

This Week in Robotics

Ժամանակահատված՝ 2026-09-07 → 2026-09-13:

Մեքենայական թարգմանություն անգլերենից, չիմբագրված · բնօրինակը՝ twir.arcanerobotics.com/2026-w37/

ԳԼԽԱՎՈՐԸ

Գործիքների շերտը գարնանից ի վեր առաջին անգամ շարժվեց՝ [LeRobot v0.6.1](#), [MuJoCo 3.13.0](#)՝ նոր դիսկրետ ինտեգրատորով, [Genesis 1.4.1](#), իսկ [OpenWAM](#)-ը անցյալ շաբաթվա առանց քարտի ստուգակետերի (checkpoints) կույտը վերածեց աշխարհ-գործողություն մոդելների (world-action models) նախատեսուցման ամբողջությամբ հրապարակված, վերահսկվող ուսումնասիրության, մինչդեռ շաբաթվա ամենամաքուր մանիպուլյացիոն արդյունքն ընդհանրապես ուսուցում չէր օգտագործում:

1 Ռազմավարությունների ուսուցում և VLA-ներ

IMLE-VLA

ՊՐԵՊՐԻՆՏ / SFU + UPenn

Ճարտարապետական առումով փոփոխությունը մեկ օբյեկտ է. $\pi 0.5$ -ի 10 քայլանոց, հոսքի համընկեցմամբ (flow matching) գործողությունների գլուխը (action head) փոխարինվում է մեկ քայլով պայմանական գեներատորով, որն ուսուցանվում է պայմանական Implicit Maximum Likelihood Estimation-ով (անուղղակի առավելագույն ճշմարտանմանության գնահատում): Յուրաքանչյուր դիտարկման համար վերցնում եք m աղմուկի վեկտոր, գեներացնում m թեկնածու փաթեթ (chunk), յուրաքանչյուր իրական փաթեթը վերագրում ամենամոտ թեկնածուին՝ առանց գրադիենտների, և հետադարձ տարածում (backprop) կատարում միայն հաղթողի միջով: Հենց ամենամոտ հարևանի այդ վերագրումն է թույլ չտալիս, որ գլուխը փլուզվի դեպի պայմանական միջինը, ինչպես L1 գլխի դեպքում է. մոդաների ծածկույթն ապահովվում է կառուցվածքով, ոչ թե օժանդակ ռեգուլյարիզատորով: Գործնականում՝ $m=2$: Հիմնական ցանցը (backbone) սառեցված է, ուստի ուսուցանվում է միայն գլուխը:

Գնահատում՝ LIBERO-ի 40 առաջադրանք $\times$ 50 դրվագ (episodes) (98.0% միջին, ամենաբարձրը համեմատված մեթոդների մեջ), LIBERO-Plus խաթարումներ՝ չորս առանցք $\times$ հինգ խստության մակարդակ, և իրական Franka-ի չորս առաջադրանք՝ յուրաքանչյուրը 20 դրվագով: Գործարկումը (inference) L40S-ի վրա 15 Հց-ից դառնում է 55 Հց, իսկ կատարման 30 հորիզոնի դեպքում թողունակությունը գումարվում է մինչև $11\times$. իրական մանիպուլյատորի վրա չափված պրոպրիոցեպտիվ ցնցումը (jerk) նվազում է $2.2\text{--}3.0\times$: Կարևոր պնդումը կայունությանն է վերաբերում. OpenVLA-OFT-ի L1 գլուխը կտրուկ վատթարանում է LIBERO-Plus-ի խստության աճի հետ, մինչդեռ IMLE-VLA-ն հետևում է $\pi 0.5$ -ին: Թույլ կողմերը՝ 98%-ով LIBERO-ն հագեցած չափիչ գործիք է, իսկ LIBERO-Plus-ի կորերը կարդացվում են գծապատկերից, ոչ թե աղյուսակից: [Կողմ հրապարակված է](#). սառեցված 3B հիմնական ցանցի վրա միայն գլխի ուսուցումը սա լիովին դնում է մեկ GPU-ի հնարավորությունների սահմաններում:

ContextFlow

ՊՐԵՊՐԻՆՏ / KAUST + Berkeley

Համատեքստային (in-context) նմանակում առանց գրադիենտային թարմացումների՝ ավտոռեգրեսիայից հրաժարվելով: Նախորդ աշխատանքը (ICRT) գործողությունները թոքենացնում և ապակողավորում է մեկ առ մեկ, ուստի վաղ սխալը կուտակվում է, և ամենից շատ հենց կոնֆիգուրացիայի այն փոփոխության ժամանակ, որի համար էլ նախատեսված է համատեքստային ուսուցումը: ContextFlow-ը ցուցադրումներն անցկացնում է perceiver ոճի սեղմիչներով (յուրաքանչյուր մոդալության համար 32 ուսուցանվող հարցում, առանձին սեղմիչներ պատկերի, պրոպրիոցեպտիկայի և գործողության համար) դեպի ֆիքսված չափի համատեքստ, և π_0 -ով սկզբնավորված գործողությունների փորձագետը դրա նկատմամբ հոսքի համընկեցմամբ գեներացնում է անընդհատ փաթեթ՝ բլոկային պատճառային ուշադրությամբ, որպեսզի համատեքստի KV-ն քեշավորվի հոսքի 10 քայլերի ընթացքում:

Վերահսկվող թիվը՝ ContextFlow-Plain-ը 10.5 տ.կ.-ով գերազանցում է համապատասխանեցված ավտոռեգրեսիվ ելակետային մոդելին (baseline), որը կառուցված է π_0 -ի նույն սկզբնավորման վրա. սեղմիչների ավելացումը տալիս է ևս 9.5 տ.կ.՝ հասնելով 73.5%-ի չտեսնված LIBERO կոնֆիգուրացիաներում՝ ընդդեմ 53.5%-ի: Այն հավասարվում է յուրաքանչյուր չտեսնված առաջադրանքից մեկ ցուցադրումով լրաուսուցված (fine-tuned) π_0 -ին (73.5 ընդդեմ 72.5-ի)՝ ընդհանրապես առանց լրաուսուցման: Հետագծի կուտակային MSE-ն նվազում է 4.7 $\rightarrow$ 3.0 $\rightarrow$ 2.2 մ²: Թույլ կողմերը՝ ընդհանրացումը նույն պրիմիտիվների ներսում չտեսնված *կոնֆիգուրացիաների* վրա է, ինչը հեղինակներն անկեղծորեն նշում են. իրական ALOHA արդյունքները 10 փորձ են յուրաքանչյուր բջջի համար (լավագույն բջիջը՝ 6/10). իսկ LIBERO-ի չորս խմբերի (suites) վրա ուսուցումը LIBERO-Goal-ը և LIBERO-10-ը թողնում է 0.0%-ի վրա, մինչև չավելացվի LIBERO-90-ը: Կող նշված չէ:

DeCAL

ՊՐԵՊՐԻՆՏ / PKU + BAAI

Mixture-of-Transformers ճարտարապետությամբ ճարպիկ VLA, որի երկու նոր մասերն էլ վերաբերում են նրան, թե *երբ* է շոշափողական (tactile) տվյալը կարևոր: Ազապտիվ տեսողական-շոշափողական միաձուլումը շոշափողական խաչաձև ուշադրությունը փականավորում է ուսուցված հպման ազդանշանով, ուստի ազատ տարածությունում շոշափողական տվյալը ճնշված է մնում և բացվում է հպման ընթացքում. հողվածը ցույց է տալիս, որ փականի արժեքն աճում է հպման սկզբում: Տեսողական-շոշափողական թաքնված համատեղ երևակայումը միասին կանխատեսում է Cosmos-VAE-ի ապագա տեսողական թաքնված վեկտորները և ապագա շոշափողական թաքնված վեկտորները՝ ուսուցման ժամանակ ապակողավորելով դրանք չմշակված շոշափողական պատկերների, ձևափոխության քարտեզների և 6-DoF ուժերի, իսկ գործարկման ժամանակ ապակողավորիչները հեռացվում են: 71% միջին հաջողություն հպումներով հագեցած վեց առաջադրանքում՝ յուրաքանչյուրը 20 փորձով (լավագույն ելակետը՝ DECO $\approx$ 56%), 0.27 վ մեկ փաթեթի համար 4090-ի վրա: Թույլ կողմերը՝ հարթակը երկու UR5 է՝ 22-DoF SharpaWave ձեռքերով և յուրաքանչյուր մատնածայրին 320×240 տեսողական շոշափողական սենսորով, ինչը փոքր մասշտաբով վերարտադրելի չէ. ուսուցումը 8×H100 է՝ 100 հազ. քայլի համար, և յուրաքանչյուր բջիջ մեկ գործարկում է $n=20$ -ով: Միայն նախագծի էջ:

Արժե նշել որպես ուղղություն. **համատեքստային հարմարվում առանց թեստի ժամանակ գրադիենտների** թեման այս ժամանակահատվածում ստացավ չորս աշխատանք՝ ContextFlow, [ICI-VLA](#) (DTW-ով որոնված միկրոցուցադրումներ՝ սառեցված տեքստ-գործողություն VLA-ի մեջ, RoboTwin՝ 60.4%, իրական աշխարհ՝ 83.2%՝ $\sim$ 1 000 փորձում), [վարքային հուշումով ստուգակետեր LIBERO-ի վրա](#) և [2AM](#): Իսկ **առցանց RL իրական սարքավորման վրա** թեման ստացավ երկու՝ [VLA-Precision](#) (բացարձակ Q-ի փոխարեն հարաբերական առավելություն, սառեցված նախաձանցի KV-ի վերօգտագործում, 98.3% ինքնաշեն ինը քիմիական առաջադրանքում. հեղինակների սեփական հենանիշ, կող չկա) և ստորև ներկայացված WHIRL-ը:

2 Մանիպուլյացիա, ճարակություն և զգայուն

Գրիչով գրել ձեռքի ներսում՝ առցանց գնահատվող առաջադրանքի յակոբյանով

ԴԻԵԴԻԻՆՏ / ETH Zurich Soft Robotics Lab. [Կող և Նախագծի էջ](#)

Շաբաթվա լավագույն սարքավորումային արդյունքը, և այն ուսուցված ռազմավարություն (policy) չի պարունակում: Զլային (tendon) կառավարմամբ 17-DoF ORCA ձեռքի վրա 10 հոդ (բութ մատ, ցուցամատ, միջնամատ) շարժում են գրիչը, որը պահվում է ֆիքսված ուժային-ճշգրիտ բռնմամբ TPU թաղանթի մեջ, որը մոտավորապես քառապատկում է դրա տրամագիծը: 2×10 հրամանային տարածության առաջադրանքի յակոբյանը (Jacobian)՝ հողի հրամանային աճից դեպի գրիչի ծայրի շարժումը թղթի հարթությունում, գնահատվում է ռեկուրսիվ կերպով՝ մոռացման գործակցով, ArUco նշիչից ~ 15 Հց հաճախականությամբ, շրջվում է մարված կեղծ հակադարձով (pseudoinverse) և կառավարվում առաջադրանքի տարածությունում PID-ով՝ ուղիղ կապով (feedforward). զրոյական տարածությունը կրում է դիրքի անդամ, որը ձգում է դեպի սկզբնական բռնումը: Ռեգրեսորը *հրամանային* աճն է, ոչ թե չափված հոդային արագությունը, ինչը ջլային ձեռքի համար ճիշտ ընտրություն է:

Թվերը՝ ~ 18 վ սցենարային գրգռում, ապա հարթության մեջ 0.62 ± 0.12 մմ միջին սխալ օդում 22 գործարկման ընթացքում և 0.67 ± 0.08 մմ թղթի վրա 16 գործարկման ընթացքում ($p95 \approx 1.4$ մմ երկու դեպքում էլ), միասին՝ 0.64 ± 0.10 մմ $n=38$ -ի վրա: Մեկ անընդհատ 31 րոպեանոց գործարկման ընթացքում գրվեցին բոլոր 26 տառերը, և պատուհանային միջինները մնացին 0.59 -ի և 0.69 մմ-ի միջև՝ առանց աճի միտման. «K» տառի ժամանակ ~ 20 մմ շեղումը վերականգնվեց նույն տառի ներսում՝ առանց վերագրգռման: Հոդվածի իրական ներդրումն աբլյացիաներն (ablations) են. յակոբյանը գրգռումից անմիջապես հետո սառեցնելը տարամիտում է 3 գործարկումից 2-ում, գրելու 12 վ-ից հետո սառեցնելը՝ 2-ից 1-ում, իսկ ~ 30 վ-ից հետո սառեցնելը պահպանում է ելակետային ճշգրտությունը 4-ից 4-ում. այսինքն՝ գրգռման փուլն ինքնին օգտագործելի քարտեզ չի տալիս, իսկ մոտ կես րոպե աշխատանքի ընթացքում հարմարվելը՝ տալիս է:

Թույլ կողմերը, որոնց մեծ մասը հեղինակներն իրենք են նշում. կառավարվում է միայն հարթության (x,y) շարժումը, իսկ 2–3 մմ չկառավարվող z շեղումը կլանվում է թուղթը միտումնավոր ուռուցիկ դարձնելով, ուստի համակարգը չի կարող գրել կոշտ մակերեսի վրա կամ բազմահարված տառեր գրել. տառերի միջև մանիպուլյատորը վերադիրքավորվում է. հպման ընդհատումով շարժումները (վերաբռնում, մատների քայլք) ձևակերպումից դուրս են. գրելը կատարվում է 8×10^{-4} մ/վ արագությամբ, և նույնիսկ $2 \times$ արագացումն անհուսալի է. և՛ ամենակարևորը, որը պետք է զսպի վերնագրի ոգևորությունը, յուրաքանչյուր հաղորդված սխալ չափվում է նույն տեսողական խողովակաշարով, որը փակում է կառավարման շղթան, առանց թղթին մնացած թանաքի անկախ չափման: Սիմուլյացված Shadow Hand-ի (0.17 մմ RMSE) և Wuji Hand 2-ի (1.48 մմ) գործարկումները կառավարում են բոլոր երեք կոորդինատները, բայց ունեն առանց աղմուկի իրական արժեքներ: Մեկ լաբորատորիա, չվերարտադրված: Փորձելու արժեքը՝ ORCA ձեռք, վեբ-տեսախցիկ և նոութբուքի պրոցեսոր:

TacPAC

ԴԻԵԴԻԻՆՏ / Fudan + NeoteAI, [ԿՈՂ](#)

Մեխանիզմը ժամանակի ուղղում է: Աշխարհ–գործողություն մոդելը պլանավորում է փաթեթ՝ պայմանավորված կանխատեսված ապագա շոշափողական տվյալներով. այդ կանխատեսումն այնուհետև սառեցվում է փաթեթի կատարման ընթացքում, այսինքն՝ հենց այն պահին, երբ տեղեկատվական շոշափողական տվյալը գալիս է: TacPAC-ն աղմուկագերծումից հետո կատարում է մեկ լրացուցիչ մաքուր անցում, քեշավորում է կանխատեսված շոշափողական տեսքերի և պլանավորված փաթեթի շերտային բանալիներն ու արժեքները և թույլ է տալիս առանձին շոշափողական փորձագետին յուրաքանչյուր նոր շոշափողական կադրը կարդալ այդ քեշի նկատմամբ՝ չկատարված վերջածանցի վրա արտադրելով քայլային ուղղում: Տեսողական բանալիները միտումնավոր չեն քեշավորվում: Ուղղումները գրանցվում են $m+d$ շեղումից սկսած, որտեղ d -ն վերջերս դիտարկված ուղղման ամենամեծ ուշացումն է՝ գործողության քայլերով. գործնականում այն մոտ է զրոյի, և ռոբոտը ուղղման կեսին քայլ է առաջանում կանչերի միայն 1.6% -ում:

Հինգ իրական առաջադրանք Flexiv Rizon 4-ի վրա՝ երկու InTac S1 մատնաձայրային սենսորով, յուրաքանչյուր առաջադրանքի համար յուրաքանչյուր մեթոդով 20 փորձ: Միայն տեսողությամբ հիմնական մոդելը՝ $22\% \rightarrow 64\%$. ամենաուժեղ ելակետը (T-Rex)՝ 48% : Մաքուր աբլյացիան քեշն է. նույն հիմնական մոդելը, նույն ուղղիչը, նույն

շոշափողական մուտքը, պարամետրերի նույն քանակը, հեռացված է միայն ուշադրությունը քեշավորված կանխատեսված շոշափողական տվյալին՝ 47% ընդդեմ 64%-ի: Միայն շոշափողական կանխատեսման ավելացումը տալիս է 22% → 37%, այսինքն՝ շահույթի մեկ երրորդը: Մեկ ուղղումը տևում է 30.4 մվ՝ ընդդեմ փաթեթը վերագեներացնելու 628.6 մվ-ի: Օժանդակ ախտորոշումը լավն է. հպման անցումները կադրերի 4.9%-ն են, բայց շոշափողական ընդհանուր փոփոխության 18.6%-ը, ուստի ապագա տեսքի վերակառուցման կորստի ֆունկցիայում գերիշխում են այն կադրերը, որտեղ ոչինչ չի կատարվում: Առանց սպասումի ռեակտիվ ուղղումն իրականում իջնում է միայն տեսողության մակարդակից ցածր վնասվածքով գնահատվող երկու առաջադրանքում. ապացույց, որ առանց նախնական գիտելիքի հպման հետադարձ կապը կարող է ճիշտ պլանը շեղել:

Երկու մասի հավաքում մեկ ձեռքում

ՊՐԵՊՐԻԵՏ / HKU + HKUST(GZ)

+ ETH): Մեկ 22-DoF Sharpa Wave ձեռքը միացնում է երկու առարկա՝ առանց երկրորդ մանիպուլյատորի և առանց հարմարանքի՝ շփ կափարիչ, ներարկիչի միոց, մարկերի կափարիչ: RL IsaacSim-ում 34-չափանի դիտարկման վրա (առարկաների երկու կենտրոն և համաչափության առանցքներ՝ գումարած 22 հոդային անկյուն. առանց արագությունների, առանց հպման ուժերի), ռեկուրենտ LSTM ռազմավարություն, նպատակին հարաբերական դիրքի պարզև՝ գումարած բազմապատկող օժանդակ անդամ, որը մեղմորեն վերագրում է բուք մատը+ցուցամատը պահվող մասին, իսկ միջնամատը/մատնեմատը/ճկույթը՝ ֆիքսված մասին, ռեգուլյարիզացված դեպի մարդու մեկ պատկերված դիրքը: Սարքավորմանն առանց լրացուցիչ ուսուցման (zero-shot) փոխանցում մեկ RealSense D435-ով և FoundationPose-ով: Փակ շղթայով՝ 15/20, 17/20, 16/20 երեք առաջադրանքում՝ ընդդեմ հաջողված սիմուլյացիոն հետազոտների բաց շղթայով կրկնարտադրման 2/20, 0/20, 0/20 արդյունքների. սա ամենամաքուր վկայությունն է, որ այստեղ սիմուլյացիայից իրական աշխարհ խզումը փակվում է հետադարձ կապով, ոչ թե ճշգրտությամբ: Ձեռքի մորֆոլոգիայի ուսումնասիրությունն է պահպանելու արժանի մասը. նույն խողովակաշարը Allegro-ի վրա ձախողում է z առանցքով տեղադրումը հինգերորդ մատի բացակայության պատճառով, իսկ XHand-ի վրա ձախողում է xy հավասարեցումը հեռացման հոդերի բացակայության պատճառով: Թույլ կողմերը՝ երկու առարկաներն էլ ձեռքի մեջ տեղադրում է մարդը և պահում, մինչև ռազմավարությունը միանա. նախնական տաքացման քայլերի համար ծանրության ուժը հանված է. և հեղինակները PhysX-ի՝ մակերեսային հպումը մոդելավորելու անկարողությունն են նշում որպես պատճառ, թե ինչու են սեղմող բռնումները դաստակի որոշ թեքությունների դեպքում վատթարանում: Աբլյացիաները կատարված են ուսուցման 3 գործարկման վրա:

WHIRL

ՊՐԵՊՐԻԵՏ / HKUST(GZ)

+ IIT, միայն նախագծի էջ): Մարդը-շղթայում RL-ը սովորաբար մարդու միջամտությունն օգտագործում է մեկ անգամ՝ որպես վարքի կլոնավորման դիմակ կամ կրկնարտադրման կշիռ: WHIRL-ը մեկ քայլով թաքնված աշխարհի մոդելին (world model) ավելացնում է չորրորդ գլուխ՝ դինամիկայի, պարզևի և ավարտի կողքին, որը կանխատեսում է հաջորդ վիճակում մարդու միջամտության հավանականությունը, և միայն այդ գլուխն է ուղղում դերակատարի (actor) օգտակարության մեջ, երբեք՝ Բելմանի թիրախի մեջ: Պատճառաբանությունը բացահայտ է և ճիշտ. միջամտության հավանականությունը կանխատեսում է օպերատորի վարքը, ոչ թե առաջադրանքի արժեքը, ուստի որպես պարզևի հավելում այն կփոխեր օպտիմալ ռազմավարությունը և օպերատորի սովորությունները կտարածեր յուրաքանչյուր թարմացման միջով: 16-DoF LEAP Hand-ի և Franka FR3-ի վրա՝ հինգ առաջադրանք. +15–30 տ.կ. համապատասխանեցված, առանց մոդելի մնացորդային HIL ելակետի նկատմամբ (96.7% Pick LEGO-ում, 100% Pick Cube-ում), և Pull Drawer-ում քայլերով կշռված միջամտության բաժինը 0.113-ից իջավ 0.018-ի, այսինքն՝ 84% կրճատում: ~8 հազ. առցանց քայլ մեկ RTX 4090-ի վրա վերցնելու և դարակի առաջադրանքների համար, յուրաքանչյուրին՝ 15–30 ցուցադրում: Թույլ կողմերը, որոնք հեղինակներն իրենք են հաղորդում՝ մեկ սերմ (seed) և մեկ օպերատոր յուրաքանչյուր բջջի համար, Ֆիշերի ճշգրիտ թեստը $p < 0.05$ -ի է հասնում միայն վերցնելու երեք առաջադրանքում, և միջամտության գլուխը ժառանգում է պիտակավորող օպերատորի զգուշության շեմը:

Թե մարդը, թե ռոբոտի ձեռքը կրում են նույն 20-DoF իզոմորֆ էկզոկմախքը, որի ընդհանուր շրջանակին կոշտ ամրացված են էնկոդերներ և դաստակի տեսախցիկներ: Նախագծային այդ մեկ որոշումը վերաթիրախավորումը (retargeting) բաց շղթայով ազատ տարածության ռեգրեսիայից վերածում է զուգակցված տվյալների վրա վերահսկվող ուսուցման. մարդու ցուցադրումը կրկնարտադրեք ռոբոտի վրա «թոթովանքով» հարմարեցված քարտեզավորման միջոցով, և մարդու ու ռոբոտի կողմի էնկոդերների հետքերի տարբերությունը հենց ուղղման ազդանշանն է: Այն նաև վերացնում է պատկերի լրացման (inpainting) խնդիրը, քանի որ երկու մարմնավորումներն էլ դաստակի տեսախցիկին ցույց են տալիս նույն արտաքին մեխանիզմը: AirPods-ի տեղադրման 52 հաջող ցուցադրում 30 րոպեում՝ ընդդեմ հեռակառավարմամբ 18-ի (3.0×). 70.0% միջին հաջողություն փորձարկումներում՝ ընդդեմ հեռակառավարմամբ հավաքված տվյալների 71.7%-ի, իսկ զուգակցված լրատուցումն արժե +12.7 տ.կ.՝ կենտրոնացած հսկման բեռով առաջադրանքների վրա (+16.7 պտուտակ պտտելու, +15.0 սեղան մաքրելու և ցողելու մեջ): Ոչինչ չի հրապարակվել, մեկ օպերատոր, մեկ թիրախային ձեռք:

Զգայման երկու արդյունք, որոնք արժանի են տեղի: **X-Hinges** (գրախոսված, [UIST '26](#), MIT CSAIL + Tianjin) միասին տալում է երկու հաղորդիչ թել, որոնց տեսակարար դիմադրությունները տարբերվում են երեք կարգով (1.23×10^4 ընդդեմ 8.75 Օմ-սմ-ի), Lamina Emergent Torsional ճկուն հողի մեջ, այնպես որ բարձր դիմադրությամբ տարրը կրում է ձևափոխության ազդանշանը, մինչդեռ հաղորդողիները աննշան աղմուկ են ավելացնում, իսկ ոչ հաղորդիչ TPU մարմինը մեկուսացնում է մաշկի հայումը: Յուրաքանչյուր առանցքի դիֆերենցիալ դասավորությունները Ուիթսթոնի կամրջի տրամաբանությունն իրականացնում են տպված երկրաչափությամբ. առանցքների միջև խաչաձև միջամտությունը՝ մինչև 8.2%. կամրջի լայնությունը կարգավորում է կոշտությունը երկու կարգի սահմաններում (մինչև $470 \times$). 420 000 ցիկլ 233 ժամում՝ 5.5%-ից պակաս դիմադրության շեղումով. շերտերի վերադրմամբ թելերի միացումը 22.3 կՕմ է՝ ընդդեմ ծայրից ծայր հսկման 952 կՕմ-ի: Անհրաժեշտ են բազմանյութ FDM տպիչ և հատուկ տրանսիմպեդանսային նախնական մշակման սխեմա (ADS1256, 50 Հց/ալիք, 0.1–100 ՄՕմ). CAD-ը և ծրագրակազմը խոստացված են բաց կոդով հրապարակել գրախոսությունից հետո: **TacClip** (ՊԵՏԱԿԱՆ, Stanford, Cutkosky lab) Fiber Bragg Grating սենսորն ամրացնում է մատնածայրի վրայով, այնպես որ *մատի բարձիկը բաց է մնում*. հսկման բեռները հյուսվածքը կողքից ուռեցնում և ծռում են ամրակը: Ուժի մեծության սխալը սովորաբար 0.5 Ն-ից պակաս է 0–8 Ն միջակայքում, 2 կՀց նմուշառում, հարթ արձագանք մինչև ~20–30 Հց (ցածր հաճախությունների ֆիլտրը մատն է, ոչ թե FBG-ն), և այն աշխատում է ջրի տակ: ~\$2 տպված ամրակի համար և ~\$20 յուրաքանչյուր մատի սենսորային գլխիկի համար՝ չհաշված օպտիկական հարցախույզ սարքը (interrogator), որտեղ էլ իրական ծախսն է: Այն գնահատում է Ifl-ը, ոչ թե նորմալ ուժը. մեկ FBG-ն առանցքներն իրարից չի անջատում:

3 Հարթակներ և մարմնավորում

Մարդանման ռոբոտներ արտադրող ոչ մի ընկերություն կրկին մեթոդ չհրապարակեց՝ երրորդ անընդմեջ ժամանակահատվածում: Ամենամոտն Unitrue-ն էր, որը Hub-ում տեղադրեց [UnifoLM-ER-1](#)-ը՝ Apache-2.0 լիցենզիայով Qwen3-VL-4B մարմնավորված դատողական մոդել (reasoner), որն ուսուցանվել է տարածական դատողության և մարմնավորված 5 մլն+ նմուշի վրա և առաջատար է բաց մոդելների մեջ 16 հենանիշից 7-ում (93.4 BLINK, 88.9 EmbSpatial, 82.0 Where2Place), գումարած օպտիկական հոսքի տիրույթների կանխատեսիչ: 6B գործողությունների մոդելը դեռ նշված է որպես սպասվող: Այսպիսով՝ դատողական մոդել, ոչ թե ռազմավարություն:

Այս շաբաթ հարթակների շերտում մնացած ամեն ինչ արտադրություն և ֆինանսներ էին: [XPENG-ը մեկնարկեց IRON-ի զանգվածային արտադրությունը՝ 2 250 TOPS սեփական արտադրության երեք Turing չիպերում, >80% գծի ավտոմատացում, մատակարարման շղթայի 85% համընկնում իր էլեկտրամեքենաների բիզնեսի հետ, 110 000 մ² Գուանչժոուում, ցուցասրահներում տեղակայում 2026 թ. 4-րդ եռամսյակում և կորպորատիվ առաքում 2027-ին: Ոչ գնահատում, ոչ մեթոդ, ոչ ռազմավարության բացահայտում. պնդումը, թե ռոբոտները գծից դուրս են գալիս ինքնուրույն, մամլո հաղորդագրության մակարդակի սպացույց է: \[Nucleus-ը Nucleus II-ում հրաժարվեց երկուսանի\]\(#\)](#)

[նութերից՝ հօգուտ անվավոր հենքի՝](#) գործարանային արձագանքից հետո. մատակարարի պատմություն, առանց տվյալների, բայց անվավոր երկձեռ ուղղության երկրորդ տվյալը երկու շաբաթում՝ Nori-ից հետո:

Մարդանման ռոբոտների ակադեմիական աշխատանքներ կան, բայց չեմ կարող երաշխավորել դրանց համար. [TANGO](#) (ամբողջ մարմնով նավիգացիայի VLA, միայն սիմուլյացիա), [CAP](#) (խորության աղմուկագերծում վատթարացած ընկալմամբ G1-ի տեղաշարժի համար, $n=5$ իրական փորձ, միայն նախագծի էջ), [ViBe](#), [տեղաշարժ հատիկավոր տեղանքում](#): Բոլորը պրեպրինտներ, բոլորը գնահատված ամփոփումներից, ոչ մեկում կոդ նշված չէ: ROS Discourse-ում 2026-09-10-ին հայտարարվեց [Bimo անունով ROS 2 բաց երկոտանի հավաքածու](#). չստուգված է, և դրա համար կայուն մշտական հղում չունեմ:

4 Բաց կոդ և գործիքներ

Վերացված խոչընդոտների առումով սա W35-ից ի վեր լավագույն շաբաթն է:

[LeRobot v0.6.1](#)՝ առաջին թողարկումը ապրիլի v0.5.1-ից ի վեր և առաջինը NVIDIA/Hugging Face գործարքից հետո: Էական փոփոխությունը վերին հոսքի (upstream) Transformers-ի վրա համախմբումն է. Wall-X-ն այժմ ժառանգում է բնիկ Qwen2.5-VL-ը, իսկ XVLA-ն՝ բնիկ Florence2-ը, և VLA-ի ընդհանուր բաղադրիչները կիսվում են $PI0/PI05/EO1/PI0$ -Fast/SmolVLA մոդելների միջև: Դիֆուզիոն ռազմավարության ուսուցումը ստանում է գրադիենտների ստուգակետավորում (gradient checkpointing). $PI0$ -Fast-ի ապակոդավորիչը կանգ է առնում գործողության ավարտի նշիչի վրա. տվյալների հավաքածուների վիճակագրությունն այժմ կադրերի ենթանմուշառում է անում՝ հիշողությունը սահմանափակելու համար. հոսքային տվյալների հավաքածուներն աջակցում են bucket-ներ և մասնավոր պահոցների թոքեններ: Սարքավորում՝ RealSense-ի ձեռքով լուսակայում/ ուժեղացում/սպիտակի հավասարակշռություն, Dynamixel-ի երեք նոր մոդել (XH540-W150, XC330-T288, XC330-T181), Feetech-ի ժամանակավոր դողի սխալների դեպքում կրկնափորձ: **Համատեղելիությունը խախտող փոփոխություն՝** `lerobot.types` → `lerobot.lerobot_types`:

[MuJoCo 3.13.0](#) (2026-09-08)՝ տարվա ամենանշանակալի սիմուլյատորային փոփոխությունը բոլոր նրանց համար, ովքեր աշխատում են ջլերի կամ կոշտ շարժաբեքերի (actuators) հետ: Նոր *discrete* ինտեգրատորը սահմանափակումների լուծումը և արագության անուղղակի թարմացումը միավորում է մեկ գործողության մեջ՝ հողերի, ջլերի և շարժաբեքերի կոշտությունն *m* մարումը ներառելով լուծիչի արդյունավետ մետրիկայի մեջ. պասիվ զսպանակները և շարժաբեքերի դիրքային ուժեղացումները կայուն են դառնում բացահայտ սահմանից շատ ավելի մեծ ժամանակային քայլերի դեպքում, իսկ հենց դա էր սերվո մոդելների վրա ստիպում 1 կՀց+ քայլեր: Առանձին՝ էլիպսային կոներով Նյուտոնի լուծիչն այժմ կատարում է մեկ վերաֆակտորիզացիա՝ յուրաքանչյուր հպման համար ռանգ-1 թարմացումների փոխարեն, որտեղ դա ավելի էժան է. միաժամանակ սահող բազմաթիվ հպումներով տեսարաններն աշխատում են $1.4\text{--}2\times$ արագ: Գումարած հարթություն-ցանց բախումների վերաշարադրում, ցանցերին կապված կետեր (sites) և Python 3.15՝ ներառյալ free-threading: Համատեղելիությունը խախտող փոփոխություն՝ 3.11.0-ի անուղղակի flex-ի արդյունավետ մետրիկայի հատուկ դեպքը հանված է. ազդված մոդելները պետք է անցնեն *discrete*-ի:

Այն լրացումն է, որը ոչ ոք չէր հրապարակել. MuJoCo շարժիչի պլագին, որը մալուխի *ուղղորդումը* դարձնում է մեխանիզմի մաս: Բնիկ տարածական ջլերը տալիս են մարմնին ամրացված կետեր, գնդային/գլանային փաթաթումներ և մեկ փոխանցվող ուժ՝ Հուկի օրենքով, որը կարող է թե՛ հրել, թե՛ քաշել: MuJoCable-ը համատեղ օպտիմալացնում է հարակից շարժվող վերլուծական և ցանցային մակերեսներով անցնող կարգավորված ուղին, կիրառում է միակողմանի, միայն քաշող առանցքային օրենք՝ թույլության պաշարով և լարման սահմանով, տարածում է ուղղորդված Կապստանի առնչությունները, որպեսզի յուրաքանչյուր հատված կրի իր սեփական լարումը, և ստացված հանգուցային ուժերը մարմիններին է փոխանցում վիրտուալ աշխատանքի միջոցով: Ծախարակային հենանիշները վերականգնում են վերլուծական փոխանցումը՝ Կապստանի հարաբերակցության 0.5%-ից պակաս սխալով: Արժանիքն ապացուցող արդյունքը՝ 18 հոդով SpiRobs-ի վրա բնիկ ջիլը և MuJoCable-ը տալիս են գրեթե նույն կայուն ծոռումը, մինչդեռ MuJoCable-ը լուծում է մալուխի գազաթնային լարման $\sim 4\times$ և 9%-ով ավելի ձգում. նման ընդհանուր շարժում, բայց փոխանցման տարբեր վիճակ: $\mu=0.60$ դեպքում ուղղորդիչի շփման ստուգումը խլում է ծոռան մեկ քառորդը և պտույտը վերաբաշխում դեպի մոտակա հոդեր: Լրացուցիչ ծախսը քայլի ժամանակի 22.5% է (41.96 ընդդեմ 34.25 մվկ մեդիանի), և այն դեռ $11.9\times$ արագ է իրական ժամանակից 0.5 մվ քայլի դեպքում: Պահոց նշված չէ, ինչը պլագինի մասին հոդվածի ամբողջ խնդիրն է:

Genesis v1.4.0 (2026-09-06) և **v1.4.1** (2026-09-12): v1.4.0-ի `gs replay`-ը տեսարանը և ամբողջական հետագիծն արտահանում է որպես մեկ ինքնուրույն արխիվ՝ մեկ այլ մեքենայի վրա բիր առ բիր ճշգրիտ կրկնարտադրմամբ. սա ամենաօգտակար բանն է սխալների մասին հաղորդումների և բոլոր նրանց համար, ովքեր փորձում են վերարտադրել ուրիշի սիմուլյացիոն արդյունքը: v1.4.1-ը մեծ տեսարանների ծախսը դարձնում է կոդիների քանակից ենթագծային թե՛ CPU-ի, թե՛ GPU-ի վրա: Նշում՝ GitHub-ի թողարկումների արտապատկերված էջերն իմ ներբեռնման ճանապարհով վերադարձնում են ակնհայտորեն սխալ ամսաթվեր (v0.6.1-ը վերադարձավ որպես «August 3, 2024», Genesis v1.4.1-ը՝ «September 12, 2024»): ստուգել եմ պիտակները, իսկ ամսաթվերը վերցված են թողարկումների API-ից:

Նաև հրապարակվեցին՝ **FARM**՝ 33 985 պարամետրով վերահսկվող ընթերցիչ (readout) սառեցված VLA-JEPA կանխատեսողական վիճակների վրա, որը ձախողումները հայտնաբերում է յուրաքանչյուր քայլում, 85.68/88.59 միավորված AUROC/AUPRC յոթ աղբյուր առաջադրանքում՝ հնգապատիկ out-of-fold գնահատմամբ, Seen դեպքում լավագույն արդյունքը 15 համապատասխանեցված ելակետի մեջ, 0.2256 մվ միջին CUDA ուշացում, և իրական ռոբոտի վրա փոխանցում՝ ստուգված PIPER X-ի, SO-101-ի և Franka-ի վրա (**ՊՐԵՊՐԻԵՏ**): Առանց լրացուցիչ ուսուցման փոխանցումը միատարր չէ. ընթերցիչի ծագումը հսկայական նշանակություն ունի (41.35 ընդդեմ 84.70 AUROC-ի PIPER X-ի նույն պոպուլյացիայի վրա՝ կախված նրանից, թե որ աղբյուր պոպուլյացիան է ուսուցանել ընթերցիչը), բայց միայն ընթերցիչի հարմարեցումը հասնում է 98.48-ի: **HuRo**-ն հրապարակում է ~ 630 հազ. ռոբոտացված դրվագ / 142 մլն կադր / ~ 1 317 ժամ, որոնք ստեղծվել են հինգ էգոկենտրոն կորպուսներից մարդկանց «ջնջելով» (inpainting) և դրանց փոխարեն Isaac-Sim-ով արտապատկերված ռոբոտ տեղադրելով՝ IK-ով վերաթիրախավորված գործողություններով. կարևոր աբլյացիան այն է, որ մարդու սկզբնական դիտարկումները պահելը (գործողությունները՝ նույնպես վերաթիրախավորված) համընկնում է բաշխման ներսում, բայց բաշխումից դուրս կորցնում է 16.5 կետ: **Xiaomi-Robotics-U0**-ն հրապարակում է Apache-2.0 լիցենզիայով 4B և Sequence ստուգակետեր 34B ընտանիքից՝ գումարած VisionTokenizer, բայց քարտում ասվում է, որ ուսուցման կողը «կավելացվի ապագա թողարկումում», ինչը հակասում է այն հաղորդումներին, թե ուսուցման կողը բացվել է: Նշված է որպես հակասական. կշիռներն իրական են, ուսուցման կողը դեռ չկա:

Տվյալների շերտի նյութեր մեկ մանիպուլյատորի մասշտաբով՝ **RobotArena360-ի ակտիվներ** (DROID-ի 22 իրական տեսարան՝ որպես փոխազդելի Gaussian-splat երկվորյակներ՝ բախման ցանցերով և չափաբերմամբ, CC-BY-4.0), **1 530 դրվագից / 440 967 կադրից բաղկացած իրական UR5 + Robotiq եռամատ հավաքածու** Apache-2.0 լիցենզիայով, **LIBERO-90-ը՝ վերափաթեթավորված LeRobot v3.0-ի** (առանց քարտի), և **LoRA աղապսներ** **lerobot/pi05_base-ի վրա՝ SO-101-ով բլոկներ դարսելու համար**:

5 Փող և գործարքներ

Maven Robotics, \$100 մլն Series A

ՈՒՈՐՏԱՅԻՆ ՄԱՄՈՒԼ

, գլխավորել է RoboStrategy-ն՝ LocalGlobe-ի, Vine Ventures-ի և XTX Markets Ventures-ի մասնակցությամբ. թիմը եկել է Apple-ի Special Projects Group-ից: Անվավոր երկձեռ ծծող մանիպուլյատոր խառը պալետավորման համար, ութ միավոր 16 ժամյա հերթափոխերով՝ պնդվող $\geq 99\%$ աշխատունակությամբ, ֆինանսավորում է երրորդ սերնդի 250 միավոր: **ԵԶՐԱՀԱՆԳՈՒՄ** → հիմնական փաստարկը մոդելը չէ, այլ ինքնավար մեքենաների (AV) ոճի տվյալների գործառնական ցիկլը (ժամերի ընթացքում հավաքել օբյեկտի եզրային դեպքերը, գործարկել ավտոմատացված աբյուսիաներ, թարմացնել կշիռները, նորից տեղակայել)՝ գումարած աքցանաձև ձեռնոցներ, որոնք հավաքում են մարդու ցուցադրումներն արդեն վերջնական գործիքին (end-effector) համապատասխանեցված: Կապիտալը գնահատում է կրկնությունների ցիկլը, ոչ թե ռազմավարությունը:

Skild AI-ը տասը ամսում անցավ \$100 մլն ARR-ի շեմը

ՈՒՈՐՏԱՅԻՆ ՄԱՄՈՒԼ

, 60+ հաճախորդ, հասույթի բաշխում՝ 86% մանիպուլյացիա / 10% շարժունակություն / 4% Fetch. S1-ի համար հաղորդվում է բաշխումից դուրս քայլային 66% հաջողություն, իսկ մեկ ցուցադրական տեսանյութը համարժեք է մոտ 380 հեռակառավարման օրինակի: **ԵԶՐԱՀԱՆԳՈՒՄ** → սա ռոբոտների ուսուցման կիրառական հասույթի առաջին թիվն է այս մասշտաբով, և դա մանիպուլյացիա է, ոչ թե մարդանման ռոբոտներ: Բոլոր թվերն ընկերության սեփական տվյալներն են, բացի անվանված հաճախորդի թողունակությունից (50 → 130–140 տող/ժամ G10 Fulfillment-ում). ոչինչ աուդիտի չի ենթարկվել:

Agility Robotics-ի S-4-ը

ՈՒՈՐՏԱՅԻՆ ՄԱՄՈՒԼ

2025-ին \$1.8 մլն զուտ վաճառք՝ ~\$140 մլն գործառնական վնասի դիմաց, \$111 մլն գործառնական ծախս (2024-ին՝ \$71 մլն), ~\$100 մլն միջոցների ծախս, >\$300 մլն բազմամյա Digit v5 պատվերներ մեկ հաճախորդից, 65 000+ աշխատանքային ժամ ինը վայրում: Churchill Capital XI-ի հետ միաձուլումն այն գնահատում է \$2.5 մլրդ, այսինքն՝ 2025-ի հասույթի մոտ 1 400×, >\$620 մլն համախառն մուտքերի դիմաց, ներառյալ Foxconn-ի գլխավորած \$200 մլն PIPE: Դրան հակադրեք **Unitree-ի բաժնետոմսերի 53% անկումը ներօրյա գազաթնակետից**. 2025-ին ¥1.70 մլրդ (\$252 մլն) հասույթ, 5 500+ առաքված մարդանման ռոբոտ, շահութաբեր, այժմ՝ ~\$30 մլրդ, կամ հասույթի ~125×: **ԵԶՐԱՀԱՆԳՈՒՄ** → հանրային շուկաներն այժմ ունեն մարդանման ռոբոտների երկու համեմատելի ընկերություն, որոնց հասույթի բազմապատկիչները տարբերվում են ավելի քան 10×, և շահութաբերն է էժանր: Սա ոլորտի ստացած առաջին իրական գնային ազդանշանն է:

OpenAI-ը հաստատում է մարդանման ռոբոտի ստեղծումը

ՈՒՈՐՏԱՅԻՆ ՄԱՄՈՒԼ

Altman-ը՝ փողքաստում, գումարած ռոբոտաշինության 19 բաց հաստիք Սան Ֆրանցիսկոյում՝ չորսը շարժաբեղերի վրա և մեկ TPM շարժաբեղերի գծով, ինչպես նաև PCB դասավորում, ջերմային սիմուլյացիա, ներկառուցված ծրագրեր, նախատիպավորում և ապրանքային մատակարարման կառավարում: Ոչ բնութագրեր, ոչ ժամանակացույց: **ԵԶՐԱՀԱՆԳՈՒՄ** → աշխատանքի ընդունումն ամբողջությամբ սարքավորման ոլորտում է: Շարժաբեղերը կազմում են մարդանման ռոբոտի նյութերի ցուցակի (BOM) 40–60%-ը, և մատակարարումը կենտրոնացված է. OpenAI-ը գնում է այն մասը, որը չի կարող ներբեռնել:

Նաև՝ ARM Institute՝ ~\$90 մլն ռազմական արտադրության ռոբոտաշինության տասը նախագծի համար, Swarmer-ը ձեռք է բերում Ratel Robotics-ը մինչև \$224 մլն-ով (պաշտպանական UGV), Enovis/eCential՝ ~\$245 մլն (վիրաբուժական), Teradyne-ը դատի է տալիս JAKA-ին Universal Robots-ի երեք արտոնագրի համար, և Boston Dynamics-ի վետերանները հիմնում են Dynamic Creatures-ը՝ կերպարային ռոբոտաշինության համար:

ԵԶՐԱՀԱՆԳՈՒՄ → շաբաթվա միաձուլումներից և ձեռքբերումներից (M&A) ոչ մեկը չի առնչվում ռոբոտների ուսուցմանը: Վիրաբուժական և պաշտպանական համախմբման երեք անընդմեջ շաբաթ:

6 Շաբաթվա խորը ընթերցումը

OpenWAM — ինչ է իրականում ժառանգում աշխարհ-գործողություն մոդելը, և չափման ինչ

ճշգրտությամբ: [arXiv:2609.07398](https://arxiv.org/abs/2609.07398), արեարիմտ, NUS + Tsinghua + PKU + HKU և ուրիշներ: [Կոդ](#), [կշիռներ և](#)

[տվյալների բաղադրատոմսեր](#): Չվերարտադրված: Մյուս թեկնածուն գրիչով գրելու հոդվածն էր, որը ներկայացված է վերևում. այն ավելի լավ արդյունք է, բայց ավելի փոքր մակերես. նրա արվեստագիտներն արդեն ամբողջական են և ազնիվ, և հետաքննելու բան քիչ է մնացել:

Փորձի կառուցվածքը: WAM-ը ժառանգում է տեսանյութ գեներացնող նախնական գիտելիքը (priors) և այն ուղղում գործողությունների: OpenWAM-ը դա բաժանում է երեք փոխարինելի մոդուլային դասի՝ տեսողական կոդավորիչ (encoder) E , հոսքերի հիմնական ցանցեր S , տեսանելիության ուշադրության դիմակ M , և կազմման կանոնի, որը հավաքում է վեց ճարտարապետություն մեկ-, երկ- և եռահամակարգային ընտանիքներում՝ հինգ տեսանյութային հիմնական ցանցերի վրա (Wan2.1-VACE-1.3B-ից մինչև Wan2.1-I2V-14B), չորս կոդավորիչի (Wan2.2-VAE, FLUX.2-VAE, DINOv3, V-JEPA 2.1՝ ըստ ցանկության S-VAE-ով սեղմված) և չորս միջմոդալ դիմակի (Isolated, Action-Sees-Video, Video-Sees-Action, Mutual): Վերահսկվող ուսումնասիրությունն անցկացվում է RoboTwin2.0-ի վրա Full և Clean2Random կարգավորումներում, որոնցից վերջինը տալիս է տիրույթի ներսի և տիրույթից դուրս առանձին բաժանումներ: Մասշտաբավորված տարբերակը՝ OpenWAM- α -ն, նախատեսված է 518.5 մլն կադրի վրա (~6 369 ϕ), որոնք ընտրված են 1.33 մլրդ կադրի չմշակված ավագանից. 30%՝ սեփական էգոկենտրոն մարդկային տեսանյութ (71.6 հազ. առաջին դեմքով ձայնագրություն, 3 006 առաջադրանք, գործողությունների ալիքներն ամբողջությամբ դիմակավորված, ուստի այն վերահսկում է միայն աշխարհի հոսքը), 40%՝ իրական ռոբոտ (AgiBotWorld-Beta, RoboCOIN, DROID) և 30%՝ սինթետիկ (InternData-A1), բոլորը 80-չափանի միասնական գործողությունների տարածության միջոցով՝ ֆիքսված պլոտային իմաստաբանությամբ և վավերականության դիմակավորմամբ: Սիմուլյացիոն ուր հենանիշ՝ գումարած իրական մեկ մանիպուլյատորով, երկձեռ և ճարպիկ ձեռքով հարթակներ, ներառյալ 21-DoF Wujii ձեռքը, որը նախատեսված է մեջ ընդհանրապես չկար:

Մեթոդը: Այստեղ գրեթե ոչինչ նոր մեխանիզմ չէ: Հոսքի համընկեցում, DiT տեսանյութային հիմնական ցանցեր, գործողությունների առանձին փորձագետ, հոսքերի միջև խառը ինքնաուշադրություն, ասինխրոն սպասարկում՝ բոլորը ժառանգված են: Նորը հենց բաժանումն է. փոխկապակցված նախագծային ընտրությունները ներկայացնել որպես վերահսկվող փոփոխականներ և յուրաքանչյուր որոշում հաղորդել նախորդի ստեղծած պայմաններում: Սա է ներդրումը, և այն իրական է. WAM-ների մասին գրականությունը մինչ օրս բաղկացած էր միաձուլյ համակարգերից, որոնցում շահույթը հնարավոր չէր վերագրել որևէ բանի:

Կարևոր թիվը: Ոչ թե որևէ հենանիշի արդյունք: Դա նախատեսված է շահույթի բաժանումն է RoboTwin2.0-Clean2Random-ում. **+0.68 տ.կ. տիրույթի ներսում (87.0 $\rightarrow$ 87.68)՝ ընդդեմ +12.12 տ.կ.-ի տիրույթից դուրս (14.5 $\rightarrow$ 26.62)**: Մարմնավորված նախատեսումն այս ապացույցով գրեթե ամբողջությամբ OOD գնում է: Երեք օժանդակ արդյունք. միայն ռոբոտային նախատեսումը տալիս է ամենաուժեղ արդյունքը տիրույթի ներսում, մինչդեռ էգո+ռոբոտ խառնուրդները՝ լավագույն OOD արդյունքը. մեկ փուլով համատեղ ուսուցումը հավասարվում է երկու փուլով հաջորդական ուսուցմանը և վերացնում ուսումնական ծրագրի անցումը. և՛ այն արդյունքը, որի վրա հոդվածը կառուցում է իր բաղադրատոմսը, նախընտրելի տեղեկատվական հոսքը նախատեսման հետ շրջվում է. RoboTwin-Full-ում գրոյից ուսուցման դեպքում Action-Sees-Video-ն առավելություն ունի 0.24 տ.կ.-ով, իսկ նախատեսումից հետո Mutual-ը՝ 0.16 տ.կ.-ով, Clean2Random-ում Mutual-ն առավելություն ունի 0.70 տ.կ.-ով (ID) և 0.72 տ.կ.-ով (OOD): Հոդվածը նաև հաղորդում է, որ WAM-ները բաշխման ներսում գերազանցում են VLA-ներին, մինչդեռ VLA-ներն ավելի լավ են ընդհանրացնում բաշխումից դուրս, և այդ անհամաչափությունը վերագրում է փոփոխված պայմաններում երկարաժամկետ պիքսելային կանխատեսման սխալների կուտակմանը:

Թույլ կողմերը: Վերջնական բաղադրատոմսը որոշող շրջումը 0.16–0.72 տ.կ. է, 5B-ի և 14B-ի միջև հիմնական ցանցի ընտրությունը՝ 1.40 տ.կ., և **ամբողջ ուսումնասիրության մեջ որևէ տեղ սերմեր, սխալի գծեր կամ կրկնություններ չեն հաղորդվում**: Մեկ շաբաթ առաջ MINERVA-ն համեմատելի հենանիշում չափեց ուսուցման սերմերից կախված ± 1 կետի տիրույթ՝ ելակետային մոդելը երեք անգամ վերաուսուցանելով (94.60–96.75): OpenWAM-ի առաջ տարած յուրաքանչյուր որոշում այդ տիրույթից փոքր է: Սա չի նշանակում, որ բաղադրատոմսը սխալ է. նշանակում է, որ հոդվածը ցույց չի տվել, որ այն ճիշտ է, և հարցը կլուծի հենց նրա իսկ

հրապարակած ենթակառուցվածքը: Երկրորդ՝ ութ հենանիշի արդյունքները բոլորը SFT-ից հետո են՝ յուրաքանչյուր հենանիշի համար առանձին խմբաքանակի չափերով և քայլերի քանակով (Աղյուսակ 10), ուստի դրանք առանց լրացուցիչ ուսուցման փոխանցման թվեր չեն: Երրորդ՝ LIBERO-Plus-ի անոմալիան վերագրվում է մեկ մանիպուլյատորով նախաուսուցման ավելի նոսր խառնուրդին, բայց հոդվածն այլ տեղում պիքսելային վերակառուցման կողավորիչներն անվանում է տեսանկյան և աղմուկի փոփոխությունների նկատմամբ անբավարար կայուն. երկու բացատրություն, մեկ տվյալների հավաքածու, ոչ մի փորձ, որը դրանք կտարանջատի: Չորրորդ՝ իրական ռոբոտի ելակետերը միայն LingBot-VA-ն և $\pi 0.5$ -ն են: Հինգերորդ, և անխուսափելի՝ 128 NVIDIA H200 ~7 օր, ~21 500 GPU-ժամ՝ ընդդեմ չհայտարարված բյուջեներով ուսուցանված ելակետերի:

Ինչ արժե փոխառել: (1) Նախաուսուցման շահույթը հաղորդեք երկու թվով՝ տիրույթի ներսում և տիրույթից դուրս, երբեք՝ մեկով. այս տվյալների վրա միջինացված մեկ թիվը կթաքցներ ամբողջ արդյունքը: (2) 80-չափանի միասնական գործողությունների տարածությունը՝ ֆիքսված սլոտային իմաստաբանությամբ և վավերականության դիմակավորմամբ, տարասեռ մարմնավորումները համատեղ ուսուցանելու էժան եղանակ է՝ առանց յուրաքանչյուր մարմնավորման համար առանձին գլուխների, և հենց դա թույլ տվեց չտեսնված 21-DoF ձեռքին ընդհանրապես հարմարվել: (3) Չպիտակավորված էգոկենտրոն տեսանյութը՝ գործողությունների և պրոպրիոցեպցիայի ամբողջությամբ դիմակավորված ալիքներով, անվճար վերահսկում է աշխարհի հոսքը. եթե առաջադրանքի որևէ տեսանյութ ունեք, այս բաղադրատոմսով այն ուսուցման համար պիտանի տվյալ է: (4) Կիրառման բաժանումը՝ ուսուցանել ամբողջ համատեղ աղմուկի հարթության վրա, գործարկել համաժամեցված անկյունագծով և արագացնել CUDA graph-ով կրկնարտադրվող ֆիքսված ձևի համատեղ աղմուկագերծման ցիկլով, հարակից քայլերի միջև արագության քեշով, հուշման ներկայացման (prompt embedding) քեշով և կառավարման ուղում VAE ապակողավորման բացակայությամբ. 170 մվ մեկ փաթեթի համար RTX 5090-ի վրա, և հենց այս թիվն է սա կիրառելի դարձնում, ոչ թե H200-ների քանակը:

Փոքրացված տարբերակի վերարտադրման արժեքը: Նախաուսուցումն անհասանելի է, բայց *ուսումնասիրությունը*՝ ոչ, և հենց դա է փոխանցելի պնդումներով մասը: Վերցրեք 1.3B VACE հիմնական ցանցը 5B-ի փոխարեն, չորս դիմակների համեմատությունը կատարեք RoboTwin2.0-Clean2Random-ի վրա (հանրային է), և դուք լրաուսուցում (fine-tuning) եք անում, ոչ թե նախաուսուցում. α բաղադրատոմսն այդ բաժանման համար օգտագործում է 10 740 օպտիմիզատորի քայլ 256 խմբաքանակով: **ԵԶՐԱՀԱՆԳՈՒՄ** $\rightarrow$ 8 խմբաքանակով և կուտակումով մեկ 80 GB քարտի վրա մեկ բջիջը կպահանջի մի քանի հարյուր GPU-ժամ, ուստի սերմերով կրկնված չորս դիմակների շարքը մեկ GPU-ի մեկ շաբաթ է, ոչ թե կլաստեր: Ոչ նոր տվյալների հավաքում, ոչ նոր սարքավորում: Դա է [§10](#)-ի փորձը:

7 Ինչ է առնչվում ձեր նախագծին

CONTEXT.md-ի `## Active build` բաժինը դեռ չլրացված տեղապահ է՝ «(empty — pending interview)», ուստի չկան աշխատանքային ուղղություններ, որոնց վրա կարելի լիներ տեղադրել այս շաբաթվա նյութերը, իսկ դրանք հորինելը ավելի վատ կլիներ, քան բաց թողնելը: Բաժինը բաց է թողնված, մինչև տեղապահը չփոխարինվի:

8 Շուկայական ազդանշան

Տեսանելիորեն աշխատողներ են ընդունում կամ ընդլայնվում են՝ OpenAI (ռոբոտաշինության 19 հաստիք Սան Ֆրանցիսկոյում, չորսը՝ շարժաբեքերի վրա, գումարած շարժաբեքերի TPM, բոլորը՝ սարքավորման ոլորտում, ռազմավարությունների գծով ոչ մի հաստիք նշված չէ), Maven Robotics (250 միավորի արտադրություն Apple-ի փակված AV խմբից հավաքված թիմով), Skild AI (60+ հաճախորդ, հասույթի 86%-ը՝ մանիպուլյացիա), XPENG (ավտոմոբիլային մասշտաբի արտադրություն), Agility (SPAC-ից $> \$620$ մլն մուտքեր՝ տարեկան $\sim \$100$ մլն ծախսի դիմաց): Ուշագրավ բացակայությունը՝ այս շաբաթ ոչ ոք տեսանելիորեն չի ընդլայնում *ռազմավարությունների հետազոտման* թիմ: Փողը գնաց շարժաբեքերի, գործարանների և կիրառման ցիկլերի վրա:

Նշված մեթոդաբանությամբ վարձատրության կամ պաշտոնային մակարդակների տվյալներ կրկին չհայտնվեցին. ռոբոտաշինության աշխատավարձերի հարցումների համար վերադարձվող միակ աղբյուրները

մնում են SEO ազդեցատուները, և ավելի լավ է ոչինչ չհաղորդել, քան դրանք «լվանալ»: Պաշտոնային մակարդակների իրական ազդանշանին ամենամոտը Agility-ի S-4-ն է, որը հանրային հաշվետվություն է, ոչ թե հարցում՝ \$111 մլն գործառնական ծախս \$1.8 մլն հասույթի դիմաց, նախորդ տարվա \$71 մլն-ից աճած: Սա այն բյուջեի ձևն է, որի ներսում են գտնվում դեռ հասույթ չունեցող մարդանման ռոբոտների ընկերության ավագ պաշտոնները, և այն այժմ բացահայտված է, ոչ թե լուրերի մակարդակում:

Ուշադրության արժանացած աշխատանքներում շարունակում են կրկնվել հմտությունների երեք ուղղություններ: **Էժան ընթերցիչներ սառեցված հիմնական ցանցերի վրա՝** FARM-ն անփոփոխ VLA-JEPA կանխատեսողական վիճակից քայլային ձախողումների հայտնաբերում է քաղում 34 հազ. ուսուցանվող պարամետրով և միլիվայրկյանից պակաս ուշացումով. [No Free Checker](#)-ը նույն ընտանիքը ձևակերպում է որպես «մոդելին ներհատուկ ստուգիչներ»: Իմանալը, թե գոյություն ունեցող ստեկի որ մասում է արդեն ազդանշանը, այժմ նոր մոդելներ ուսուցանելուց տարբեր իրավասություն է: **Կատարման ժամանակացույցի ճարտարագիտություն՝** W36-ի շարունակությունը. IMLE-VLA-ի մեկ քայլով գլուխը, TacPAC-ի՝ քեշավորող և $m+d$ -ից ետ գրող ասինխրոն ուղղիչը, OpenWAM-ի համաժամեցված աղմուկազերծման անկյունագիծն ու գործարկման արագացման ստեկը: **Գնահատման նախագծումը որպես ներդրում՝** No Free Checker-ն ուսումնասիրում է ~150 ստուգիչ, պնդում է, որ բոլոր ընտանիքներում հասանելիության աճին զուգընթաց վստահելիությունը նվազում է, և առաջարկում է հաղորդման ենթակա ինը չափանիշ, այդ թվում՝ «փոխանցվող միջնորդավորված շահույթի» հարաբերակցությունը ($\Delta_{\text{panel}}/\Delta_{\text{proxy}}$, որը նույն ստուգակետերը գնահատում է ուսուցման ստուգիչով և մի խմբով, որը ռազմավարությունը երբեք չի տեսել). այն նաև անկեղծորեն նշում է, որ ռոբոտաշինության գրականության բոլոր համաձայնության ցուցանիշներն ու հետագա շահույթները չափվել են ոչ հակառակորդային թեկնածուների վրա, ուստի խոցելիությունն ըստ էության չափված չէ: [MEMOBench](#)-ը գտնում է, որ $\pi 0$ -ն հիշողության պահպանման մեջ 94.0% է, իսկ առաջադրանքում՝ 4.4%: [EgoGenEval](#)-ը 16 գեներատորում տարանջատում է տեսախցիկի շարժման հիմնավորումը տեսարանի պահպանումից՝ լավագույնը 0.662՝ ընդդեմ օրակույի 0.940-ի: Եվ [Northwestern-ի ակնարկը](#) պնդում է, որ բազմամատ ձեռքերի ճարակության պնդումներն այժմ ընդհանրապես համեմատելի չեն:

Պարտադիր նվազագույնն այժմ՝ կող և կշիռներ հողվածը ներկայացնելիս. ուշացումը՝ հաջողության կողքին. հենանիշի գործարկման միջավայր, որը մեկ ուրիշը կարող է գործարկել. և՛ նորովի, OpenWAM-ից հետո, ID/OOD բաժանում՝ միջինացված մեկ թվի փոխարեն: Դեռևս տարբերակող են՝ սերմերով կրկնված աբլյացիաները (դեռ գրեթե բացակայում են. OpenWAM-ը չունի), իրական սարքավորման վրա փակվող RL ցիկլերը (WHIRL, VLA-Precision), հրապարակված սարքավորման ֆայլերը և ձեր սեփական ստուգիչի հակառակորդային գնահատումը, որը, ըստ ակնարկի, ռոբոտաշինության մեջ ոչ ոք չի արել:

9 Թեմաներ

ԹԵՄԱ	ԿԱՐԳԱՎԻՃԱԿ	ՀԱՋՈՐԴ ԻՐԱԿԱՆ ԹԱՐՄԱՑՈՒՄԸ
Արդյո՞ք VLA-ի մասշտաբն է ապահովում լեզվով պայմանավորված կառավարումը	ԿԱՆԳՆԱԾ	Այս շաբաթ 10 մլն-ից փոքր մոդելների շարունակություն չկա: Դեռ՝ 10 մլն-ից փոքր ռազմավարություն իրական սարքավորման վրա, բազմառաջադրանքային, հրապարակված ելակետի համեմատ
Հենանիշների վավերականությունը ֆիզիկական ԱԲ-ում	ԱՌԱՋԸՆԹԱՑ (No Free Checker-ի ինը չափանիշ. MEMOBench՝ $\pi 0$ 94.0% պահպանում / 4.4% հաջողություն. EgoGenEval. IMLE-VLA՝ 98.0% LIBERO)	Խոշոր լաբորատորիա, որն աբլյացիոն աղյուսակում հաղորդում է սերմերի տիրույթներ, կամ որևէ խումբ, որը գնահատում է ստուգիչի խոցելիությունը որոնման պայմաններում
Հպման զգայում առանց առանձին շոշափողական սարքավորման	ՎԻՃԵԼԻ	Շաբաթը պաշտպանեց հակառակ կողմը. TacPAC-ը, DeCAL-ը, TacClip-ը և X-Hinges-ը բոլորն էլ սարքավորում են ավելացում: W35-ի VISTA-ն դեռ կող չունի

ԹԵՄԱ	ԿԱՐԳԱՎԻՃԱԿ	ՀԱԶՈՐԴ ԻՐԱԿԱՆ ԹԱՐՄԱՑՈՒՄԸ
Վճարել ուսուցման ժամանակ, ոչ թե աշխատանքի ժամանակ (միավորված)	ԱՌԱՋԸՆԹԱՏ (IMLE-VLA՝ 15→55 Հg, 11x թողունակություն. TacPAC՝ 30.4 մվ ընդդեմ 628.6 մվ-ի. OpenWAM՝ 170 մվ/փաթեթ 5090-ի վրա. թաքնված իմաստային կմախք)	Ուշացումն ու ID/OOD բաժանումը՝ երկուսն էլ որպես լռելյայն սյունակներ
Առցանց ուղղումը որպես վերջին քայլ	ԱՌԱՋԸՆԹԱՏ (WHIRL՝ +15-30 տ.կ. և օպերատորի ժամանակի 84% կրճատում 16-DoF LEAP ձեռքի վրա. երրորդ խումբը՝ այլ սարքավորման վրա)	Սերմերով կամ օպերատորներով կրկնություն. մինչ այժմ բոլոր արդյունքները մեկ սերմ են, մեկ օպերատոր, մեկ սարք
Աշխարհ-գործողություն մոդելները որպես միասնական ռազմավարություն + սիմուլյատոր	ԱՌԱՋԸՆԹԱՏ (OpenWAM-ը հրապարակում է ենթակառուցվածքը, կշիռները և տվյալների բաղադրատոմսերը. FARM-ը ձախողումն ապակողավորում է սառնեցված WAM վիճակից)	Հրապարակված աշխարհի մոդել, որն օգտագործվում է ռազմավարությունը բարելավելու համար մեկի կողմից, ով այն չի ստեղծել. այժմ առաջին անգամ հնարավոր է
Ում է պատկանում ռոբոտների ուսուցման բաց ստեկը	ԱՌԱՋԸՆԹԱՏ (LeRobot v0.6.1՝ գործարքից հետո առաջին թողարկումը)	Այն համախմբվեց HF Transformers-ի վերին հոսքի դասերի, ոչ թե NVIDIA-ի ստեկերի վրա: Հետևեք տվյալների ձևաչափի կառավարմանը և արդյոք GR00T-ը արտոնյալ վերաբերմունք կստանա
Սիմուլյացիայի ճշգրտությունը փոխանցման ու հպման համար	ՆՈՐ (MuJoCo 3.13.0-ի discrete ինտեգրատոր. MuJoCable-ի Capstan ուղղորդում. Aero Hand-ի ստուգված ջլային մոդելը W36-ում)	Ջլային ձեռքի սիմուլյացիայից իրական աշխարհ արդյունք, որը նշում է նոր ինտեգրատորի դերը, և MuJoCable-ի պահոց
Էժան բաց սարքավորումը որպես ուսուցման հենք	ԱՌԱՋԸՆԹԱՏ (X-Hinges՝ գրախոսված ինքնագզայող ճկուն հոդեր. TacClip՝ ~\$20/մատ սենսորային գլխիկ. ORCA ձեռքը՝ շաբաթվա լավագույն ճարակության արդյունքով)	Դեռևս Aero Hand-ի, BRIDGE-ի կամ Peg-in-Bench-ի որևէ երրորդ կողմի հավաքում չկա
Հայտարարված բաց թողարկումներն ընդդեմ իրական արտեֆակտների	ՎԻՃԵԼԻ , ավելի սուր	Կողմ՝ OpenWAM-ը հրապարակեց ամեն ինչ. FARM, IMLE-VLA, HuRo, գրիչով գրելու կողը: Դեմ՝ Xiaomi U0-ի քարտը հետաձգում է ուսուցման կողը. Unitree-ի գործողությունների մոդելը սպասվում է. MuJoCable-ը պահոց չի նշում. StreamPI-ի կշիռները W35-ից ի վեր երեք շաբաթ է՝ սպասվում են
Մարդանման ռոբոտների կապիտալն ընդդեմ ցուցադրված կարողության	ԱՌԱՋԸՆԹԱՏ (Agility-ի աուդիտված \$1.8 մլն/\$140 մլն՝ ընդդեմ \$2.5 մլրդ SPAC-ի. Unitree՝ -53%. XPENG-ի զանգվածային արտադրություն. OpenAI-ի մուտքը)	Մարդանման ռոբոտներ արտադրող ընկերություններից դեռ զրո մեթոդ: Unitree-ի դատողական մոդելի կշիռներն առայժմ ամենամոտն են. գործողությունների մոդելի հրապարակումը կփոխեր սա

Մանիպուլյացիան գնվում է, ոչ թե կառուցվում

ԿԱՆԳԵԱԾ

W35-ի՝ եռամսյակի ընթացքում իրականանալու կանխատեսմանը մնացել է ~6 շաբաթ: Այս շաբաթվա M&A-ը կրկին պաշտպանական UGV և վիրաբուժական էր: Եթե ժամկետն անցնի, փակել թեման

Այս շաբաթ հանված է՝ «Ուսուցման ժամանակ վերահսկողություն՝ գործարկման զրո ծախսով» թեման, որը միավորվել է «Վճարել ուսուցման ժամանակ, ոչ թե աշխատանքի ժամանակ» թեմային, որն ընդգրկում է նույն սկզբունքը թե օժանդակ կորստի, թե ուշացման կողմից:

10 Բաց հարց

Արդյո՞ք աշխարհ-գործողություն մոդելների նախագծման տարածությունը լուծելի է այն ճշգրտությամբ, որով ոլորտն այժմ չափում է այն: OpenWAM-ի առաջ տարած բաղադրատոմսը հենվում է ուշադրության դիմակի համար 0.16–0.72 տ.կ. և հիմնական ցանցի համար 1.40 տ.կ. տարբերությունների վրա՝ առանց հաղորդված սերմերի, և դա այն բանից մեկ շաբաթ անց, երբ MINERVA-ն համեմատելի հենանիշում չափեց ուսուցման սերմերից կախված ± 1 կետի տիրույթ՝ պարզապես երեք անգամ վերաուսուցանելով: Կամ WAM-ների նախագծման տարածությունն իսկապես ավելի նուրբ ազդանշան ունի, քան ռազմավարությունների մասշտաբինը, կամ լայնորեն հրապարակված բաղադրատոմսը քամվել է աղմուկից, և հրապարակվածից ես չեմ կարող ասել՝ որն է ճիշտ:

Ինչը կլուծեր հարցը՝ դիմակների երեք բջիջը վերաուսուցանել երեքական սերմով RoboTwin2.0-Clean2Random-ի վրա և հրապարակել տարածվածությունը: OpenWAM-ը հրապարակել է ենթակառուցվածքը, գնահատման արձանագրությունը և տվյալների բաղադրատոմսերը հենց դա անելու համար, և RoboTwin2.0-ն հանրային է. ուստի սա արժե լրաուսուցման մի քանի գործարկում մեկ GPU-ի վրա և ոչինչ կառուցել չի պահանջում: Սա ռոբոտների ուսուցման մեջ այժմ հասանելի ամենաեժան և բարձր արժեքով փորձն է, և այն հետադարձ ուժով կիրառելի է այս թողարկման աբլյացիոն աղյուսակների մեծ մասի նկատմամբ:

Աղբյուրները նշված են տեքստում: arXiv-ի նյութերը պրեպրինտներ են և չգրախոսված, բացի X-Hinges-ից, որն ընդունված է UIST '26-ում: Մեկ լաբորատորիայի պնդումները նշված են, որտեղ դա տեղի է: Ամբողջական տեքստի փոխարեն ամփոփումներից գնահատված նյութերը նկարագրված են առանց թվերի:

ԱՅՍ ՇԱԲԱԹՎԱ ԱՅԼ ՆՅՈՒԹԵՐ

Հերթում 4+ գնահատական ստացած բոլոր նյութերը, որոնք վերևում առանձին բաժին չստացան՝

- **SyncWorld** — չափաբերման դրվագը համատեքստում սահմանում է գործողություն-պատկեր համապատասխանությունը՝ աշխարհի մոդելը վերածելով առանց լրացուցիչ ուսուցման աշխատող սիմուլյատորի՝ թեստի ժամանակ ռազմավարությունը բարելավելու համար: Բաց թողնված է՝ կող նշված չէ, և շաբաթն արդեն ունի OpenWAM-ը նույն առանցքի վրա՝ հրապարակված ստեղծով:
- **Show-Harness** — դիսկրետ իմաստային գործողությունները թույլ են տալիս առաջատար VLM-ներին ուղղակիորեն կառավարել ռոբոտներ՝ ուղեկցող սարքով, որը ցուցադրումներ է հավաքում առանց հեռակառավարման սարքավորման: Բաց թողնված է՝ նախագծի էջ, թվերը չստուգված:
- **VLA-Precision** — հարաբերական առավելությամբ առցանց RL՝ սառեցված նախաձեռնի KV-ի վերօգտագործմամբ, 98.3% ինը քիմիական առաջադրանքում: Բաց թողնված է տեղի սղության պատճառով՝ հեղինակների սեփական առաջադրանքների խումբ, կող չկա:
- **ICI-VLA** — DTW-ով որոնված միկրոցուցադրումները թեստի ժամանակ հարմարեցնում են սառեցված տեքստ-գործողություն VLA-ն. RoboTwin՝ 60.4%, իրական աշխարհ՝ 83.2%՝ ~1 000 փորձում: Բաց թողնված է՝ նույն համատեքստային տեղն է զբաղեցնում, ինչ ContextFlow-ը, 8×A100 օֆլայն ուսուցմամբ և առանց նշված կողի:

- [EgoGenEval](#) — էգոկենտրոն տեսանյութային աշխարհի մոդելների համար տարանջատում է տեսախցիկի շարժման հիմնավորումը տեսարանի պահպանումից. 16 առանց դիրքի գեներատոր, լավագույնը 0.662՝ ընդդեմ օրակուլի 0.940-ի, մարդկանց կողմից վավերացված $\rho=0.943$ -ով: Ներկայացված է [§8](#)-ում.
հրապարակումը չստուգված է: