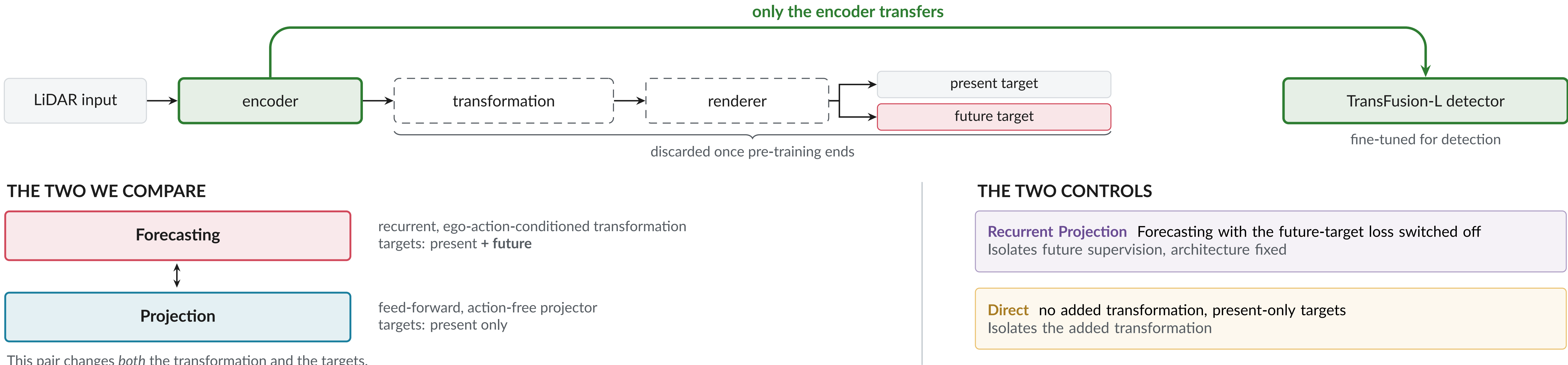


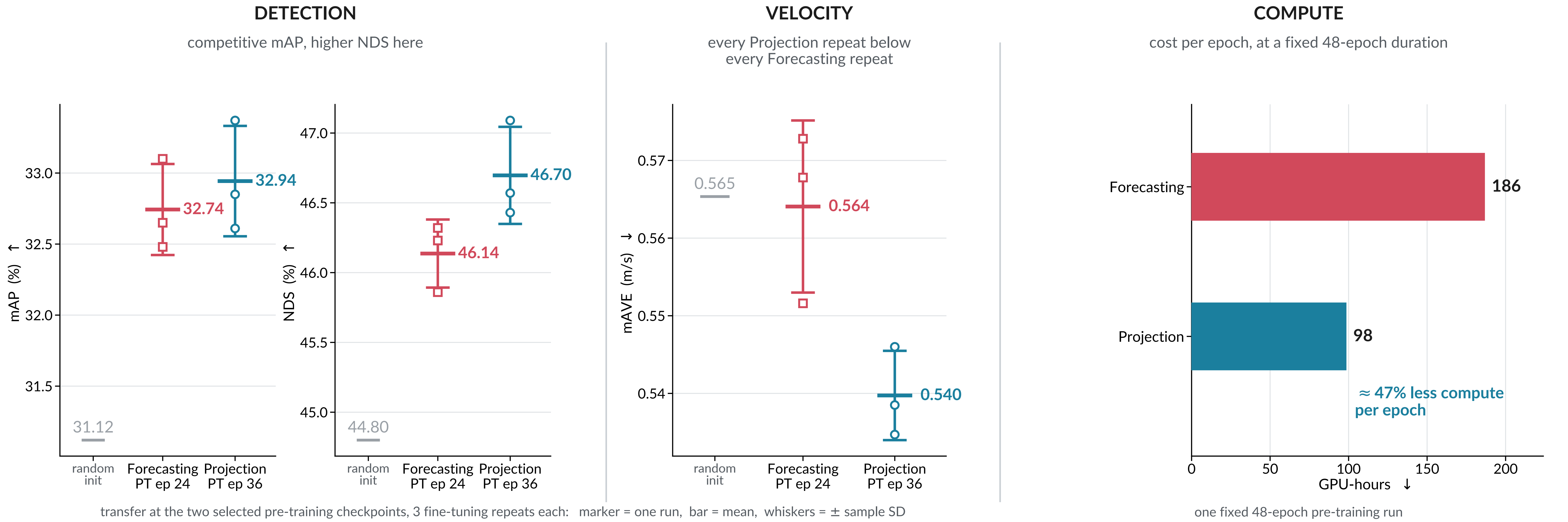
What Does the Future Buy?

Separating projection from prediction in LiDAR pre-training

Against full forecasting, present-only pre-training achieves competitive detection transfer on nuScenes and ONCE, higher NDS on nuScenes, and lower pre-training cost.

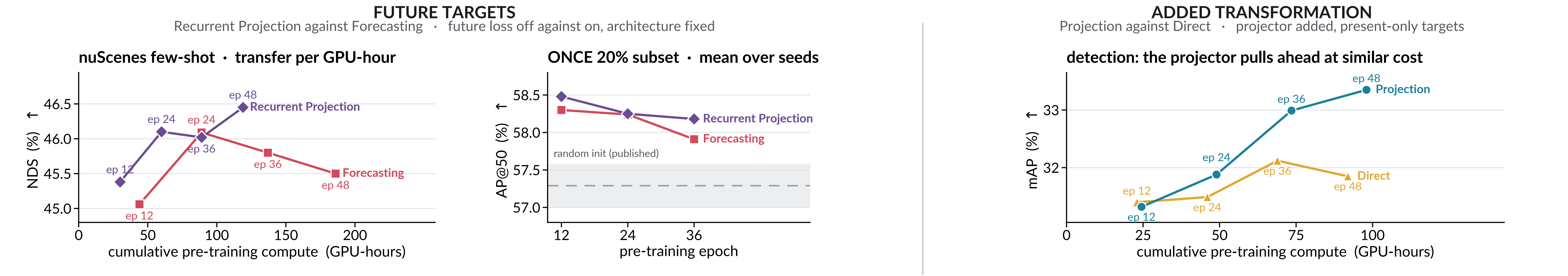


Competitive detection, lower compute, and lower mAVE.

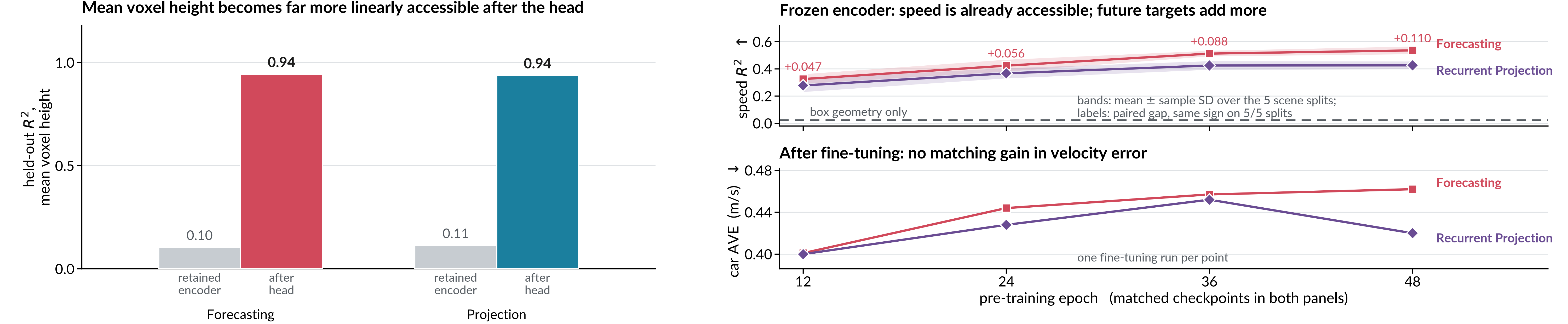


Checkpoints selected after examining detection results; one pre-training run per branch.

What supports this simplification?



What changes in the representation?



What the head changes. The head adapts the encoder's features to the rendering objective, turning them into a geometrically explicit code that fine-tuning then discards.

What the future changes. Future targets improve the frozen speed readout, but do not reduce the detector's velocity error.

Probe: linear regression of car speed from frozen encoder features pooled within ground-truth car boxes

A present-only projector supports competitive detection transfer with less pre-training compute. The two ingredients come apart: future targets change what the frozen encoder exposes, not what the detector achieves.

Open question: which driving tasks benefit from future-target supervision beyond what present-only pre-training provides?

Scope. One pre-training run per branch; selected checkpoints; one few-shot regime and one detector per dataset. ONCE is scored with AP@50, which carries no velocity term, so it speaks to detection only.

[1] TREND: Unsupervised 3D Representation Learning via Temporal Forecasting for LiDAR Perception, arXiv:2412.03054. A discarded projection before a rendering decoder is established practice (UniPAD, CVPR 2024); what is decomposed here is the *temporal* recipe.



details and code