is saved both locally and backed up to the UBC Sockeye research cloud. It has been made publicly available through Hugging Face. [4]

## Methods

### Inference



### Data Collection

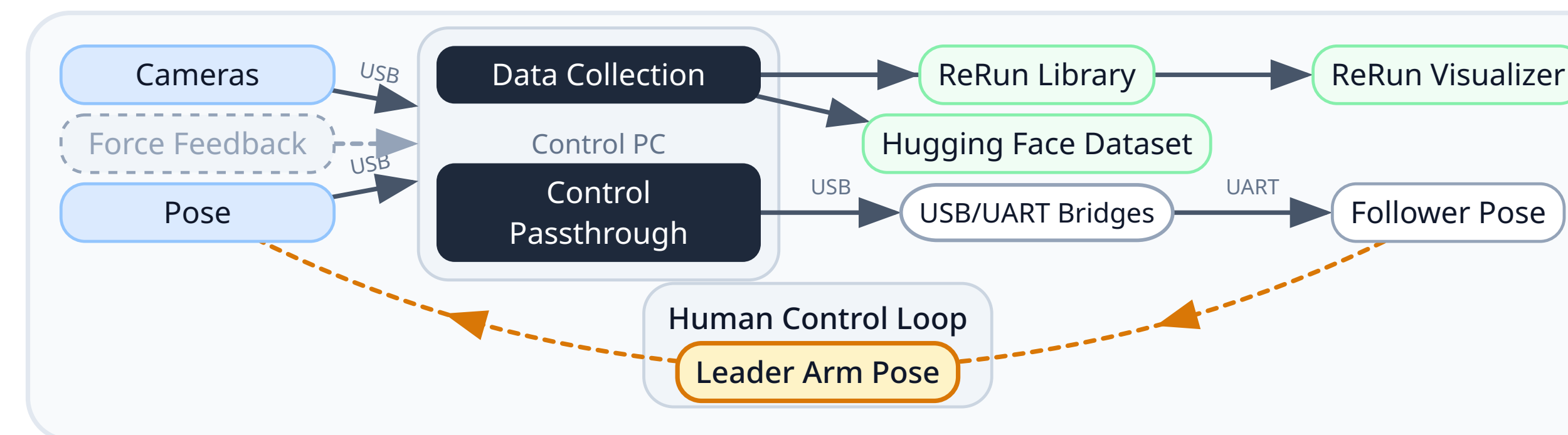


Figure 2: System Architecture



Figure 3: Example Visual Information Collected

In order to provide sufficient data, 499 episodes were collected, totalling over 550k images at 30 frames per second. This data was split between full-task data (400 episodes, ~4 hours) and pickup data (99 episodes, ~1 hour) as initial models showed that the policy struggled to learn the pickup phase.

## Policy Training

Training was conducted locally, on the public cloud, and on a UBC HPC cluster, but local training was shown to have the fastest development cycle due to reduced data transmission and job initialization overhead.

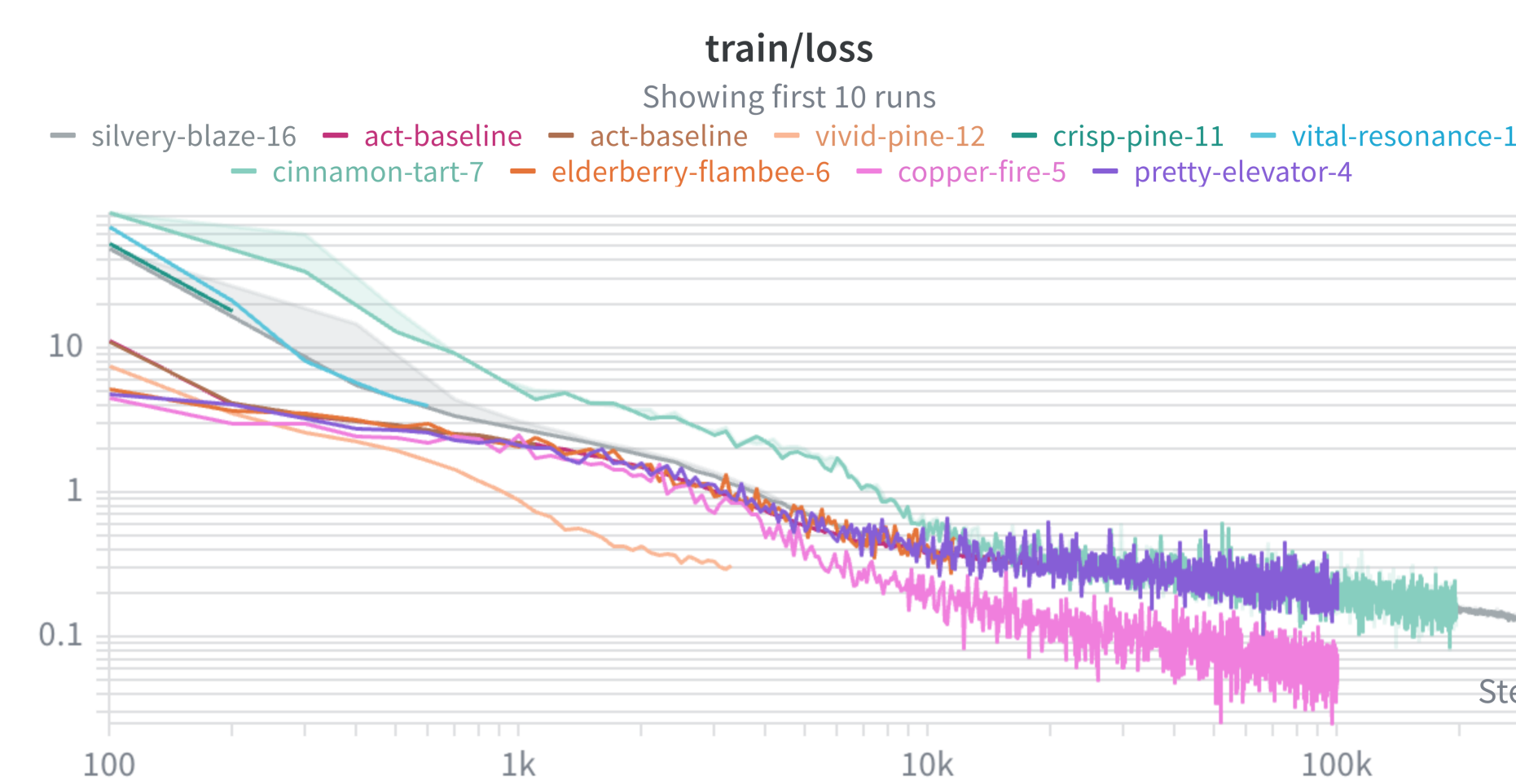


Figure 4: Training loss over 100k steps (15 h)

The loss curve, while difficult to mechanistically interpret, indicates that the policy is learning and that it is unlikely that there are any critical errors in the data collection or training scripts that prevent the model from navigating the loss landscape.

## Results and Analysis

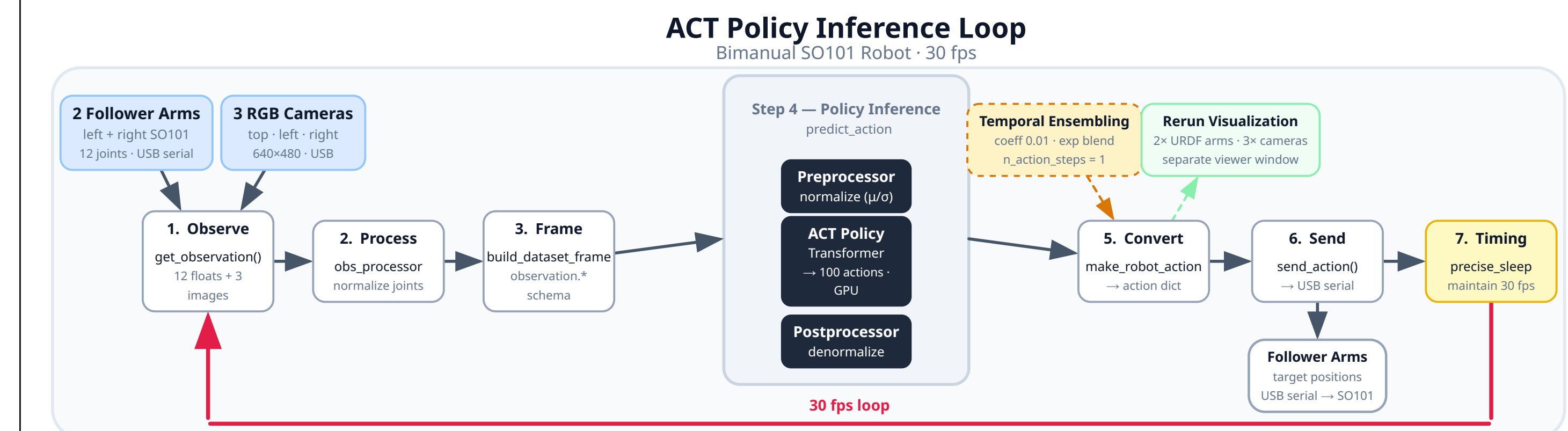


Figure 5: Inference Diagram

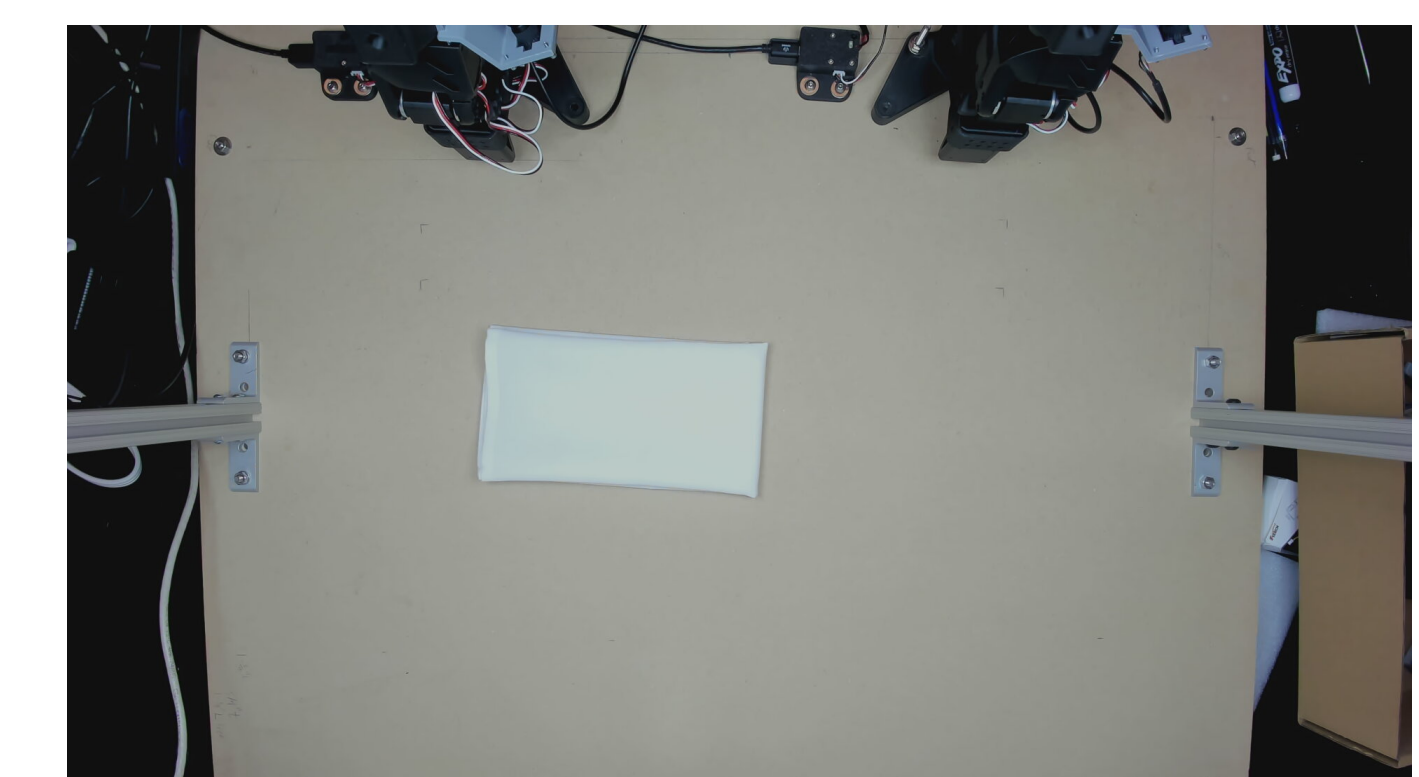


Figure 6: Example Successful Fold

Initial results show that the learned policy can execute the folding sequence once a stable grasp has been achieved, but performance degrades significantly when the napkin enters an out-of-distribution pose. Critically, the napkin rarely reaches the fully-folded state because the policy struggles during pickup. To isolate this bottleneck, we trained separate policies on pickup-only and fold-only data subsets. The pickup-only policy grasps successfully in nearly all trials, and the fold-only policy (given a human-assisted grasp) completes the fold reliably. However, a single policy trained on the full task fails to compose these stages, suggesting that the combined distribution is too broad for the ACT model to capture at this dataset scale, or that distributional shift between the pickup and folding phases creates conflicting action modes.

## Conclusion

This work shows that low-cost imitation learning can produce partial success on robotic napkin folding, but robust end-to-end performance is still limited by the pickup stage. Since training loss plateaued without corresponding task success, the main bottleneck is unlikely to be additional compute alone; dataset quality, distribution coverage, and stage-specific supervision are more probable limitations.

Future work will explore stage-aware policy architectures such as SARM to explicitly decompose the task, SmoVLA to leverage pretrained vision-language representations for better generalization, and learned reward modeling to provide a continuous napkin fold quality metric beyond binary success. Together, these steps should clarify whether improved data and task decomposition are sufficient to achieve reliable full-fold execution.

## Bibliography

- [1] T. Z. Zhao, V. Kumar, S. Levine, and C. Finn, "Learning Fine-Grained Bimanual Manipulation with Low-Cost Hardware," 2023, doi: 10.48550/arXiv.2304.13705.
- [2] Q. Chen, J. Yu, M. Schwager, P. Abbeel, Y. Shentu, and P. Wu, "SARM: Stage-Aware Reward Modeling for Long Horizon Robot Manipulation," 2025, doi: 10.48550/arXiv.2509.25358.
- [3] DYNA, "DYNA - Production-Ready Foundation Model Robot." Accessed: Mar. 27, 2026. [Online]. Available: <https://www.dyna.co/dyna-1/research>
- [4] J. Himmens, "linique-v2-fold-pickup (Revision fe33cb4)." [Online]. Available: <https://huggingface.co/datasets/jhimmens/linique-v2-fold-pickup>