

This Week in Robotics

Window: 2026-08-31 → 2026-09-06.

THE LINE

NVIDIA agreed to buy Hugging Face for \$12.9B, putting LeRobot's maintainer inside the company that ships Isaac GR00T — and in the same week a 0.54M-parameter policy with no language encoder and no pretrained vision hit 95.1% on LIBERO, which says more about the benchmark than about the policy.

1 Policy learning & VLAs

MINERVA

PREPRINT / U. Tokyo

Scales a from-scratch CNN plus flow-matching chunk head *down* until LIBERO breaks. The architecture is subtraction: language becomes a 40-entry task-ID table, action-token self-attention becomes an MLP mixer, and the freed 26% of parameters goes to the encoder. [Code](#), [recipes](#), [all checkpoints](#), Apache-2.0. Deep read in [§6](#).

Knowing When to Stop

PREPRINT / Tsinghua

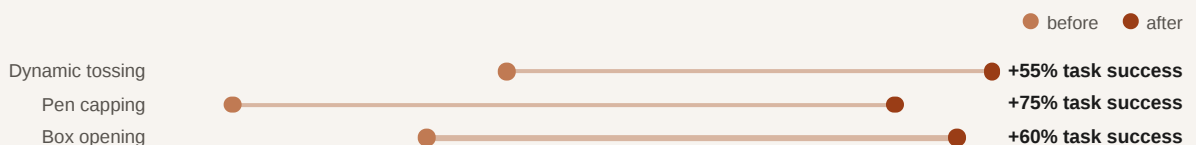
Training-free adaptive chunking: cross-attention from action queries to VLM tokens disperses with the action index and saturates near 95% of its log-N bound, so truncate where entropy is both high and flat. The mechanism check earns it: entropy tracks held-out action MSE at Spearman $\rho = 0.432$, and the partial correlation after removing the index is still 0.214. LIBERO 94.88 → 97.25 over the best fixed horizon and real dual-arm 47.5 → 61.7 (n=20/task), but RoboTwin $\pi_{0.5}$ only 61.6 → 62.9 with no standard deviations reported — inside plausible seed noise. The real claim is matching a post-hoc-tuned horizon without knowing it, at a *longer* average horizon (18.1 vs 10) and under 7 ms overhead. No code stated.

SmoothRL

PREPRINT / Atribot

Under asynchronous inference each chunk is only partly executed, so $\nabla_a Q$ is truncated to the execution region $[n, 2n)$ while the critic keeps the committed prefix as state augmentation for the concurrent decision process. Frozen base plus a TD3-style attached network in *raw* action space, which lets human interventions serve directly as regression targets — no latent inversion. 250 on-robot episodes per task take dynamic tossing from 39% to 94% and pen capping from 8% to 83%, with box opening reaching 90% only after dipping to 20% — below the frozen base — at 150 episodes. No code.

SmoothRL: frozen base policy → 250 on-robot rollout episodes



One episode per initial-state configuration, single robot, no seed repeats, authors' own protocol. Box opening dipped to 20% — below the frozen base — at 150 episodes before recovering.

arXiv 2608.29768, preprint, unreproduced, no code released

ZETA

PREPRINT / Galbot + PKU

Controlled cross-embodiment ablations on a $\pi 0.5$ -class backbone, 14 held-out embodiments, 6,300 sim rollouts per model over 3 runs. EEf-delta state *and* actions beat absolute-EEf/world-delta 75.7% vs 60.3% in sim and 89.9% vs 56.0% real, concentrated on arm-only and full-embodiment shifts; 5% target-embodiment data in pretraining is worth 13.4 pp, which is why they insist “strict” and “pretrain-exposed” zero-shot are different claims. No code stated.

XR-2

PREPRINT / PrimeBot + PKU

Releases 1,500 hours of bimanual household data — 531.7 h of real-robot teleoperation over 32,518 trajectories on a 25-DoF mobile bimanual platform, plus roughly 1,000 h of in-the-wild UMI collection across 200+ households and 5,000 garments — and trains a 5B VLA on it. The number worth having is the scaling shape on clothes folding: 34% at 30 h, 84% at 120 h, then flat (82% at 160 h). Three rounds of DAGger under a failure-weighted budget take a 58% checkpoint to 74%, 82% and 93%. Expert data buys coverage of the expert’s own distribution; only on-policy corrections reach the states a deployed policy visits. Rig, model and pipeline are proprietary, so what transfers is the curve, not the result.

Worth naming: four independent papers add supervision that is **free or discarded at inference** — [PHR-VLA](#), [Temporal Forcing](#), [GIFT](#) (+4.6 to +12.6 pp), [EGR](#) (code out). Nobody wants to pay for their inductive bias at run time any more.

2 Manipulation, dexterity & sensing

Aero Hand Open

PREPRINT / Chestnut Robotics

Sixteen joints, seven motors, 374 g, **\$314 bill of materials**, fully 3D-printed, covering all 33 GRASP-taxonomy grasps on one hardware configuration. The contribution is that the *cable transmission* is modeled in MuJoCo — every cable and return spring a spatial tendon wrapping CAD-derived pulley cylinders, wrap sides pinned so moment arms follow from geometry rather than being fitted. Agreement is therefore a prediction, not a fit: 0.04–1.40 mm cable-excursion error across the four fingers, 0.29–0.45 mm RMS tracking. The thumb is the weak channel at 31–33% excursion mismatch, its abduction–flexor coupling carried only to first order. In-hand cube rotation trained purely in sim runs zero-shot on seven motor encoders and nothing else. Design files, firmware, model, RL package and deployment node released.

N0-Foundation

PREPRINT / NeoteAI + Fudan

A tactile-UMI handheld rig, 30,000+ hours of synchronized visuo-tactile demonstrations, and [OpenNeoData](#), a **5,000-hour open subset** over six embodiments and 250+ tasks. The bet is that dense three-axis force fields, not raw sensor images, are the hardware-agnostic interface between tactile sensors. Their own simulated suite is the honest part: $\pi 0.5$ leads at 45.8% mean over twelve tasks, and sustained-contact tasks stay near-unsolved (18% on gear placement).

ChainSplat

PREPRINT / KTH, [CODE](#)

Models a cable as an eight-joint screw-theoretic revolute chain with link-wise Gaussians, recovering geometry, dynamics and *external force* from multi-view RGB alone — no depth, no F/T sensor, no particle set — and the compact joint-space state is what makes trajectory optimization and real-time state estimation affordable. Soft: three ropes, 24 five-second trajectories, one model fit per trajectory and evaluated on the base trajectory it was fit to. System identification, not generalization.

Facet-0

PREPRINT

Makes force a *generated* quantity rather than a monitored one: flow matching emits an action chunk and the predicted wrist wrench jointly, and a distributional action-wrench critic scores contact outcomes for RL post-training. The reported contact-rich assembly gap is large — 82% against 15% — but the evidence is a project page with no weights and no dataset release, and the result rests on a proprietary 1,000-hour force-synchronized corpus. Unreproduced single-lab claim.

Three contact-control results worth knowing, all preprints, none with code stated: [torque-space MPPI](#) with explicit rigid-body dynamics at a 166 Hz solver rate for compliant 7-DoF force control; [behavior-realization separation](#) for physical HRI, splitting desired contact-port acceleration from its receding-horizon QP realization, with workspace overshoot falling from centimetres to sub-millimetre on a Franka FR3; and a [finger with a passive continuously-variable transmission](#) whose moving-pulley routing buys 3.6× mean force amplification with abduction unchanged — no CAD released.

Contact without tactile hardware advanced twice more, both preprints: [Blind Dexterity](#) drives whole-body G1 manipulation from joint encoders alone, and [distributed whole-arm interaction](#) reaches 71% real grasping success from four IMUs in a soft arm.

3 Platforms & embodiment

BRIDGE

PREPRINT / CMU + HUST + Joyln

Co-designs morphology and controller instead of choosing actuators after the fact: candidate topologies are screened by kinematic retargeting error, then by closed-loop tracking error under a trained policy, with calibrated torque–speed limits enforced during training. An 88 cm, 21-DoF humanoid at **\$1.5K** with open design *and* open policy — against \$3K/\$5.7K/\$6K for Bumi, K1 and ToddlerBot, none of which open both. Tracking success 94.83% vs 91.87/92.66/88.23 in MuJoCo on a merged LaFAN1 benchmark; the backflip and dance clips are video evidence only.

[WM-LOCO](#) (preprint) reports 93.3% real G1 success on foothold-constrained terrain, single lab, no code stated, unreproduced. No humanoid *company* published a method this window — Galbot’s [ET1 pre-order](#) is a press release with no evaluation.

Four more legged preprints, none with code stated. [ADAPT](#) fits a diffusion action prior over text-labelled trajectories and refines it with RL for closed-loop text-steered whole-body control. [Agile traversal of sparse 3D structures](#) uses monkey bars to force perception of thin overhanging geometry — a harder perceptual problem than terrain, because the thing you must not miss is mostly empty space. [Stay Seated](#) puts a humanoid on a passive caster chair so the actuators stop paying for weight support. And [humanoid safe stop](#) casts emergency stop as a learned stoppability value rather than one fixed maneuver executed regardless of feasibility — the most deployment-relevant of the four. Separately, [nonlinear body oscillations](#) finds six normal modes of a compliant quadruped that each develop into a distinct gait; hardware-validated, and a reminder that morphology still does control work.

4 Open source & tooling

By friction removed. **MINERVA**: Apache-2.0, four checkpoints, ~2 h to train one model on a single consumer GPU, ~55 min for the full four-suite eval. **Aero Hand Open**: a dexterous hand for \$314 whose MuJoCo model ships through Menagerie, ~13 h of printing. **Peg-in-Bench** (preprint, AIST): [STLs plus a scenario generator](#) — five peg shapes × three tolerances (0.1/1/3 mm), magnet-mounted hole pieces with eight orientations, machine-readable task JSON. A reproducible contact-rich physical eval for the price of resin; no baselines yet. **SolarWM**: data engine, metric annotations for 1.43M clips, code and weights.

The data layer moved more than the tooling layer. [TacVerse released handheld-UMI bimanual visuo-tactile sets](#) in LeRobot v3.0 with rectified tactile video streams, CC-BY-SA. [RoboMIND’s Agilex bimanual split was reprocessed](#) to v3.0 at 8,638 episodes, 5.4M frames and 54 tasks. [Apple’s EgoDex was converted](#) to v3.0 with train/test subdirectories — somebody else’s conversion week you no longer spend. [MIT Media Lab published a tactile and hand-pose contrastive encoder](#) with three seeds, retrieval mAP, training code *and* the exact preprocessed cache, a level of release completeness almost nobody manages. A [v3.0 → v2.1 downgrade script](#) now exists for openpi stacks that have not migrated.

Nothing in window from LeRobot (still v0.5.1, April), MuJoCo/MJX, Isaac Lab, ManiSkill, Genesis, RoboCasa or ROS 2.

[OpenWAM](#) dropped ~35 world-action checkpoints with no card or license — unusable as a baseline until someone says what they are.

5 Money & moves

NVIDIA → Hugging Face, \$12.9B

, 2026-09-03 ([TRADE PRESS](#)); NVIDIA states the platform stays open and its compute won't be required. **INFERENCE** \2192 LeRobot's maintainer — the default stack for one-arm builds — now sits inside the vendor of Isaac GR00T, and an open-platform commitment in a press release is a statement, not a governance structure.

Figure ↔ Nscale, \$3.5B compute

Scaling past \$6B, up to 100,000 Vera Rubin GPUs, Nscale taking equity ([TRADE PRESS](#)). **INFERENCE** \2192 humanoid capital is buying *training* capacity rather than fleets, and equity-for-compute means the supplier co-underwrites the bet.

Lyte, \$165M Series C at \$1.6B

For 4D coherent vision silicon ([TRADE PRESS](#)). **INFERENCE** \2192 perception hardware is a funded category again, on the premise that depth quality rather than policy architecture is the binding constraint. Also: [Medtronic put \\$700M into Cornerstone Robotics](#) (surgical, no robot-learning content); [Wandercraft is seeking €100M at €750M](#), not closed.

Below the headline numbers, the data layer is commercialising. [Visko closed a \\$10M pre-seed](#) for a streaming world model; [Kinetic Blocks opened a gated-beta marketplace](#) pricing teleop and robot-execution datasets with LeRobot v3.0 support and A+ to F quality grading; [X Square unveiled TwinDEX](#), an isomorphic wearable rig claiming 5.3× collection throughput and a 24-step chemistry task acquired with zero on-robot teleoperation — vendor claim, nothing released. **INFERENCE** \2192 once demonstrations are priced with quality grades, the scarce asset stops being episode volume and becomes the grading rubric.

6 Deep read of the week

One this week. Aero Hand Open and the adaptive-chunking paper were the other candidates, covered above; neither independently earns this depth — the hand's contribution is a released artifact rather than a result to interrogate, and the chunking gains sit inside unreported seed noise on its own primary benchmark.

MINERVA — how small a policy standard LIBERO actually demands. [arXiv:2609.03715](#), preprint, U. Tokyo. Unreproduced; code and checkpoints released.

Setup. All four LIBERO suites, `lerobot/libero` (1,693 demos, 273k frames, 40 tasks pooled), Franka in sim, 2,000 rollouts per point with hard resets. Held out in the standard protocol: nothing — that is the point. Two extensions do hold things out: LIBERO-90 (89 tasks) and LIBERO-Plus (10,030 perturbed variants over 7 factors).

Method. Inherited: ACT-style chunking and temporal ensembling, a flow-matching DiT head, spatial-softmax keypoints, FiLM conditioning, velocity distillation. New is the design principle — contain nothing the benchmark exercises, then shrink until it breaks.

LIBERO four-suite average vs. total parameters

MINERVA-0.5M (0.54M params)		95.05% success
MINERVA-1M (0.99M)		96.75% success
MINERVA-10M (9.66M)		97.45% success
π 0.5, LeRobot impl. (4.1B)		97.50% success

MINERVA rows are 2,000 rollouts at a single training seed. The π 0.5 row is the LeRobot implementation's reported result at 10 episodes/task (400 rollouts) — a reference point, not a protocol-matched comparison. A three-seed rerun of the 1M model spans 94.60–96.75, so gaps under ~1 point are not resolvable.

[arXiv 2609.03715](#), preprint, unreproduced; code and checkpoints released

The number that matters

Is not 95.05%. It is the ± 1 -point **training-seed band**, from retraining the 1M baseline three times (94.60–96.75; its long-suite score alone spans 87.6–92.8). Read the scaling curve against it: 96.75 at 1M to 97.45 at 9.66M is 0.7 points for 10 $\times$ the parameters, and the 97.50 reference is inside the same band. Against that instrument only two choices survive: chunk length (–3.16 at chunk 8, –2.20 at chunk 32, and the failure is interpretable — short chunks break the goal suite by letting the policy dither between branches, long chunks break the long suite by not reacting at subgoal boundaries) and vision allocation (–1.41 mean, –4.6 on long). Flow matching versus one-pass L1 regression is +0.34 across three seeds — no measurable difference — while regression is up to 3.8 $\times$ faster. The shipped 0.54M model closes its loop in ≈ 5 ms on eight laptop CPU threads, against 1.0 s for SmolVLA and 12.8 s for $\pi 0.5$.

The recipe also survives 2.25 $\times$ the task count: 94.6% average over 89 LIBERO-90 tasks at 0.995M parameters, single seed, 10 episodes per task, with a thin failure tail and no collapsed scene family.

Where it's soft. One simulator, one embodiment, one eval seed; the headline rows, LIBERO-90 and the distillation comparison are single runs inside that band. The $\pi 0.5$ comparison is against a differently-measured number. No real-robot validation. The permutation probe cuts both ways: freezing the released 1M checkpoint and rotating only its task-ID mapping collapses 96.75% to 6.5%, with Spatial, Goal and Long landing almost exactly on the 1-in-10 chance line. LIBERO-Plus prices the rest — 46–56% under perturbation, photometric robustness $\leq 6.5\%$ at every scale tested.

Steal list. (1) Measure your seed band before believing any ablation; three retrains is cheap here and invalidates most one-point deltas in the literature, several in this issue included. (2) Run the permutation probe — freeze weights, rewrite only the conditioning channel; if performance falls to chance, your conditioning is an index, not language. (3) Spend parameters on the encoder, not the action head: a token-mixing MLP matched self-attention at 26% fewer parameters, and moving that capacity into the CNN moved the long suite ~ 10 points at constant total. (4) Replan every step with ACT-style temporal ensembling (96.3%) rather than BID-style candidate selection (95.0%) or RTC soft inpainting (91.8%), which over-constrains an already-noisy velocity field; for sub-1M models drop the initial-noise temperature to 0.85, worth up to +2.25 points where the model is weakest and nothing at all above 1M.

Cost to reproduce. ~ 2 h per model on one consumer GPU, public LIBERO demos, no pretraining, no data collection. A seed-replicated ablation sweep is a weekend, not a cluster.

7 What this changes for a small build

Assume one arm, one GPU. Four things moved.

1. **Your LIBERO baseline is now a two-hour job.** MINERVA's checkpoints and recipe make seed-replicated ablations affordable, and hand you a ± 1 -point yardstick that most single-run deltas fall inside. Use a perturbation protocol before claiming robustness.
2. **A dexterous hand entered the printable price band.** \$314, 374 g, validated MuJoCo transmission model, zero-shot sim-to-real from encoders alone — with commissioning time budgeted for the thumb, which carries the residual error.
3. **Physical evaluation got cheap**, via Peg-in-Bench, if your contact-rich results currently live on a bespoke fixture and you have been unable to compare them to anyone.
4. **Two inference-loop changes are implementable straight from the papers.** Entropy-plateau truncation needs no retraining and removes a tuning knob. One-pass L1 regression instead of ten Euler steps was 3.8 $\times$ faster at no measured accuracy cost on this benchmark — worth testing on your own action head before assuming flow matching is buying you anything.

An existence proof at this scale also landed, in trade press only: a [DIY SO-101 on a tracked base running \$\pi 0.5\$ through LeRobot](#), Raspberry Pi on-robot with inference offloaded to a PC, reporting 96% on sock picking. Single build, no stated protocol — an anecdote, but on exactly the hardware envelope in question.

Nothing relevant appeared in hardware: no actuator, tactile-sensor, low-cost-arm, depth-camera or edge-compute release in window, and no releases at all from the major sim and middleware stacks.

8 Market signal

Visibly scaling: NVIDIA (buying the platform layer), Figure (buying training compute, supplier taking equity), Lyte (perception silicon), Wandercraft (raising). Last week's manipulation-specific buyers did not reappear; capital went to compute, sensing silicon and platform ownership rather than end-effectors or policy teams.

Three skill clusters recur in the work that got attention. **Execution-schedule engineering** — async inference loops, chunk stitching, one-step generation, adaptive horizons; knowing where a value gradient is even valid under async execution is now a specialization. **Training-time supervision at zero inference cost. Evaluation design as the contribution** — ZETA's protocol split, MINERVA's seed band, [FailBench](#) finding the best of 13 VLM success-judges reaches 0.77, and a [physics-consistent assistive-care benchmark](#) where zero-shot $\pi 0.5$ scores 0.7% with no physics-valid successes.

One structural note under the funding: with graded, priced demonstration marketplaces appearing alongside purpose-built collection rigs, the scarce skill in data work shifts from collecting episodes to specifying what counts as a good one — the acceptance criteria and IK-feasibility filters corpora like NeoData now document as a pipeline stage rather than an afterthought.

Table stakes: code and checkpoints at submission, action chunking, a flow or diffusion action head, async inference, LIBERO $\geq 95\%$. Still scarce: seed-replicated ablations, RL loops that close on real hardware, tactile data at scale, and hardware whose files you can print. No credible compensation or levelling datapoints surfaced — the only sources that appear for robotics comp queries are SEO aggregator pages with no stated methodology, and reporting nothing beats laundering those.

9 Threads

THREAD	STATUS	NEXT REAL UPDATE
Does VLA scale buy language-conditioned control?	ADVANCING (MINERVA)	A sub-10M policy on <i>real</i> hardware, multi-task, vs. a published baseline
Benchmark validity in physical AI	ADVANCING (MINERVA capacity floor + permutation probe; FailBench)	A major lab reporting a perturbation protocol alongside standard LIBERO by default
Contact sensing without dedicated tactile hardware	ADVANCING (Blind Dexterity, soft-arm IMUs, ChainSplat)	Independent reproduction on different hardware; VISTA (W35) still has no code
Training-time supervision at zero inference cost	NEW (PHR-VLA, Temporal Forcing, GIFT, EGR)	One replicated on a second backbone independently, with a seed band
Inference latency as deployment bottleneck	ADVANCING (adaptive chunking, DriftingVLA, MINERVA at ≈ 5 ms CPU)	Latency as a default column beside success rate
Online correction as the last mile	NEW (SmoothRL 250 episodes; XR-2 DAgger 58 $\rightarrow$ 93%)	A third group showing correction data beats more expert data past saturation
World-action models as unified policy + simulator	ADVANCING (WCD, SolarWM release)	A <i>released</i> world model used for policy improvement by someone who didn't build it
Who owns the open robot-learning stack	NEW (NVIDIA/Hugging Face, \$12.9B)	LeRobot's next release, and whether governance or format decisions visibly change
Low-cost open hardware as training substrate	ADVANCING (Aero Hand \$314, BRIDGE \$1.5K, Peg-in-Bench)	A third party training and deploying a policy on any of them

THREAD	STATUS	NEXT REAL UPDATE
Claimed open releases vs. actual artifacts	CONTESTED	For: MINERVA, Aero Hand, BRIDGE, SolarWM shipped. Against: OpenWAM has no card or license; StreamPI weights still pending since W35
Humanoid capital vs. demonstrated capability	ADVANCING (Figure \$3.5B; Wandercraft; Galbot pre-order)	Still zero methods from humanoid <i>companies</i> ; CMU's BRIDGE is the academic counterexample
Manipulation acquired rather than built	STALLED	W35's prediction of a second logistics buyer this quarter has ~7 weeks left

ALSO SEEN

VLAAct — 20% of the data beats full-data GR00T-N1.6 on an unseen humanoid. Skipped: 16-GPU continued-pretraining budget, out of reach despite open weights. Motus2 — shared-weight policy, simulator and evaluator in a closed improvement loop. Skipped: proprietary-scale data and no artifact, so the loop can't be inspected. ZimaBlue — video-to-action curriculum, 30 Hz control from egocentric video, code released. Skipped: fleet-scale pretraining means the code reproduces the recipe, not the result. Zeva — frozen policy retrieves action-effect transitions as an in-context causal prompt. Skipped: same gradient-free adaptation slot as WCD, with a less clean protocol. DriftingVLA — one-step noise-to-chunk generation, 3.36× faster. Skipped: same axis MINERVA's L1-vs-flow result attacks more cheaply and with three seeds. SA-WAM — frozen VAE tokenizes metric depth; geometry error flags 80% of failures. Skipped: no code stated, and the failure-flagging half has no detection baseline. World-Coherent Decoding — ranks sampled futures by flow surprisal and action-path effort, frozen backbone. Skipped: no code stated, in a week already carrying a training-free method with broader eval.	MAGP — scale-equivariant augmentation recovers metric scale into an existing VLA. Skipped: plug-and-play claim with no released adapter. Seeing the World and the Self — ego-adapted geometry backbone conditions a diffusion motion head. Skipped: code and dataset promised, not released; same unshipped metric-depth cluster as MAGP and VI3. Contact-Guided Exploration — exploration critic seeded by grasp-sampler contacts, annealed to zero mid-training. Skipped: locomanipulation-specific, sim-to-real only, no released environment. Non-Prehensile Throwing — RL over joint-jerk trajectories, zero-shot to a UR5e via jerk system identification. Skipped: single task, explicitly friction-sensitive, needs re-identification on your surface. DemoMimic — contact-local geometry rewards, 71% over 16 objects from one demonstration. Skipped: no code stated, reward still hand-specified per task family. Temporal robustness of imitation learning — ACT degrades far faster than its scripted expert as execution speed rises; code and data released. Skipped for space only; the week's cleanest negative result. RoboCurve / GPT-6 Astra on real arms (trade press) — 95% bowl / 10% puzzle on dual YAM arms. Skipped: n=20 per task, no paper, so the task spread is the only usable signal.
---	--

10 Open question

Is the ± 1 -point LIBERO seed band MINERVA measured a property of *that* policy family below 10M parameters, or of the benchmark at every scale? Nearly every ablation delta in this issue — GIFT's +4.6 pp, adaptive chunking's +1.3 on RoboTwin, StreamPI's +2.6 from W35 — is a single run, and how many survive depends entirely on the answer. Settled by any group retraining a $\pi 0.5$ - or GR00T-class policy three times on identical data and publishing the spread. It costs three training runs, and nobody has.

Sources labeled inline. All arXiv items are preprints, unreviewed. Single-lab claims flagged where relevant.