

This Week in Robotics

Window: 2026-09-07 → 2026-09-13.

THE LINE

The tooling layer moved for the first time since spring — [LeRobot v0.6.1](#), [MuJoCo 3.13.0](#) with a new discrete integrator, [Genesis 1.4.1](#) — and [OpenWAM](#) turned last week's uncarded checkpoint dump into a fully released, controlled study of world-action-model pretraining, while the week's cleanest manipulation result used no learning at all.

1 Policy learning & VLAs

IMLE-VLA

PREPRINT / SFU + UPenn

Architecturally the change is one object: the 10-step flow-matching action head of $\pi 0.5$ is replaced by a single-step conditional generator trained with conditional Implicit Maximum Likelihood Estimation. Per observation you draw m noise vectors, generate m candidate chunks, assign each ground-truth chunk to its nearest candidate without gradients, and backprop only through the winner. That nearest-neighbour assignment is what keeps the head from collapsing to the conditional mean the way an L1 head does — mode coverage by construction rather than by auxiliary regularizer. $m=2$ in practice. The backbone is frozen, so only the head trains.

Eval: LIBERO 40 tasks $\times$ 50 episodes (98.0% average, the highest among compared methods), LIBERO-Plus perturbations at four axes $\times$ five severities, and four real Franka tasks at 20 episodes each. Inference goes 15 Hz $\rightarrow$ 55 Hz on an L40S, and at execution horizon 30 the throughput compounds to 11 $\times$; measured proprioceptive jerk drops 2.2–3.0 $\times$ on the real arm. The robustness claim is the one that matters: OpenVLA-OFT's L1 head degrades sharply under LIBERO-Plus severity while IMLE-VLA tracks $\pi 0.5$. Soft: LIBERO at 98% is a saturated instrument, and the LIBERO-Plus curves are read off a figure rather than a table. [Code released](#); head-only training on a frozen 3B backbone puts this squarely in one-GPU range.

ContextFlow

PREPRINT / KAUST + Berkeley

In-context imitation without gradient updates, moved off autoregression. Prior work (ICRT) tokenizes actions and decodes them one at a time, so an early error compounds — and compounds worst exactly under the configuration shift in-context learning is for. ContextFlow puts the demonstrations through perceiver-style compressors (32 learnable queries per modality, separate compressors for image, proprioception and action) into a fixed-size context, and a $\pi 0$ -initialised action expert flow-matches a continuous chunk against it under block-wise causal attention so the context KV can be cached across the 10 flow steps.

The controlled number: ContextFlow-Plain beats a matched autoregressive baseline built on the same $\pi 0$ initialisation by 10.5 pp; adding the compressors adds another 9.5 pp, to 73.5% on unseen LIBERO configurations against 53.5%. It matches $\pi 0$ fine-tuned on one demonstration from each unseen task (73.5 vs 72.5) with no fine-tuning at all. Cumulative trajectory MSE falls 4.7 $\rightarrow$ 3.0 $\rightarrow$ 2.2 m². Soft: generalisation is to unseen *configurations* within the same primitives, which the authors state plainly; real ALOHA results are 10 trials per cell (best cell 6/10); and training four LIBERO suites leaves LIBERO-Goal and LIBERO-10 at 0.0% until LIBERO-90 is added. No code stated.

DeCAL

PREPRINT / PKU + BAAI

A Mixture-of-Transformers dexterous VLA whose two new pieces are both about *when* tactile counts. Adaptive visuo-tactile fusion gates tactile cross-attention on a learned contact signal, so tactile stays suppressed in free space and opens during contact — the paper shows the gate value rising at contact onset. Visuo-tactile latent co-imagination predicts future Cosmos-VAE visual latents and future tactile latents jointly, decoding to raw tactile images, deformation maps and 6-DoF forces at training time and dropping the decoders at inference. 71% average success across six contact-rich tasks at 20 trials each (best baseline DECO $\approx$ 56%), 0.27 s per chunk on a 4090. Soft: the platform is two UR5s with 22-DoF SharpaWave hands and per-fingertip 320×240 vision-based tactile — not reproducible at small scale — training is 8×H100 for 100k steps, and every cell is a single run at n=20. Project page only.

Worth naming as a cluster: **in-context adaptation without test-time gradients** got four entries this window — ContextFlow, [ICI-VLA](#) (DTW-retrieved micro-demos into a frozen text-action VLA, RoboTwin 60.4%, real 83.2% over ~1,000 trials), [behavior-prompting checkpoints on LIBERO](#), and [2AM](#). And **online RL on real hardware** got two: [VLA-Precision](#) (relative-advantage instead of absolute Q, frozen-prefix KV reuse, 98.3% on nine self-built chemistry tasks — authors’ own benchmark, no code) and WHIRL below.

2 Manipulation, dexterity & sensing

In-hand pen writing from an online task Jacobian

PREPRINT / ETH Zurich Soft Robotics Lab; [code and project page](#)

The week’s best hardware result, and it contains no learned policy. On the 17-DoF tendon-driven ORCA hand, 10 joints (thumb, index, middle) articulate a pen held in a fixed power-precision grip inside a TPU sleeve that roughly quadruples its diameter. A 2×10 command-space task Jacobian — commanded joint increment to pen-tip motion in the paper plane — is estimated recursively with a forgetting factor from an ArUco marker at ~15 Hz, inverted with a damped pseudoinverse, and driven by a task-space PID with feedforward; the nullspace carries a posture term pulling back toward the initial grip. The regressor is the *commanded* increment, not measured joint velocity, which is the right call on a tendon hand.

Numbers: ~18 s of scripted excitation, then 0.62 ± 0.12 mm mean in-plane error over 22 in-air runs and 0.67 ± 0.08 mm over 16 on-paper runs (p95 $\approx$ 1.4 mm both), pooled 0.64 ± 0.10 mm over n=38. A single continuous 31-minute run wrote all 26 letters with per-window means held between 0.59 and 0.69 mm and no upward trend; a ~20 mm excursion during “K” recovered inside the same letter without re-excitation. The ablations are the paper’s real contribution: freezing the Jacobian right after excitation diverges in 2 of 3 runs, freezing after 12 s of writing diverges in 1 of 2, and freezing after ~30 s holds baseline in 4 of 4 — so the excitation phase alone does not produce a usable map, and about half a minute of on-the-job adaptation does.

Where it’s soft, and the authors say most of it: only the in-plane (x,y) is controlled, with 2–3 mm of uncontrolled z drift absorbed by deliberately bulging the paper, so it cannot write on a rigid surface or do multi-stroke letters; the arm repositions between letters; contact-discontinuous motions (regrasp, finger gaiting) are outside the formulation; writing runs at 8×10^{-4} m/s and even 2× is unreliable; and — the one that should temper the headline — every reported error is measured through the same vision pipeline that closes the loop, with no independent measurement of the deposited ink. The simulated Shadow Hand (0.17 mm RMSE) and Wuji Hand 2 (1.48 mm) runs control all three coordinates but have noise-free ground truth. Single lab, unreproduced. Cost to try: an ORCA hand, a webcam, and a laptop CPU.

TacPAC

PREPRINT / Fudan + NeoteAI, [CODE](#)

The mechanism is a timing fix. A world-action model plans a chunk conditioned on predicted future tactile; that prediction is then frozen while the chunk executes, which is exactly when the informative tactile arrives. TacPAC runs one extra clean pass after denoising, caches the per-layer keys and values of the predicted tactile views and of the planned chunk, and lets a separate tactile expert read each incoming tactile frame against that cache, emitting a per-step delta over the unexecuted suffix. Visual keys are deliberately not cached. Corrections are written back from offset $m+d$, where d is the largest recently observed correction latency in action steps — in practice near zero, and the robot advances a step mid-correction in only 1.6% of calls.

Five real tasks on a Flexiv Rizon 4 with two InTac S1 fingertip sensors, 20 trials per method per task. Vision-only base 22% → 64%; strongest baseline (T-Rex) 48%. The clean ablation is the cache: same base model, same corrector, same tactile input, same parameter count, only attention to the cached predicted tactile removed — 47% vs 64%. Adding tactile prediction alone gets 22% → 37%, a third of the gain. One correction is 30.4 ms against 628.6 ms to regenerate a chunk. The supporting diagnostic is good: contact transitions are 4.9% of frames but 18.6% of total tactile change, so a future-view reconstruction loss is dominated by frames where nothing happens. Reactive correction without an expectation actually falls *below* vision-only on the two damage-scored tasks — evidence that contact feedback with no prior can pull a correct plan off course.

Assembling two parts in one hand

PREPRINT / HKU + HKUST(GZ)

+ ETH). A single 22-DoF Sharpa Wave hand mates two objects with no second arm and no fixture — bottle cap, syringe plunger, marker cap. RL in IsaacSim over a 34-D observation (two object centroids and symmetry axes plus 22 joint angles; no velocities, no contact forces), recurrent LSTM policy, goal-relative-pose reward plus a multiplicative auxiliary that softly assigns thumb+index to the held part and middle/ring/pinky to the fixed one, regularised toward a single human snapshot pose. Zero-shot to hardware from one RealSense D435 with FoundationPose. Closed-loop 15/20, 17/20, 16/20 on the three tasks against open-loop replay of successful sim trajectories at 2/20, 0/20, 0/20 — the cleanest available statement that the sim-to-real gap here is closed by feedback, not by fidelity. The hand-morphology study is the part worth keeping: the same pipeline on Allegro fails the z-axis insertion for lack of a fifth finger, and on XHand fails xy alignment for lack of abduction joints. Soft: both objects are placed in the hand by a human and held until the policy engages; gravity is removed for the warm-up steps; and the authors name PhysX's inability to model patch contact as why pinch grasps degenerate at some wrist tilts. Ablations are over 3 training runs.

WHIRL

PREPRINT / HKUST(GZ)

+ IIT, project page only). Human-in-the-loop RL usually spends a takeover once, as a behaviour-cloning mask or a replay weight. WHIRL adds a fourth head to a one-step latent world model — beside dynamics, reward and termination — that predicts the *probability a human takes over at the next state*, and routes only that head into the actor's utility, never into the Bellman target. The reasoning is explicit and correct: takeover probability predicts operator behaviour, not task cost, so as a reward bonus it would change the optimal policy and propagate operator habits through every backup. On a 16-DoF LEAP Hand plus Franka FR3, five tasks: +15–30 pp over a matched model-free residual-HIL baseline (96.7% on Pick LEGO, 100% on Pick Cube), and step-weighted intervention fraction on Pull Drawer down from 0.113 to 0.018, an 84% cut. ~8k online steps on one RTX 4090 for the pick and drawer tasks, 15–30 demonstrations each. Soft, and the authors report it: one seed and one operator per cell, Fisher's exact reaching $p < 0.05$ on the three pick tasks only, and the intervention head inherits the labelling operator's caution threshold.

SEED-UMI

PREPRINT / PKU + Delta Intelligence

Both the human *and* the robot hand wear the same 20-DoF isomorphic exoskeleton, with encoders and wrist cameras rigidly mounted to the shared frame. That single design decision converts retargeting from an open-loop free-space regression into supervised learning on paired data: replay a human demonstration on the robot through a babbling-fit mapping, and the discrepancy between the human-side and robot-side encoder traces is the correction signal. It also removes the inpainting problem, since both embodiments present the same outer mechanism to the wrist camera. 52 successful AirPods-insertion demos in 30 minutes against 18 by teleoperation (3.0×); 70.0% mean rollout success against 71.7% for teleoperation-collected data, with paired fine-tuning worth +12.7 pp concentrated on the contact-load tasks (+16.7 screw driving, +15.0 table cleaning and spray). Nothing released, single operator, one target hand.

Two sensing results worth the shelf space. **X-Hinges** ([peer-reviewed](#), [UIST '26](#), MIT CSAIL + Tianjin) co-prints two conductive filaments three orders of magnitude apart in resistivity (1.23×10^4 vs $8.75 \Omega\text{-cm}$) into a Lamina Emergent Torsional compliant joint, so the high-resistivity element carries the deformation signal while the traces contribute negligible noise, with a non-conductive TPU body isolating skin contact. Differential layouts per axis implement Wheatstone logic in printed geometry; cross-axis interference down to 8.2%; bridge width tunes stiffness over two orders of magnitude (up to 470×); 420,000 cycles over 233 hours with resistance drift under 5.5%; the layer-overlap filament interface is 22.3 k Ω against 952 k Ω for end-to-end contact. It needs a multi-material FDM printer and a custom transimpedance front end (ADS1256, 50 Hz/channel, 0.1–100 M Ω); CAD and

software are promised open-source after review. **TacClip** ([PREPRINT](#)), Stanford, Cutkosky lab) clips a Fiber Bragg Grating over the fingertip so the *fingertip stays bare*: contact loads bulge the tissue laterally and bend the clip. Force-magnitude error typically below 0.5 N over 0–8 N, 2 kHz sampling, flat response to ~20–30 Hz (the finger is the low-pass filter, not the FBG), and it works underwater. ~\$2 for the printed clip and ~\$20 per finger for the sensor head, excluding the optical interrogator — which is where the real cost is. It estimates $\|f\|$ rather than normal force; the single FBG does not decouple axes.

3 Platforms & embodiment

No humanoid company published a method again — third consecutive window. The closest anyone came is Unitree putting [UnifoLM-ER-1](#) on the Hub: an Apache-2.0 Qwen3-VL-4B embodied reasoner trained on 5M+ spatial-reasoning and embodied samples, leading open models on 7 of 16 benchmarks (93.4 BLINK, 88.9 EmbSpatial, 82.0 Where2Place), plus an optical-flow region predictor. The 6B action model is still listed as coming. So: a reasoner, not a policy.

Everything else at the platform layer this week was manufacturing and finance. [XPENG launched mass production of IRON](#) — 2,250 TOPS across three in-house Turing chips, >80% line automation, 85% supply-chain overlap with its EV business, 110,000 m² in Guangzhou, showroom deployment Q4 2026 and enterprise delivery 2027. No evaluation, no method, no policy disclosure; the claim that units walk off the line unassisted is press-release evidence. [Nucleus dropped bipedal legs for a wheeled base on Nucleus II](#) after factory-floor feedback — vendor account, no data, but the second wheeled-bimanual datapoint in two weeks after Nori.

Academic humanoid work exists but I cannot vouch for it: [TANGO](#) (whole-body navigation VLA, sim-only), [CAP](#) (depth denoising for perception-degraded G1 locomotion, n=5 real trials, project page only), [ViBe](#), [granular-terrain locomotion](#). All preprints, all scored from abstracts, none with code stated. A [ROS 2 open biped kit called Bimo](#) was announced on ROS Discourse on 2026-09-10; unverified, and I do not have a stable permalink for it.

4 Open source & tooling

By friction removed, this is the best week since W35.

[LeRobot v0.6.1](#) — the first release since v0.5.1 in April, and the first since the NVIDIA/Hugging Face deal. The substantive change is consolidation onto upstream Transformers: Wall-X now subclasses native Qwen2.5-VL and XVLA subclasses native Florence2, with shared VLA components across PI0/PI05/EO1/PI0-Fast/SmolVLA. Diffusion-policy training gets gradient checkpointing; PI0-Fast’s decoder stops at the end-of-action marker; dataset statistics now subsample frames to bound memory; streaming datasets support buckets and private-repo tokens. Hardware: RealSense manual exposure/gain/white-balance, three new Dynamixel models (XH540-W150, XC330-T288, XC330-T181), retry on transient Feetech bus errors. **Breaking:** `lerobot.types` → `lerobot.lerobot_types`.

[MuJoCo 3.13.0](#) (2026-09-08) is the most consequential simulator change of the year for anyone doing tendon or stiff-actuator work. A new **discrete** integrator merges the constraint solve and the implicit velocity update into one operation, folding joint, tendon and actuator stiffness *and* damping into the solver’s effective metric — passive springs and actuator position gains become stable at timesteps far beyond the explicit limit, which is the thing that has been forcing 1 kHz+ steps on servo models. Separately, the Newton solver with elliptic cones now does a single refactorization instead of per-contact rank-1 updates where that is cheaper: scenes with many simultaneously sliding contacts run 1.4–2× faster. Plus a plane-mesh collision rewrite, mesh-associated sites, and Python 3.15 including free-threading. Breaking: the implicit-flex effective-metric special case from 3.11.0 is gone; affected models must move to **discrete**.

MuJoCable

PREPRINT

Is the complement nobody shipped: a MuJoCo engine plugin that makes cable *routing* part of the mechanism. Native spatial tendons give body-fixed sites, sphere/cylinder wraps and one transmitted force with a Hookean law that can push as well as pull. MuJoCable jointly optimises an ordered path across adjacent moving analytic and mesh surfaces, applies a unilateral pull-only axial law with slack reserve and tension limit, propagates directional Capstan relations so each segment carries its own tension, and maps the resulting nodal forces to bodies by virtual work. Pulley benchmarks recover analytic transmission with Capstan-ratio error below 0.5%. The finding that earns it: on the 18-joint SpiRobs, the native tendon and MuJoCable produce nearly the same steady bend while MuJoCable resolves $\sim 4\times$ the peak cable tension and 9% more take-up — similar global motion, different transmission state. A guide-friction sweep at $\mu=0.60$ costs a quarter of the bend and redistributes rotation proximally. Overhead is 22.5% on step time (41.96 vs 34.25 μ s median), still 11.9 $\times$ real time at a 0.5 ms step. No repository stated, which is the whole problem with a plugin paper.

[Genesis v1.4.0](#) (2026-09-06) and [v1.4.1](#) (2026-09-12). v1.4.0's `gs_replay` exports a scene plus full trajectory as one standalone archive with bit-exact replay on another machine — the single most useful thing for bug reports and for anyone trying to reproduce someone else's sim result. v1.4.1 makes large-scene cost sub-linear in island count on both CPU and GPU. Note: GitHub's rendered release pages return visibly wrong date fields through my fetch path (v0.6.1 came back as "August 3, 2024", Genesis v1.4.1 as "September 12, 2024"); the tags are what I verified, dates come from the release API.

Also released: [FARM](#) — a 33,985-parameter supervised readout over frozen VLA-JEPA predictive states that detects failure per step, 85.68/88.59 pooled AUROC/AUPRC over seven source tasks in five-fold out-of-fold evaluation, best Seen performance among 15 matched baselines, 0.2256 ms mean CUDA latency, and real-robot transfer tested on PIPER X, SO-101 and Franka ([PREPRINT](#)). Zero-shot transfer is not uniform — the source-readout lineage matters enormously (41.35 vs 84.70 AUROC on the same PIPER X population depending on which source population trained the readout) — but readout-only adaptation reaches 98.48. [HuRo](#) releases ~ 630 K robotized episodes / 142M frames / $\sim 1,317$ hours built by inpainting humans out of five egocentric corpora and compositing an Isaac-Sim-rendered robot back in with IK-retargeted actions; the ablation that matters is that keeping the original human observations (actions still retargeted) matches in-distribution but loses 16.5 points out-of-distribution.

[Xiaomi-Robotics-U0](#) ships Apache-2.0 4B and Sequence checkpoints from a 34B family plus a VisionTokenizer — but the card says training code "will be added in a future release," which contradicts the reporting that training code was opened. Flagged as conflicting; the weights are real, the training code is not there yet.

Data-layer items at one-arm scale: [RobotArena360 assets](#) (22 real DROID scenes as interactable Gaussian-splat twins with collision meshes and calibration, CC-BY-4.0), a [1,530-episode / 440,967-frame real UR5 + Robotiq three-finger set](#) under Apache-2.0, [LIBERO-90 repacked to LeRobot v3.0](#) (uncarded), and a [LoRA adapter on lerobot/pi05_base for SO-101 block stacking](#).

5 Money & moves

Maven Robotics, \$100M Series A

TRADE PRESS

, led by RoboStrategy with LocalGlobe, Vine Ventures and XTX Markets Ventures; team out of Apple's Special Projects Group. Wheeled bimanual suction manipulator for mixed palletizing, eight units on 16-hour shifts at claimed $\geq 99\%$ uptime, funding 250 third-generation units. **INFERENCE** \2192 the pitch is not the model — it is the AV data-ops loop (collect facility edge cases within hours, run automated ablations, update weights, redeploy) plus pincer gloves that capture human demonstrations already mapped to the end-effector. Capital is pricing the iteration loop, not the policy.

Skill AI crosses \$100M ARR in ten months

TRADE PRESS

, 60+ customers, revenue split 86% manipulation / 10% mobility / 4% Fetch; S1 reported at 66% out-of-distribution per-step success with one demonstration video valued at roughly 380 teleoperation examples. **INFERENCE \2192** this is the first applied robot-learning revenue number at this scale, and it is manipulation, not humanoids. All figures are first-party except the named customer throughput (50 → 130–140 lines/hour at G10 Fulfillment); nothing is audited.

Agility Robotics' S-4

TRADE PRESS

\$1.8M 2025 net sales against a ~\$140M operating loss, \$111M opex (from \$71M in 2024), ~\$100M cash burn, >\$300M multi-year Digit v5 orders from a single customer, 65,000+ operating hours across nine sites. The Churchill Capital XI merger values it at \$2.5B — roughly 1,400× 2025 revenue — for >\$620M gross proceeds including a \$200M Foxconn-led PIPE. Set against **Unitree down 53% from its intraday peak**: ¥1.70B (\$252M) 2025 revenue, 5,500+ humanoids shipped, profitable, now ~\$30B or ~125× revenue. **INFERENCE \2192** public markets now have two humanoid comparables whose revenue multiples differ by more than 10×, and the profitable one is the cheap one. That is the first real price signal the sector has had.

OpenAI confirms a humanoid

TRADE PRESS

Altman on a podcast, plus 19 open San Francisco robotics roles — four on actuators and one actuator TPM, alongside PCB layout, thermal simulation, firmware, prototyping, and commodity management. No specs, no timeline. **INFERENCE \2192** the hiring is entirely hardware. Actuators are 40–60% of humanoid BOM and the supply is concentrated; OpenAI is buying the part it cannot download.

Also: [ARM Institute ~\\$90M across ten military-manufacturing robotics projects](#), [Swarmer to acquire Ratel Robotics for up to \\$224M](#) (defence UGV), [Enovis/eCential ~\\$245M](#) (surgical), [Teradyne sues JAKA over three Universal Robots patents](#), and [Boston Dynamics veterans launch Dynamic Creatures](#) for character robotics. **INFERENCE \2192** none of the week's M&A touches robot learning. Three consecutive weeks of surgical and defence consolidation.

6 Deep read of the week

OpenWAM — what a world-action model actually inherits, and at what measurement precision. [arXiv:2609.07398](#), preprint, NUS + Tsinghua + PKU + HKU and others. [CODE](#), [weights and data recipes](#). Unreproduced. The other candidate was the pen-writing paper, covered above; it is the better result but the smaller surface — its ablations are already complete and honest, and there is little left to interrogate.

Setup. A WAM inherits video-generative priors and channels them into actions. OpenWAM factorises that into three interchangeable module classes — visual encoder E , stream backbones S , visibility attention mask M — and a composition rule assembling six architectures across single-, dual- and tri-system families, over five video backbones (Wan2.1-VACE-1.3B through Wan2.1-I2V-14B), four encoders (Wan2.2-VAE, FLUX.2-VAE, DINOv3, V-JEPA 2.1, optionally compressed by an S-VAE), and four cross-modality masks (Isolated, Action-Sees-Video, Video-Sees-Action, Mutual). The controlled study runs on RoboTwin2.0 in Full and Clean2Random settings, the latter giving separate in-domain and out-of-domain splits. The scaled instantiation, OpenWAM- α , pretrains on 518.5M frames (~6,369 h) curated from a 1.33B-frame raw pool: 30% own egocentric human video (71.6K first-person recordings, 3,006 tasks, action channels fully masked so it supervises only the world stream), 40% real robot (AgiBotWorld-Beta, RoboCOIN, DROID) and 30% synthetic (InternData-A1), all through an 80-D unified action space with fixed slot semantics and validity masking. Eight simulation benchmarks plus real single-arm, bimanual and dexterous-hand platforms — including a 21-DoF Wuji hand absent from pretraining entirely.

Method. Almost nothing here is a new mechanism. Flow matching, DiT video backbones, a dedicated action expert, mixed self-attention across streams, asynchronous serving — all inherited. What is new is the factorisation itself: exposing coupled design choices as controlled variables and reporting each decision under the conditions created by the previous one. That is the contribution, and it is a real one; the WAM literature to date has been monolithic systems with no way to attribute a gain.

The number that matters. Not any benchmark score. It is the split of the pretraining gain on RoboTwin2.0-Clean2Random: **+0.68 pp in-domain (87.0 → 87.68) against +12.12 pp out-of-domain (14.5 → 26.62)**. Embodied pretraining is, on this evidence, almost entirely an OOD purchase. Three supporting findings: robot-only pretraining gives the strongest in-domain while ego+robot mixtures give the best OOD; one-stage co-training matches two-stage sequential and removes a curriculum transition; and — the finding the paper builds its recipe on — the preferred information flow *reverses* with pretraining, RoboTwin-Full favouring Action-Sees-Video by 0.24 pp from scratch and Mutual by 0.16 pp after pretraining, with Clean2Random favouring Mutual by 0.70 pp (ID) and 0.72 pp (OOD). The paper also reports that WAMs beat VLAs in-distribution while VLAs generalise better out-of-distribution, attributing the asymmetry to error accumulation in long-horizon pixel prediction under shift.

Where it's soft. The reversal that fixes the final recipe is 0.16–0.72 pp, the backbone choice between 5B and 14B is 1.40 pp, and **no seeds, error bars or repeats are reported anywhere in the study**. One week ago MINERVA measured a ± 1 -point training-seed band on a comparable benchmark by retraining a baseline three times (94.60–96.75). Every decision OpenWAM carries forward is smaller than that band. This does not mean the recipe is wrong; it means the paper has not shown it is right, and its own released infrastructure is what would settle it. Second: the eight-benchmark results are all post-SFT with per-benchmark batch sizes and step counts (Table 10), so these are not zero-shot transfer numbers. Third: the LIBERO-Plus anomaly is attributed to a thinner single-arm pretraining mixture, but the paper elsewhere names pixel-reconstruction encoders as insufficiently robust to viewpoint and noise variation — two explanations, one dataset, no experiment separating them. Fourth: real-robot baselines are LingBot-VA and $\pi 0.5$ only. Fifth, and unavoidable: 128 NVIDIA H200s for ~7 days, ~21,500 GPU-hours, against baselines trained under unstated budgets.

Steal list. (1) Report pretraining gains as two numbers, in-domain and out-of-domain, never one — on this data the single averaged number would have hidden the entire finding. (2) The 80-D unified action space with fixed slot semantics and validity masking is a cheap way to co-train heterogeneous embodiments without per-embodiment heads, and it is what let an unseen 21-DoF hand adapt at all. (3) Unlabelled egocentric video with action and proprioception channels fully masked supervises the world stream for free — if you have any video of the task, it is trainable data under this recipe. (4) The deployment split: train across the whole joint noise plane, infer along the synchronised diagonal, and accelerate with a CUDA-graph-replayed fixed-shape joint denoising loop, a velocity cache across adjacent steps, a prompt-embedding cache, and no VAE decode in the control path — 170 ms per chunk on an RTX 5090, which is the number that makes this deployable rather than the H200 count.

Cost to reproduce a scaled-down version. The pretrain is out of reach; the *study* is not, and it is the part with the transferable claims. Take the 1.3B VACE backbone rather than the 5B, run the four-mask comparison on RoboTwin2.0-Clean2Random (public), and you are fine-tuning, not pretraining — the α recipe uses 10,740 optimizer steps at batch 256 for that split.

INFERENCE [v2192](#) at batch 8 with accumulation on one 80 GB card, one cell lands in the low hundreds of GPU-hours, so a seed-replicated four-mask sweep is a week of one GPU, not a cluster. No new data collection, no new hardware. That is the experiment in [§10](#).

7 Relevant to what you're building

CONTEXT.md's `## Active build` section is still the unfilled placeholder — “(empty — pending interview)” — so there are no workstreams to map this week's items onto, and inventing them would be worse than skipping. Section skipped until the placeholder is replaced.

8 Market signal

Visibly hiring or scaling: OpenAI (19 SF robotics roles, four on actuators plus an actuator TPM — all hardware, no policy roles named), Maven Robotics (250-unit build-out on a team assembled from Apple's shuttered AV group), Skild AI (60+ customers, revenue 86% manipulation), XPENG (manufacturing at automotive scale), Agility (>\$620M of SPAC proceeds against ~\$100M annual burn). The notable absence: nobody is visibly scaling a *policy research* team this week. The money went to actuators, factories, and deployment loops.

No compensation or levelling datapoints with a stated methodology surfaced again — the only sources that return for robotics comp queries remain SEO aggregators, and reporting nothing beats laundering those. The nearest thing to a real levelling signal is the Agility S-4, which is a public filing rather than a survey: \$111M of operating expense against \$1.8M of revenue, up from \$71M the year before. That is the shape of the budget the senior roles at a pre-revenue humanoid company sit inside, and it is now disclosed rather than rumoured.

Three skill clusters keep recurring in the work that got attention. **Cheap readouts on frozen backbones** — FARM extracts step-wise failure detection from an unmodified VLA-JEPA predictive state with 34K trainable parameters and sub-millisecond latency; [No Free Checker](#) formalises the same family as “model-intrinsic verifiers.” Knowing where in an existing stack the signal already lives is now a distinct competence from training new models. **Execution-schedule engineering**, continuing from W36 — IMLE-VLA’s single-step head, TacPAC’s cache-and-write-back-from- $m+d$ asynchronous corrector, OpenWAM’s synchronised denoising diagonal and inference acceleration stack. **Evaluation design as the contribution** — No Free Checker surveys ~150 verifiers, argues that credibility falls as availability rises across every family, and proposes nine reportable metrics including a “proxy gain that transfers” ratio ($\Delta_{\text{panel}}/\Delta_{\text{proxy}}$, scoring the same checkpoints under the training verifier and under a panel the policy never saw); it also states plainly that every agreement rate and downstream gain in the robotics literature was measured on non-adversarial candidates, so exploitability is essentially unmeasured. [MEMOBench](#) finds π_0 at 94.0% on memory storage and 4.4% on the task. [EgoGenEval](#) separates camera-motion grounding from scene preservation across 16 generators, best 0.662 against a 0.940 oracle. And a [Northwestern review](#) argues multifingered-hand dexterity claims are not currently comparable at all.

Table stakes now: code and weights at submission; latency reported beside success; a benchmark harness someone else can run; and — newly, after OpenWAM — an ID/OOD split rather than one averaged number. Still differentiating: seed-replicated ablations (still nearly absent; OpenWAM has none), RL loops that close on real hardware (WHIRL, VLA-Precision), released hardware files, and adversarial evaluation of your own verifier, which the survey says nobody in robotics has done.

9 Threads

THREAD	STATUS	NEXT REAL UPDATE
Does VLA scale buy language-conditioned control?	STALLED	No sub-10M follow-up this week. Still: a sub-10M policy on real hardware, multi-task, vs a published baseline
Benchmark validity in physical AI	ADVANCING (No Free Checker’s nine metrics; MEMOBench π_0 94.0% storage / 4.4% success; EgoGenEval; IMLE-VLA at 98.0% LIBERO)	A major lab reporting seed bands in an ablation table, or any group scoring a verifier’s exploitability under search
Contact sensing without dedicated tactile hardware	CONTESTED	The week argued the other side: TacPAC, DeCAL, TacClip and X-Hinges all add hardware. W35’s VISTA still has no code
Paying at training time, not run time (merged)	ADVANCING (IMLE-VLA 15→55 Hz, 11× throughput; TacPAC 30.4 ms vs 628.6 ms; OpenWAM 170 ms/chunk on a 5090; latent semantic scaffolding)	Latency and an ID/OOD split both appearing as default columns
Online correction as the last mile	ADVANCING (WHIRL: +15–30 pp and 84% less operator time on a 16-DoF LEAP hand — the third group, on different hardware)	Seed or operator replication; every result so far is one seed, one operator, one rig

THREAD	STATUS	NEXT REAL UPDATE
World-action models as unified policy + simulator	ADVANCING (OpenWAM releases infra, weights and data recipes; FARM decodes failure from a frozen WAM state)	A released world model used for policy improvement by someone who didn't build it — now possible for the first time
Who owns the open robot-learning stack	ADVANCING (LeRobot v0.6.1, first post-deal release)	It consolidated onto upstream HF Transformers classes, not onto NVIDIA stacks. Watch dataset-format governance and whether GR00T gets privileged treatment
Simulation fidelity for transmission and contact	NEW (MuJoCo 3.13.0 discrete integrator; MuJoCable Capstan routing; Aero Hand's validated tendon model in W36)	A tendon-hand sim-to-real result that credits the new integrator, and a MuJoCable repository
Low-cost open hardware as training substrate	ADVANCING (X-Hinges peer-reviewed self-sensing flexures; TacClip ~\$20/finger sensor head; ORCA hand carrying the week's best dexterity result)	Still no third-party build of Aero Hand, BRIDGE or Peg-in-Bench
Claimed open releases vs actual artifacts	CONTESTED , sharper	For: OpenWAM shipped everything; FARM, IMLE-VLA, HuRo, the pen-writing code. Against: Xiaomi U0's card defers training code; Unitree's action model pending; MuJoCable states no repo; StreamPI weights pending three weeks since W35
Humanoid capital vs demonstrated capability	ADVANCING (Agility's audited \$1.8M/\$140M against a \$2.5B SPAC; Unitree —53%; XPENG mass production; OpenAI entering)	Still zero methods from humanoid companies. Unitree's reasoner weights are the closest yet — the action model landing would move this
Manipulation acquired rather than built	STALLED	W35's within-the-quarter prediction has ~6 weeks left. This week's M&A was defence UGV and surgical again. If it lapses, kill

Dropped this week: “Training-time supervision at zero inference cost,” folded into “Paying at training time, not run time,” which covers the same principle on both the auxiliary-loss and the latency side.

10 Open question

Is the world-action-model design space resolvable at the precision the field is currently measuring it? OpenWAM's carried-forward recipe rests on deltas of 0.16–0.72 pp for the attention mask and 1.40 pp for the backbone, with no seeds reported — one week after MINERVA measured a ± 1 -point training-seed band on a comparable benchmark by the simple expedient of retraining three times. Either the WAM design space has a genuinely finer signal than the policy-scale one, or a widely-released recipe has been distilled from noise, and I cannot tell which from what is published.

What would settle it: retrain three of the mask cells at three seeds each on RoboTwin2.0-Clean2Random and publish the spread. OpenWAM released the infrastructure, the evaluation protocol and the data recipes to do exactly this, and RoboTwin2.0 is public — so this costs a handful of fine-tuning runs on one GPU and requires building nothing. It is the cheapest high-value experiment available in robot learning right now, and it applies retroactively to most of the ablation tables in this issue.

Sources labeled inline. arXiv items are preprints and unreviewed except X-Hinges, which is accepted at UIST '26. Single-lab claims flagged where relevant. Items scored from abstracts rather than full text are described without numbers.

ALSO THIS WEEK

Everything scoring 4+ in the queue that did not get a section above:

- [**SyncWorld**](#) — a calibration episode specifies the action-visual mapping in context, turning a world model into a zero-shot simulator for test-time policy improvement. Skipped: no code stated, and the week already carries OpenWAM on the same axis with a released stack.
- [**Show-Harness**](#) — discrete semantic actions let frontier VLMs drive robots directly, with a companion rig that collects demonstrations without teleoperation hardware. Skipped: project page, numbers unverified.
- [**VLA-Precision**](#) — relative-advantage online RL with frozen-prefix KV reuse, 98.3% on nine chemistry tasks. Skipped for space: authors' own task suite, no code.
- [**ICI-VLA**](#) — DTW-retrieved micro-demonstrations adapt a frozen text-action VLA at test time; RoboTwin 60.4%, real 83.2% over ~1,000 trials. Skipped: same in-context slot as ContextFlow, with 8×A100 offline training and no code stated.
- [**EgoGenEval**](#) — separates camera-motion grounding from scene preservation for egocentric video world models; 16 pose-free generators, best 0.662 vs a 0.940 oracle, human-validated at $p=0.943$. Covered in [§8](#); release unverified.