

Visual context



Multimodal conditioning

Text

Cross-attention

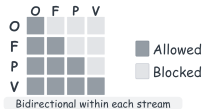
Place the blue bottle in the open drawer.

Robot actions

AdaLN

e.g. $[\Delta x, \Delta y, \Delta \theta, \text{grip}]$

Causal attention



Diffusion Transformer

Causal Chain-of-Thought

1

Predict motion

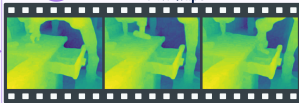
Optical flow F



2

Predict geometry

Pointmaps P



3

Predict future

Future RGB V



Stage 1 Pre-training

31K h collected $\rightarrow$ 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

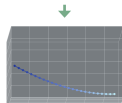
Test-time: In-context Control

Action



Robot Rendering

Camera



Camera Trajectory

Diffusion Transformer



Future RGB