

VTAM: Video-Tactile-Action Models for Complex Physical Interaction Beyond VLAs

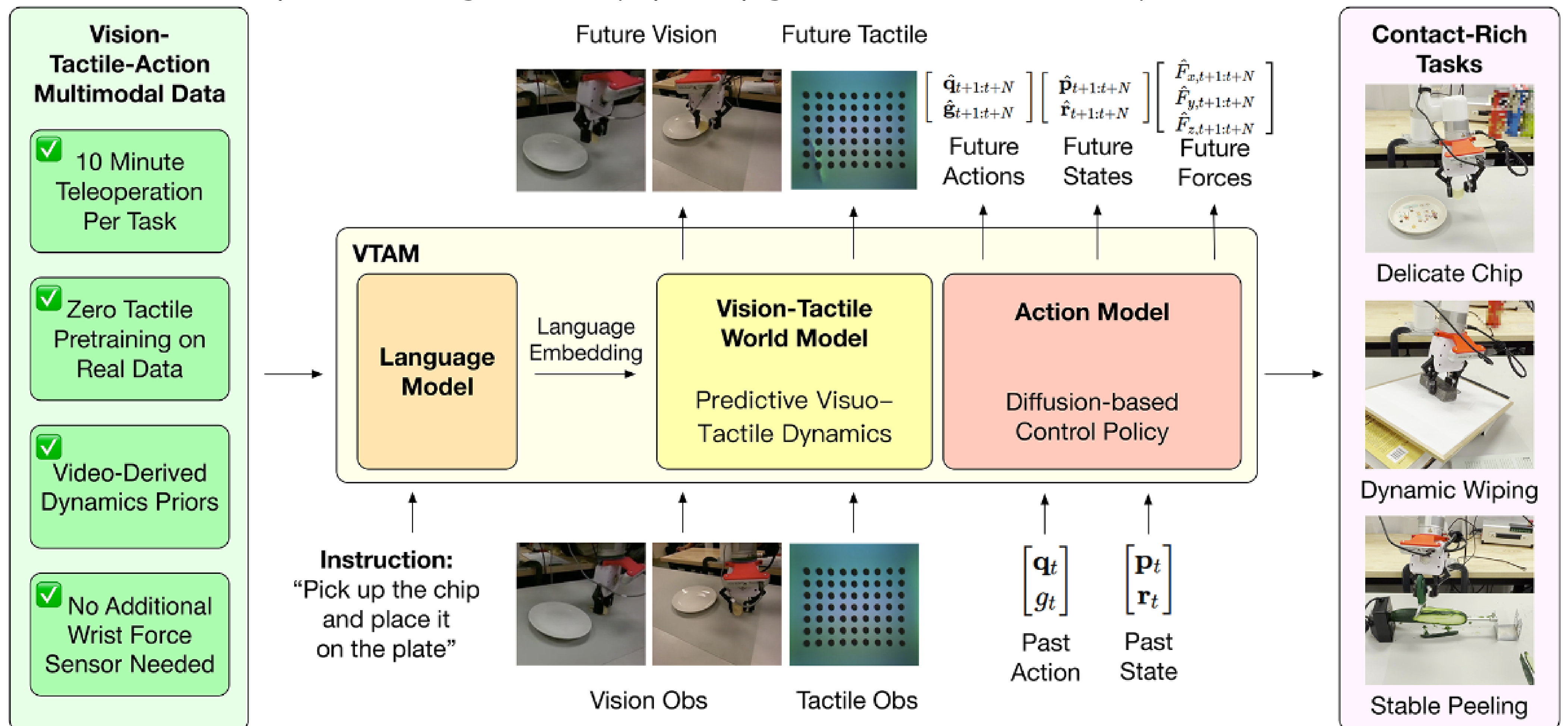
Haoran Yuan^{1,*}, Weigang Yi^{1,*}, Zhenyu Zhang^{2,*}, Wendi Chen^{3,*}, Yuchen Mo¹, Jiashi Yin¹, Xinzhuo Li¹, Xiangyu Zeng¹, Chuan Wen³, Cewu Lu³, Katherine Driggs-Campbell¹, Ismini Lourentzou¹

plan-lab.github.io/vtam



VTAM

We propose a new **video-tactile world action model** for contact-rich robotic manipulation. Given multi-view visual observations, tactile signals, and robot state/action context, VTAM predicts interaction dynamics and generates physically grounded actions for complex real-world tasks.



Contributions

- ✓ Predictive visuo-tactile world model
- ✓ Lightweight modality-transfer finetuning
- ✓ Tactile regularization for balanced fusion
- ✓ Action + state + virtual-force prediction

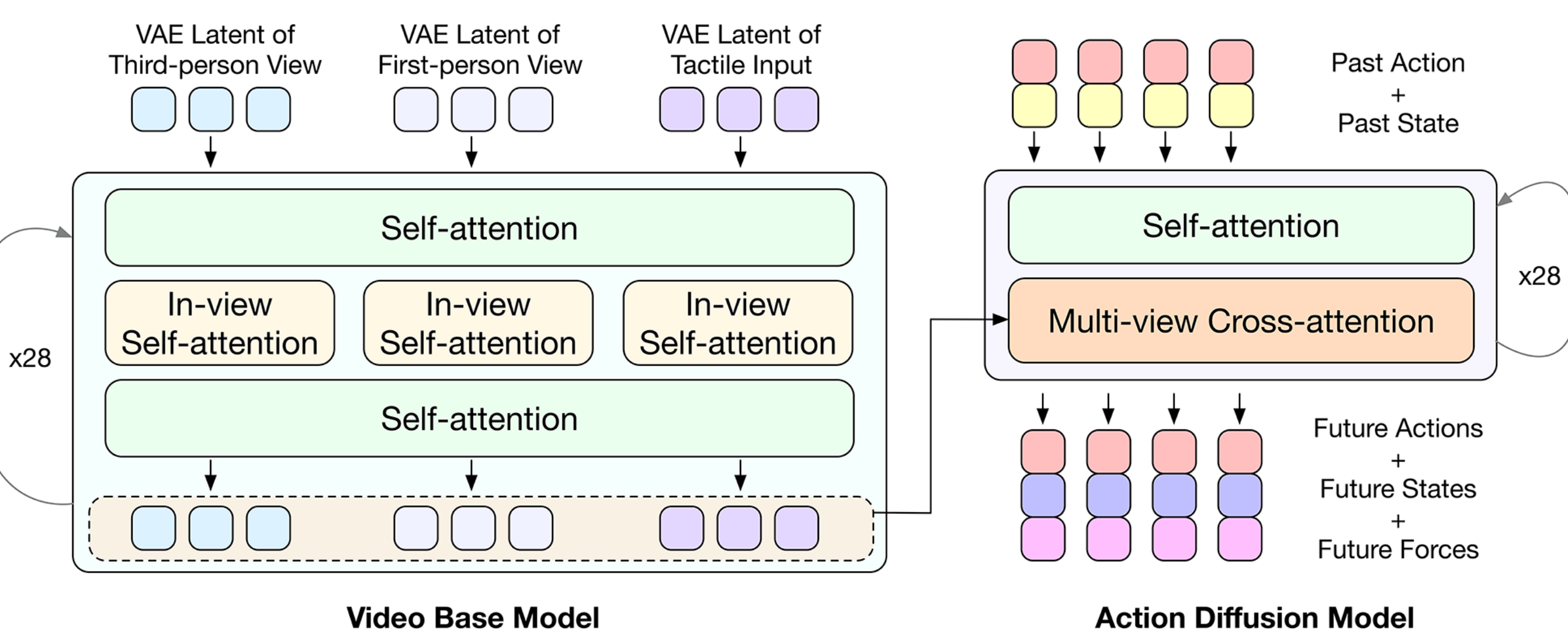
90%
average success
across contact-rich tasks

10 min
teleoperation / task
efficient real-world data

Why tactile?

- Vision misses subtle contact, slip, deformation, and force transitions.
- Contact-rich manipulation needs precise force modulation in addition to semantic understanding.
- Simply adding tactile input can still fail when visual features dominate multimodal attention.

Method



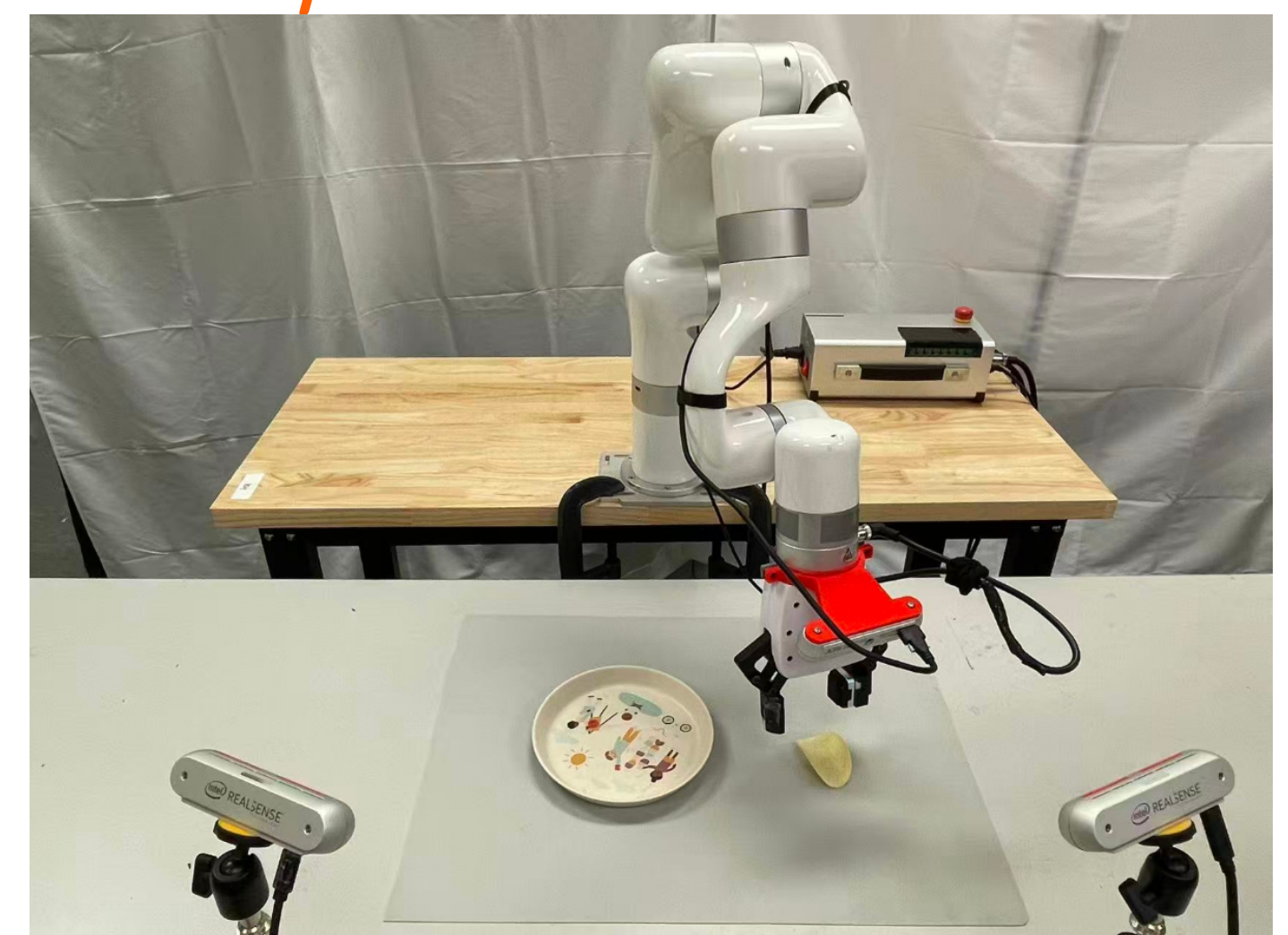
Video Base Model

Action Diffusion Model

Stage I: joint visuo-tactile predictive dynamics

Stage II: conditional diffusion predicts future actions, states, and virtual forces

Teleoperation data collection

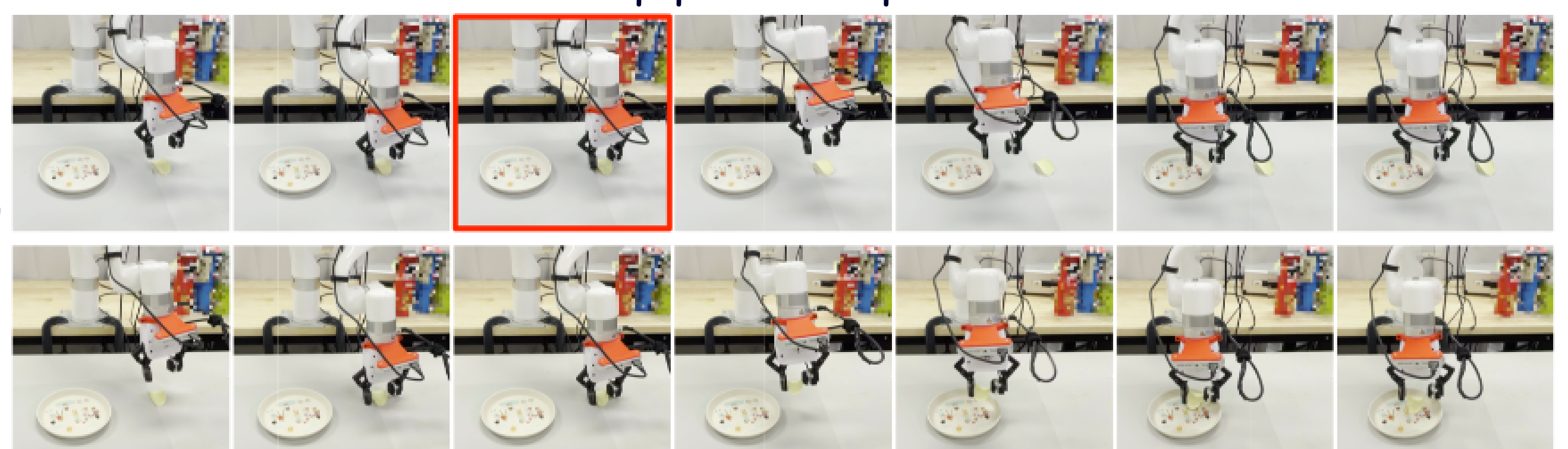


Experiments

80 trials/model · 20 trials/task setting

Model	Chip	Peel	Wipe
Genie Envisioner	0%	0%	2.5%
$\pi_{0.5}$ (Vision)	10%	0%	0%
$\pi_{0.5}$ + Tactile	5%	0%	0%
VTAM (Ours)	90%	85%	95%

Baseline
VTAM



Chip pick-and-place