

This Week in Robotics

Ժամանակահատված՝ 2026-08-31 → 2026-09-06:

Մեքենայական թարգմանություն անգլերենից, չիմբագրված · բնօրինակը՝ twir.arcanerobotics.com/2026-w36/

ԳԼԽԱՎՈՐԸ

NVIDIA-ն համաձայնեց գնել Hugging Face-ը \$12.9 մլրդ-ով, և LeRobot-ը պահպանողն այժմ հայտնվում է Isaac GR00T-ը թողարկող ընկերության ներսում: Նույն շաբաթ 0.54 մլն պարամետրով ռազմավարությունը (policy)՝ առանց լեզվական կոդավորիչի (encoder) և առանց նախաուսուցված տեսողության, LIBERO-ում հասավ 95.1%-ի, ինչն ավելին է ասում հենանիշի (benchmark), քան ռազմավարության մասին:

1 Ռազմավարությունների ուսուցում և VLA-ներ

MINERVA

ՊՐԵՊՐԻԵՏ / U. Tokyo

Զրոյից ուսուցանվող CNN-ը և հոսքի համընկեցմամբ (flow matching) գործողությունների փաթեթի (action chunk) գլուխը *փոքրացնում* է այնքան, մինչև LIBERO-ն «կոտրվի»: Ճարտարապետությունը հանման վրա է կառուցված. լեզուն դառնում է առաջադրանքների 40 տողով ID աղյուսակ, գործողությունների թոքենների ինքնաուշադրությունը (self-attention) դառնում է MLP խառնիչ, և ազատված 26% պարամետրերը տրվում են կոդավորիչին: [Կոդ, բաղադրատոմսեր, բոլոր ստուգակետերը](#), Apache-2.0: Խորը ընթերցում՝ [\\$6-ում](#):

Knowing When to Stop

ՊՐԵՊՐԻԵՏ / Tsinghua

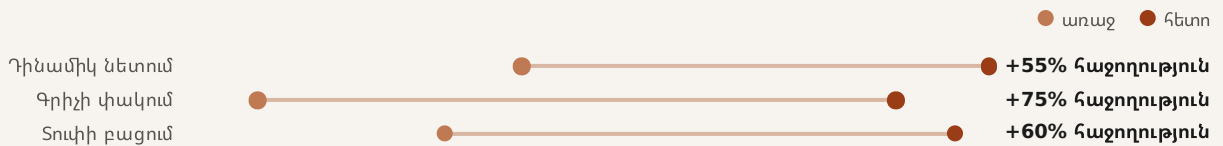
Առանց ուսուցման ադապտիվ փաթեթավորում (chunking). գործողությունների հարցումներից (queries) դեպի VLM թոքենները խաչաձև ուշադրությունը ցրվում է գործողության ինդեքսի աճին զուգընթաց և հազեճում իր log-N սահմանի մոտ 95%-ի վրա, ուստի փաթեթը կտրվում է այնտեղ, որտեղ էնտրոպիան և՛ բարձր է, և՛ հարթ: Մեխանիզմի ստուգումն արդարացնում է մոտեցումը. էնտրոպիան հետևում է չտեսնված տվյալների վրա գործողությունների MSE-ին Spearman $\rho = 0.432$ համահարաբերակցությամբ, և ինդեքսի ազդեցությունը հանելուց հետո մասնակի համահարաբերակցությունը դեռ 0.214 է: LIBERO՝ 94.88 → 97.25 լավագույն ֆիքսված հորիզոնի համեմատ, իրական երկձեռ ռոբոտ՝ 47.5 → 61.7 ($n=20$ /առաջադրանք), բայց RoboTwin-ում $\pi 0.5$ -ը՝ ընդամենը 61.6 → 62.9, առանց ստանդարտ շեղումների, ինչը սերմերի (seeds) հավանական աղմուկի սահմաններում է: Իրական պնդումն այն է, որ մեթոդը համընկնում է հետագայում կարգավորված հորիզոնին՝ առանց այն իմանալու, *ավելի երկար* միջին հորիզոնով (18.1 ընդդեմ 10-ի) և 7 մվ-ից պակաս լրացուցիչ ծախսով: Կոդ նշված չէ:

SmoothRL

ՊՐԵՊՐԻԵՏ / Astribot

Ասինխրոն գործարկման (inference) դեպքում յուրաքանչյուր փաթեթ կատարվում է միայն մասամբ, ուստի ∇_{θ} -ն կտրվում է մինչև կատարման տիրույթը՝ $[n, 2n)$, մինչդեռ քննադատը (critic) կատարման համար արդեն հանձնված նախածանցը պահում է որպես վիճակի լրացում զուգահեռ որոշումների գործընթացի համար: Սառեցված հիմնական մոդել՝ զումարած TD3 ռճի կցված ցանց *չխախտված* գործողությունների տարածությունում, ինչը թույլ է տալիս մարդու միջամտություններն ուղղակիորեն օգտագործել որպես ռեգրեսիայի թիրախներ՝ առանց թաքնված տարածության շրջման: Յուրաքանչյուր առաջադրանքի համար իրական ռոբոտի վրա 250 դրվագ (episodes) դինամիկ նետումը 39%-ից հասցնում է 94%-ի, գրիչի փակումը՝ 8%-ից 83%-ի, իսկ տուփի բացումը 90%-ի հասնում է միայն 150 դրվագում 20%-ի անկումից հետո, ինչը սառեցված հիմնական մոդելից ցածր է: Կոդ չկա:

SmoothRL. սառեցված հիմնական ռազմավարություն → իրական ռոբոտի վրա 250 փորձարկման դրվագ



Մեկ դրվագ՝ սկզբնական վիճակի յուրաքանչյուր կոնֆիգուրացիայի համար, մեկ ռոբոտ, առանց սերմերի կրկնությունների, հեղինակների սեփական արձանագրությամբ: Տուփի բացումը 150 դրվագում իջավ 20%-ի՝ սառեցված հիմնական մոդելից ցածր, նախքան վերականգնվելը: arXiv 2608.29768, պրեպրինտ, չվերարտադրված, կող չի հրապարակվել

ZETA

ՊՐԵՊՐԻՆՏ / Galbot + PKU

Վերահսկվող աբլյացիաներ (ablations) տարբեր մարմնավորումների միջև $\pi 0.5$ դասի հիմնական ցանցի (backbone) վրա՝ 14 չտեսնված մարմնավորում (embodiment), յուրաքանչյուր մոդելի համար 6 300 սիմուլյացիոն փորձարկում (rollouts)՝ 3 կրկնությամբ: EEF-delta վիճակն m գործողությունները գերազանցում են բացարձակ EEF/աշխարհային delta-ին՝ 75.7% ընդդեմ 60.3%-ի սիմուլյացիայում և 89.9% ընդդեմ 56.0%-ի իրական աշխարհում, հիմնականում միայն մանիպուլյատորի և ամբողջ մարմնավորման փոփոխությունների դեպքում: Նախատեսված մեջ թիրախային մարմնավորման 5% տվյալն արժե 13.4 տ.կ., և հենց դրա համար են նրանք պնդում, որ «խիստ» և «նախատեսված ընթացքում տեսած» zero-shot (առանց նախնական օրինակների) արդյունքները տարբեր պնդումներ են: Կող նշված չէ:

XR-2

ՊՐԵՊՐԻՆՏ / PrimeBot + PKU

Հրապարակում է տնային երկձեռ աշխատանքների 1 500 ժամ տվյալ՝ 531.7 ժ իրական ռոբոտի հեռակառավարում (teleoperation) 32 518 հետագծով 25-DoF շարժական երկձեռ հարթակի վրա, գումարած մոտ 1 000 ժ UMI հավաքագրում իրական միջավայրում՝ 200+ տնային տնտեսությունում և 5 000 հազուստի վրա, և դրա վրա ուսուցանում է 5B VLA: Արժեքավոր թիվը հագուստ ծալելու մասշտաբավորման կորի ձևն է՝ 34%՝ 30 ժ-ի դեպքում, 84%՝ 120 ժ-ի դեպքում, այնուհետև հարթ (82%՝ 160 ժ-ի դեպքում): Զախողումներով կշռված բյուջեով DAGger-ի երեք փուլը 58%-անոց ստուգակետը (checkpoint) հասցնում են 74%-ի, 82%-ի և 93%-ի: Փորձագետի տվյալները ծածկում են փորձագետի սեփական բաշխումը: միայն սեփական ռազմավարությամբ (on-policy) ուղղումներն են հասնում այն վիճակներին, որոնց հանդիպում է կիրառված ռազմավարությունը: Սարքավորումը, մոդելը և խողովակաշարը փակ են, ուստի փոխանցվողը կորն է, ոչ թե արդյունքը:

Արժե նշել, չորս անկախ հոդված ավելացնում են վերահսկողություն (supervision), որը **գործարկման ժամանակ անվճար է կամ դեռ է նետվում՝ PHR-VLA, Temporal Forcing, GIFT** (+4.6-ից +12.6 տ.կ.), **EGR** (կողը հրապարակված է): Այլևս ոչ ոք չի ուզում իր ինդուկտիվ կողմնակալության համար վճարել աշխատանքի ժամանակ:

2 Մանիպուլյացիա, ճարպկություն և զգայուն

Aero Hand Open

ՊՐԵՊՐԻԵՏ / Chestnut Robotics

Տասնվեց հոդ, յոթ շարժիչ, 374 գ, **\$314 նյութերի ցուցակ**, ամբողջությամբ 3D տպված, և մեկ սարքավորումային կոնֆիգուրացիայով ընդգրկում է GRASP դասակարգման բոլոր 33 բռնումները: Ներդրումն այն է, որ *մալուխային փոխանցումը* մոդելավորված է MuJoCo-ում. յուրաքանչյուր մալուխ և հետադարձ զսպանակ տարածական ջիլ (tendon) է, որը փաթաթվում է CAD-ից ստացված ճախարակների գլաններին, իսկ փաթաթման կողմերը ֆիքսված են, որպեսզի մոմենտի բազուկները բխեն երկրաչափությունից, այլ ոչ թե հարմարեցվեն: Ուստի համապատասխանությունը կանխատեսում է, ոչ թե հարմարեցում. մալուխի տեղաշարժի սխալը չորս մասներում 0.04–1.40 մմ է, հետևման RMS սխալը՝ 0.29–0.45 մմ: Բութ մատը թույլ օղակն է՝ տեղաշարժի 31–33% անհամապատասխանությամբ, քանի որ նրա հեռացման–ծալման կապը մոդելավորված է միայն առաջին կարգի ճշտությամբ: Ձեռքի ներսում խորանարդի պտտումը, որն ուսուցանվել է միայն սիմուլյացիայում, առանց լրացուցիչ ուսուցման աշխատում է միայն յոթ շարժիչների էնկոդերներով: Հրապարակված են նախագծային ֆայլերը, ներկառուցված ծրագիրը (firmware), մոդելը, RL փաթեթը և տեղակայման հանգույցը:

N0-Foundation

ՊՐԵՊՐԻԵՏ / NeoteAI + Fudan

Շոշափողական (tactile) UMI ձեռքի սարք, 30 000+ ժամ համաժամեցված տեսողական-շոշափողական ցուցադրումներ և **OpenNeoData՝ 5 000 ժամանոց բաց ենթաբազմություն**՝ վեց մարմնավորումով և 250+ առաջադրանքով: Խաղաղությանն այն է, որ շոշափողական սենսորների միջև սարքավորումից անկախ միջերեսը ուժի խիտ եռառանցք դաշտերն են, ոչ թե սենսորների չմշակված պատկերները: Ազնիվ մասը նրանց սեփական սիմուլյացիոն առաջադրանքների խումբն է. $\pi 0.5$ -ն առաջատար է՝ տասներկու առաջադրանքում 45.8% միջինով, իսկ տևական հպում պահանջող առաջադրանքները գրեթե չլուծված են մնում (ատամնանիվի տեղադրում՝ 18%):

ChainSplat

ՊՐԵՊՐԻԵՏ / KTH, ԿՈԴ

Մալուխը մոդելավորում է որպես պտուտակների տեսությամբ նկարագրված ութ հոդով պտտական շղթա՝ յուրաքանչյուր օղակին կցված գաուսյաններով, և միայն բազմատեսանկյուն RGB-ից վերականգնում է երկրաչափությունը, դինամիկան և *արտաքին ուժը*՝ առանց խորության, առանց ուժի/մոմենտի սենսորի, առանց մասնիկների բազմության. հենց հոդերի տարածությունում կոմպակտ վիճակն է հետագծի օպտիմալացումն ու իրական ժամանակում վիճակի գնահատումը մատչելի դարձնում: Թույլ կողմերը՝ երեք պարան, 24 հինգվայրկյանային հետագիծ, յուրաքանչյուր հետագծի համար առանձին հարմարեցված մոդել, որը գնահատվում է հենց այն հետագծի վրա, որին հարմարեցվել է: Համակարգի նույնականացում, ոչ թե ընդհանրացում:

Facet-0

ՊՐԵՊՐԻԵՏ

Ուժը դարձնում է *գեներացվող* մեծություն, ոչ թե վերահսկվող. հոսքի համընկեցումը միասին արտադրում է գործողությունների փաթեթը և դաստակի կանխատեսված ուժ-մոմենտը (wrench), իսկ գործողություն–ուժ-մոմենտի բաշխումային քննադատը գնահատում է հպման արդյունքները RL հետուսուցման համար: Հպումներով հագեցած հավաքման մեջ հաղորդված տարբերությունը մեծ է՝ 82% ընդդեմ 15%-ի, բայց ապացույցը նախագծի էջ է՝ առանց կշիռների և տվյալների հավաքածուի հրապարակման, և արդյունքը հենվում է 1 000 ժամանոց փակ, ուժի հետ համաժամեցված կորպուսի վրա: Մեկ լաբորատորիայի չվերարտադրված պնդում:

Հպման կառավարման երեք արդյունք, որոնց մասին արժե իմանալ, բոլորը պրեպրինտներ, ոչ մեկում կող նշված չէ. **պտտող մոմենտի տարածությունում MPPI՝** կոշտ մարմնի բացահայտ դինամիկայով և լուծիչի 166 Հց հաճախականությամբ՝ 7-DoF ճկուն ուժային կառավարման համար. **վարքի և իրականացման բաժանում** ֆիզիկական HRI-ի համար, որը հպման կետի ցանկալի արագացումն առանձնացնում է դրա QP իրականացումից՝ շարժվող հորիզոնով, և աշխատանքային տարածքում գերազանցումը նվազում է սանտիմետրերից մինչև միլիմետրից պակաս Franka FR3-ի վրա. և **պասիվ անընդհատ փոփոխվող փոխանցումով մատ**, որի շարժական

ճախարակով մալուխների ուղղորդումը տալիս է ուժի $3.6\times$ միջին ուժեղացում՝ անփոփոխ հեռացմամբ. CAD ֆայլեր չեն հրապարակվել:

Առանց շոշափողական սարքավորման հպման զգայուն առաջ գնաց ևս երկու անգամ, երկուսն էլ պրեպրինտներ. [Blind Dexterity](#)-ն G1-ի ամբողջ մարմնով մանիպուլյացիան կառավարում է միայն հոդերի էնկոդերներով, իսկ [ամբողջ թևով բաշխված փոխազդեցությունը](#) իրական բռնման 71% հաջողության է հասնում փափուկ մանիպուլյատորի չորս IMU-ով:

3 Հարթակներ և մարմնավորում

BRIDGE

ՊՐԵՊՐԻՆՏ / CMU + HUST + JoyIn

Միաժամանակ նախագծում է մորֆոլոգիան և կառավարիչը՝ շարժաբերերը (actuators) հետազայում ընտրելու փոխարեն. թեկնածու տոպոլոգիաները նախ զտվում են կինեմատիկ վերաթիրախավորման (retargeting) սխալով, ապա ուսուցանված ռազմավարության ներքո փակ շղթայով հետևման սխալով, իսկ ուսուցման ընթացքում կիրառվում են չափաբերված պտտող մոմենտ–արագություն սահմանները: 88 սմ, 21-DoF մարդանման ռոբոտ **\$1.5 հազ.** արժեքով՝ բաց նախագծով և բաց ռազմավարությամբ, ընդդեմ Bumi-ի, K1-ի և ToddlerBot-ի \$3 հազ./\$5.7 հազ./\$6 հազ. գների, որոնցից ոչ մեկը երկուսն էլ բաց չի դնում: Հետևման հաջողություն՝ 94.83% ընդդեմ 91.87/92.66/88.23-ի MuJoCo-ում՝ միավորված LaFAN1 հենանիշի վրա. հետադարձ սալտոյի և պարի տեսահոլովակները միայն տեսանյութային ապացույց են:

[WM-LOCO](#)-ն (պրեպրինտ) հաղորդում է իրական G1-ի 93.3% հաջողություն ոտքի հենման սահմանափակ կետերով տեղանքում, մեկ լաբորատորիա, կող նշված չէ, չվերարտադրված: Այս ժամանակահատվածում մարդանման ռոբոտներ արտադրող ոչ մի [ընկերություն](#) մեթոդ չի հրապարակել. Galbot-ի [ET1-ի նախնական պատվերը](#) մամլո հաղորդագրություն է՝ առանց գնահատման:

Եվս չորս պրեպրինտ ոտքավոր ռոբոտների մասին, ոչ մեկում կող նշված չէ: [ADAPT](#)-ը տեքստով պիտակավորված հետազոտի վրա դիֆուզիոն գործողությունների նախնական բաշխում է հարմարեցնում և այն ճշգրտում RL-ով՝ տեքստով ուղղորդվող ամբողջ մարմնի փակ շղթայով կառավարման համար: [Agile traversal of sparse 3D structures](#) աշխատանքը մարզական սանդուղքներով (monkey bars) ստիպում է ընկալել բարակ, վերնից կախված երկրաչափություն. դա ավելի բարդ ընկալման խնդիր է, քան տեղանքը, քանի որ այն, ինչ չպետք է բաց թողնել, հիմնականում դատարկ տարածություն է: [Stay Seated](#)-ը մարդանման ռոբոտին նստեցնում է պասիվ անվավոր աթոռին, որպեսզի շարժաբերերն այլևս չծախսվեն քաշը պահելու վրա: Եվ [մարդանման ռոբոտի անվտանգ կանգառը](#) վթարային կանգառը ձևակերպում է որպես կանգ առնելու ուսուցված արժեք, այլ ոչ թե մեկ ֆիքսված մանևր, որը կատարվում է անկախ իրազործելիությունից. չորսից սա ամենից շատ է առնչվում իրական կիրառմանը: Առանձին՝ [ոչ գծային մարմնական տատանումների](#) աշխատանքը ճկուն քառոտանի ռոբոտի մոտ գտնում է վեց նորմալ մոդ, որոնցից յուրաքանչյուրը վերածվում է առանձին քայլվածքի. ստուգված է սարքավորման վրա և հիշեցնում է, որ մորֆոլոգիան դեռ կառավարման աշխատանք է կատարում:

4 Բաց կող և գործիքներ

Ըստ վերացված խոշորոտների: **MINERVA**՝ Apache-2.0, չորս ստուգակետ, ~2 ժ մեկ մոդելի ուսուցման համար մեկ սպառողական GPU-ի վրա, ~55 ր չորս խմբերի ամբողջական գնահատման համար: **Aero Hand Open**՝ \$314 արժողությամբ ճարպիկ ձեռք, որի MuJoCo մոդելը տարածվում է Menagerie-ի միջոցով, ~13 ժ տպագրություն: **Peg-in-Bench** (պրեպրինտ, AIST)՝ [STL ֆայլեր և սցենարների գեներատոր](#). ցցիկի հինգ ձև $\times$ երեք թույլատրելի շեղում (0.1/1/3 մմ), մագնիսով ամրացվող անցքով մասեր՝ ութ կողմնորոշմամբ, մեքենայական ընթերցման համար առաջադրանքների JSON: Վերարտադրելի, հպումներով հագեցած ֆիզիկական գնահատում խեժի գնով. ելակետային արդյունքներ դեռ չկան: **SolarWM**՝ տվյալների շարժիչ, չափողական անոտացիաներ 1.43 մլն տեսահատվածի համար, կող և կշիռներ:

Տվյալների շերտն ավելի շատ շարժվեց, քան գործիքների շերտը: [TacVerse-ը հրապարակեց ձեռքի UMI-ով հավաքված երկձեռ տեսողական-շոշափողական հավաքածուներ](#) LeRobot v3.0 ձևաչափով՝ ուղղված

շոշափողական տեսահոսքերով, CC-BY-SA: [RoboMIND-ի Agilex երկձեռն բաժինը վերամշակվեց v3.0-ի՝ 8 638 դրվագ, 5.4 մլն կադր և 54 առաջադրանք: Apple-ի EgoDex-ը փոխարկվեց v3.0-ի՝ train/test ենթապանակներով. ուրիշի փոխարկման մի շաբաթ, որն այլևս պետք չէ ծախսել: MIT Media Lab-ը հրապարակեց շոշափողական տվյալների և ձեռքի դիրքի կոնտրաստիվ կողավորիչ՝ երեք սերմով, որոնման mAP-ով, ուսուցման կողով և նախապես մշակված ճշգրիտ քեշով. հրապարակման ամբողջականության մի մակարդակ, որին գրեթե ոչ ոք չի հասնում: Այժմ կա \[v3.0 → v2.1 հետադարձման սկրիպտ\]\(#\) այն openpi ստեղծիչի համար, որոնք դեռ չեն անցել նոր ձևաչափին:](#)

Այս ժամանակահատվածում ոչինչ չի եղել LeRobot-ից (դեռ v0.5.1, ապրիլ), MuJoCo/MJX-ից, Isaac Lab-ից, ManiSkill-ից, Genesis-ից, RoboCasa-ից կամ ROS 2-ից: [OpenWAM-ը](#) հրապարակեց աշխարհ-գործողություն մոդելների ~35 ստուգակետ՝ առանց քարտի կամ լիցենզիայի. որպես ելակետ անօդաչուների են, քանի դեռ ինչ-որ մեկը չի ասել, թե ինչ են դրանք:

5 Փող և գործարքներ

NVIDIA → Hugging Face, \$12.9 մլրդ

, 2026-09-03 ([ՈՒՈՐՏԱՅԻՆ ՄԱՄՈՒԼ](#)): NVIDIA-ն հայտարարում է, որ հարթակը բաց կմնա, և նրա հաշվողական ռեսուրսները պարտադիր չեն լինի: [ԵԶՐԱՀԱՆԳՈՒՄ →](#) LeRobot-ը պահպանողը՝ մեկ մանիպուլյատորով նախագծերի համար լռելյայն ստեղծ, այժմ գտնվում է Isaac GR00T-ի մատակարարի ներսում, իսկ մամուլ հաղորդագրության մեջ բաց հարթակի խոստումը հայտարարություն է, ոչ թե կառավարման կառուցվածք:

Figure ↔ Nscale, \$3.5 մլրդ հաշվողական ռեսուրս

' \$6 մլրդ-ից ավել աճի հեռանկարով, մինչև 100 000 Vera Rubin GPU, Nscale-ը ստանում է բաժնեմաս ([ՈՒՈՐՏԱՅԻՆ ՄԱՄՈՒԼ](#)): [ԵԶՐԱՀԱՆԳՈՒՄ →](#) մարդանման ռոբոտների կապիտալը գնում է *ուսուցման* հզորություններ, ոչ թե ռոբոտների պարկեր, իսկ հաշվողական ռեսուրսի դիմաց բաժնեմասը նշանակում է, որ մատակարարը նույնպես երաշխավորում է խաղաղությունը:

Lyte, \$165 մլն Series C՝ \$1.6 մլրդ գնահատմամբ

' 4D կոհերենտ տեսողության չիպերի համար ([ՈՒՈՐՏԱՅԻՆ ՄԱՄՈՒԼ](#)): [ԵԶՐԱՀԱՆԳՈՒՄ →](#) ընկալման սարքավորումները կրկին ֆինանսավորվող կատեգորիա են՝ այն ենթադրությամբ, որ հիմնական սահմանափակումը խորության որակն է, ոչ թե ռազմավարության ճարտարապետությունը: Նաև՝ [Medtronic-ը \\$700 մլն ներդրեց Cornerstone Robotics-ում](#) (վիրաբուժական, առանց ռոբոտների ուսուցման բովանդակության). [Wandercraft-ը փնտրում է €100 մլն՝ €750 մլն գնահատմամբ](#), գործարքը չի փակվել:

Գլխավոր թվերից ներքև տվյալների շերտն առևտրայնանում է: [Visko-ն փակեց \\$10 մլն pre-seed փուլ](#) հոսքային աշխարհի մոդելի համար. [Kinetic Blocks-ը բացեց փակ բետա շուկա](#), որտեղ գնակարգվում են հեռակառավարման և ռոբոտի կատարման տվյալների հավաքածուները՝ LeRobot v3.0 աջակցությամբ և A+-ից F որակի գնահատականներով. [X Square-ը ներկայացրեց TwinDEX-ը՝](#) իզոմորֆ կրելի սարք, որի մասին պնդվում է հավաքագրման 5.3× թողունակություն և քիմիական 24 քայլանոց առաջադրանքի յուրացում՝ առանց իրական ռոբոտի հեռակառավարման. մատակարարի պնդում, ոչինչ չի հրապարակվել: [ԵԶՐԱՀԱՆԳՈՒՄ →](#) երբ ցուցադրումները գնակարգվում են որակի գնահատականներով, սակավ ակտիվ դադարում է դրվագների ծավալը լինել և դառնում է գնահատման չափանիշը:

6 Շաբաթվա խորը ընթերցումը

Այս շաբաթ՝ մեկը: Մյուս թեկնածուներն էին Aero Hand Open-ը և ադապտիվ փաթեթավորման հոդվածը, որոնք ներկայացված են վերևում. դրանցից ոչ մեկն ինքնին չի արժանանում այս խորությանը. ձեռքի ներդրումը հրապարակված արտեֆակտ է, ոչ թե հետաքննելու արդյունք, իսկ փաթեթավորման շահույթներն իր իսկ հիմնական հենանիշում սերմերի չհաղորդված աղմուկի սահմաններում են:

MINERVA — որքան փոքր ռազմավարություն է իրականում պահանջում ստանդարտ LIBERO-ն:

[arXiv:2609.03715](https://arxiv.org/abs/2609.03715), արեպրինտ, Ս. Tokyo: Չվերարտադրված. կողը և ստուգակետերը հրապարակված են:

Փորձի կառուցվածքը: LIBERO-ի բոլոր չորս խմբերը (suites), LeRobot/libero (1 693 ցուցադրում, 273 հազ. կադր, 40 առաջադրանք միասին), Franka սիմուլյացիայում, յուրաքանչյուր կետի համար 2 000 փորձարկում՝ լրիվ վերագործարկումներով: Ստանդարտ արձանագրությունում ստուգման համար առանձնացված է՝ ոչինչ, և հենց դա է էությունը: Երկու ընդլայնում իսկապես առանձնացնում են տվյալներ՝ LIBERO-90 (89 առաջադրանք) և LIBERO-Plus (10 030 խաթարված տարբերակ՝ 7 գործոնով):

Մեթոդը: Ժառանգված՝ ACT ոճի փաթեթավորում և ժամանակային համույթավորում, հոսքի համընկեցմամբ DiT գլուխ, տարածական softmax-ով առանցքային կետեր, FiLM պայմանավորում, արագության թորում (distillation): Նորը նախագծման սկզբունքն է՝ ներառել միայն այն, ինչ հենանիշն իրականում ստուգում է, ապա փոքրացնել, մինչև այն «կոտրվի»:

LIBERO-ի չորս խմբերի միջին ընդդեմ պարամետրերի ընդհանուր թվի

MINERVA-0.5M (0.54 մլն պարամետր)		95.05% հաջողություն
MINERVA-1M (0.99 մլն)		96.75% հաջողություն
MINERVA-10M (9.66 մլն)		97.45% հաջողություն
$\pi 0.5$, LeRobot-ի իրականացում (4.1 մլրդ)		97.50% հաջողություն

MINERVA-ի տողերը 2 000 փորձարկում են՝ ուսուցման մեկ սերմով: $\pi 0.5$ -ի տողը LeRobot-ի իրականացման հաղորդված արդյունքն է՝ 10 դրվագ/առաջադրանք (400 փորձարկում), ինչը հղման կետ է, ոչ թե արձանագրությամբ համապատասխանեցված համեմատություն: 1M մոդելի երեք սերմով վերաուսուցումը տալիս է 94.60–96.75 միջակայք, ուստի ~ 1 կետից փոքր տարբերությունները տարբերակելի չեն:

[arXiv 2609.03715](https://arxiv.org/abs/2609.03715), արեպրինտ, չվերարտադրված. կողը և ստուգակետերը հրապարակված են

Կարևոր թիվը

95.05%-ը չէ: Դա **ուսուցման սերմերից կախված ± 1 կետի տիրույթն է**, որը ստացվել է 1M ելակետային մոդելը երեք անգամ վերաուսուցանելով (94.60–96.75. միայն long խմբի արդյունքը տատանվում է 87.6–92.8 միջակայքում): Մասշտաբավորման կորը կարդացեք դրա համեմատ. 1 մլն-ի դեպքում 96.75-ից մինչև 9.66 մլն-ի դեպքում 97.45-ը 0.7 կետ է $10\times$ պարամետրերի դիմաց, և 97.50 հղման արդյունքը նույն տիրույթում է: Այդ չափիչ գործիքի համեմատ պահպանվում է միայն երկու ընտրություն՝ փաթեթի երկարությունը (-3.16^8 երկարության դեպքում, -2.20^3 32 երկարության դեպքում, և ձախողումը մեկնաբանելի է. կարճ փաթեթները «կոտրում» են goal խումբը, քանի որ թույլ են տալիս ռազմավարությանը տատանվել ճյուղերի միջև, իսկ երկար փաթեթները «կոտրում» են long խումբը, քանի որ չեն արձագանքում ենթանպատակների սահմաններում) և տեսողությանը հատկացված ռեսուրսը (-1.41 միջինում, -4.6 long-ում): Հոսքի համընկեցումն ընդդեմ մեկ անցումով L1 ռեգրեսիայի՝ $+0.34$ երեք սերմով, այսինքն՝ չափելի տարբերություն չկա, մինչդեռ ռեգրեսիան մինչև $3.8\times$ արագ է: Հրապարակված 0.54 մլն պարամետրով մոդելը փակում է իր կառավարման շղթան ≈ 5 մլ-ում՝ նոսրեբուքի պրոցեսորի ութ հոսքով, ընդդեմ SmolVLA-ի 1.0 վ-ի և $\pi 0.5$ -ի 12.8 վ-ի:

Բաղադրատոմսը դիմանում է նաև առաջադրանքների $2.25\times$ ավելի մեծ թվին՝ LIBERO-90-ի 89 առաջադրանքում 94.6% միջին՝ 0.995 մլն պարամետրով, մեկ սերմ, յուրաքանչյուր առաջադրանքի համար 10 դրվագ, ձախողումների բարակ «պոչով» և առանց փլուզված տեսարանների ընտանիքի:

Թույլ կողմերը: Մեկ սիմուլյատոր, մեկ մարմնավորում, գնահատման մեկ սերմ. հիմնական տողերը, LIBERO-90-ը և թորման համեմատությունը մեկական գործարկումներ են այդ տիրույթի ներսում: $\pi 0.5$ -ի հետ համեմատությունն այլ կերպ չափված թվի հետ է: Իրական ռոբոտի վրա ստուգում չկա: Տեղափոխությունների ստուգումը երկսայրի է. հրապարակված 1M ստուգակետը սառեցնելը և միայն առաջադրանքների ID-ների համապատասխանությունը պտտելը 96.75%-ն իջեցնում է 6.5%-ի, ընդ որում Spatial, Goal և Long խմբերը գրեթե ճշգրիտ ընկնում են պատահականության 1-ը-10-ից գծի վրա: Մնացածի գինը ցույց է տալիս LIBERO-Plus-ը՝ 46–56% խաթարումների ներքո, ֆոտոմետրիկ կայունություն $\leq 6.5\%$ ստուգված բոլոր մասշտաբներում:

Ինչ արժե փոխառել: (1) Չափեք ձեր սերմերի տիրույթը, նախքան որևէ աբլյացիայի հավատալը. երեք վերաուսուցումն այստեղ էժան է և անվավեր է դարձնում գրականության մեկ կետի տարբերությունների մեծ մասը, այդ թվում՝ այս թողարկումում մի քանիսը: (2) Կատարեք տեղափոխությունների ստուգումը. սառեցրեք

կշիռները, վերագրեք միայն պայմանավորման ալիքը. եթե արդյունքն ընկնում է պատահականության մակարդակի, ձեր պայմանավորումը ինդեքս է, ոչ թե լեզու: (3) Պարամետրերը ծախսեք կոդավորիչի, ոչ թե գործողությունների գլխի վրա. թոքենները խառնող MLP-ն 26%-ով պակաս պարամետրերով հավասարվեց ինքնաուշադրությանը, և այդ հզորությունը CNN-ին տեղափոխելը long խումբը ~10 կետով բարելավեց՝ ընդհանուր նույն չափի դեպքում: (4) Վերապլանավորեք յուրաքանչյուր քայլում ACT ոճի ժամանակային համույթավորմամբ (96.3%), ոչ թե BID ոճի թեկնածուների ընտրությամբ (95.0%) կամ RTC փափուկ լրացմամբ (91.8%), որը գերսահմանափակում է առանց այդ էլ աղմկոտ արագության դաշտը. 1 մլն-ից փոքր մոդելների համար սկզբնական աղմուկի ջերմաստիճանն իջեցրեք 0.85-ի, ինչը տալիս է մինչև +2.25 կետ այնտեղ, որտեղ մոդելն ամենաթույլն է, և ոչինչ՝ 1 մլն-ից բարձր:

Վերարտադրման արժեքը: ~2 ժ մեկ մոդելի համար մեկ սպառողական GPU-ի վրա, LIBERO-ի հանրային ցուցադրումներ, առանց նախաուսուցման (pretraining), առանց տվյալների հավաքման: Սերմերով կրկնված աբլյացիաների շարքը մեկ հանգստյան օրվա գործ է, ոչ թե կլաստերի:

7 Ինչ է սա փոխում փոքր նախագծի համար

Ենթադրենք՝ մեկ մանիպուլյատոր, մեկ GPU: Չորս բան փոխվեց:

1. **Ձեր LIBERO ելակետն այժմ երկու ժամվա գործ է:** MINERVA-ի ստուգակետերն ու բաղադրատոմսը սերմերով կրկնված աբլյացիաները դարձնում են մատչելի և ձեզ տալիս ± 1 կետի չափանիշ, որի ներսում են ընկնում մեկ գործարկումով տարբերությունների մեծ մասը: Կայունություն պնդելուց առաջ կիրառեք խաթարումների արձանագրություն:
2. **Ծառայիկ ձեռքը մտավ տպելի գնային միջակայք:** \$314, 374 գ, ստուգված MuJoCo փոխանցման մոդել, սիմուլյացիայից իրական աշխարհ առանց լրացուցիչ ուսուցման՝ միայն էնկոդերներով, բայց բուրձ մատի կարգաբերման համար ժամանակ նախատեսեք, քանի որ մնացորդային սխալը նրանում է:
3. **Ֆիզիկական գնահատումն էժանացավ՝** Peg-in-Bench-ի շնորհիվ, եթե ձեր հպումներով հագեցած արդյունքներն այժմ հենվում են ինքնաշեն հարմարանքի վրա, և դուք չեք կարողացել դրանք որևէ մեկի հետ համեմատել:
4. **Գործարկման ցիկլի երկու փոփոխություն կարելի է իրականացնել ուղիղ հոդվածներից:** Էնտրոպիայի հարթակով կտրումը վերաուսուցում չի պահանջում և վերացնում է կարգավորման մեկ պարամետր: Էյլերի տասը քայլի փոխարեն մեկ անցումով L1 ռեգրեսիան այս հենանիշում $3.8\times$ արագ էր՝ առանց չափելի ճշգրտության կորստի. արժե ստուգել ձեր սեփական գործողությունների գլխի վրա, նախքան ենթադրելը, թե հոսքի համընկեցումը ձեզ ինչ-որ բան է տալիս:

Այս մասշտաբում հայտնվեց նաև գոյության ապացույց, միայն ոլորտային մամուլում. [ինքնաշեն SO-101՝ թրթուրավոր հենքի վրա, որը \$\pi 0.5\$ է աշխատեցնում LeRobot-ի միջոցով](#)՝ ռոբոտի վրա Raspberry Pi-ով և համակարգչին փոխանցված գործարկմամբ, գուլպաներ վերցնելու 96% հաջողությամբ: Մեկ սարք, արձանագրություն նշված չէ. անեկդոտ է, բայց հենց խնդրո առարկա սարքավորման սահմաններում:

Սարքավորումների ոլորտում առնչվող ոչինչ չհայտնվեց. այս ժամանակահատվածում շարժաբերի, շոշափողական սենսորի, էժան մանիպուլյատորի, խորության տեսախցիկի կամ եզրային հաշվարկի թողարկում չեղավ, և հիմնական սիմուլյատորների ու միջանկյալ ծրագրային ստեղծելից ընդհանրապես թողարկումներ չեղան:

8 Շուկայական ազդանշան

Տեսանելիորեն ընդլայնվում են՝ NVIDIA-ն (գնում է հարթակի շերտը), Figure-ը (գնում է ուսուցման հաշվողական ռեսուրս, մատակարարը ստանում է բաժնեմաս), Lyte-ը (ընկալման չիպեր), Wandercraft-ը (ներգրավում է միջոցներ): Անցյալ շաբաթվա մանիպուլյացիային հատուկ գնորդները չվերադարձան. կապիտալը գնաց հաշվողական ռեսուրսների, զգայման չիպերի և հարթակների սեփականության, ոչ թե վերջնական գործիքների կամ ռազմավարությունների թիմերի վրա:

Ուշադրության արժանացած աշխատանքներում կրկնվում են հմտությունների երեք ուղղություններ: **Կատարման ժամանակացույցի ճարտարագիտություն**՝ ասինխրոն գործարկման ցիկլեր, փաթեթների կարում, մեկ քայլով գեներացում, ադապտիվ հորիզոններ. իմանալը, թե ասինխրոն կատարման դեպքում արժեքի գրադիենտն ընդհանրապես որտեղ է վավեր, այժմ առանձին մասնագիտացում է: **Ուսուցման ժամանակ վերահսկողություն՝ գործարկման զրո ծախսով: Գնահատման նախագծումը որպես ներդրում**՝ ZETA-ի արձանագրության բաժանումը, MINERVA-ի սերմերի տիրույթը, [FailBench](#)-ը, որը գտնում է, որ հաջողությունը գնահատող 13 VLM-ներից լավագույնը հասնում է 0.77-ի, և [ֆիզիկայի հետ համաձայնեցված խնամքի հենանիշը](#), որտեղ zero-shot π 0.5-ը ստանում է 0.7%՝ առանց ֆիզիկապես վավեր որևէ հաջողության:

Ֆինանսավորման տակ մի կառուցվածքային դիտարկում. քանի որ գնակարգված ցուցադրումների շուկաները հայտնվում են հատուկ կառուցված հավաքագրման սարքերի կողքին, տվյալների հետ աշխատանքում սակավ հմտությունը դրվագներ հավաքելուց տեղափոխվում է դեպի այն սահմանելը, թե որն է լավ դրվագը. ընդունման չափանիշներն ու IK-իրագործելիության ֆիլտրերը, որոնք NeoData-ի պես կորպուսներն այժմ փաստաթղթավորում են որպես խողովակաշարի փուլ, ոչ թե հետագա մտահոգություն:

Պարտադիր նվազագույնը՝ կող և ստուգակետեր հոդվածը ներկայացնելիս, գործողությունների փաթեթավորում, հոսքային կամ դիֆուզիոն գործողությունների գլուխ, ասինխրոն գործարկում, LIBERO $\geq 95\%$: Դեռևս սակավ են՝ սերմերով կրկնված աբլյացիաները, իրական սարքավորման վրա փակվող RL ցիկլերը, մեծածավալ շոշափողական տվյալները և սարքավորումը, որի ֆայլերը կարող եք տալ: Վարձատրության կամ պաշտոնային մակարդակների վերաբերյալ վստահելի տվյալներ չհայտնվեցին. ռոբոտաշինության աշխատավարձերի հարցումների համար հայտնվող միակ աղբյուրները SEO ագրեգատոր էջերն են՝ առանց նշված մեթոդաբանության, և ավելի լավ է ոչինչ չհաղորդել, քան դրանք «վանալ»:

9 Թեմաներ

Թեմա	ԿԱՐԳԱՎԻՃԱԿ	ՀԱԶՈՐԴ ԻՐԱԿԱՆ ԹԱՐՄԱՑՈՒՄԸ
Արդյո՞ք VLA-ի մասշտաբն է ապահովում լեզվով պայմանավորված կառավարումը	ԱՌԱՋԸՆԹԱՑ (MINERVA)	10 մլն-ից փոքր ռազմավարություն <i>իրական</i> սարքավորման վրա, բազմաառաջադրանքային, հրապարակված ելակետային մոդելի համեմատ
Հենանիշների վավերականությունը ֆիզիկական ԱԲ-ում	ԱՌԱՋԸՆԹԱՑ (MINERVA-ի հզորության նվազագույն շեմ + տեղափոխությունների ստուգում. FailBench)	Խոշոր լաբորատորիա, որը ստանդարտ LIBERO-ի կողքին լռելյայն հաղորդում է խաթարումների արձանագրություն
Հպման զգայում առանց առանձին շոշափողական սարքավորման	ԱՌԱՋԸՆԹԱՑ (Blind Dexterity, փափուկ մանիպուլյատորի IMU-ներ, ChainSplat)	Անկախ վերարտադրություն այլ սարքավորման վրա. VISTA-ն (W35) դեռ կող չունի
Ուսուցման ժամանակ վերահսկողություն՝ գործարկման զրո ծախսով	ՆՈՐ (PHR-VLA, Temporal Forcing, GIFT, EGR)	Դրանցից մեկի անկախ վերարտադրությունը երկրորդ հիմնական ցանցի վրա՝ սերմերի տիրույթով
Գործարկման ուշացումը որպես կիրառման նեղ տեղ	ԱՌԱՋԸՆԹԱՑ (ադապտիվ փաթեթավորում, DriftingVLA, MINERVA՝ ≈ 5 մվ CPU-ի վրա)	Ուշացումը որպես լռելյայն սյունակ հաջողության մակարդակի կողքին
Առցանց ուղղումը որպես վերջին քայլ	ՆՈՐ (SmoothRL՝ 250 դրվագ. XR-2 DAgger՝ 58→93%)	Երրորդ խումբ, որը ցույց է տալիս, որ հազեցումից հետո ուղղումների տվյալներն ավելի արդյունավետ են, քան լրացուցիչ փորձագիտական տվյալները

ԹԵՄԱ	ԿԱՐԳԱՎԻՃԱԿ	ՀԱԶՈՐԴ ԻՐԱԿԱՆ ԹԱՐՄԱՏՈՒՄԸ
Աշխարհ-գործողություն մոդելները որպես միասնական ռազմավարություն + սիմուլյատոր	ԱՌԱՋԸՆԹԱՑ (WCD, SolarWM-ի հրապարակում)	Հրապարակված աշխարհի մոդել, որն օգտագործվում է ռազմավարությունը բարելավելու համար մեկի կողմից, ով այն չի ստեղծել
Ում է պատկանում ռոբոտների ուսուցման բաց ստեղծ	ՆՈՐ (NVIDIA/Hugging Face, \$12.9 մլրդ)	LeRobot-ի հաջորդ թողարկումը, և արդյոք կառավարման կամ ձևաչափի որոշումները տեսանելիորեն փոխվում են
Էժան բաց սարքավորումը որպես ուսուցման հենք	ԱՌԱՋԸՆԹԱՑ (Aero Hand՝ \$314, BRIDGE՝ \$1.5 հազ., Peg-in-Bench)	Երրորդ կողմ, որը դրանցից որևէ մեկի վրա ռազմավարություն է ուսուցանում և կիրառում
Հայտարարված բաց թողարկումներն ընդդեմ իրական արտեֆակտների	ՎԻՃԵԼԻ	Կողմ՝ MINERVA, Aero Hand, BRIDGE, SolarWM՝ հրապարակված: Դեմ՝ OpenWAM-ը չունի քարտ կամ լիցենզիա. StreamPI-ի կշիռները W35-ից ի վեր դեռ սպասվում են
Մարդանման ռոբոտների կապիտալն ընդդեմ ցուցադրված կարողության	ԱՌԱՋԸՆԹԱՑ (Figure՝ \$3.5 մլրդ. Wandercraft. Galbot-ի նախնական պատվեր)	Մարդանման ռոբոտներ արտադրող ընկերություններից դեռ զրո մեթոդ. CMU-ի BRIDGE-ը ակադեմիական հակաօրինակն է
Մանիպուլյացիան գնվում է, ոչ թե կառուցվում	ԿԱՆԳԵԱԾ	W35-ի կանխատեսմանը, թե այս եռամսյակում կհայտնվի լոգիստիկայի երկրորդ գնորդ, մնացել է ~7 շաբաթ

ԱՅԼ ՈՒՇԱԳՐԱՎ ՆՅՈՒԹԵՐ

VLAct — տվյալների 20%-ը գերազանցում է ամբողջ տվյալներով GR00T-N1.6-ին չտեսնված մարդանման ռոբոտի վրա: Բաց թողնված է՝ շարունակական նախաուսուցման 16-GPU բյուջե, անհասանելի՝ չնայած բաց կշիռներին: Motus2 — ընդհանուր կշիռներով ռազմավարություն, սիմուլյատոր և գնահատիչ՝ բարելավման փակ ցիկլում: Բաց թողնված է՝ խոշոր ընկերությունների մասշտաբի փակ տվյալներ և ոչ մի արտեֆակտ, ուստի ցիկլը հնարավոր չէ ստուգել: ZimaBlue — տեսանյութից գործողություն ուսումնական ծրագիր, 30 Հց կառավարում է գոկենտրոն տեսանյութից, կողը հրապարակված է: Բաց թողնված է՝ ռոբոտների պարկի մասշտաբի նախաուսուցումը նշանակում է, որ կողը վերարտադրում է բաղադրատոմսը, ոչ թե արդյունքը: Zeva — սառեցված ռազմավարությունը որոնում է գործողություն-արդյունք անցումներ որպես համատեքստային պատճառային հուշում: Բաց թողնված է՝ զբաղեցնում է առանց գրադիենտի հարմարեցման նույն տեղը, ինչ WCD-ն, ավելի պակաս մաքուր արձանագրությամբ: DriftingVLA — մեկ քայլով աղմուկից փաթեթ գեներացում, 3.36× արագ: Բաց թողնված է՝ նույն առանցքը, որը MINERVA-ի L1-ընդդեմ-հոսքի արդյունքն ավելի էժան և երեք սերմով է ուսումնասիրում: SA-WAM — սառեցված VAE-ն թրքենացնում է չափողական խորությունը. երկրաչափական սխալը նշում է ձախողումների 80%-ը: Բաց թողնված է՝ կող նշված չէ, իսկ ձախողումները նշող մասը հայտնաբերման ելակետ չունի:	World-Coherent Decoding — նմուշառված ապագաները դասակարգում է ըստ հոսքի անակնկալության և գործողությունների ուղու ջանքի, սառեցված հիմնական ցանցով: Բաց թողնված է՝ կող նշված չէ, իսկ այս շաբաթ արդեն կա առանց ուսուցման մեթոդ՝ ավելի լայն գնահատմամբ: MAGP — մասշտաբային համարժեքությամբ (scale-equivariant) տվյալների աճեցումը գոյություն ունեցող VLA-ին վերադարձնում է չափողական մասշտաբը: Բաց թողնված է՝ «միացրո և օգտագործիր» պնդում՝ առանց հրապարակված աղապտերի: Seeing the World and the Self — սեփական մարմնին հարմարեցված երկրաչափական հիմնական ցանցը պայմանավորում է դիֆուզիոն շարժման գլուխը: Բաց թողնված է՝ կողը և տվյալների հավաքածուն խոստացված են, բայց չեն հրապարակվել. նույն չհրապարակված չափողական խորության խումբը, ինչ MAGP-ն և VI3-ը: Contact-Guided Exploration — հետազոտման քննադատ, որը սկզբնավորվում է բռնումների նմուշառիչի հպումներով և ուսուցման կեսին աստիճանաբար նվազեցվում է մինչև զրո: Բաց թողնված է՝ հատուկ է տեղաշարժ-մանիպուլյացիային, միայն սիմուլյացիայից իրական աշխարհ, միջավայրը չի հրապարակվել: Non-Prehensile Throwing — RL հոդերի ցնցման (jerk) հետագծերի վրա, առանց լրացուցիչ ուսուցման փոխանցում UR5e-ին՝ ցնցման համակարգի նույնականացման միջոցով: Բաց թողնված է՝ մեկ առաջադրանք, բացահայտորեն զգայուն շփման նկատմամբ, պահանջում է վերանույնականացում ձեր մակերեսի վրա:
--	---

DemoMimic — համան տեղային երկրաչափության պարզեներ, 71% 16 առարկայի վրա մեկ ցուցադրումից: Բաց թողնված է՝ կող նշված չէ, պարզեր դեռ ձեռքով է սահմանվում առաջադրանքների յուրաքանչյուր ընտանիքի համար:

Temporal robustness of imitation learning — կատարման արագության աճին զուգընթաց ACT-ը շատ ավելի արագ է վատթարանում, քան իր սցենարային փորձագետը. կողը և տվյալները հրապարակված են: Բաց թողնված է միայն տեղի սղության պատճառով. շաբաթվա ամենամաքուր բացասական արդյունքը:

RoboCurve / GPT-6 Astra իրական մանիպուլյատորների վրա (ոլորտային մամուլ) — 95% ամանով առաջադրանքում / 10% գլուխկոտրուկում՝ երկու YAM մանիպուլյատորի վրա: Բաց թողնված է՝ $n=20$ յուրաքանչյուր առաջադրանքի համար, հոդված չկա, ուստի առաջադրանքների միջև տարբերությունը միակ օգտագործելի ազդանշանն է:

10 Բաց հարց

MINERVA-ի չափած LIBERO-ի ± 1 կետի սերմային տիրույթը 10 մլն-ից փոքր պարամետրերով *այդ* ռազմավարությունների ընտանիքի՝ հատկությունն է, թե՛ հենանիշի՝ բոլոր մասշտաբներում: Այս թողարկման գրեթե բոլոր աբլյացիոն տարբերությունները՝ GIFT-ի +4.6 տ.կ.-ը, ադապտիվ փաթեթավորման +1.3-ը RoboTwin-ում, StreamPI-ի +2.6-ը W35-ից, մեկական գործարկումներ են, և թե դրանցից քանիսը կպահպանվեն, ամբողջությամբ կախված է պատասխանից: Հարցը կլուծի ցանկացած խումբ, որը $\pi 0.5$ կամ GR00T դասի ռազմավարությունը երեք անգամ վերաուսուցանի նույն տվյալների վրա և հրապարակի տարածվածությունը: Դա արժե երեք ուսուցման գործարկում, և ոչ ոք դա չի արել:

Աղբյուրները նշված են տեքստում: arXiv-ի բոլոր նյութերը պրեպրինտներ են՝ չգրախոսված: Մեկ լաբորատորիայի պնդումները նշված են, որտեղ դա տեղին է: