



point tracking

Adam Harley
Meta

Point tracking



Given any pixel coordinate, track the corresponding world point.

The “particle video” problem

including matting, rotoscoping, labelling, and object removal. Goldman *et al.* [17] describe particle-based video annotation applications that would be difficult to create using standard motion representations.

1.1 Particle Video Representation

Our approach represents video motion using a set of particles that move through time. Each particle denotes an interpolated image point sample, in contrast to a feature patch that represents a neighborhood of pixels [30]. Particle density is adaptive, so that the algorithm can model detailed motion with substantially fewer particles than pixels.

The algorithm optimizes particle trajectories using an objective function that combines point-based appearance matching and inter-particle distortion. The algorithm extends and truncates particle trajectories to model motion near occlusion boundaries.

Our contributions include posing the particle video problem, defining the particle video representation, and presenting an algorithm for particle motion estimation. We provide a new motion optimization scheme that combines variational techniques with an adaptive motion representation. The algorithm uses weighted links between particles to implicitly represent grouping, providing an alternative to discrete layer-based representations.

1.2 Design Goals

The particle video problem can be described as dense feature tracking or long-range optical flow. **We want to track the trajectory of each pixel through a given video. Ideally each trajectory would correspond to the motion of a physical real-world surface point.**

A primary goal is the ability to model complex motion and occlusion. We want the algorithm to handle general video, which may include close-ups of people talking, hand-held camera motion, multiple independently moving objects, textureless regions, narrow fields of view, and complicated geometry (e.g. trees or office clutter).

A particle approach provides this kind of flexibility. Particles can represent complicated geometry and motion because they are small; a particle’s appearance will not change as rapidly as the appearance of a large feature patch, and it is less likely to straddle an occlusion boundary. Particles represent motion in a non-parametric manner; they do not assume that the scene consists of planar or rigid components.

This flexibility in modelling complex motion can also be achieved by optical flow, but the optical flow representation is best suited to successive pairs of frames, not to long sequences. Frame-to-frame flow fields can be concatenated to obtain longer-range correspondences, but the resulting multi-frame flow must be refined at each step to avoid drift.

In contrast, the particle representation allows a form of random-access motion evaluation: given a set of particles, we can easily find correspondences between any pair of frames (assuming the frames have a sufficient number of particles in common). Furthermore, unlike a sequence of motion fields, the particle representation provides discrete motion primitives, which are valuable for subsequent use of the motion information, such as interactive video matting [26] and interactive video annotation [17].

1.3 Overview

Section 2 describes related work in motion estimation. We combine several of these methods to create an optical flow algorithm described in Section 3 (with details in Appendix A). Optical flow is used as an input to our particle-based motion estimation.

The particle algorithm is described in Section 4, which explains how particles are added, propagated, linked, optimized, and pruned. These steps are performed as the algorithm sweeps back and forth across a video, constructing a complete particle representation of the video’s motion.

Section 5 includes an evaluation of the particle video algorithm on a variety of real-world videos. We quantify the performance of the algorithm and possible alternatives. We provide mechanisms for visualizing the

“We want to track the trajectory of each pixel through a given video.

Ideally each trajectory would correspond to the motion of a physical real-world surface point.”

Sand & Teller, Particle Video: Long-Range Motion Estimation using Point Trajectories, 2006.

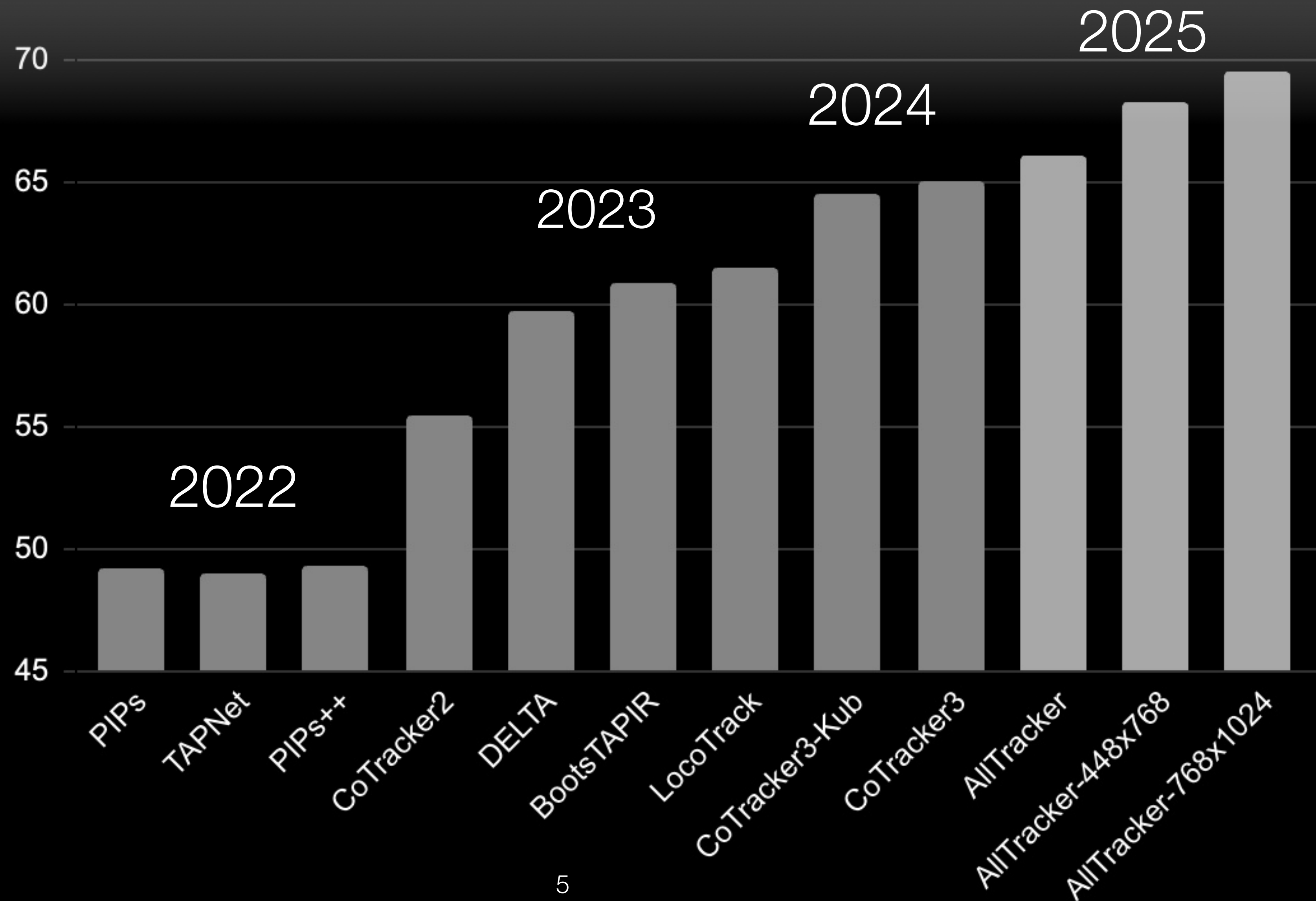
Before: Point tracking through optical flow



Sand & Teller. "Particle video: Long-range motion estimation using point trajectories." IJCV 80.1 (2008): 72.

Now: A problem all of its own

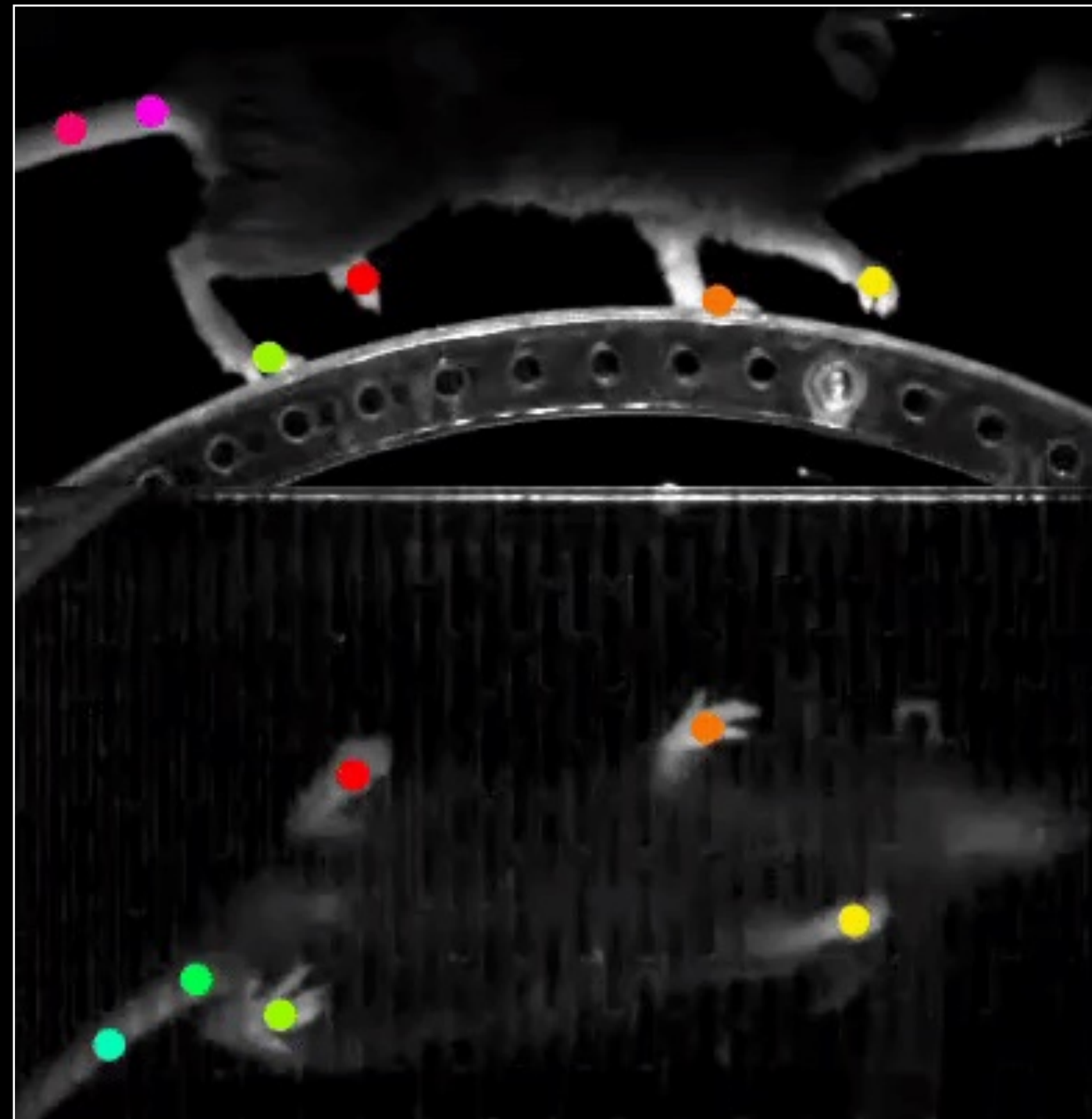
Average
accuracy
across
multiple
datasets



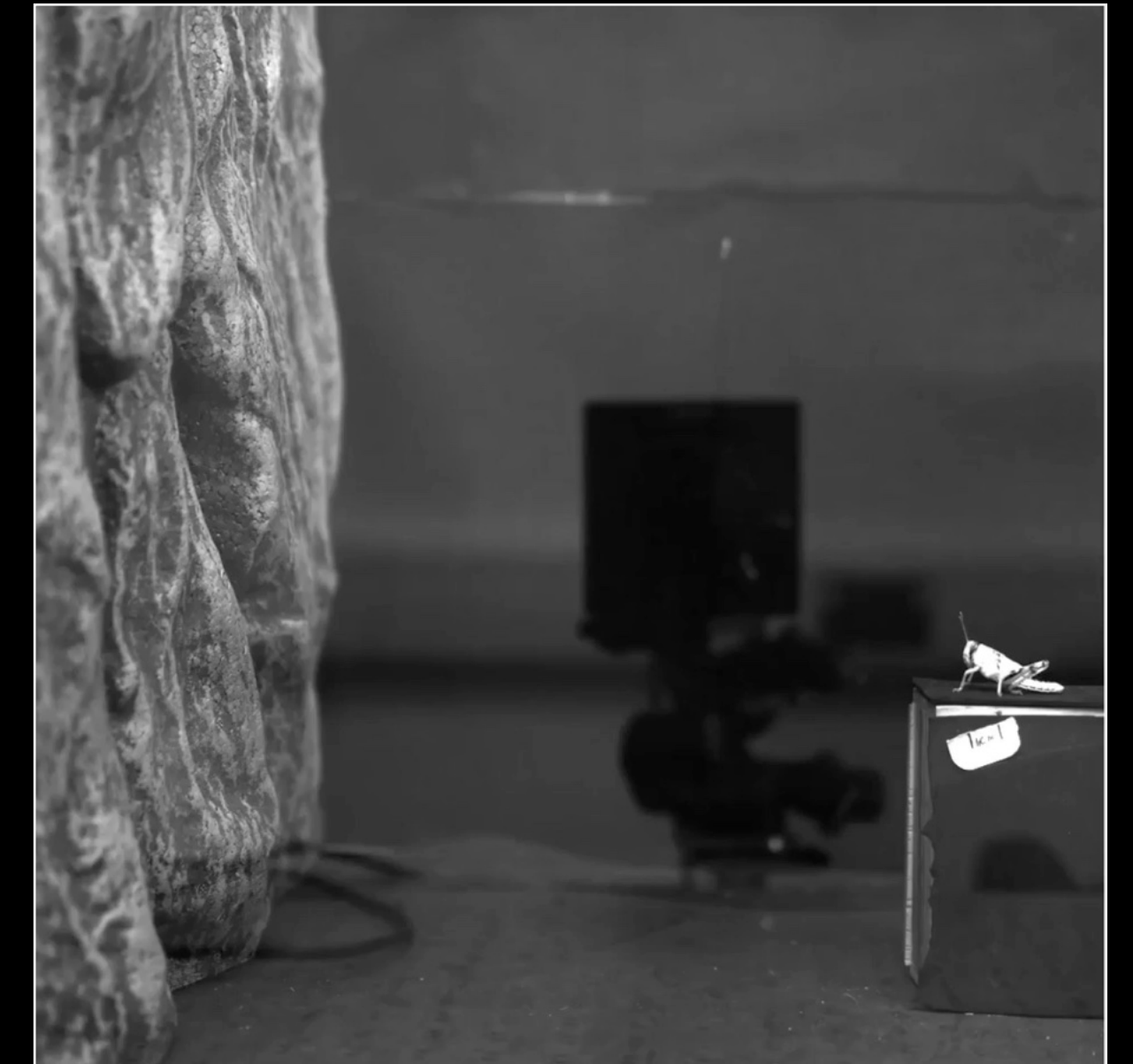
Point tracks make science easier



Polajnar, Kvinikadze, **Harley**, Malenovsky. "Wing buzzing as a mechanism for generating vibrational signals in psyllids (Hemiptera: Psylloidea)". Insect Science 2024.



Mathis et al., "DeepLabCut: Markerless pose estimation of user-defined body parts with deep learning." Nature neuroscience 21.9 (2018): 1281-1289



Goode, C. K., and Gregory P. Sutton. "Control of high-speed jumps: the rotation and energetics of the locust (*Schistocerca gregaria*)." Journal of Comparative Physiology B 193.2 (2023): 145-153.

Point tracks make robotics easier

Flow as the Cross-Domain Manipulation Interface

Mengda Xu^{1, 2, 3} Zhenjia Xu^{1, 2} Yinghao Xu¹ Cheng Chi^{1, 2}

Any-point Trajectory Modeling for Policy Learning

Chuan Wei^{1, 2, 5} Kai Chen⁶ Qi
¹UC Berkeley ²III
⁴Shanghai AI Labor

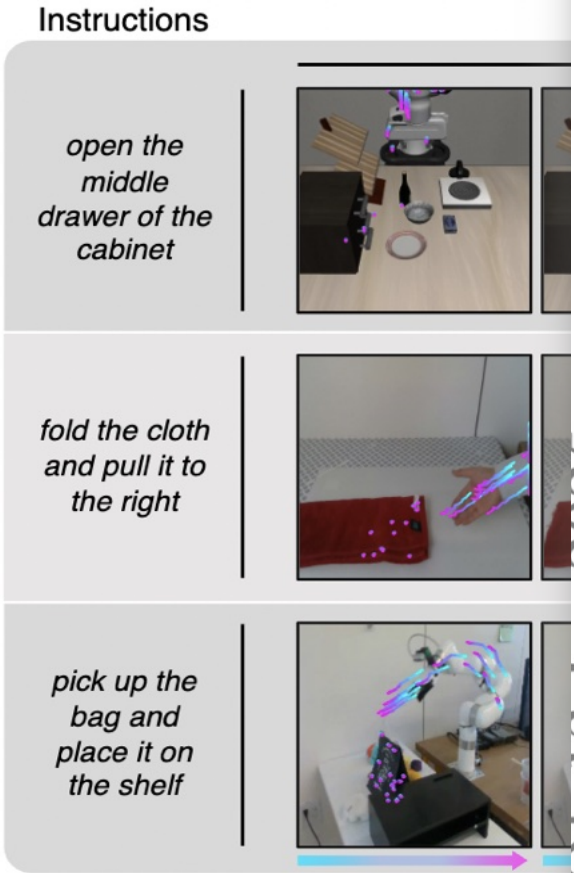
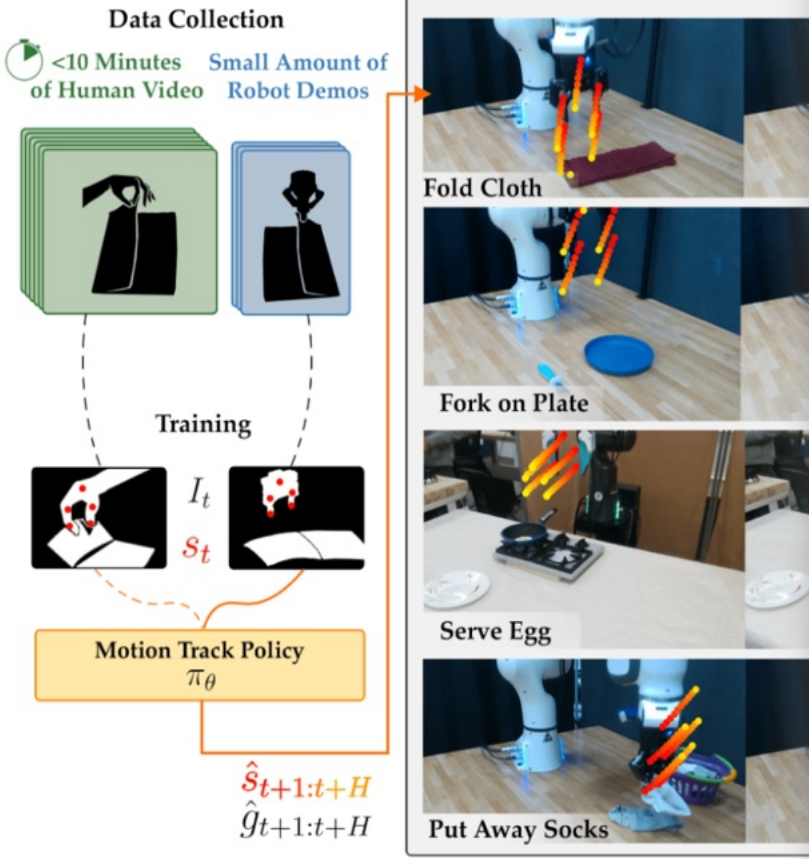


Fig. 1: Given a task instruction and the initial pos

MOTION TRACKS: A Unified Representation for Human-Robot Transfer in Few-Shot Imitation Learning

Juntao Ren¹, Priya Sundaesan², Dorsa S



Track2Act: Predicting Point Tracks from Internet Videos enables Diverse Zero-shot Robot Manipulation

Homanga Bharadhwaj¹ Roozbeh Mottaghi² Abhinav Gupta^{1,*} Shubham Tulsiani^{1,*}

¹Carnegie Mellon University and ²University of Washington, Meta

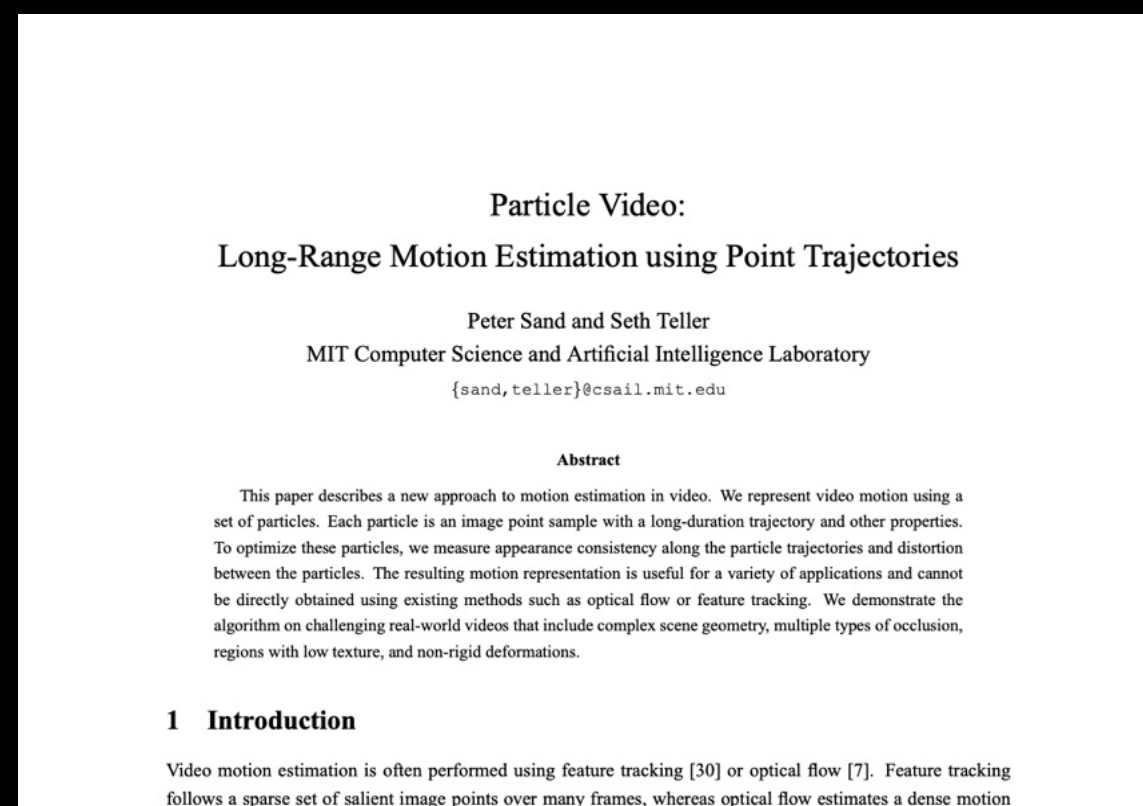
<https://homangab.github.io/track2act/>

Abstract: We seek to learn a generalizable goal-conditioned policy that enables *zero-shot* robot manipulation — interacting with unseen objects in novel scenes without test-time adaptation. While typical approaches rely on a large amount of demonstration data for such generalization, we propose an approach that leverages web videos to predict plausible interaction plans and learns a task-agnostic transformation to obtain robot actions in the real world. Our framework, *Track2Act* predicts tracks of how points in an image should move in future time-steps based on a goal, and can be trained with diverse videos on the web including those of humans and robots manipulating everyday objects. We use these 2D track predictions to infer a sequence of rigid transforms of the object to be manipulated, and obtain robot end-effector poses that can be executed in an open-loop manner.

The VOTSp2026 Challenge

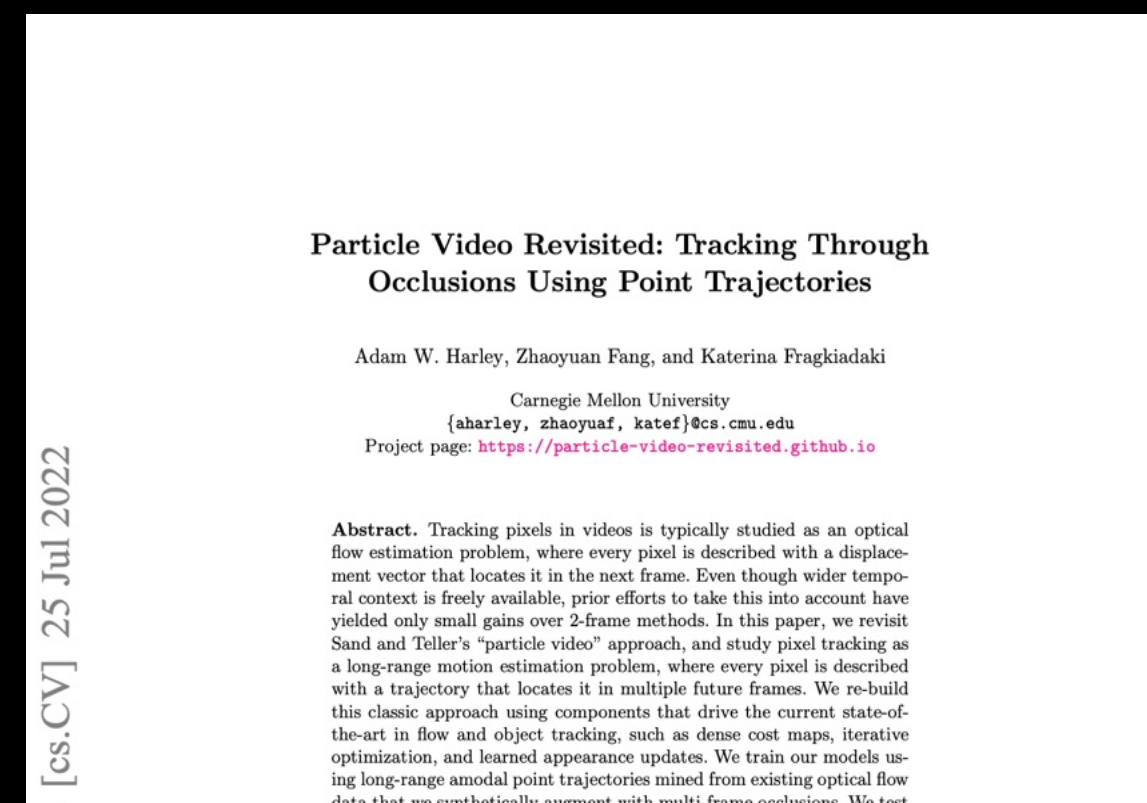
How do we evaluate point tracks?

Sand & Teller,
IJCV 2006.



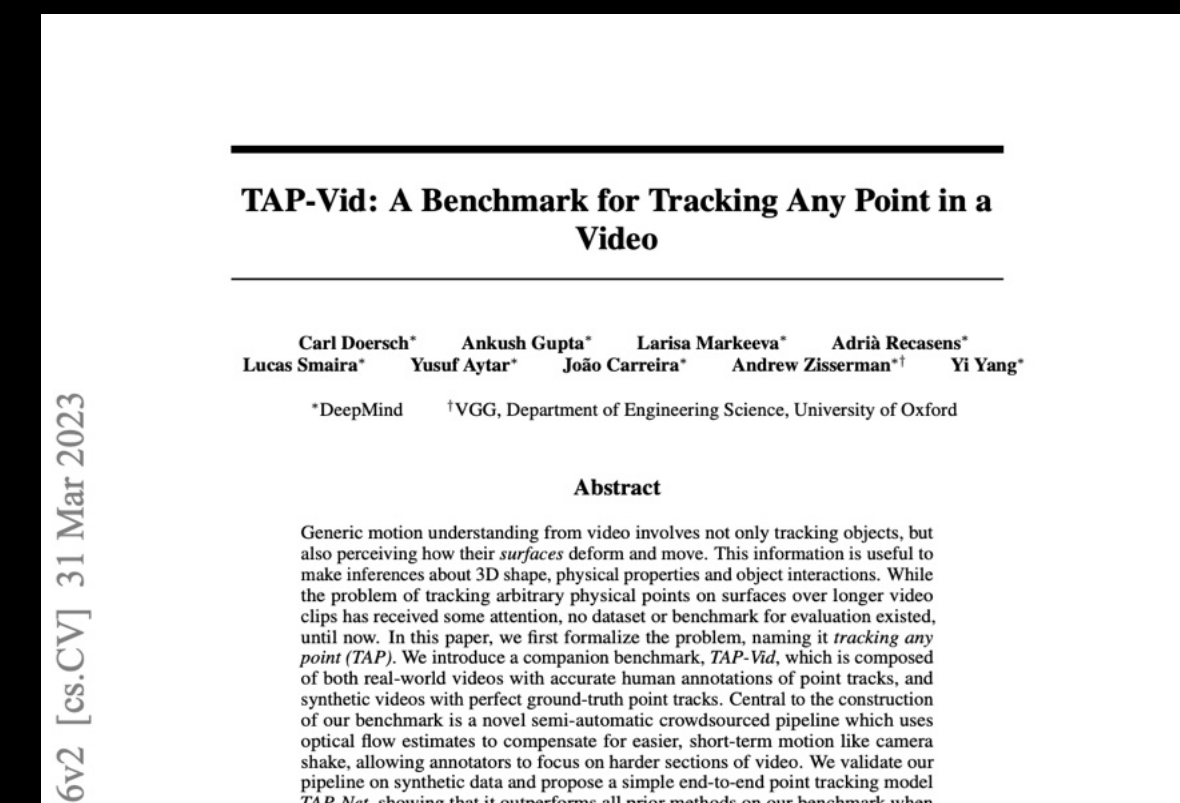
Cycle consistency

Harley et al.,
ECCV 2022.



Median trajectory error,
Percentage of Correct
Keypoints (PCK)

Doersch et al., NeurIPS Datasets &
Benchmarks, 2022.



d_avg (PCK),
Occlusion Accuracy,
Average Jaccard



original frame - 1280x720

δ_{avg}

δ_{avg}

Euclidean distance thresholds





δ_{avg}

79.2%

80%

Where do we evaluate point tracks?

Driving: KITTI, Waymo (adapted by PIPs, DriveTrack)

Animals: BADJA, Horse10 (adapted by PIPs, AllTracker)

Robotics: RGB-Stacking, RoboTAP (adapted by TAPVid)

Egocentric: EgoPoints, Aria Digital Twin (adapted by EgoPoints, TAPVid-3D)

Capture studio: Panoptic Studio (adapted by TAPVid-3D)

Youtube: DAVIS, Kinetics (adapted by TAPVid)

Surveillance: CroHD (adapted by PIPs)

Biology: Cell Tracking, SpaceAnimal, Tri-Mouse, Schooling-Fish

Full-Suite Evaluation is Getting Difficult

oTAP) we only track points from the first available query frame, which we find gives similar results to using all possible queries and full video lengths. We argue that using a wide array of benchmarks helps flatten the effects of label noise, and gives a better sense of the overall performance in practice. We encourage future work to follow our example, but we note that our full evaluation is time-consuming (with some baselines requiring multiple days to finish their pass).

Baselines We benchmark a variety of recent point trackers and optical flow models in our tests, including the optical flow models RAFT [47] and SEAF-RAFT [54], the long-term



Annotation 20 – horsejump–high | frame 000/049 | visible

track xy=(196.47, 176.96) nearest pixel=(196, 177) RGB=(28, 30, 17)



101x61 source pixels at 8x; hollow box = nearest pixel

Annotation 14 – car-roundabout | source frame 000/074 | visible

track xy=(465.48, 214.60) nearest pixel=(465, 215) RGB=(31, 33, 32)



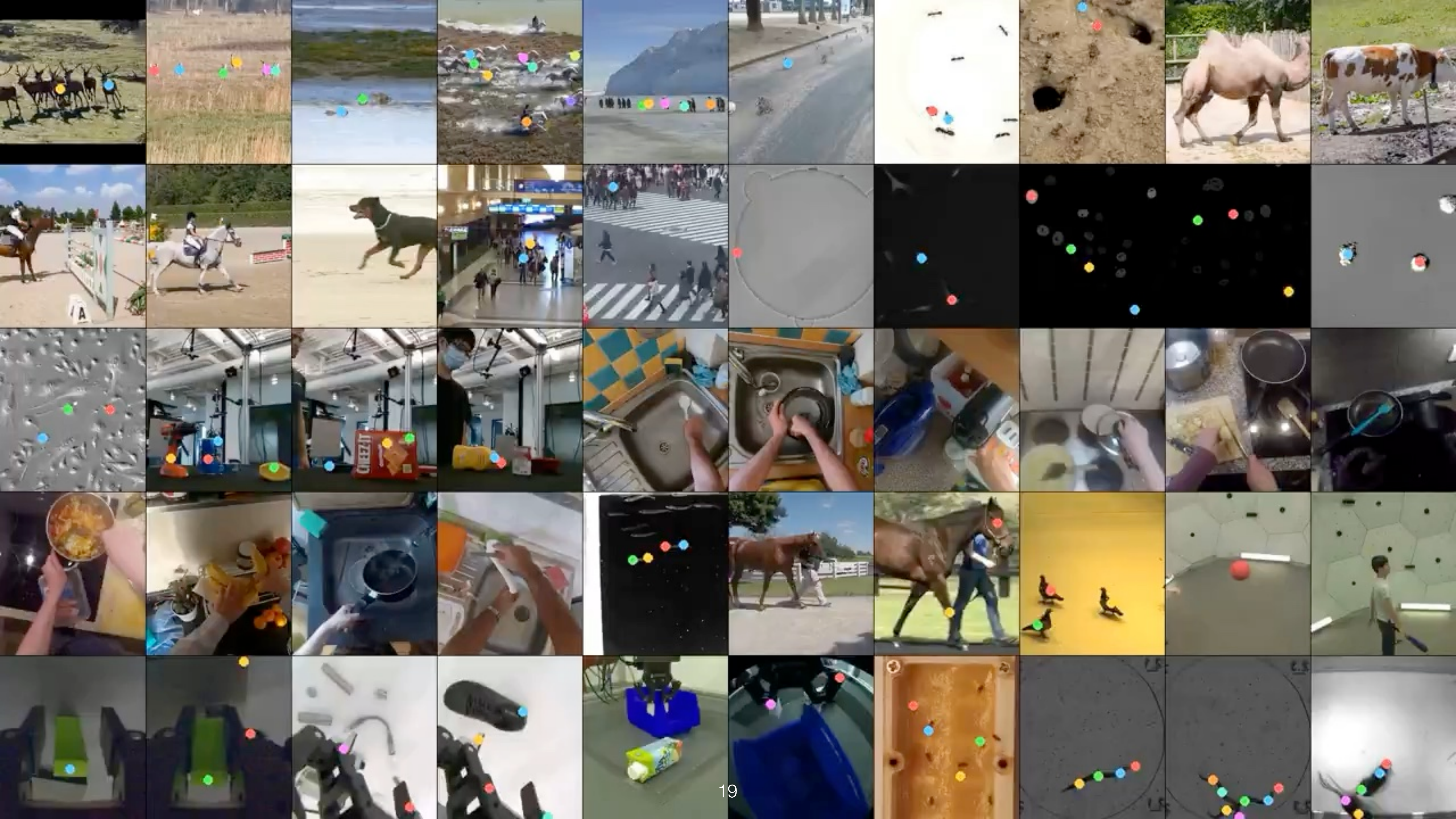
61x61 px at 6x; box = nearest pixel

VOTSp goals

Selective Track Curation
(Hand-picked)

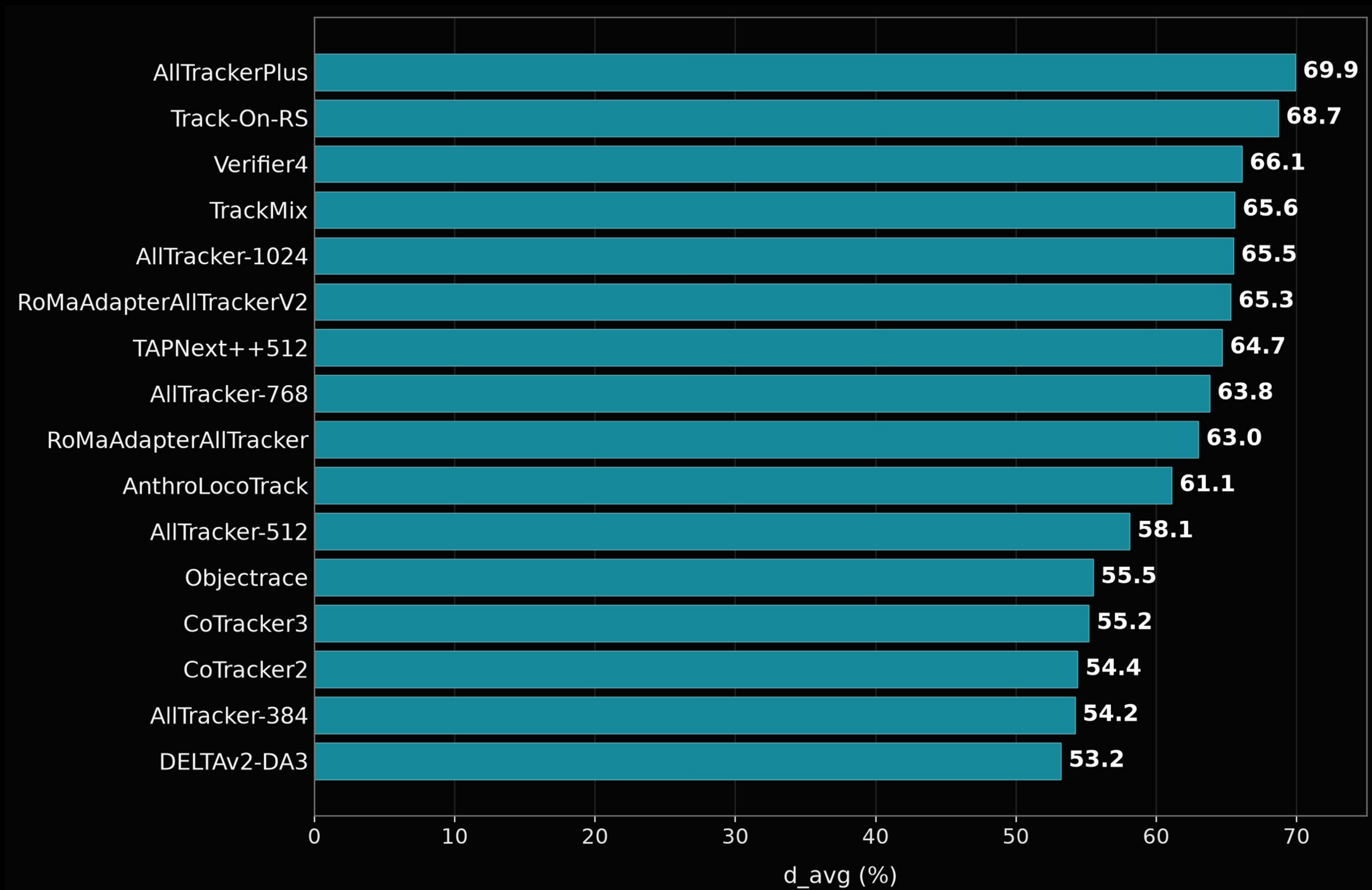
Domain diversity
(15 source datasets)

Lightweight
(50 sequences total)

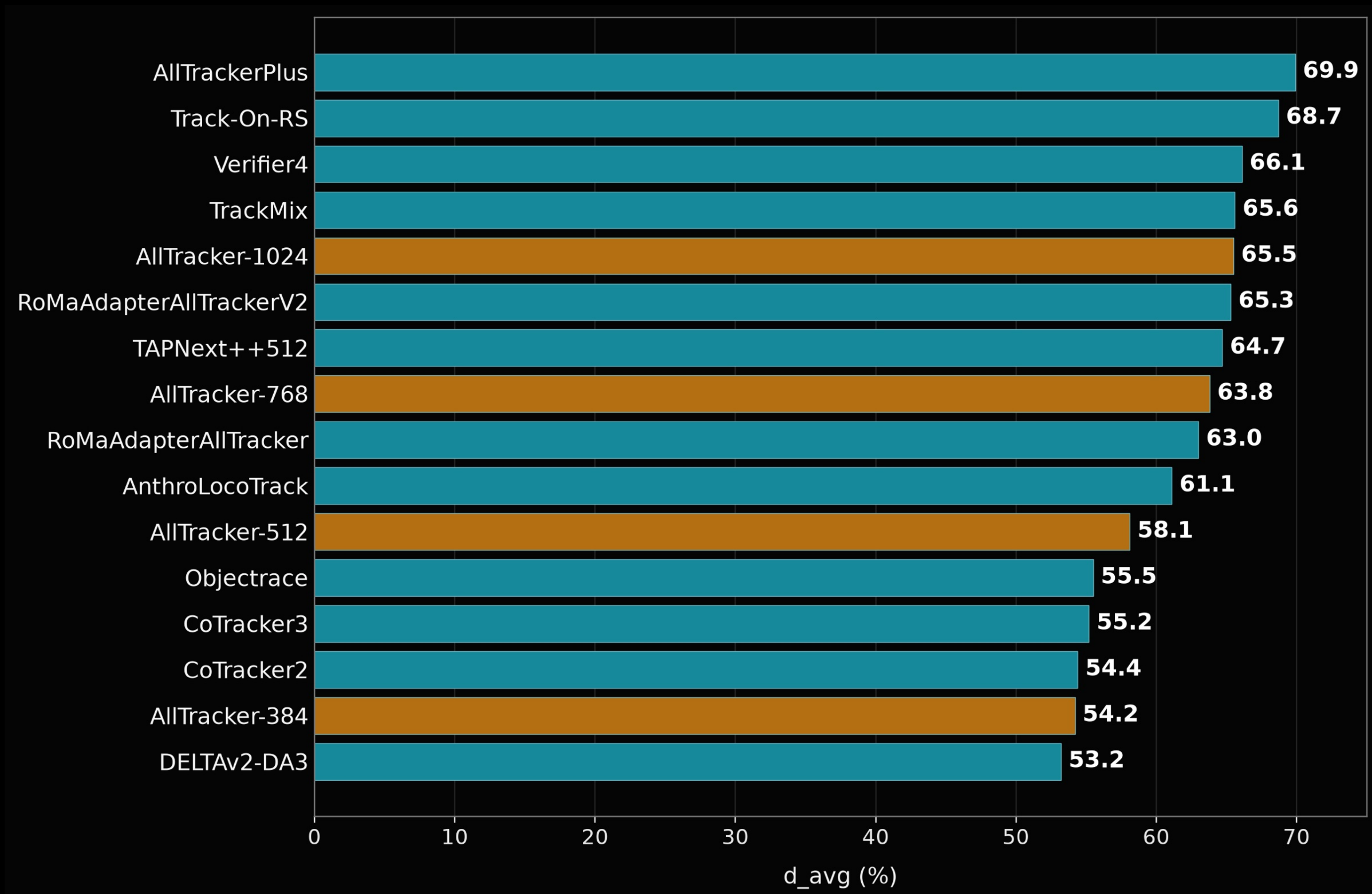


Results & Analysis

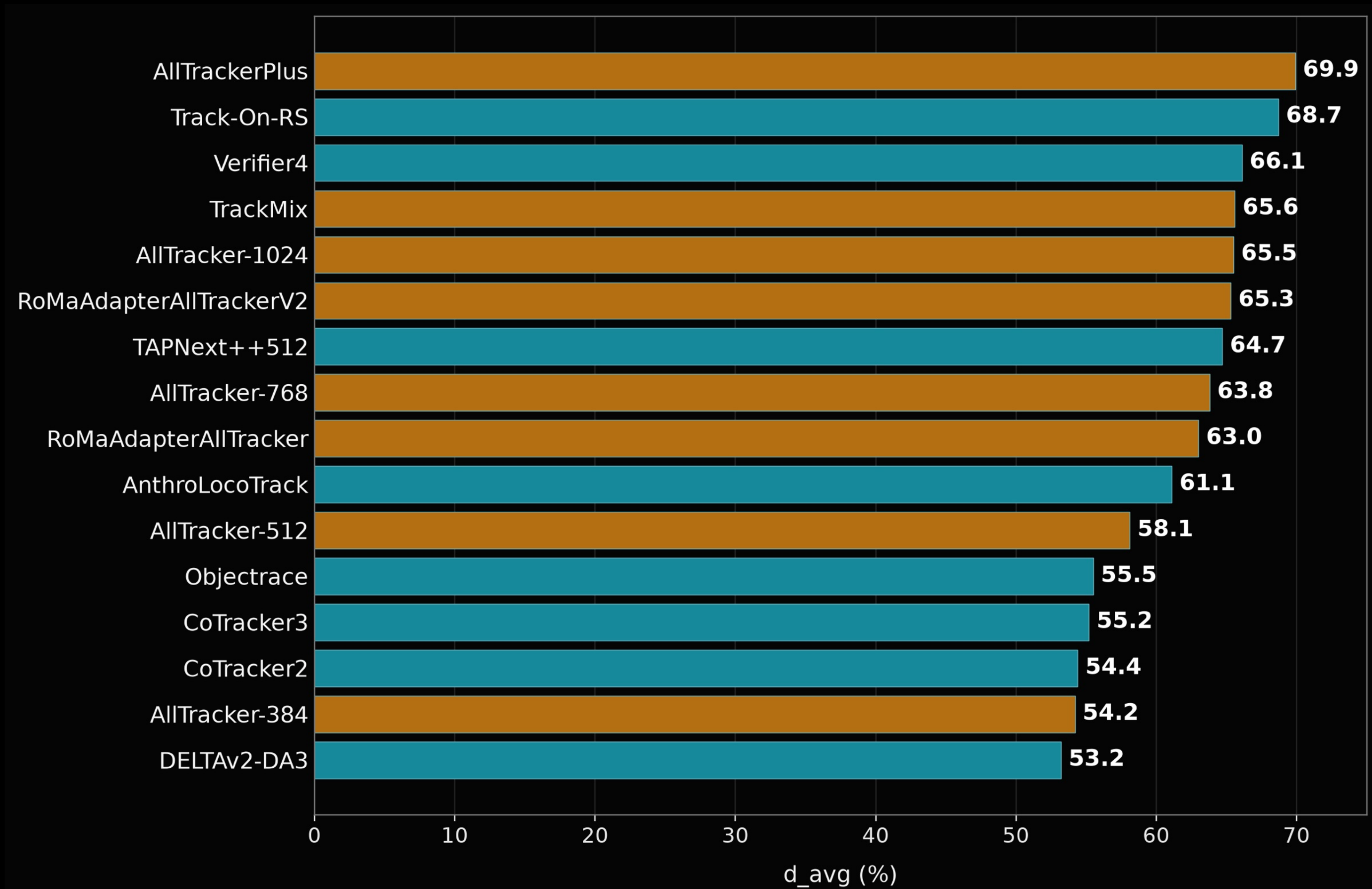
VOTSp Results: 16 submissions



Baseline: AllTracker (at multiple resolutions)

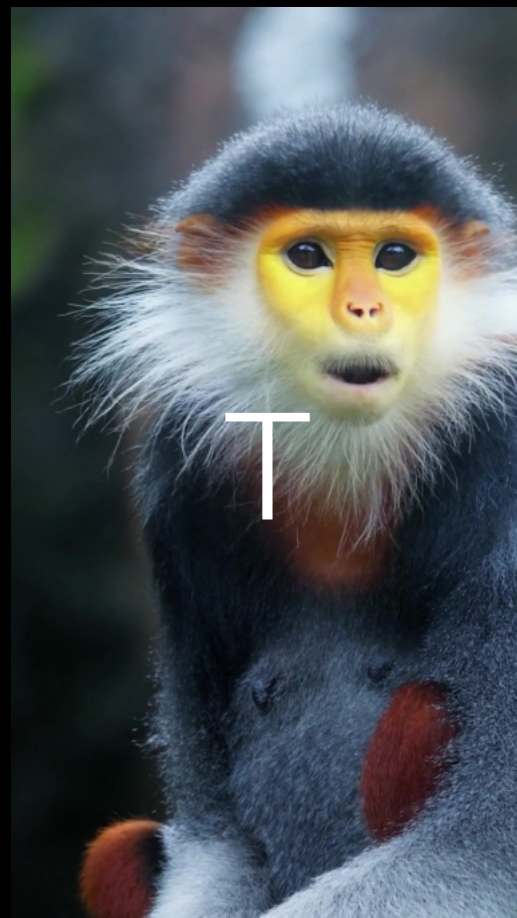
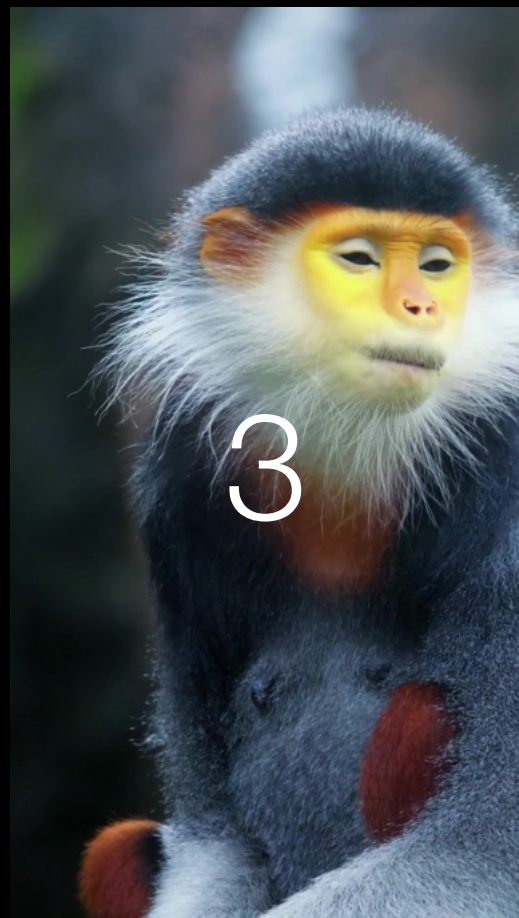
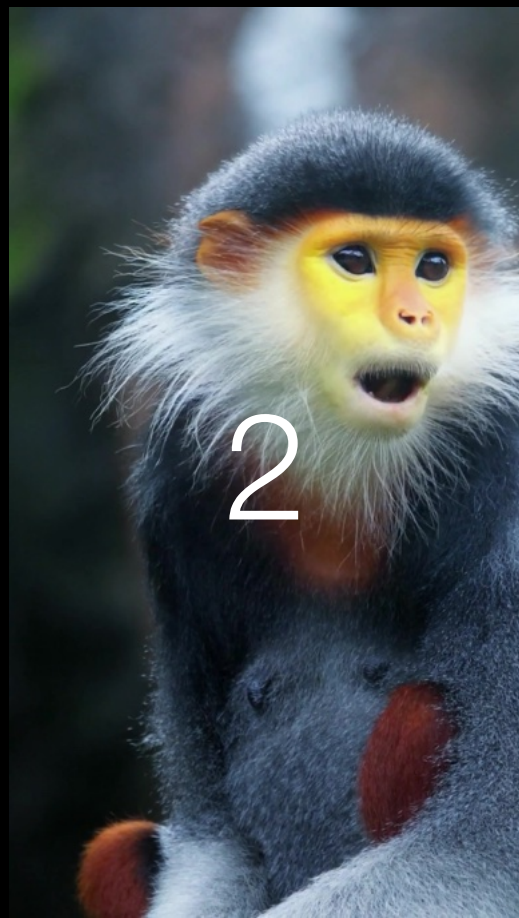
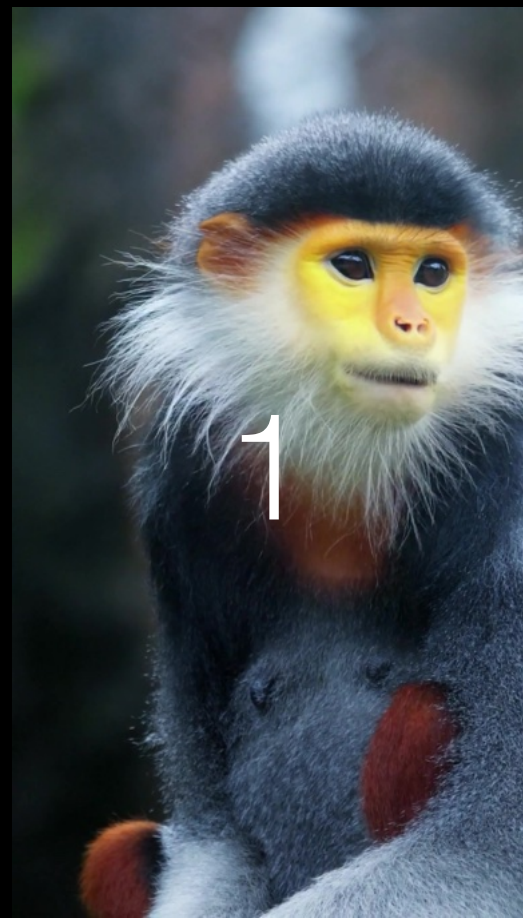
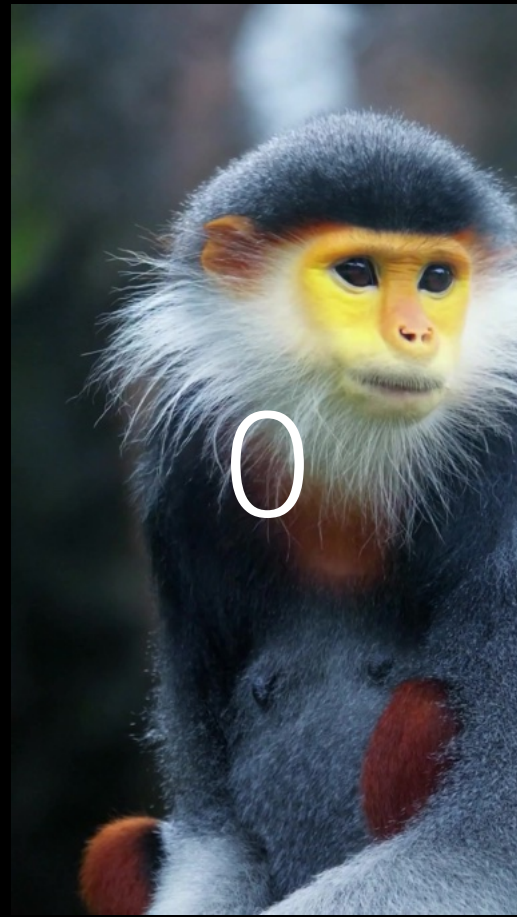
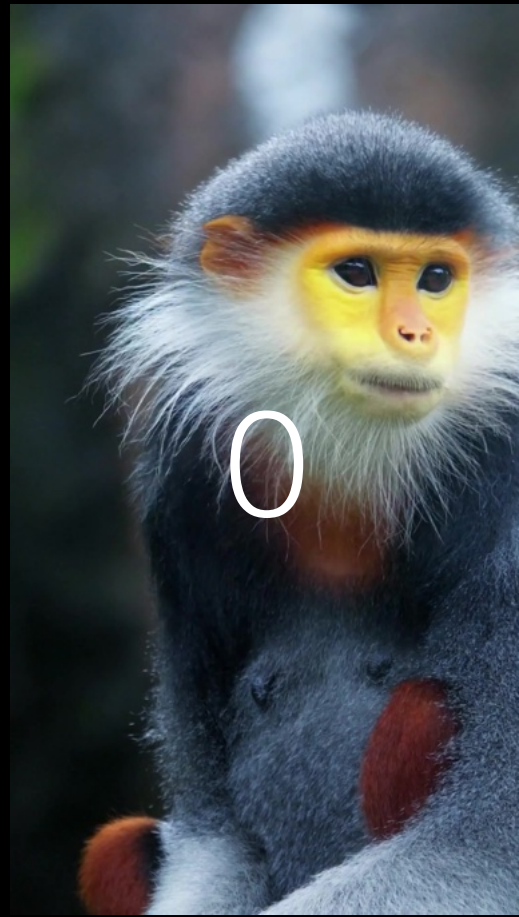
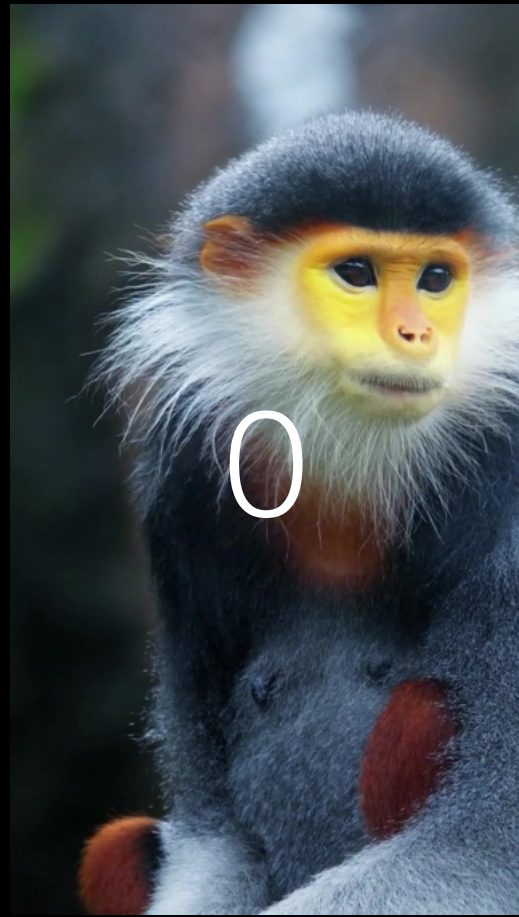
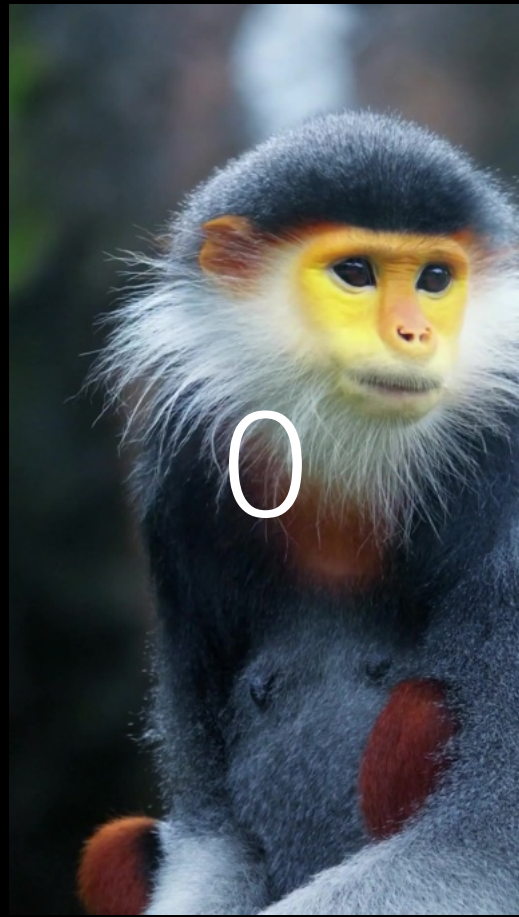
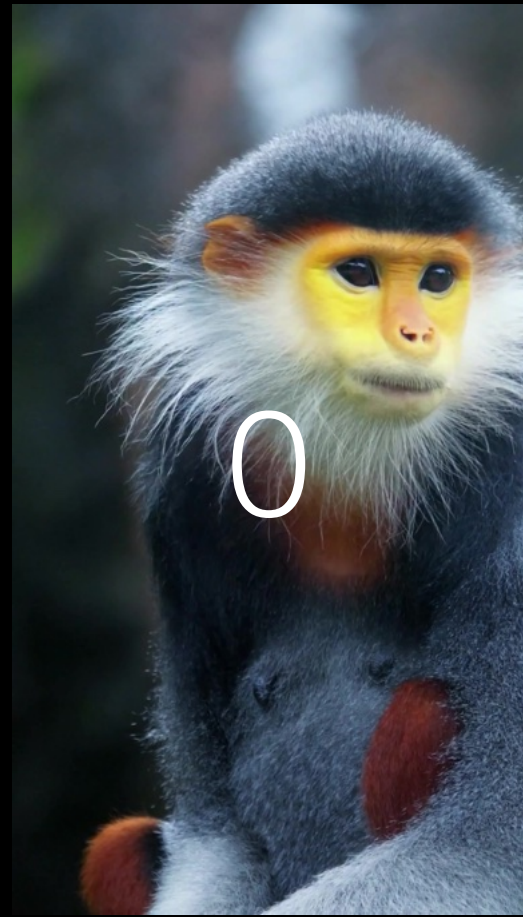


Many submissions based on AllTracker

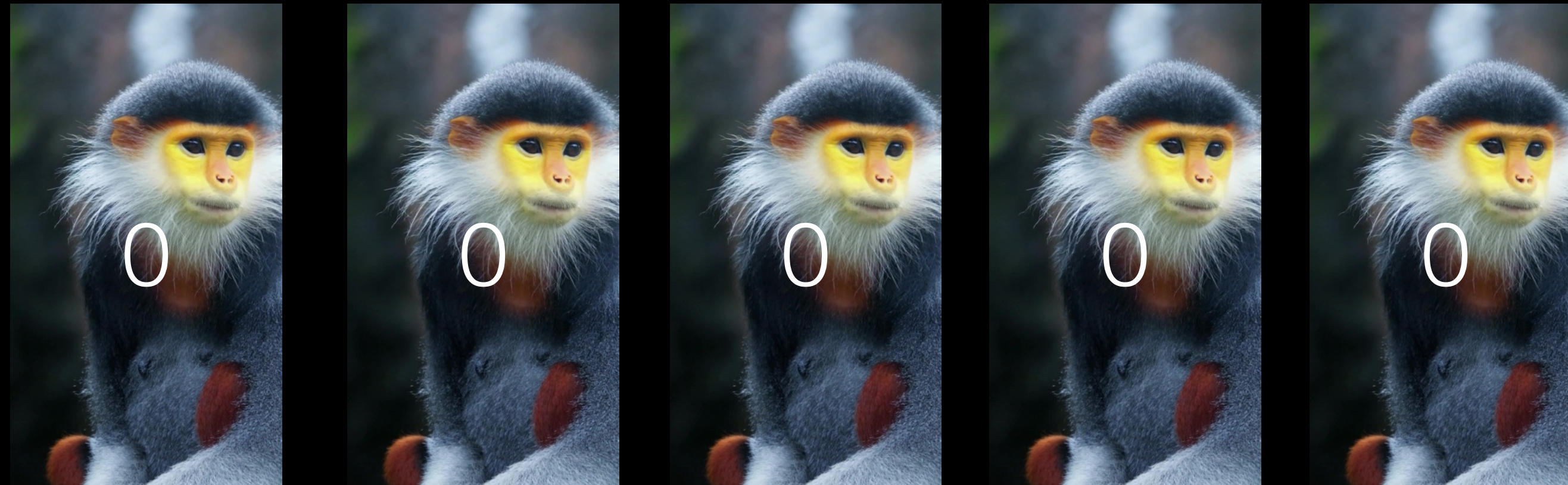


AllTracker Overview

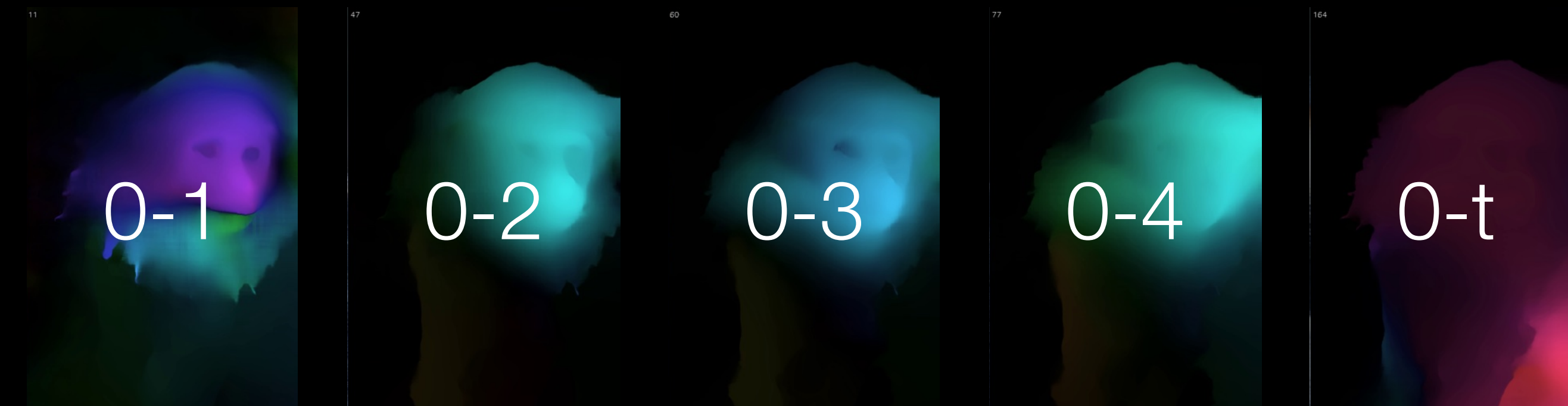




Source
frame



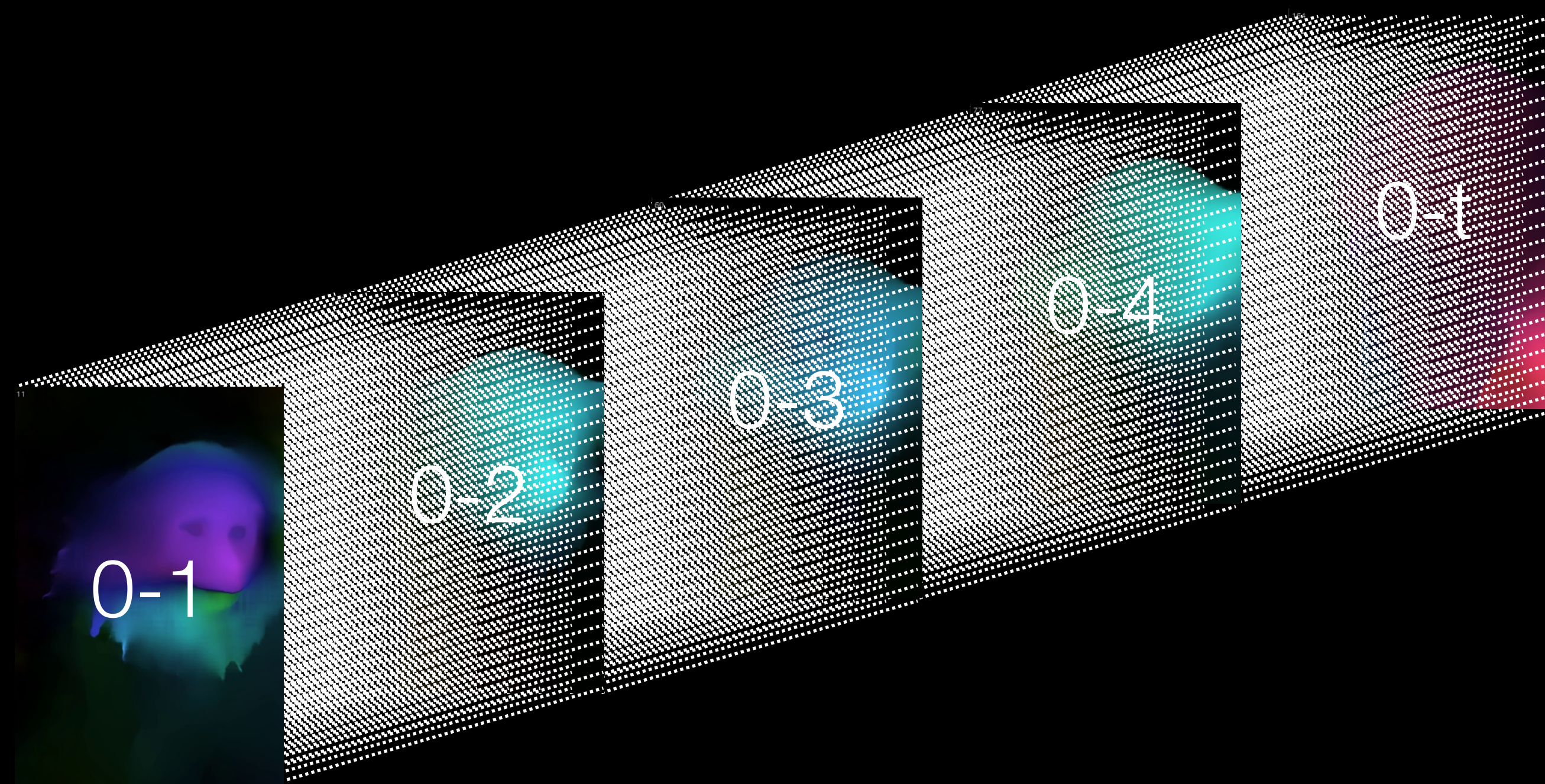
Optical
flow



Target
frame

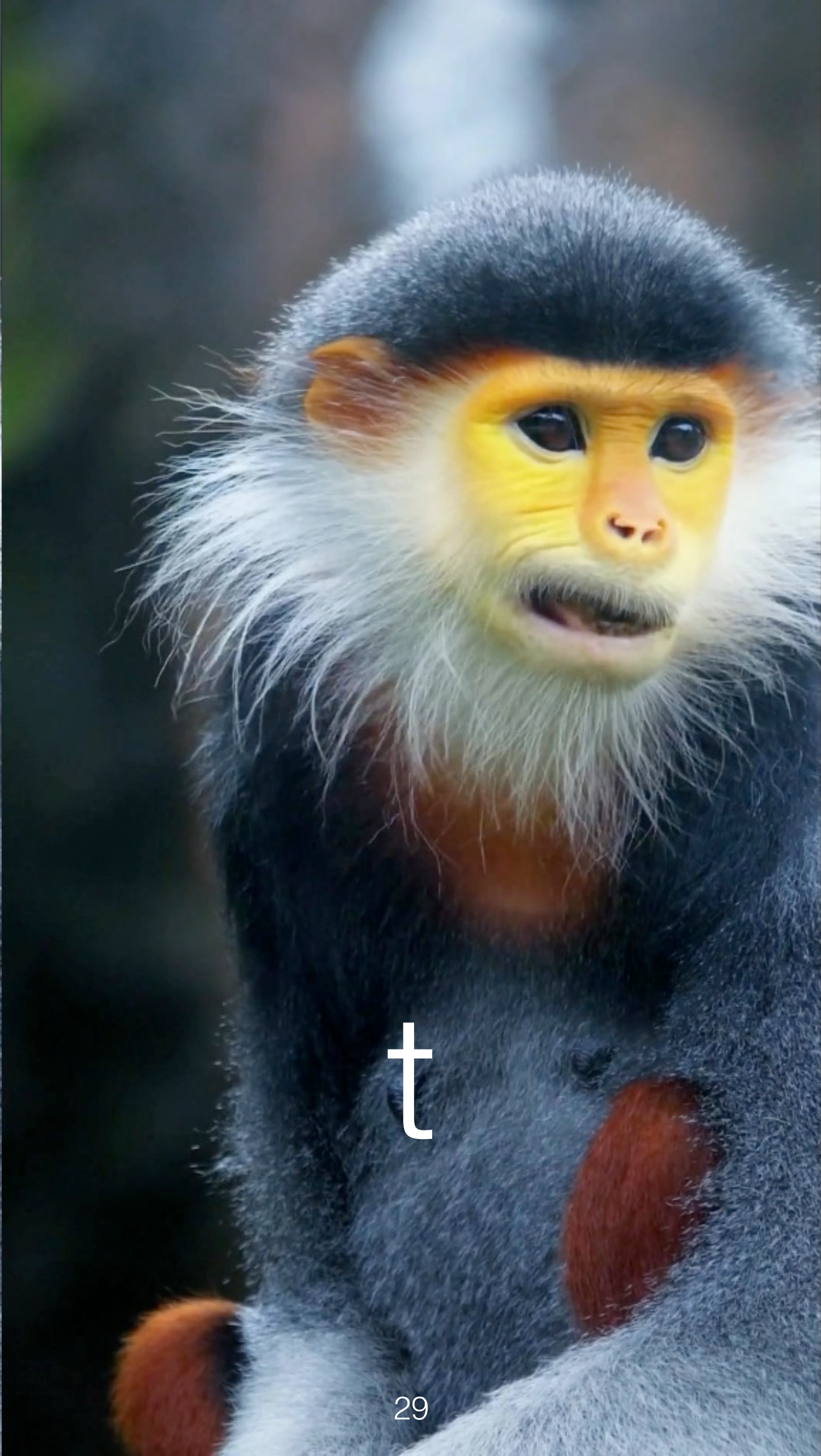


Pixel-aligned temporal attention

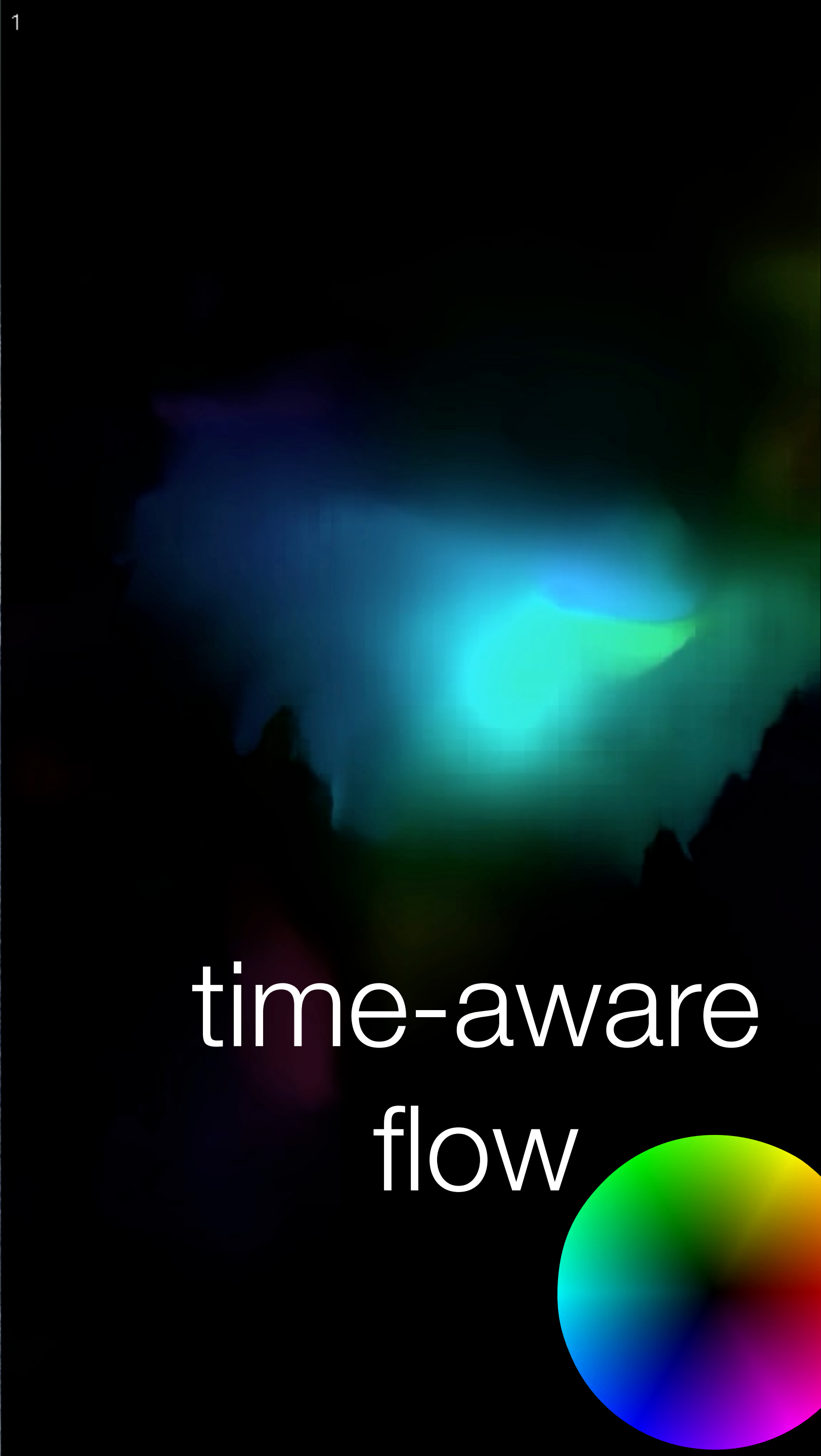




s

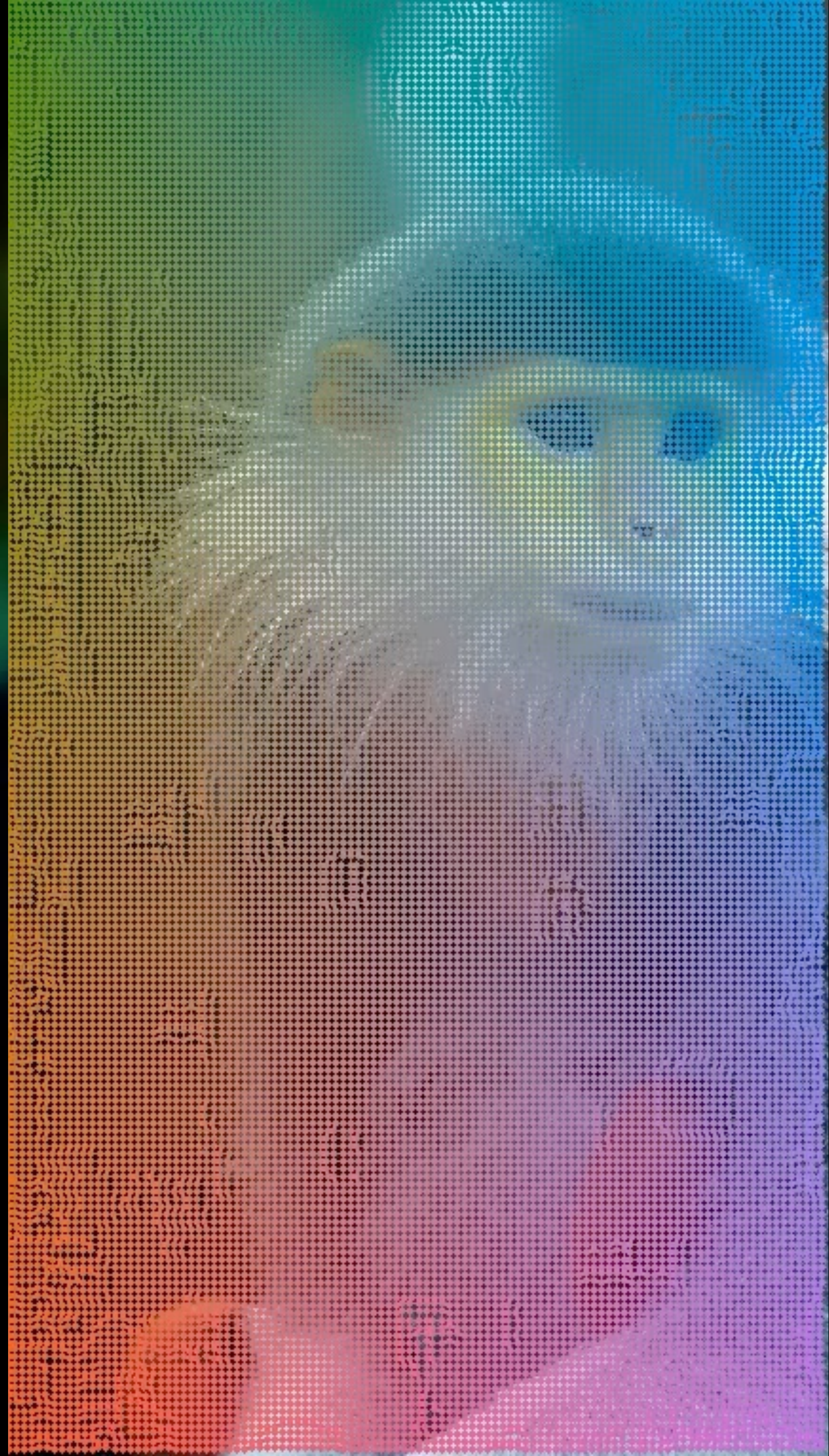
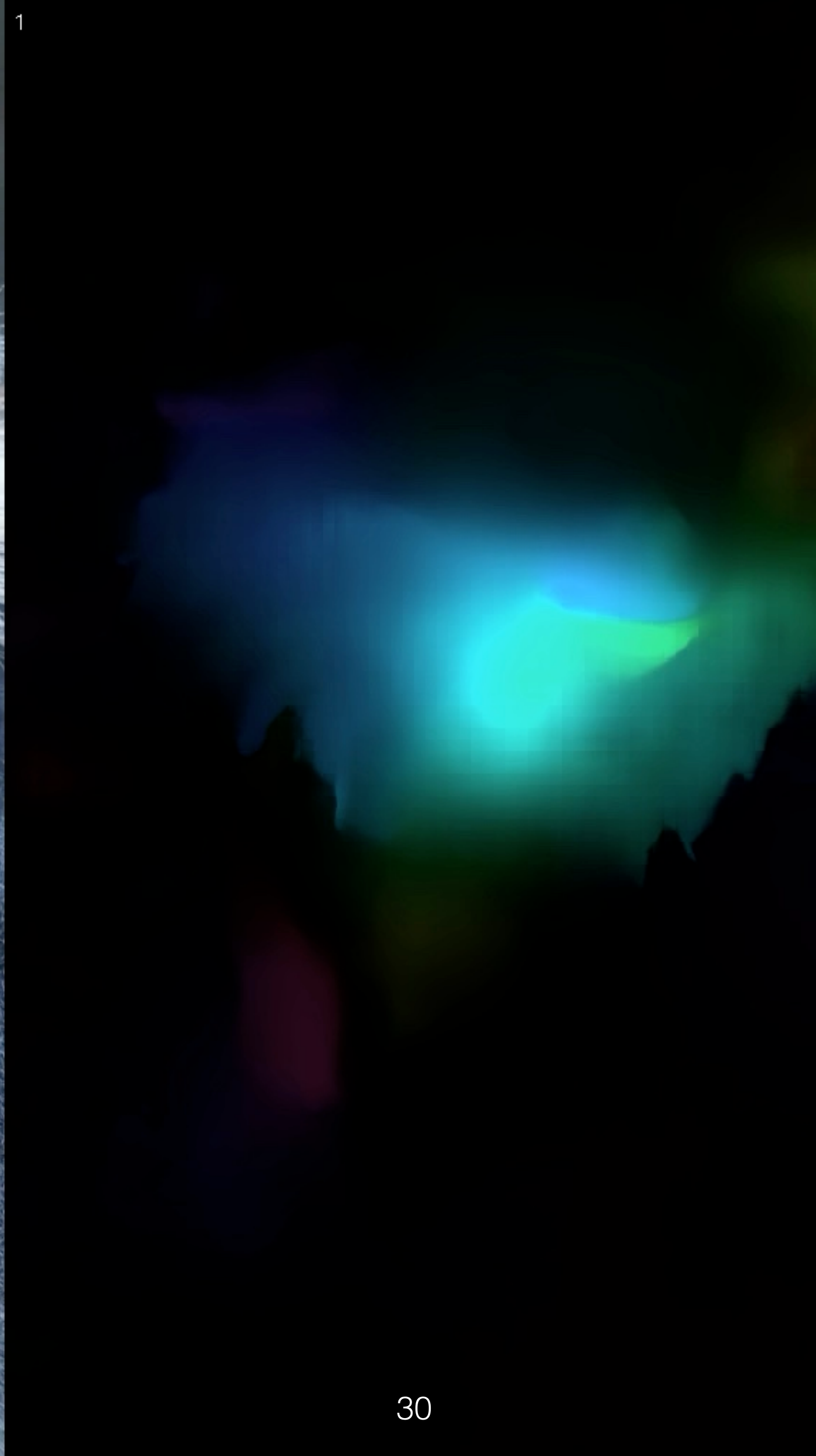


t

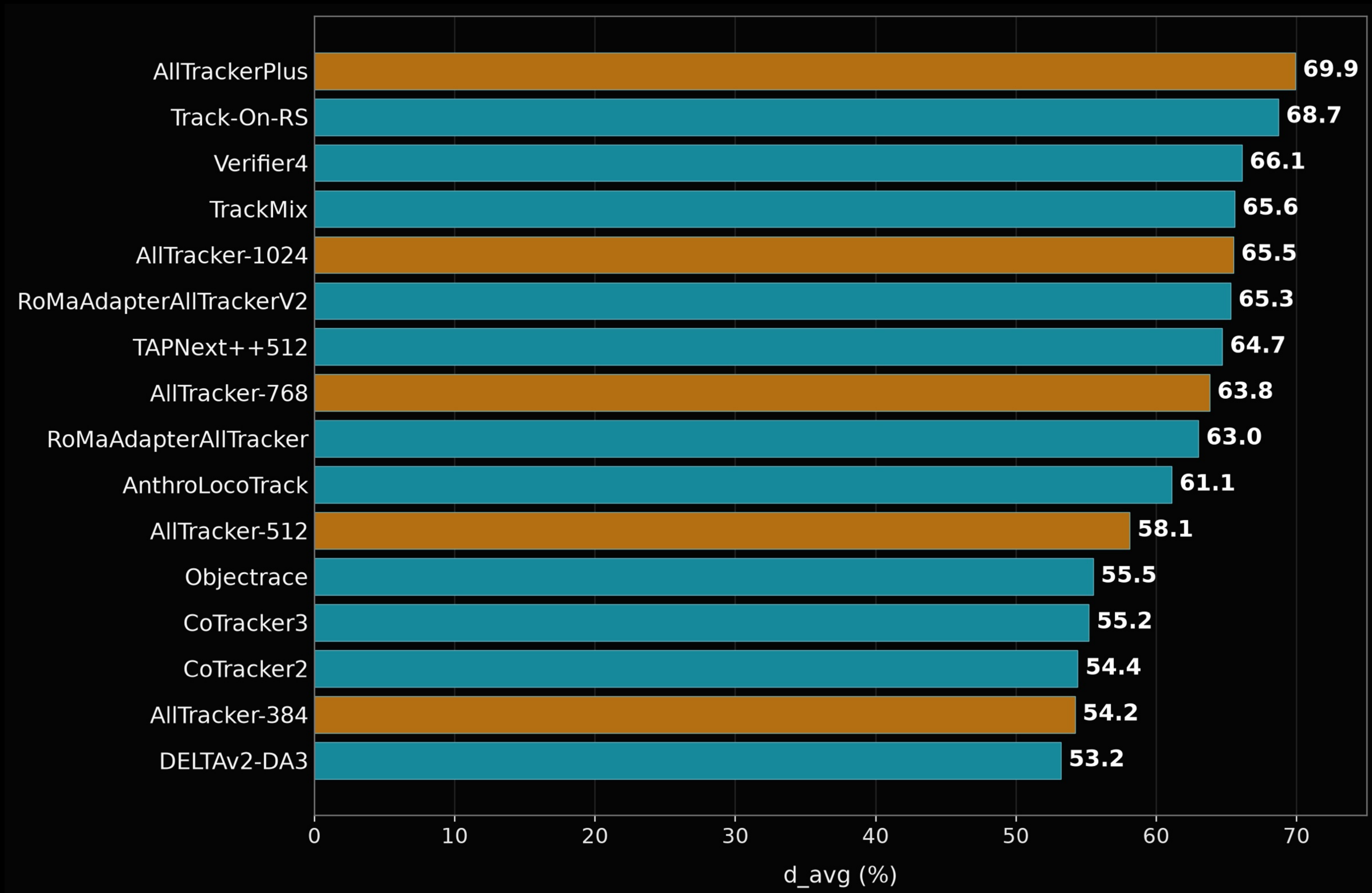


time-aware
flow



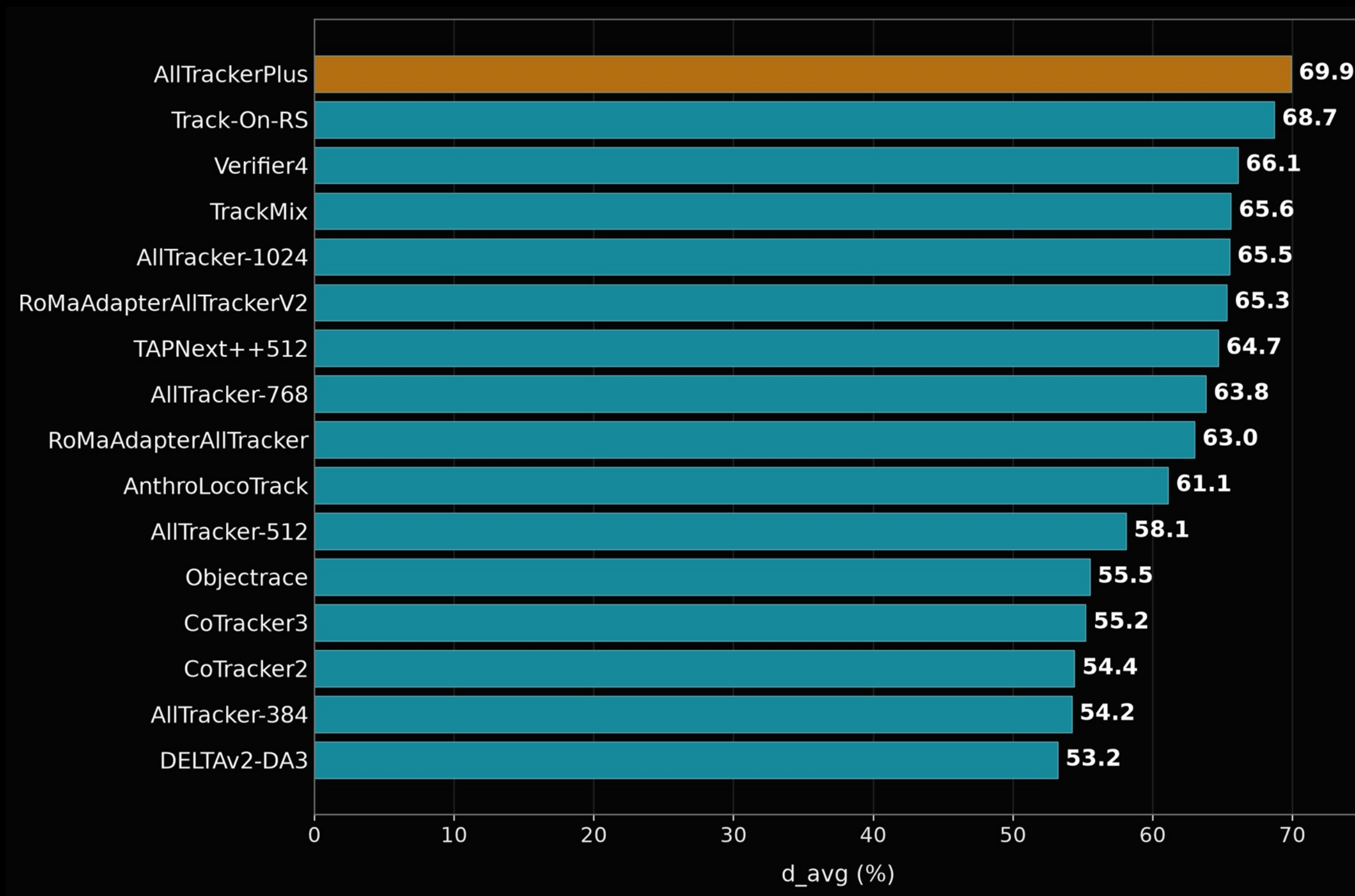


AllTracker scales well with resolution



Winner: AllTrackerPlus

Zhenglin Du, Zhengyang Li, Yi Wen (Xidian University)



Keys to victory

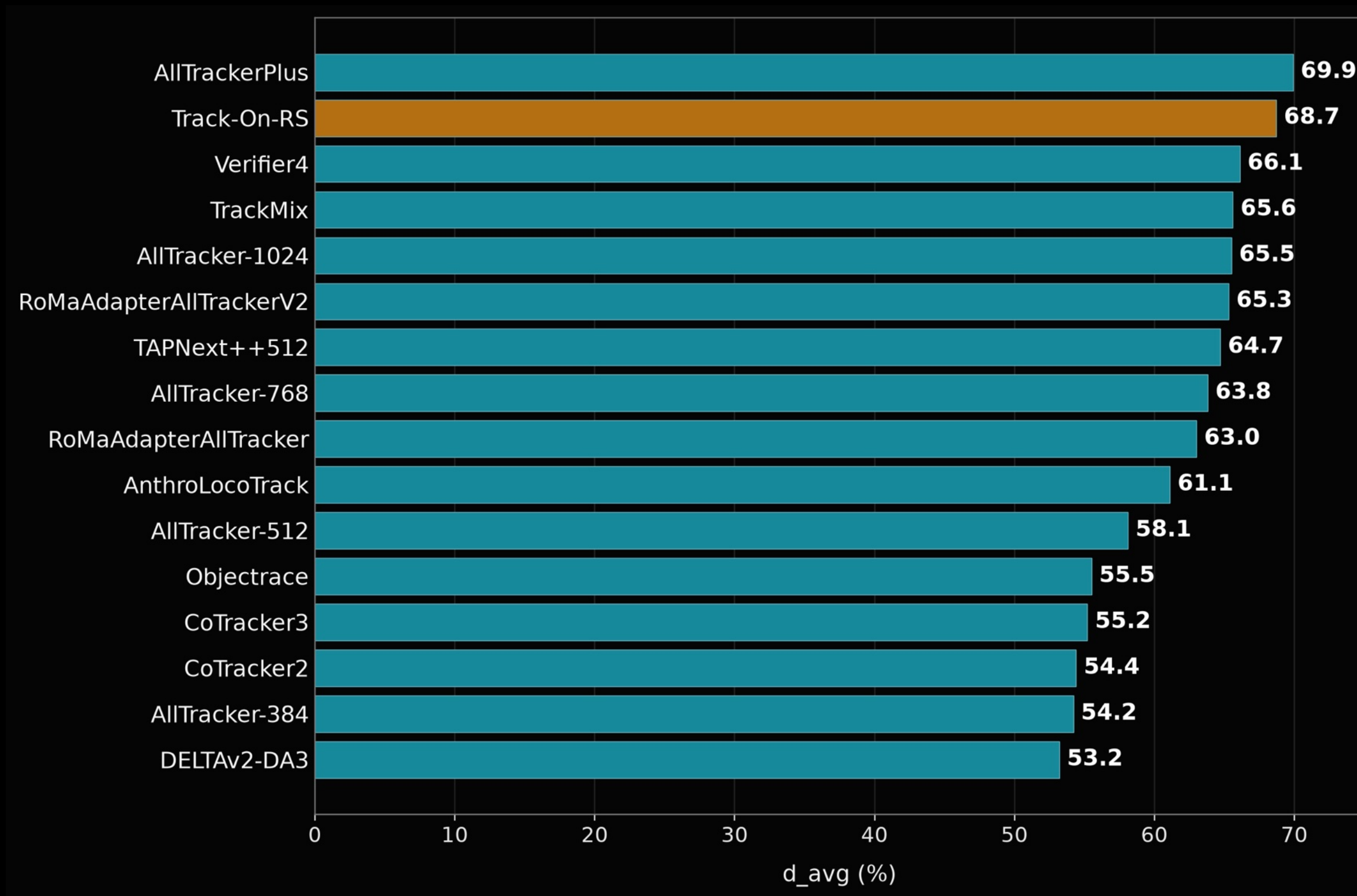
better memory engineering

forced max-resolution inference

base `alltracker.pth` weights

Runner-Up: Track-On-RS

Görkay Aydemir, Yunus Kurt, Volkan Aydingül, Nasrin Rahimi, Misra Yavuz, Burak Biner, Ahmet Emirdagi, Suleyman Aslan, M. Akin Yilmaz (Codeway AI Research & fal.ai)



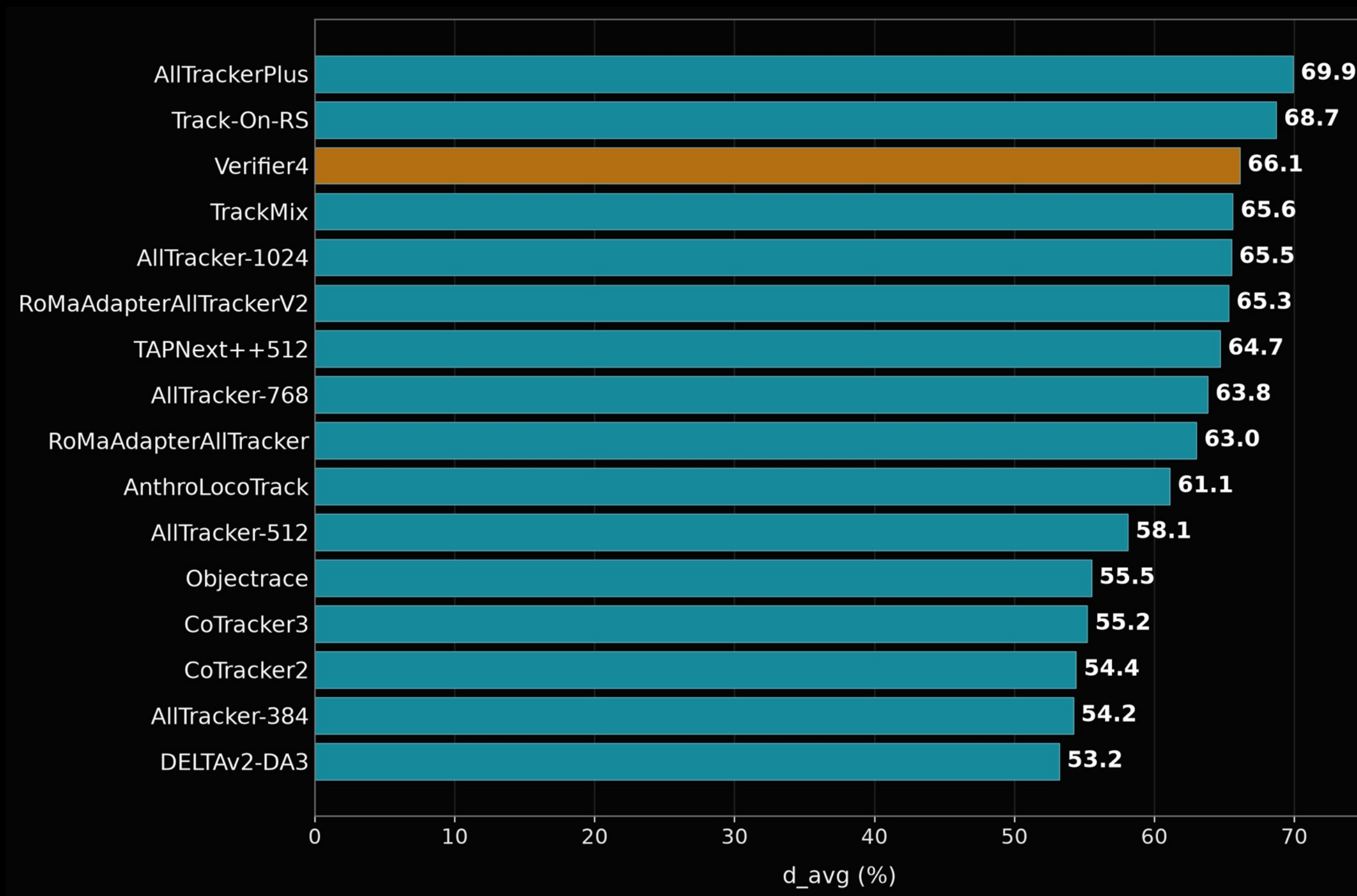
Key Mechanics

causal tracker (!)

SAM3 guidance

Highlight: Verifier4

Jinxia Xie, Liangtao Shi, Daoheng Li, Zixu Li, Jianjun Qian



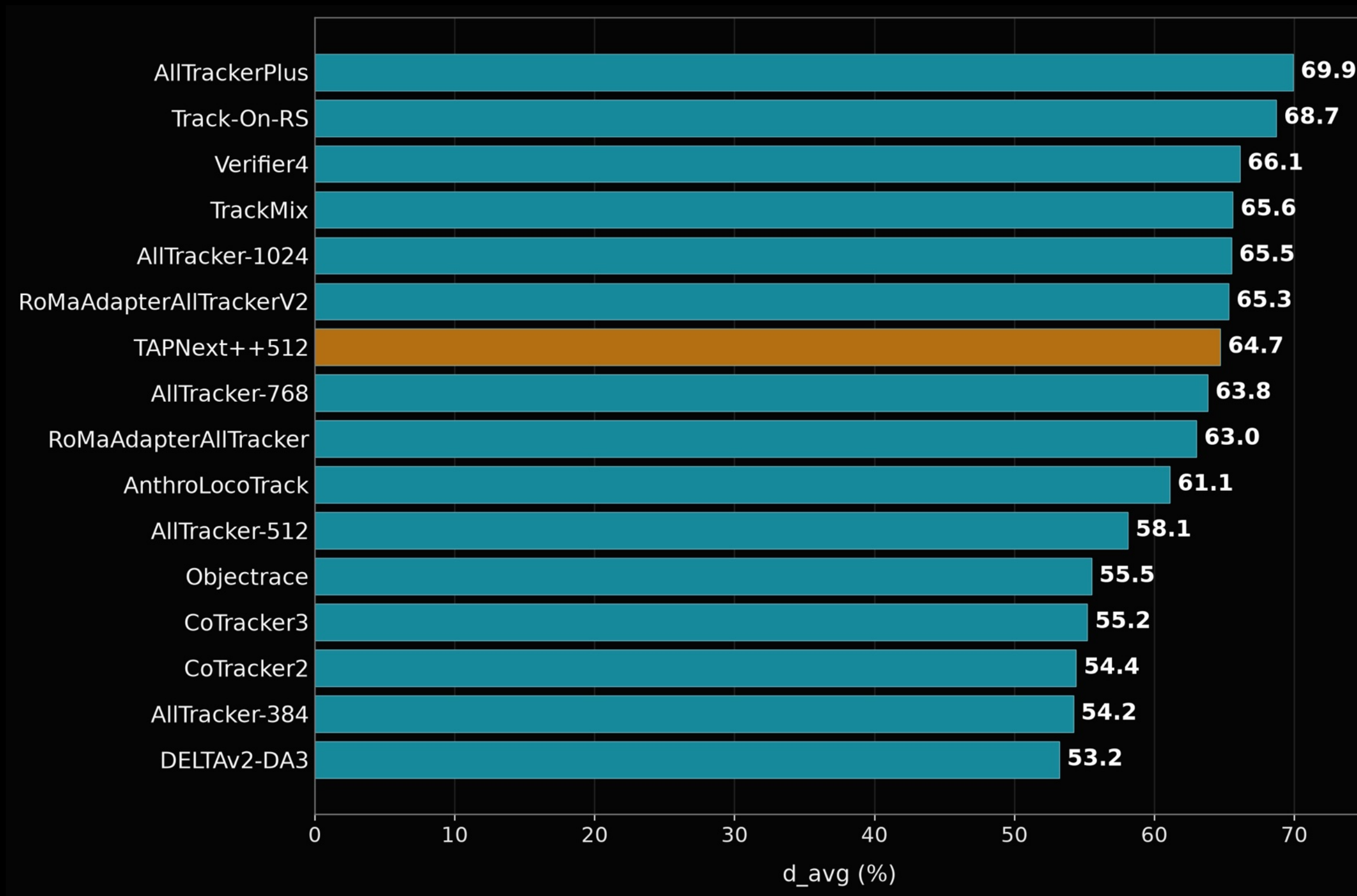
Key Mechanics

4 base trackers in parallel

verifier-guided frame arbitration

Highlight: TAPNext++512

Sebastian Jung, Martin Sundermeyer, Carl Doersch



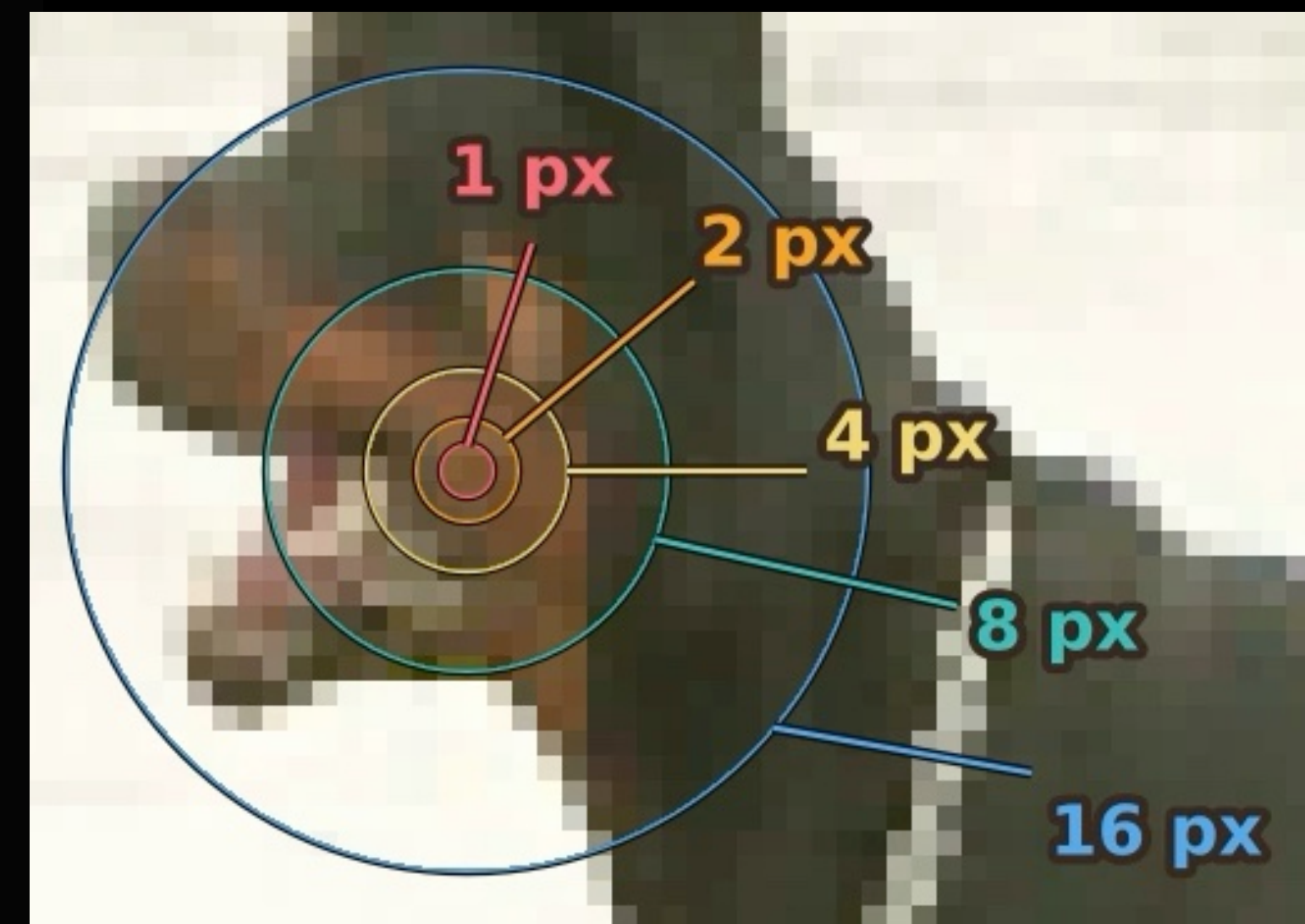
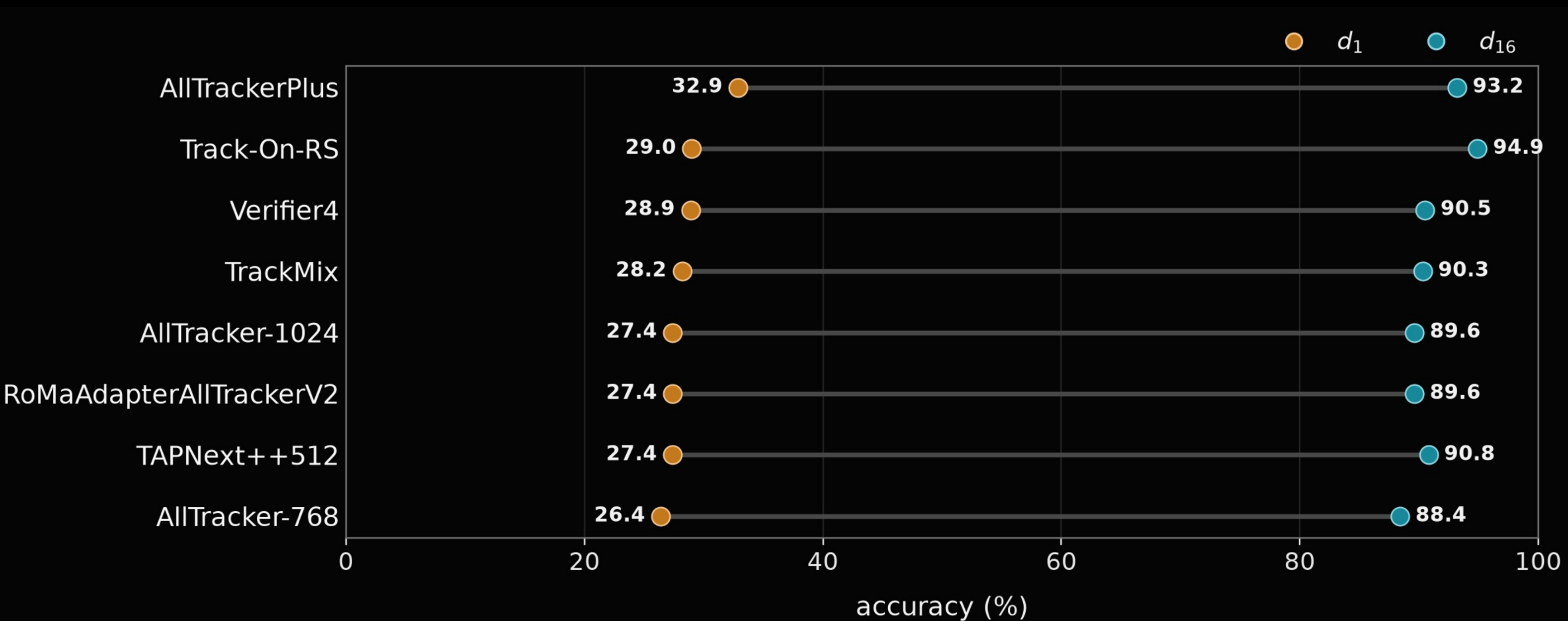
Key Mechanics

online recurrent state-space model

causal inference (!)

minimal inductive biases

All methods struggle with sub-pixel precision



Conclusion

Takeaways

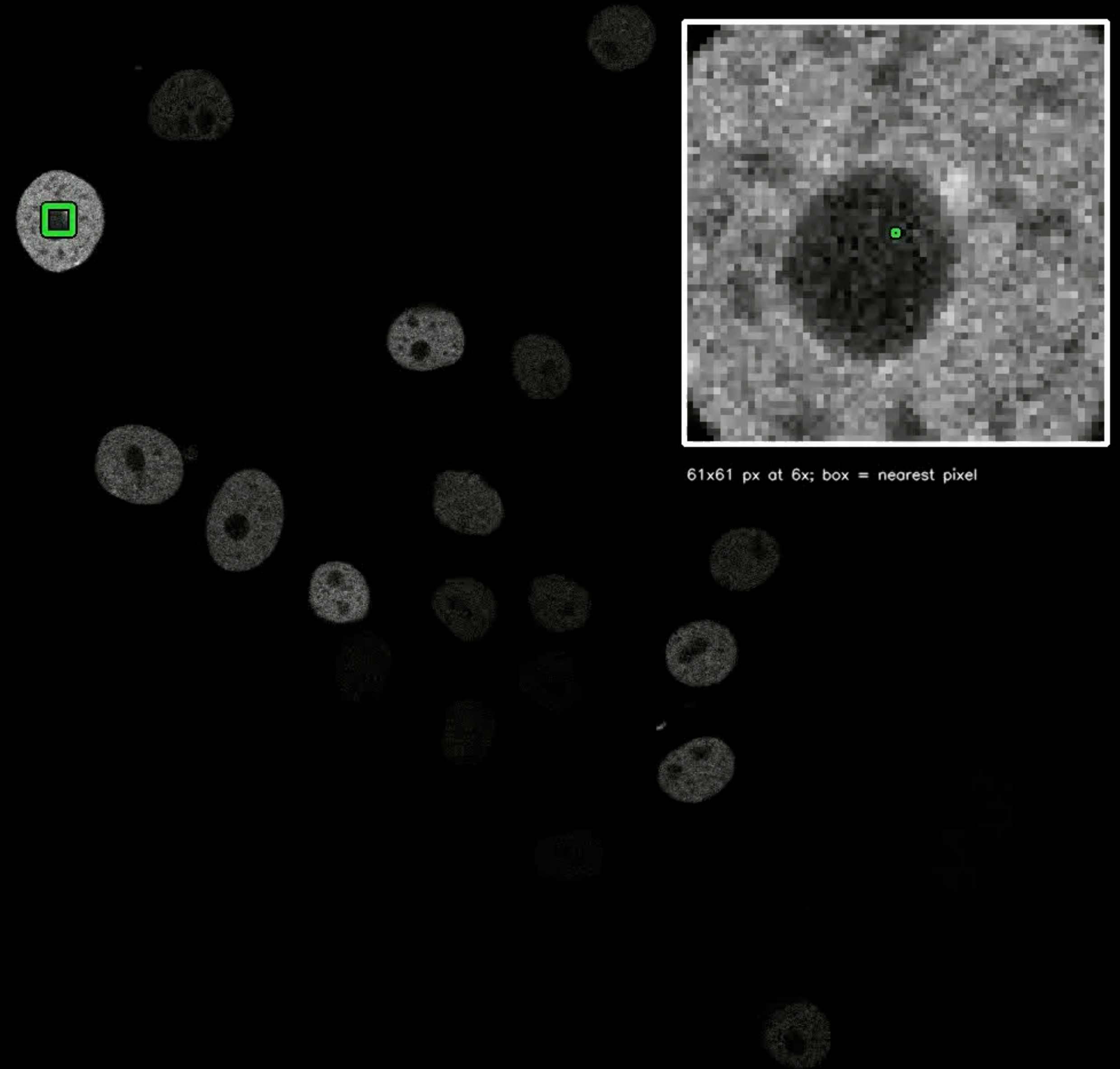
Resolution matters a lot

Causal methods are catching up

“Pure” data-driven methods are catching up

Limitation 1:

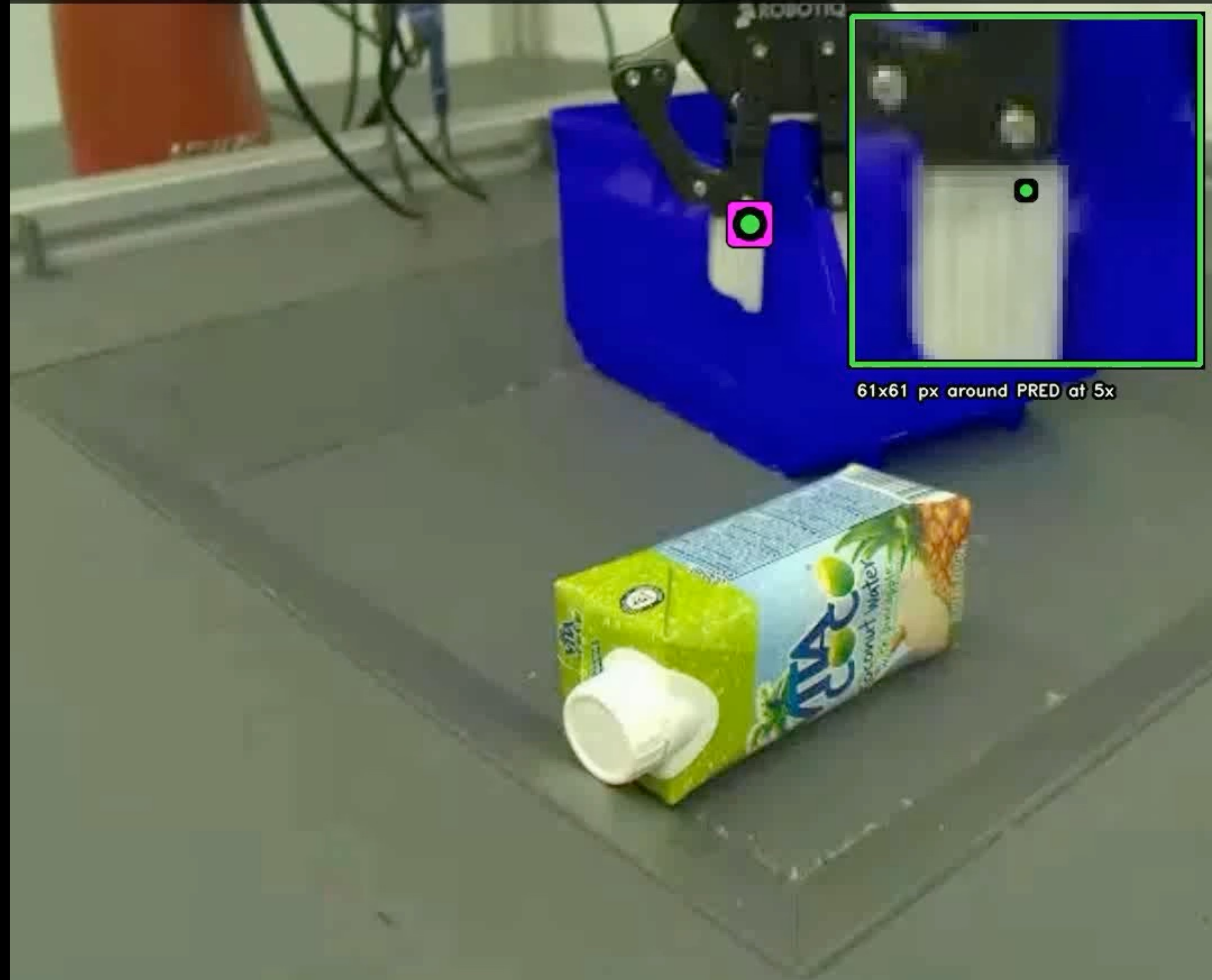
not all annotations
are pixel-precise



61x61 px at 6x; box = nearest pixel

Limitation 2:

position accuracy
was the only metric



61x61 px around PRED at 5x



point tracking

Adam Harley
Meta