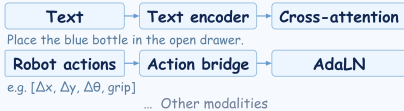


Visual context



Multimodal conditioning

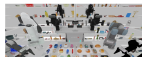


Causal attention

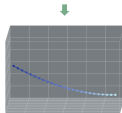


In-context control

Action



Camera



Intermediate stream



Future RGB

Causal Chain-of-Thought

Flow context

Diffusion Transformer

1

Predict motion

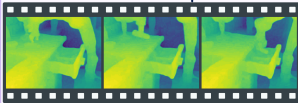
Optical flow F



2

Predict geometry

Pointmaps P



3

Predict future

Future RGB V



Stage 1 Pre-training

29K h collected → 20K h curated



Data



RGB only

Stage 2 Mid-training

Random stage as target

Flow



given

Pointmap



denoise target

RGB



masked

Stage 3 Reinforcement learning

Rollouts



Reward

- Task adherence
- Physical plausibility
- Temporal coherence
- Visual quality