



# Compositional Video Synthesis with Action Graphs

Amir Bar\*, Roei Herzig\*, Xiaolong Wang, Anna Rohrbach,  
Gal Chechik, Trevor Darrell, Amir Globerson

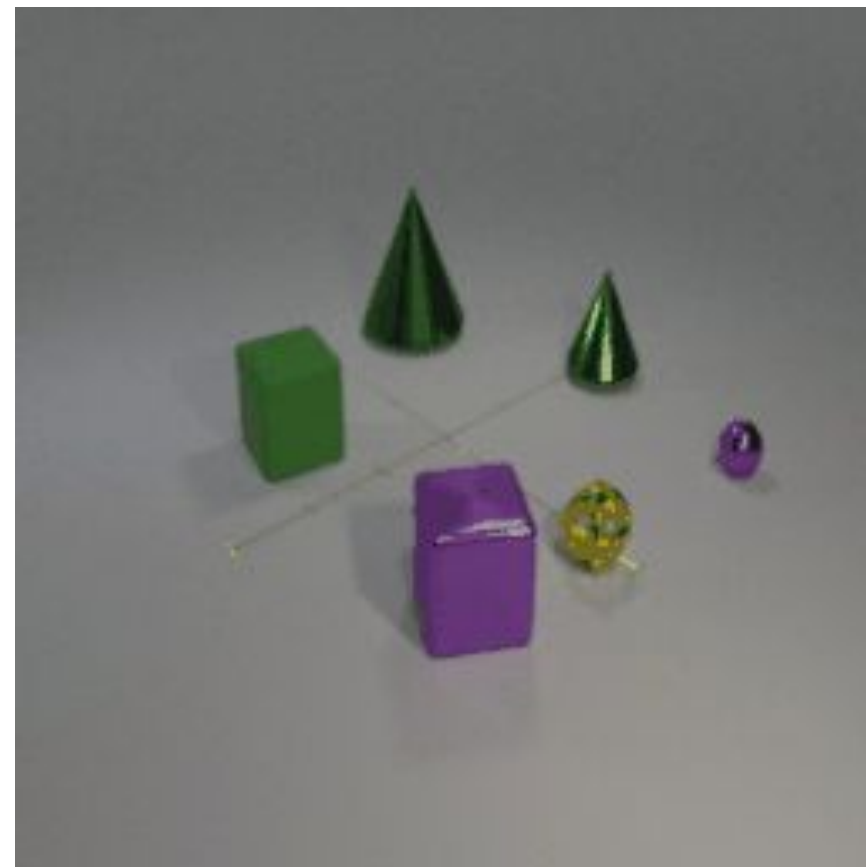
ICML 2021

\* Equally contributed.

SRVU, ICCV 2021

# Our Goal

Synthesize videos of  
actions



# Our Goal

Learn to synthesize videos of actions

---

**Our model should be able to synthesize:**

- Multiple actions and objects
- Potentially simultaneous actions
- Coordinated and timed actions

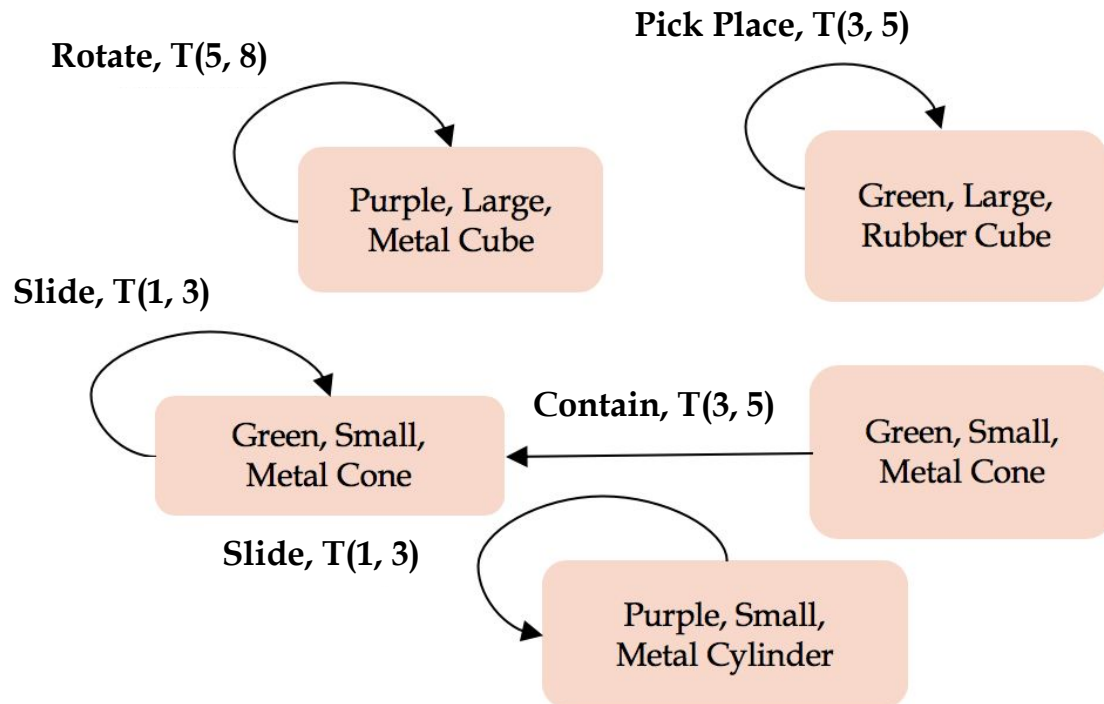


**How should we model actions?**

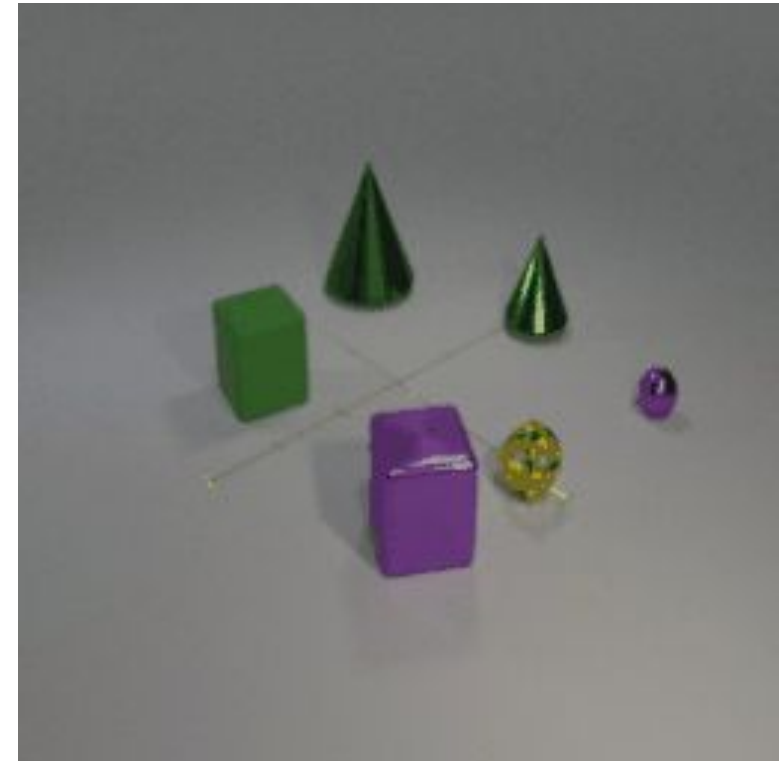
# The Action Graph Representation

- Nodes are objects
- Edges are timed actions
- Each action is annotated with a start and end time

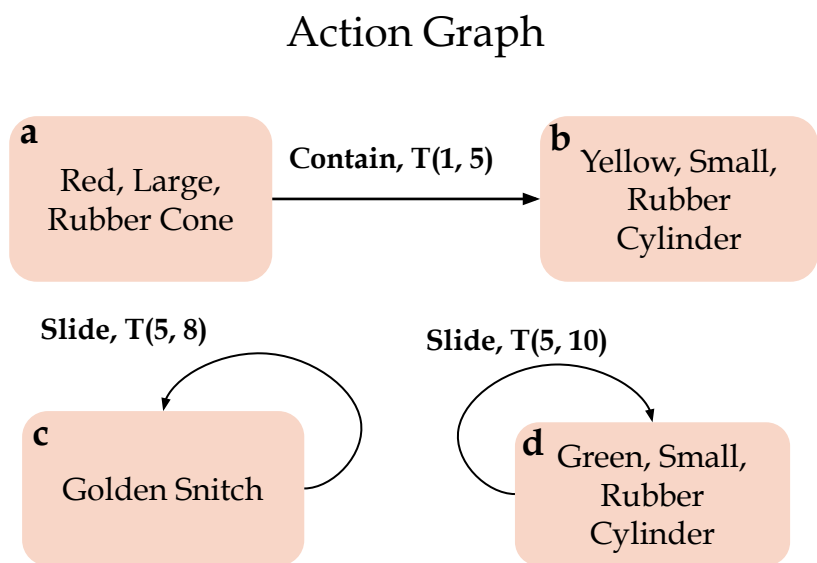
*Action Graph*



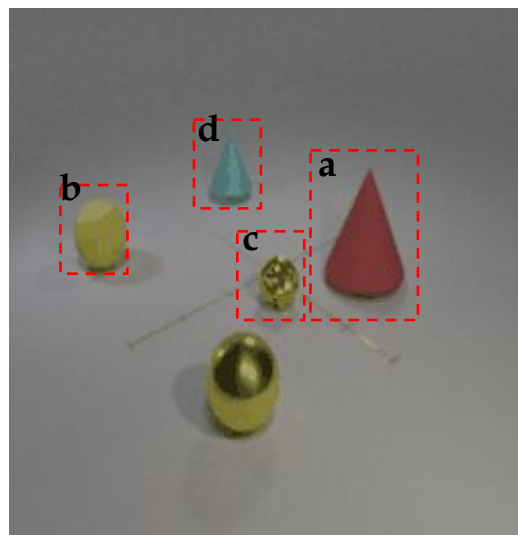
*Video*



# New Task Setting: Action-Graph-to-Video



Initial image & Layout



Input



Video

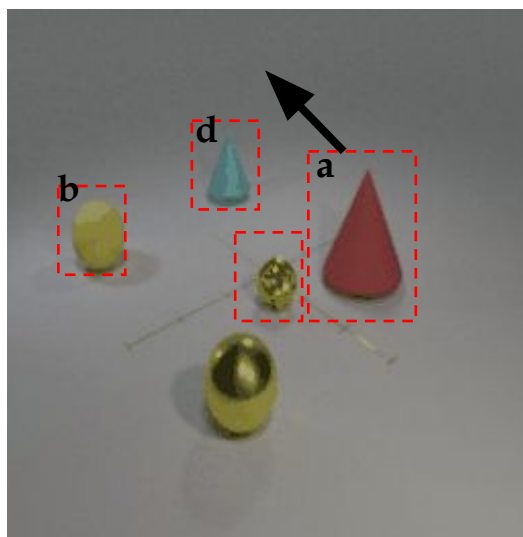


Output

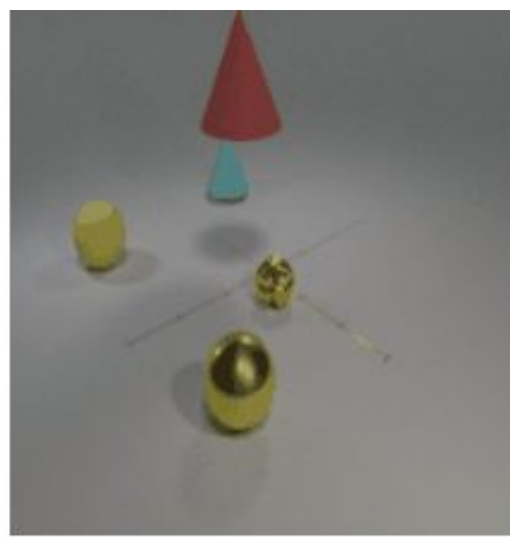
# The Action Graph to Video Model

Synthesize next frame in a coarse-to-fine manner

- Action execution schedule, given Action Graph
- Given the schedule, predict how should object moves
- Then, predict how should pixels move



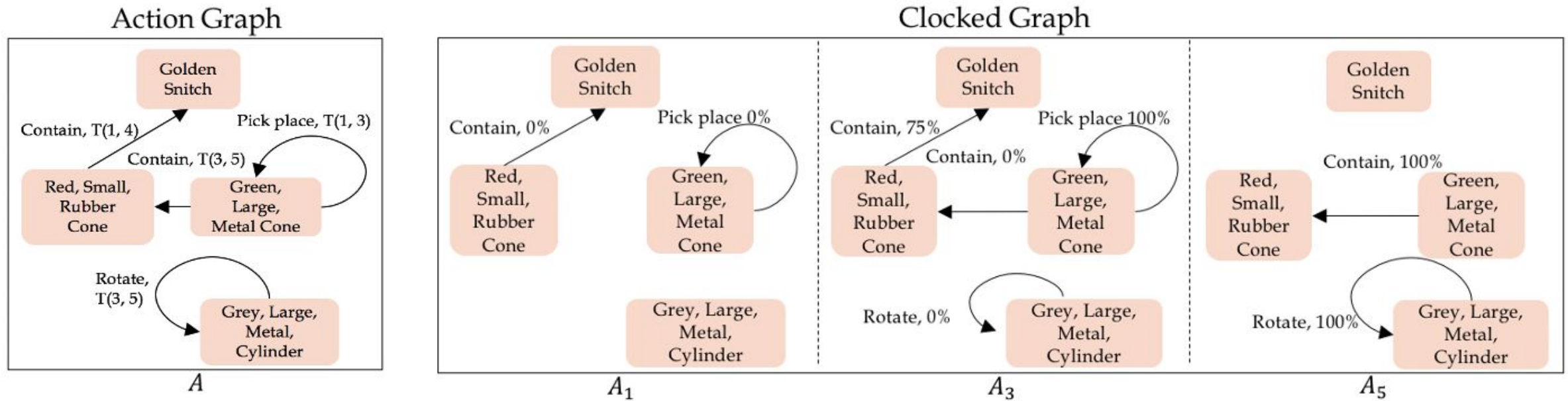
Previous image and layout



Next frame

# Scheduling Actions via “Clocked Edges”

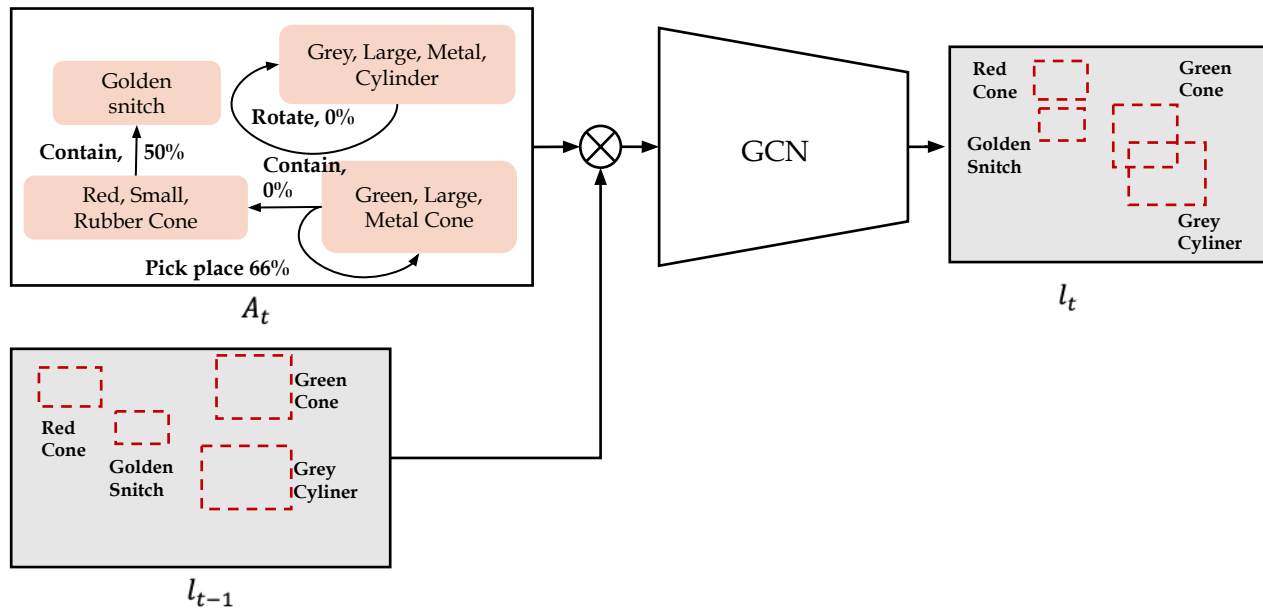
How to synchronize and schedule multiple actions?



Time specific Action Graphs

# Action Graph to Video

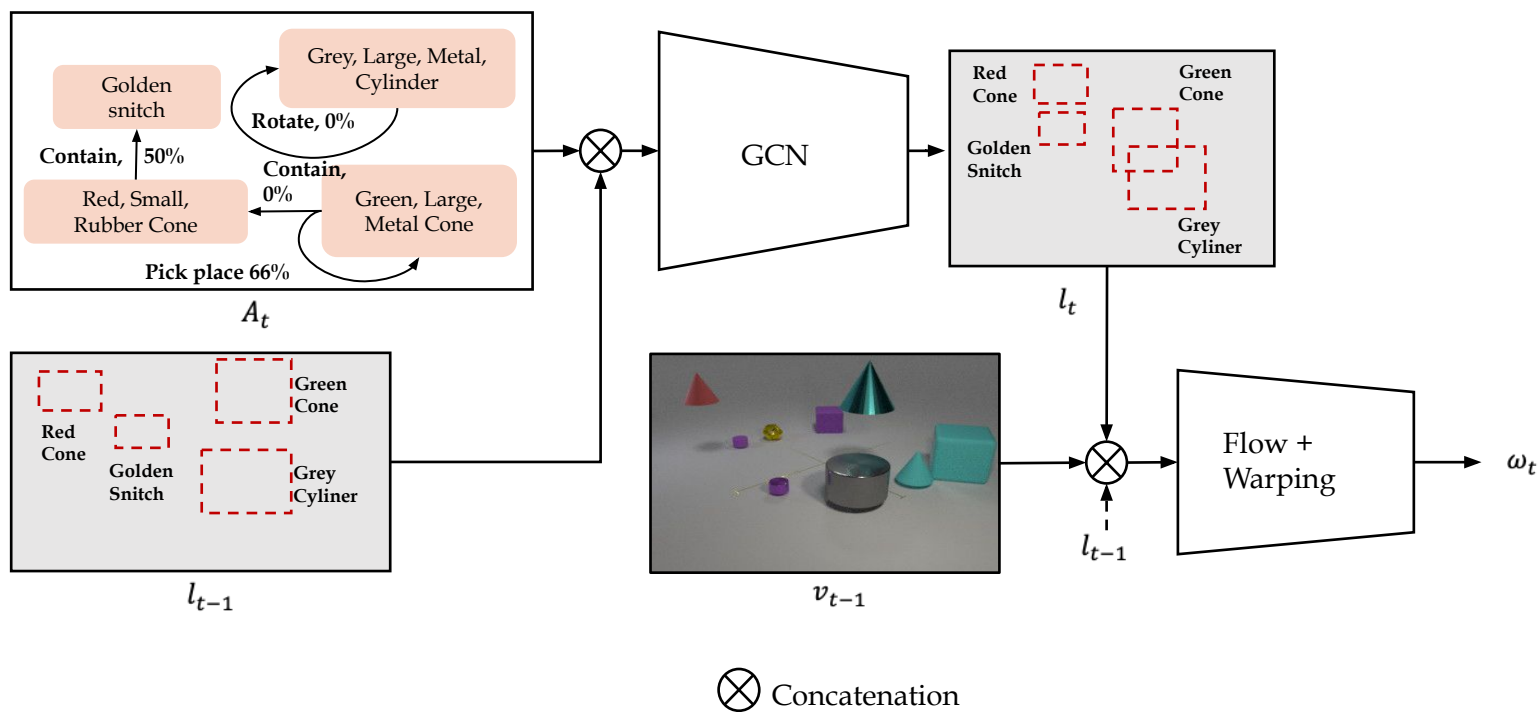
- Predict new scene layout given previous layout and Clocked Action Graph



$\otimes$  Concatenation

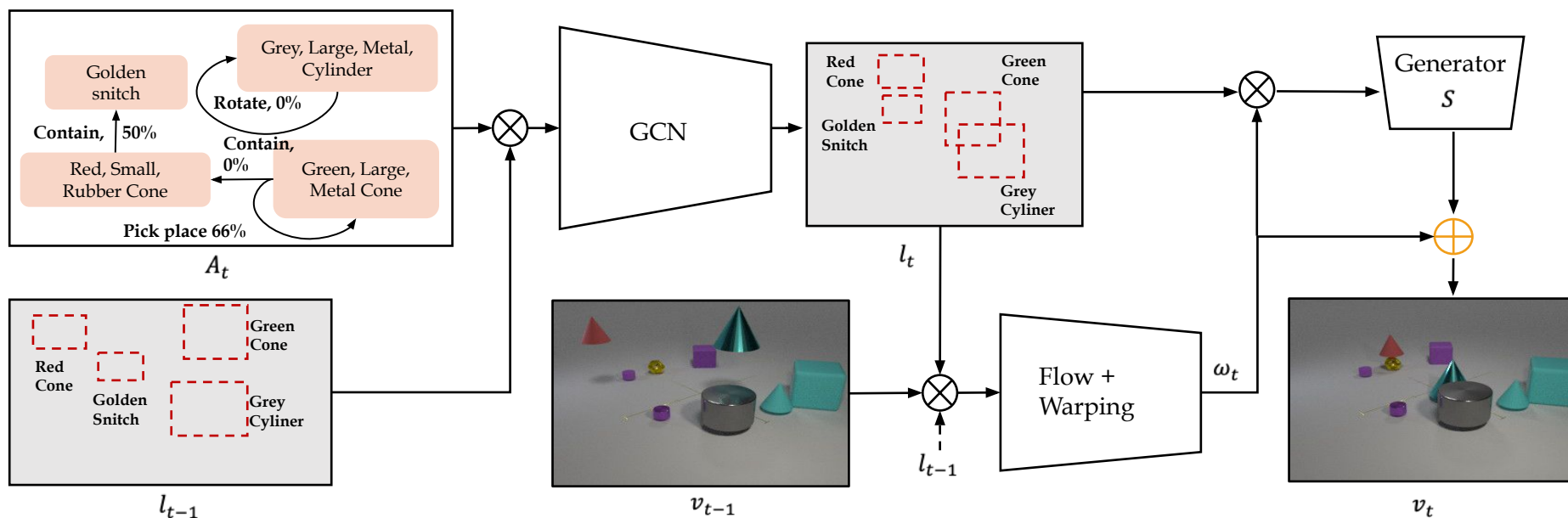
# Action Graph to Video

- Predict new scene layout given previous layout and Clocked Action Graph
- Predict the future pixels flow, and warp the previous image



# Action Graph to Video

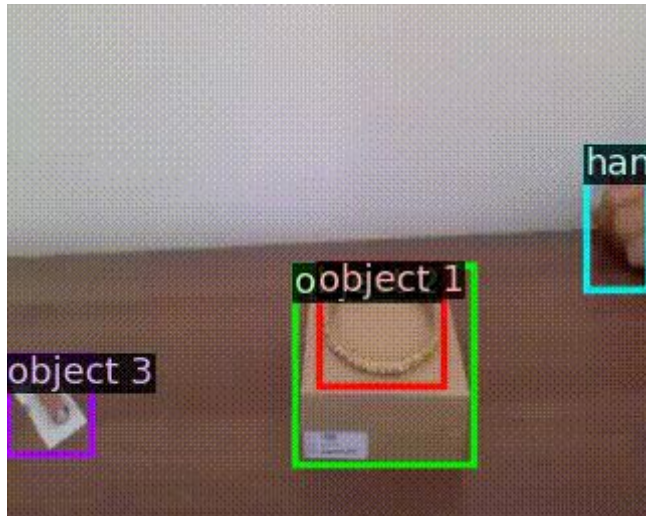
- Predict new scene layout given previous layout and Clocked Action Graph
- Predict the future pixels flow, and warp the previous image
- Refine the warped image via a SPADE Generator



$\otimes$  Concatenation  $\oplus$  Addition

# Datasets of atomic actions

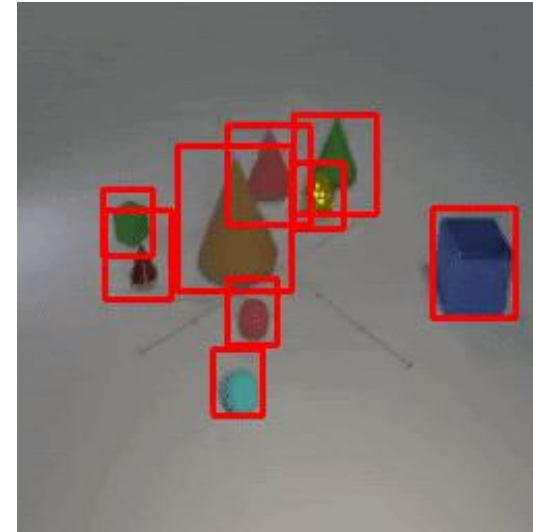
## Something Something V2



Move down  
Move up  
Push right  
Push left  
Take  
Cover  
Uncover  
Put

Videos of *single* actions

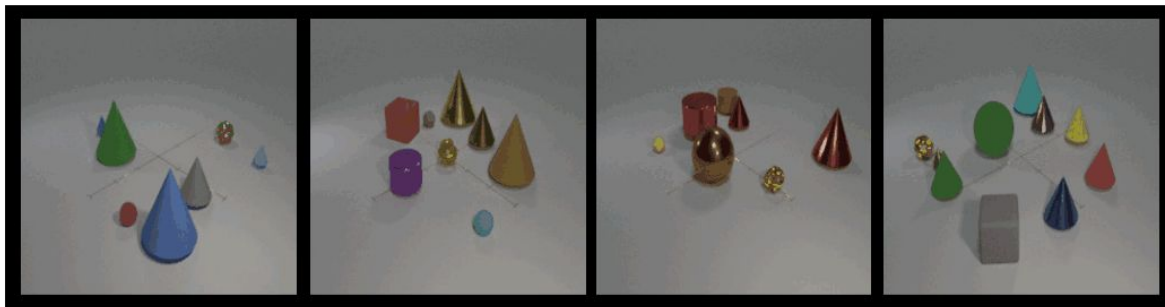
## CATER



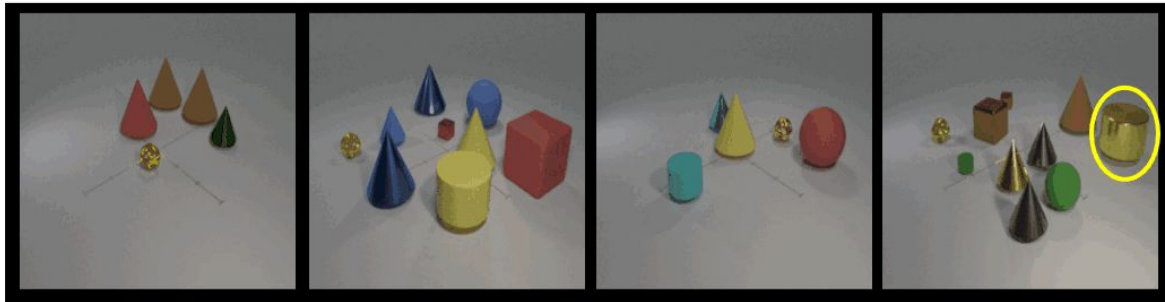
Slide  
Contain  
Rotate  
Pick Place

Videos of *multiple* actions

## Actions in CATER



*Multiple Simultaneous Actions*



*Slide*

*Contain*

*Pick Place*

*Rotate*

## Actions in Something Something



*Push Left*

*Move Down*

*Uncover*

*Push*



*Push Right*

*Move Up*

*Cover*

*Take*

# Human evaluation of the synthesized videos

| AG2Vid vs. Baseline       | Semantic Accuracy |             | Visual Quality |             |
|---------------------------|-------------------|-------------|----------------|-------------|
|                           | CATER             | SmthV2      | CATER          | SmthV2      |
| CVP (Ye et al., 2019)     | <b>85.7</b>       | <b>90.6</b> | <b>76.2</b>    | <b>93.8</b> |
| HG (Nawhal et al., 2020a) | —                 | <b>84.6</b> | —              | <b>88.5</b> |
| V2V (Wang et al., 2018a)  | <b>68.8</b>       | <b>84.4</b> | <b>68.8</b>    | <b>96.9</b> |
| RNN                       | <b>56.0</b>       | <b>80.6</b> | <b>52.0</b>    | <b>77.8</b> |
| AG2Vid-GTL                | 48.6              | 46.2        | 42.9           | 50.0        |

Each cell is the %times AG2Vid model is better

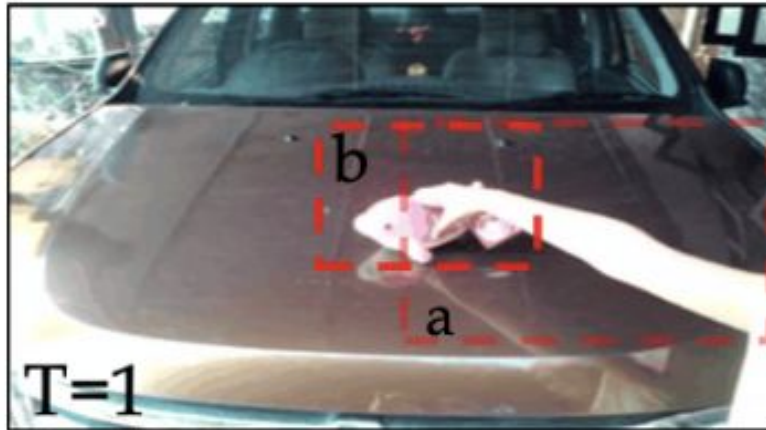
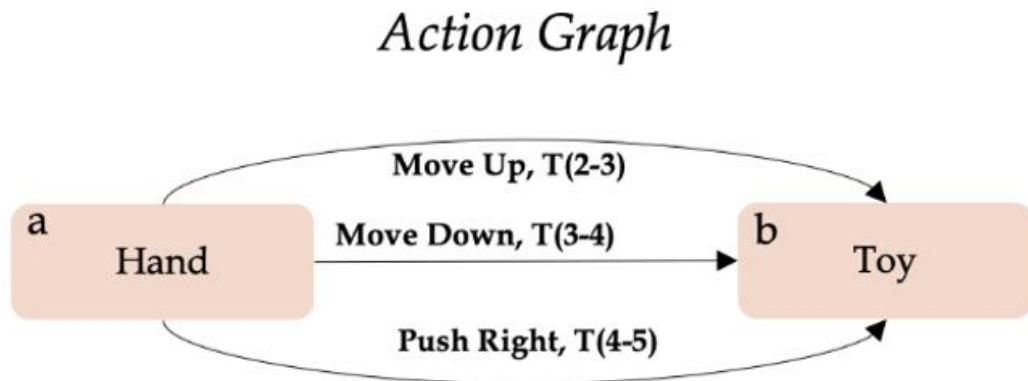
Our model outperforms baselines by synthesizing videos that are **more semantically correct** and have **better visual quality**.

# Zero-shot synthesis

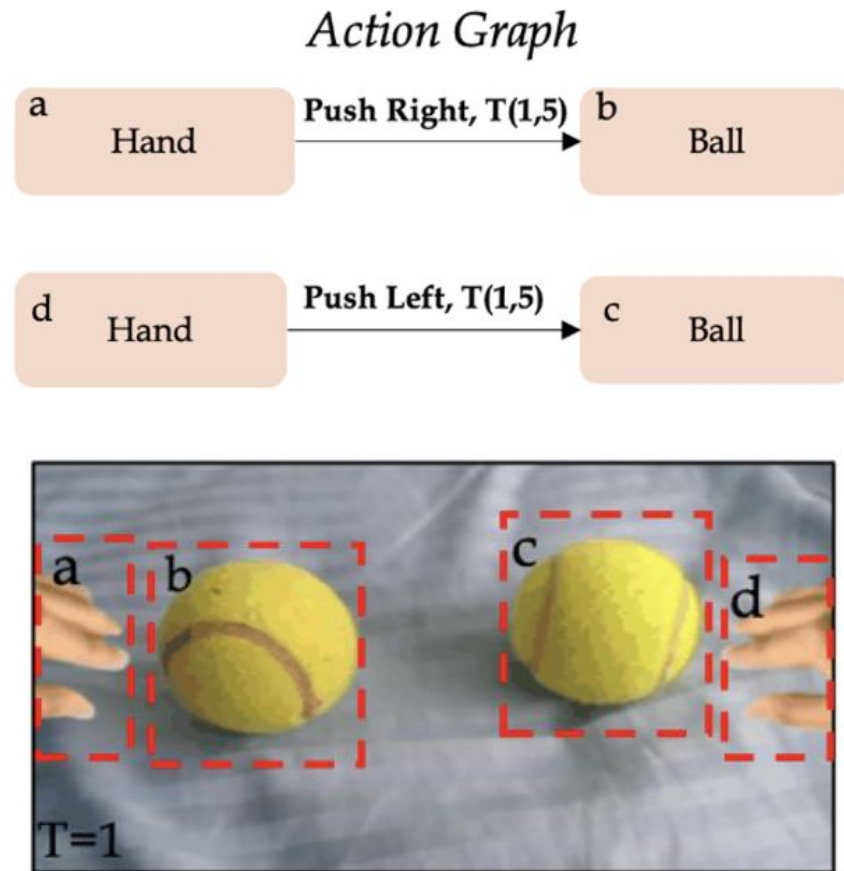
So far, we've showed that our model can synthesize the atomic actions present in the training data.

Can we use this approach to synthesize more complex videos?

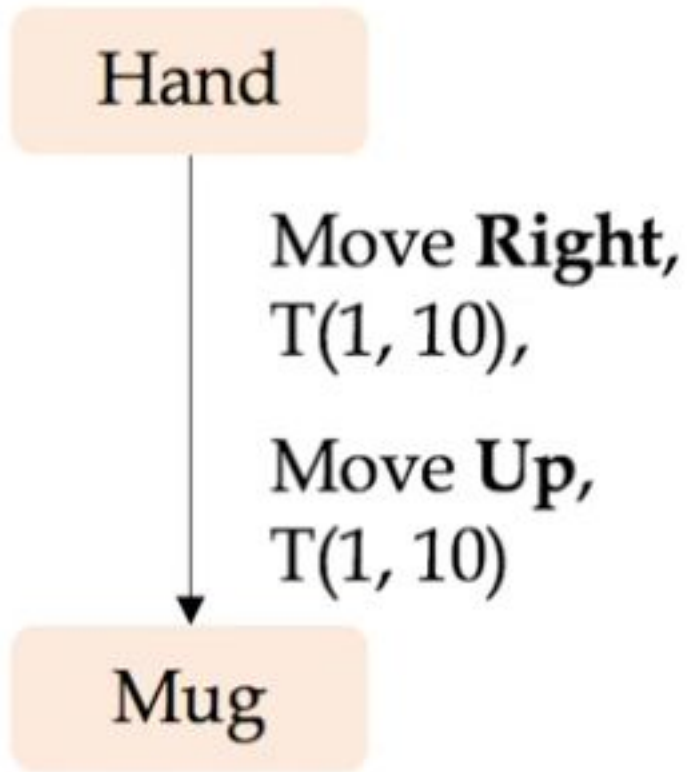
# Synthesizing zero-shot *sequential* actions



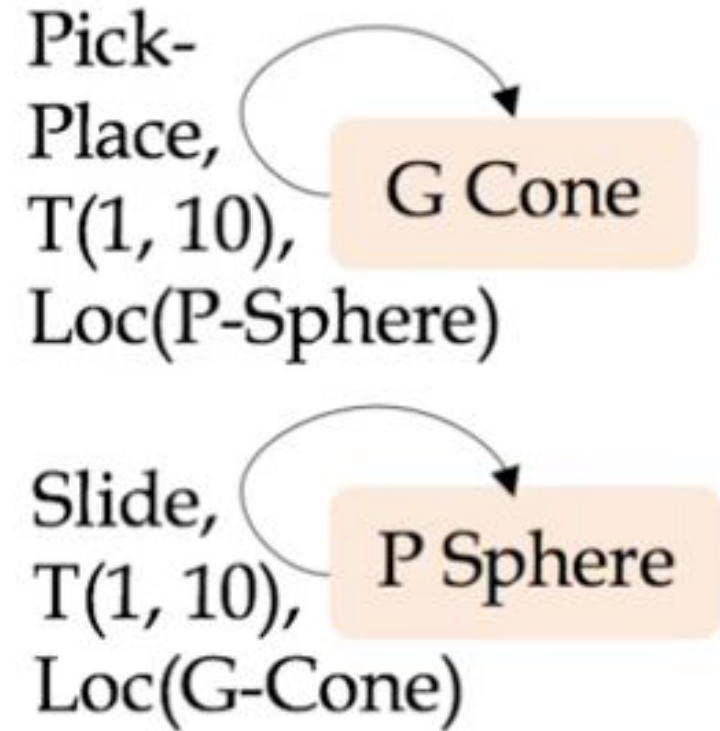
# Synthesizing zero-shot *simultaneous* actions



# Synthesizing new *action composites*

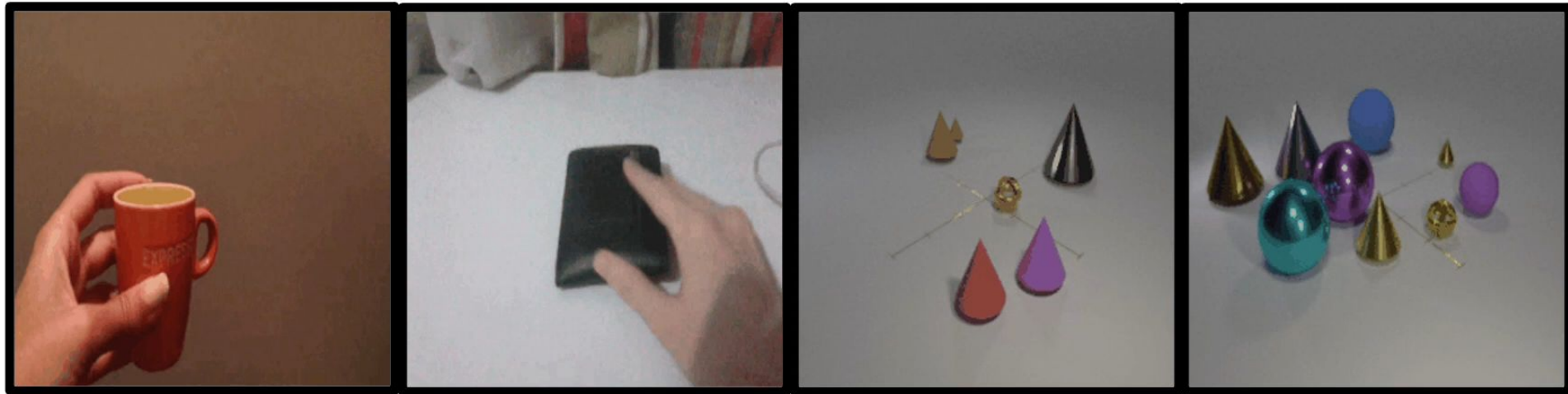
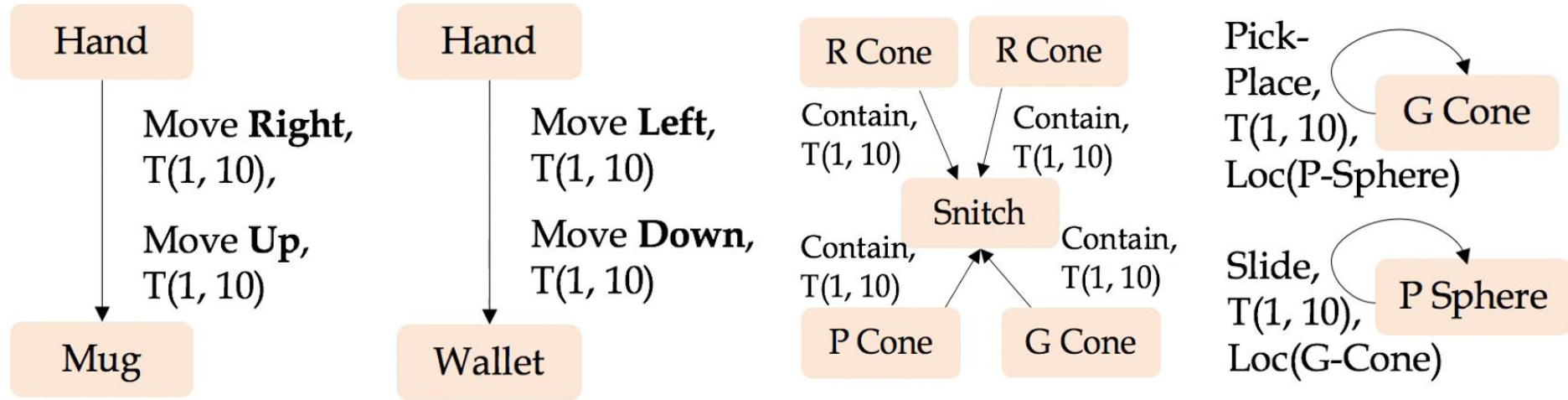


New action: **Move Right-Up**



New action: **Swap Place**

# Synthesizing new action composites

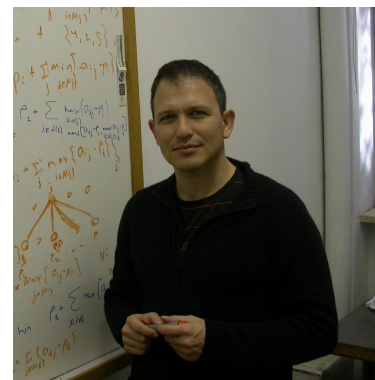


*Right Up*

*Left Down*

*Huddle*

*Swap*



See project page for code and models:  
[roeiherz.github.io/AG2Video](https://roeiherz.github.io/AG2Video)

# Thank you!