

This Week in Robotics

Ժամանակահատված՝ 2026-08-24 → 2026-08-30: Առաջին թողարկում:

Մեքենայական թարգմանություն անգլերենից, չիմբագրված՝ բնօրինակը՝ twir.arcanerobotics.com/2026-w35/

ԳԼԽԱՎՈՐԸ

Sony CSL-ի 0.68 մլն պարամետրով ռազմավարությունը (policy), որը հիմնված է կանխատեսողական կոդավորման (predictive coding) վրա, LIBERO-ում 4–7× գերազանցեց պարամետրերով համապատասխանեցված Transformer և LSTM ելակետային մոդելներին (baselines): Դա ուղղակի հարված է այն ենթադրությանը, թե լեզվով պայմանավորված կառավարումը պահանջում է VLA-ի մասշտաբի պարամետրեր: Մինչդեռ կապիտալը շարժվեց հակառակ ուղղությամբ. հաղորդվում է, որ SoftBank-ը բանակցում է 1X-ի վերահսկողությունը ձեռք բերելու շուրջ՝ ~\$6 մլրդ գնահատմամբ:

1 Ռազմավարությունների ուսուցում և VLA-ներ

PredVLA

ՊՐԵՊՐԻԵՆՏ / Sony CSL Tokyo

Ճարտարապետական նորույթը՝ դիտարկումները երբեք չեն մտնում ռեկուրենտ դինամիկայի մեջ: Հիերարխիկ, բազմաժամանակային մասշտաբով RNN-ը (PV-RNN-ի ժառանգորդ) կանխատեսում է տեսողական հատկանիշները և պրոպրիոցեպցիան (proprioception), իսկ դիտարկումներն ազդում են միայն թեստի ժամանակ՝ սահող պատուհանում թաքնված ազատ փոփոխականների վրա գրադիենտային վայրէջքի միջոցով. դա սխալի ռեգրեսիա է, ոչ թե մուտքային ուղի: Թաքնված վիճակի գնահատման (inference) քայլերի թիվը զրո դնելով՝ նույն կշիռների վրա ստացվում է *ճշգրիտ* բաց շղթայով արլյացիա (ablation): Նախնական մշակման մասը (front end) սառեցված է և պատրաստի (ResNet18 + MiniLM + ֆիքսված PCA), ուստի 675 732 ուսուցանվող պարամետրերն ամբողջ ուսուցված կառավարիչն են: Գնահատում՝ միայն LIBERO, միայն սիմուլյացիա, մեկ մարմնավորում (embodiment), յուրաքանչյուր խմբի (suite) համար առանձին ուսուցում: Կող չկա: Մեկ լաբորատորիայի չվերարտադրված պնդում: Մանրամասն՝ [§6](#)-ում:

StreamPI

ՊՐԵՊՐԻԵՆՏ / HKU + ACE Robotics

Մեկ կադրով աշխատող VLA-ին ($\pi 0.5$) ավելացնում է ժամանակային համատեքստ՝ առանց նոր պարամետրերի. դա հրահանգին կապված հոսքային մշակում է T կադրից բաղկացած պատուհանում, ոչ թե առանձին հիշողության մոդուլ: «Զրո լրացուցիչ պարամետր» պնդումը նշանակում է, որ ժամանակային ազդեցական օգտագործում է ուշադրության (attention) արդեն եղած ռեսուրսը՝ ռեկուրենտ գլուխ ավելացնելու փոխարեն: LIBERO՝ 98.3% T=5-ի դեպքում ($\pi 0.5$ -ի համեմատ Goal-ում +2.8 տ.կ., Long-ում +2.6 տ.կ., ինչը հազեցած հենանիշում աղմուկի սահմաններում է), բայց փաստարկն իրական ռոբոտի արդյունքներն են՝ +26.7-ից +36.6 տ.կ. հիշողությունից կախված և ճշգրտություն պահանջող չորս առաջադրանքներում, օրինակ՝ բաժակի տեղադրումը 92% ընդդեմ 60%-ի՝ 25 փորձում: Փոքր նմուշ և մեկ լաբորատորիա, բայց այդ տարբերությունը դուրս է սերմերի (seeds) հավանական տատանումից: [ՊԱՀՈՑ](#)՝ իրական կոդով (openpi/LeRobot-ի հիմքով). 2026-08-30-ի դրությամբ README-ն ասում է, որ կշիռները դեռ սպասվում են: Մինչ դրանց հրապարակումը սա մեկ GPU-ով վերարտադրելի չէ. $\pi 0.5$ -ի լրաուսուցումը (fine-tuning) փոքր մոդելը զրոյից ուսուցանելու պես է ժան չէ:

Riemann-1.0

ՊՐԵՊՐԻՆՏ

Մեն պատճառային ավտոռեգրեսիվ մոդել, որը միաժամանակ ռազմավարություն է և գործողությամբ պայմանավորված սիմուլյատոր՝ բազմատեսանկյուն տեսանյութի, վիճակի և գործողության վրա. հայտարարված է 200 հազ. + ժամ տվյալ, որը ներառում է մարդու էգոկենտրոն տեսանյութ, ձեռքի բռնիչով արված ցուցադրումներ և տարասեռ ռոբոտային տվյալներ: Հաղորդված արդյունքները՝ RoboTwin2.0՝ 94.3%, LIBERO՝ 99.0%, RoboCasa-365՝ 62.6%, իրական աշխարհում՝ 85.0% հաջողություն / 94.4% առաջխաղացում, ~15%-ով ավելի, քան ամենաուժեղ բաց ելակետային մոդելը: Հերթում սա նշված էր որպես «բաց կոդով հրապարակված», բայց ես չկարողացա հաստատել. abs էջում պահոց նշված չէ, և համապատասխան HF պահոց գոյություն չունի: **Հրապարակման պնդումը համարեք չհաստատված:** Ամեն դեպքում՝ սա ռոբոտների պարկի (fleet) տվյալների արդյունք է:

CLAP

ՊՐԵՊՐԻՆՏ

ՆԱԽԱԳԻԾ

Գործողությամբ պայմանավորված տեսանյութային աշխարհի մոդել (world model)՝ տարբեր մարմնավորումների միջև. թաքնված գործողությունները սովորվում են չափափակված տեսանյութից և կապվում վերջնական գործիքի (end-effector) դիրքերին՝ ընդգրկելով DROID, Bridge, երկձեռ YAM, G1 մարդանման ռոբոտ և EgoDex: Թաքնված գործողությունների տարածությունն ընդհանուր միջերեսն է, իսկ վերջնական գործիքի դիրքը՝ մարմնավորմանը հատուկ միակ կապը. հենց դրա համար է պնդվում, որ նոր մարմնավորմանը հարմարվելը տեղի է ունենում մեկ գրադիենտային քայլում: SVD հիմնական ցանց (backbone), OXE + EgoDex, 100 հազ. քայլ: Հղված են՝ [ԿՈՂ](#) և [ԿԵԻՆԵՏԻՐ](#): Կարևոր թիվը՝ 12 GB-ից պակաս VRAM, աշխատում է 3060-ի վրա:

2 Մանիպուլյացիա, ճարպկություն և զգայում

VISTA

ՊՐԵՊՐԻՆՏ

Հպումը (contact) որոշում է ճկուն բռնիչի (gripper) *տեսանելի ձևափոխությունից*՝ սովորական տեսախցիկներով ստացված 3D տեղաշարժի դաշտից, և այն վերադարձնում է որպես բռնիչի աստիճանական ուղղումներ: Իմաստը շոշափողական (tactile) մատրիցի համեմատ ճշգրտությունը չէ, այլ շոշափողական սարքավորումն ու դրա լարերը վերացնելը: Երեք առաջադրանք (տարբեր չափերի առարկաների բռնում, կափարիչի պտուտակաբացում, գեղազրություն), որոնցում գերազանցում է 3D Diffusion Policy-ին և շոշափողական ելակետային մոդելին: Երեք առաջադրանքը քիչ է, և կոդ նշված չէ. նախագծի նյութերը Google Sites էջ են: [ԵԶՐԱՀԱՆԳՈՒՄ](#) → փոքր նախագծի համար հպման ազդանշան ստանալու ամենաէփան ճանապարհը, եթե ձեր բռնիչը ձևափոխվում է տեսանելի և կրկնելի կերպով:

TacForcing

ՊՐԵՊՐԻՆՏ

Շոշափողական հետադարձ կապը տեղադրում է գործողությունների հոսքային գեներացման ներսում, ոչ թե դրա շուրջ փաթաթված ռեակտիվ շերտում: Նույն ճարտարապետական բնագոյը, ինչ StreamPI-ում, կիրառված հպման նկատմամբ:

Լիի խմբով սահմանափակված MeanFlow բռնում

ՊՐԵՊՐԻՆՏ

Հոսքի համընկեցում (flow matching) $SO(3) \times \mathbb{R}^3$ -ի վրա՝ կիսախմբի համաձայնեցվածության հանրահաշվական պայմանով. ≤ 5 ցանցի գնահատում, միլիվայրկյանային բռնումներ, ACRONYM-ում դիֆուզիոն/հոսքային ելակետային մոդելների նկատմամբ մինչև $39\times$ գերազանցություն, հայտարարված ուղիղ փոխանցում իրական աշխարհ: Կոդ չկա: Փոխանցելի գաղափարը՝ կիսախմբի սահմանափակումը որպես բազմաձևության վրա քիչ քայլերով նմուշառման հնարք, հատուկ չէ միայն բռնմանը:

Նաև (բոլորը պրեպրինտներ)՝ [ջլային կառավարմամբ հինգմատանի ձեռք](#)՝ յուրաքանչյուր մատի վրա երկու տեսակի շոշափողական զգայմամբ, [PRISM](#) GPU-ի վրա նմուշառմամբ MPC QP պրոյեկցիայով՝ երկձեռ հպման

համար, և ռոպեների ընթացքում սովորված [ձեռնածություն իրական ռոբոտի վրա](#)։ վերջինը հիշեցնում է, որ սարքավորման վրա առցանց հարմարվելը կարող է ավելի արդյունավետ լինել, քան sim2real խզումը փակելը:

3 Հարթակներ և մարմնավորում

Սակավ նորություններով շաբաթ: **Microduck** ([ԿԱՄԴՈՒԼ](#))՝ \$399, 25 սմ, 15-DOF երկոտանի ռոբոտ Pollen/Hugging Face-ից, RK3566, ToF լիդար, երկու IMU. SDK-ն, սիմուլյացիան և RL ուսուցման ստեկը բաց են GitHub-ում, սարքավորման ֆայլերը՝ ոչ: Վաճառքը՝ մինչև 2026 թ. Սուրբ Ծննդյան տոները: Կարողությունները ցույց են տրված միայն տեսանյութով (օրորվող քայլք, կտուցով առարկա վերցնել, ընկնելուց վերականգնում, անվաշմուշկներով սահք)։ դեմո տեսանյութի մակարդակի ապացույց, անհայտ է՝ բեմադրված է, թե իրական:

[SOLO](#)-ն (մարդանման ռոբոտի տեղաշարժ ընկալման հիման վրա) և ցածր գնով բաց սարքավորման վրա [Poppy-ի LQR քայլքը](#) հարթակային միակ պրեպրինտներն են, որոնց հետևում իրական մեթոդ կա: Այս ժամանակահատվածում մարդանման ռոբոտներ արտադրող ոչ մի ընկերություն տեխնիկական արդյունք չի հրապարակել:

4 Բաց կող և գործիքներ

Այս ամիս աշխատանքը հեշտացնող նյութերը՝ ըստ կարևորության. **CLAP**-ի 12 GB աշխարհի մոդելը (վերևում), որը թույլ է տալիս ուսուցված դինամիկան ստուգել ձեր սեփական սարքավորման վրա: [CRESSim-Neo](#) (պրեպրինտ)՝ GPU-ի վրա փաթեթներով աշխատող սիմուլյատոր ձևափոխվող և վիրաբուժական առարկաների համար՝ PyTorch կապերով: [Sharpa-ի շոշափողական ձևափոխության կողավորիչը](#)՝ Apache-2.0 լիցենզիայով ConvNeXtV2 կշիռներ սենսորներ արտադրողից. փոքր է, բայց իրական տվյալներ ունեցողի կողմից նախաուսուցված շոշափողական հատկանիշները հազվադեպ են: [Muse-Robotics-1](#)՝ 121.7 մլն պարամետրով, MIT լիցենզիայով VLA, առանց հոդվածի, առանց քարտի, չստուգված:

Այս ժամանակահատվածում ոչինչ չի եղել LeRobot-ից, MuJoCo/MJX-ից, Isaac Lab-ից, ManiSkill-ից, RoboCasa-ից կամ ROS 2-ի միջուկից: Genesis-ի թողարկումների էջը երեք անընդմեջ ստուգումների ժամանակ 404 է վերադարձրել. կարգավիճակն անհայտ է, և դա չի նշանակում «թողարկում չկա»:

Գնահատման թեմայով ձեր ժամանակին արժանի նյութը՝ [Ֆիզիկական ԱԲ-ի հենանիշների ավելորդության վիճակագրական աուդիտ](#) (պրեպրինտ). 51 մոդել $\times$ 12 հենանիշ (benchmark). փոխարինելի երկու զույգերը միավորելը 51 մոդելից 22-ի դիրքը փոխում է երեք և ավելի տեղով, իսկ չորս հենանիշը պահպանում են բոլոր տասներկուսի դասակարգման օգտակարության 78.5%-ը: Կարդացեք այսպես. ձեր տեսած հենանիշային աղյուսակների մեծ մասը նույն առանցքը մի քանի անգամ է հաղորդում:

5 Փող և գործարքներ

GENERALIST

\$3 մլրդ

+50%՝ 10 շաբաթում

\$200 մլն լրացուցիչ ներդրում՝ 8VC-ի գլխավորությամբ

Աղբյուրները՝ ստորև՝ ըստ նյութերի. մամուլ և ոլորտային մամուլ:

1X

~\$6 մլրդ

SoftBank-ի բանակցությունները վերահսկիչ բաժնեմասի շուրջ, չհաստատված

NSF ԿԵՆՏՐՈՆ

\$30 մլն

տասը տարվա ընթացքում, UT Austin

LOCUS / NEXERA

անհայտ

պայմանները չեն հրապարակվել

Generalist → \$3 մլրդ

ՄԱՄՈՒԼ

8VC-ի գլխավորությամբ \$200 մլն լրացուցիչ ներդրում՝ հունիսի \$2 մլրդ-ից հետո. 50% աճ տասը շաբաթում՝ Gen 1.5-ի այն պնդման հիման վրա, որ այն առաջադրանքներ է սովորում 3–12 վայրկյան տևողությամբ տեսացուցադրումներից: **ԵԶՐԱՀԱՆԳՈՒՄ** → կապիտալը գնահատում է *ցուցադրումների արդյունավետությունը*, ոչ թե առաջադրանքների ծածկույթը. համոզիչ դարձած փաստարկը «մեկ հմտության համար ավելի քիչ դրվագներն» են:

SoftBank ↔ 1X, ~\$6 մլրդ, վերահսկիչ բաժնեմաս

ՄԱՄՈՒԼ / հաղորդել է The Information-ը. SoftBank-ը հրաժարվել է մեկնաբանել, 1X-ը չի պատասխանել, չհաստատված

ԵԶՐԱՀԱՆԳՈՒՄ → *վերահսկիչ* բաժնեմասը աճի փուլ չէ: Եթե գործարքը կայանա, այն կնշանակի, որ ռազմավարական սեփականատերը կլանում է մարդանման ռոբոտի հարթակը, այլ ոչ թե շուկան է բարձրացնում դրա արժեքը:

Locus-ը ձեռք է բերում Nexera-ն

ՈՒՈՐՏԱՅԻՆ ՄԱՄՈՒԼ

NeuraGrasp-ը՝ ծծող և սեղմող բռնումը համատեղող փափուկ վերջնական գործիք՝ արտադրանքի վեց սերնդով, տեղադրվում է Locus-ի Array շարժական մանիպուլյատորի վրա: **ԵԶՐԱՀԱՆԳՈՒՄ** → լոգիստիկայի հաստատված խաղացողները եզրակացրել են, որ մանիպուլյացիան գնում են, ոչ թե կառուցում, և այն, ինչ նրանք գնում են, վերջնական գործիքի սարքավորումն է, ոչ թե ռազմավարությունը:

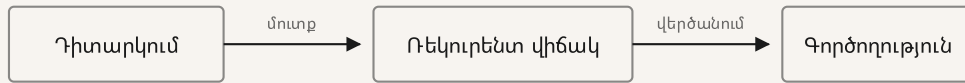
EXL/iMerit. վերլուծական ընկերությունը կլանում է ֆիզիկական ԱԲ-ի տվյալների պիտակավորման ընկերությանը. **NSF**-ը տասը տարվա ընթացքում \$30 մլն է հատկացրել UT Austin-ի կենտրոնին, որն ուսումնասիրում է ռոբոտների և մարդկանց համատեղ ուսուցումը: **ԵԶՐԱՀԱՆԳՈՒՄ** → տվյալների պիտակավորման շերտը միավորվում է կորպորատիվ ծառայությունների մեջ, մինչդեռ երկարաժամկետ HRI-ն նահանջում է դեպի դաշնային ֆինանսավորում:

6 Խորը ընթերցում՝ PredVLA

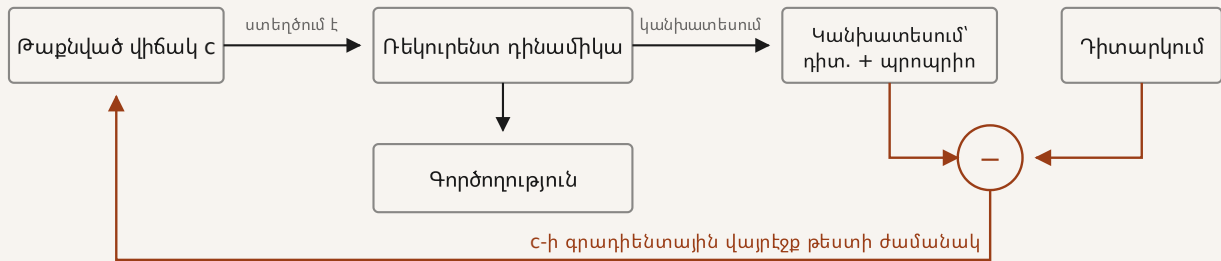
[arXiv:2608.26673](https://arxiv.org/abs/2608.26673)՝ արեպրինտ, Sawada և Kasahara, Sony Computer Science Laboratories:

Փորձի կառուցվածքը: LIBERO, բոլոր չորս խմբերը, պաշտոնական ցուցադրումներ, Franka սիմուլյացիայում, 14 սերմ, յուրաքանչյուր խումբ ուսուցանված է առանձին: Ստուգման համար առանձնացված է՝ ոչինչ LIBERO-ի սեփական բաժանումներից բացի. ոչ այլ մարմնավորումներ, ոչ սարքավորում:

Մեթոդը: Ժառանգված՝ բազմաժամանակային մասշտաբով RNN և PV-RNN-ի նախնական/հետին (prior/posterior) ազատ փոփոխականների ձևակերպումը, սառեցված ResNet18/MiniLM նախնական մշակման մասեր, զառույյան խառնուրդով գործողությունների գլուխ (action head): Նոր՝ (ա) այն գործարկել որպես *լեզվով պայմանավորված* ռազմավարություն ժամանակակից VLA արձանագրությամբ, ինչը կանխատեսողական կողավորման սենսոմոտոր աշխատանքները չէին արել. (բ) խիստ կանոնը, որ դիտարկումները մոդելին դիպչում են միայն թաքնված փոփոխականների վրա ազատ էներգիայի նվազարկման միջոցով, և հենց դա է ապահովում ճշգրիտ բաց շղթայով աբլյացիան. (գ) երկու կողային ուղի՝ ցածր չափողականության տեսողություն → գործողություն նեղ կոկորդ և մոդելի սեփական կանխատեսած գործողության ու պրոպրիոցեպցիայի էֆերենս պատճենը (efference copy), որը վերադարձվում է տեսողական ճյուղին, այնպես որ ոչ մի չափում չի մտնում ուղիղ դինամիկայի մեջ:



PREDVLA



Ճարտարապետական կանոնը, որի վրա կառուցված է հոդվածը: Վարքի կլոնավորման (behaviour cloning) դեպքում դիտարկումը ռեկուրենտ վիճակի մուտքն է: PredVLA-ում դինամիկան կանխատեսում է ստեղծում, և դիտարկումը թաքնված վիճակին հասնում է միայն որպես կանխատեսման սխալի գրադիենտ: Գրադիենտային քայլերի թիվը զրո դնելը ամբողջությամբ կտրում է նարնջագույն ուղին, և հենց դա է հնարավոր դարձնում ճշգրիտ բաց շղթայով արկայցիան անփոփոխ կշիռների վրա:

Կարևոր թիվը: Ոչ թե չորս խմբերի 75.35% միջինը, այլ վերահսկվող համեմատությունը: Համապատասխանեցված պարամետրերով և նույն նախնական մշակմամբ, ցուցադրումներով, գործողությունների գլխով ու արձանագրությամբ՝ BC-Transformer՝ 19.73%, BC-RNN՝ 10.26%: Ավելի շատ պարամետրերով և գործողությունների փաթեթավորմամբ (action chunking) Transformer-ներն ավելի վատ են (16.70%՝ 8 երկարությամբ փաթեթի դեպքում, 3.44%՝ 16 երկարությամբ փաթեթի դեպքում): Այս բյուջեի դեպքում փաթեթավորումը փլուզվում է:

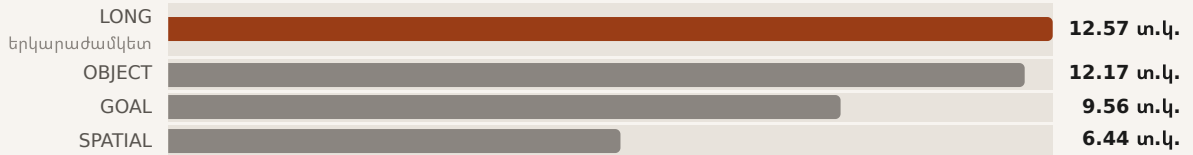
LIBERO-ի չորս խմբերի միջին հաջողության մակարդակը՝ համապատասխանեցված պարամետրային բյուջեներով



Բոլոր հինգն ունեն նախնական մշակման նույն սառեցված մասը, նույն ցուցադրումները, գործողությունների գլուխը և գնահատման արձանագրությունը, ուստի տարբերությունը վերագրելի է կառավարիչին: Գործողությունների փաթեթավորմամբ երկու տարբերակներն ավելի շատ պարամետր ունեն, քան PredVLA-ն, և այնուամենայնիվ վերջին տեղերում են, բայց սրանք հեղինակների սեփական ելակետային մոդելներն են, և այս մասշտաբում հիպերպարամետրերի որոնում չի հաղորդվում:

arXiv:2608.26673, աղյուսակ 1 և աղյուսակ 2՝ պրեպրինտ, չվերարտադրված

Կորցրած տոկոսային կետերը, երբ առցանց սխալի ռեգրեսիան անջատված է



Թաքնված վիճակի գնահատման քայլերի թիվը զրո դնելը նույն կշիռների վրա տալիս է ճշգրիտ բաց շղթայով կառավարում. սա այն աբյուսիվն է, որը հնարավոր է դարձնում ճարտարապետությունը: Փակ շղթայով ուղղումը կորցնելու գինը մոտավորապես կրկնապատկվում է ամենահեշտ խմբից դեպի երկարաժամկետ խումբը: Սրանք հոգվածի՝ նույն սերմերով համեմատված տարբերություններն են, ոչ թե 14 սերմով հիմնական միջինների տարբերությունները, որոնք հաշվարկված են սերմերի այլ բազմության վրա:

arXiv:2608.26673՝ պրեպրինտ, չվերարտադրված

Թույլ կողմերը:

- Միայն սիմուլյացիա, մեկ մարմնավորում, մեկ հենանիշ. LONG-ում 40.57%-ը երկարաժամկետ առաջադրանքների արդյունք չէ:
- Յուրաքանչյուր խմբի առանձին ուսուցումը նշանակում է, որ բազմաառաջադրանքային ոչ մի պնդում չի պահպանվում:
- Ելակետային մոդելները հեղինակների սեփական՝ պարամետրերով համապատասխանեցված իրականացումներն են, ոչ թե հրապարակված համակարգեր, և հիպերպարամետրերի որոնում չի հաղորդվում: 3.4% արդյունքով 16-փաթեթային Transformer-ը կասկածելիորեն անսարք ելակետ է. հավանաբար դա 1 մլն-ից փոքր ռեժիմում կարգավորման արտեֆակտ է, ոչ թե փաթեթավորման դեմ ապացույց:
- Համեմատությունը փոքր BC ռազմավարությունների հետ է, ոչ թե $\pi 0.5$ -ի կամ OpenVLA-ի: Հոգվածը չի պնդում, թե գերազանցում է դրանց:
- Հաշվողական ռեսուրսների անհամաչափությունը *հեղինակների օգտին* է, ինչը ազնիվ ուղղությունն է:

Ինչ արժե փոխառել:

- Ճշգրիտ բաց շղթայի անջատիչը. եթե դիտարկման ուղին առանձնացվող է, կարող եք նույն կշիռների վրա չափել, թե իրականում ինչ է անում փակ շղթայով ուղղումը: Այստեղ դրա անջատումն արժեցել է 6–13 տ.կ.: Ռազմավարությունների մեծ մասն ընդհանրապես չի կարող պատասխանել այս հարցին:
- Սառեցված կոդավորիչ (encoder) + ձեր սեփական ցուցադրումների վրա հարմարեցված ֆիքսված PCA բազիս. էժան է, դետերմինիստիկ, և ուսուցանվող մասը դարձնում է այնքան փոքր, որ փորձերը կարելի է կրկնել բոլորների ընթացքում:
- Էֆերենտ պատճենը որպես հետադարձ կապի ուղի. վերադարձրեք մոդելի *կանխատեսած* գործողությունն ու պրոպրիոցեպցիան, ոչ թե չափումները՝ ուղիղ դինամիկան պահելով գեներատիվ:
- Ըստ ալիքների աղավաղման ստուգումներ: Նրանց ամենասուր արդյունքն անհամաչափ է. պրոպրիոցեպտիվ տվյալների 25% աղավաղումը OBJECT-ում արժեցել է 49.7 տ.կ., իսկ տեսողական տվյալների 50% աղավաղումը՝ 5 տ.կ.-ից պակաս: Արժե ստուգել ցանկացած համակարգի վրա, որը կիրառում եք:

Վերարտադրման արժեքը: 0.68 մլն ուսուցանվող պարամետր, սառեցված նախնական մշակման մասեր, LIBERO-ի հանրային ցուցադրումներ. գնահատմամբ՝ մեկ խմբի համար միանիշ թվով GPU-ժամ մեկ սպառողական քարտի վրա, առանց տվյալների հավաքման, առանց նախաուսուցման (pretraining): Միակ անսովոր ծախսը թեստի ժամանակ սխալի ռեգրեսիան է, և 46 մվ/քայլ արժեքն այն ներառում է: **ԵԶՐԱՀԱՆԳՈՒՄ** → մեկ հանգստյան օր, և ամբողջ ռիսկը նրանում է, թե արդյոք ելակետային մոդելներն արդարացիորեն են կարգավորված:

7 Ինչ է սա փոխում փոքր նախագծի համար

Ենթադրենք՝ մեկ մանիպուլյատոր, մեկ GPU, ռոբոտների պարկ չկա: Այս շաբաթ երեք նյութ ինչ-որ բան է փոխում:

PredVLA-ն այստեղ միակ հոդվածն է, որը կարող եք ամբողջությամբ վերարտադրել: 0.68 մլն ուսուցանվող պարամետր, սառեցված պատրաստի նախնական մշակման մասեր, LIBERO-ի հանրային ցուցադրումներ, մեկ խմբի համար միանիշ թվով GPU-ժամ: Եթե ուզում եք ռազմավարություն, որը կարող եք իսկապես կարդալ սկզբից մինչև վերջ և հագեցնել չափիչ գործիքներով, սա է կլոնավորելու ճարտարապետությունը՝ [\\$6](#)-ի այն վերապահումով, որ ելակետային մոդելները կարող են թերի կարգավորված լինել:

CLAP-ը աշխարհի մոդելները դարձնում է ստուգելի առանց կլաստերի: 3060-ի վրա 12 GB-ից պակաս VRAM-ը «կարդալ հոդվածը» և «գործարկել այն» տարբերությունն է: Եթե ուսուցված դինամիկայի մոդելը նախկինում բյուջեից դուրս էր, այս շաբաթ այդպես չէ:

VISTA-ն վերացնում է սարքավորումային կախվածությունը: Եթե շոշափողական մատրիցը ձեր ցանկում էր միայն հսկայական ազդանշանը կառավարման շղթա բերելու համար, ճկուն բռնիչը և ձեր արդեն ունեցած տեսախցիկները կարող են բավարար լինել: Կող չի հրապարակվել, ուստի սա նախագծային ակնարկ է, ոչ թե ներբեռնելու բան, և այն պահանջում է բռնիչ, որի ձևավորությունը տեսանելի և կրկնելի է:

Այս շաբաթ ոչինչ չի առնչվում տվյալների հավաքման սարքերին կամ հեռակառավարման (teleoperation) արժեքին:

8 Շուկայական ազդանշան

Որտեղ է տեսանելիորեն աճում աշխատակազմը՝ Generalist (ընդհանուր առմամբ ներգրավված \$600 մլն), 1X, եթե SoftBank-ի գործարքը կայանա, Locus (որը կլանում է Nexera-ի մանիպուլյացիայի թիմը), և NSF-ի ֆինանսավորմամբ նոր կենտրոն UT Austin-ում: Հարկ է նշել, որ այս շաբաթ մանիպուլյացիային հատուկ պահանջարկը եկավ *լոգիստիկ* գնորդից, ոչ թե մարդանման ռոբոտների ընկերությունից:

Ուշադրության արժանացած աշխատանքներում կրկնվում են երեք տեխնիկական ուղղություններ: Հոսքային և ասինխրոն գործարկումը (inference) (StreamPI, [FlashVLA](#), քիչ քայլերով հոսքային նմուշառում) հանդիպում է երեք առանձին անգամ. ոլորտը «աշխատում է ռազմավարությունը» հարցից անցել է «աշխատում է այն կառավարման հաճախականությամբ» հարցին: Երկրորդը հպումն ու ուժն են որպես ռազմավարության առաջնային մոտեցեր, ոչ թե անվտանգության թաղանթներ (VISTA, TacForcing): Երրորդը տարբեր մարմնավորումների միջև թաքնված գործողությունների միջերեսներն են (CLAP, տեսախցիկակենտրոն նախաուսուցում):

Պարտադիր նվազագույնը, այն իմաստով, որ դրանք այլևս ոչինչ չեն տարբերակում՝ LIBERO $\geq 95\%$, գործողությունների փաթեթավորում, հոսքային կամ դիֆուզիոն գործողությունների գլուխներ, հոդվածը ներկայացնելիս հղված պահոց: Դեռևս սակավ են՝ պարամետրերով համապատասխանեցված վերահսկվող աբլյացիաները (գրեթե ոչ ոք դրանք չի անում, և հենց դրանց շնորհիվ է PredVLA-ն հասկանալի), իրական ռոբոտի վրա գնահատումները՝ նշված փորձերի քանակով, և որևէ ապացույց, որ հեղինակները գիտեն, թե որ հենանիշներն են ավելորդ:

Այս շաբաթ վարձատրության կամ պաշտոնային մակարդակների վերաբերյալ վստահելի տվյալներ չհայտնվեցին: «robotics engineer salary 2026» որոնման համար բարձր դիրքեր զբաղեցնող ագրեգատոր էջերը SEO բովանդակություն են՝ առանց նշված մեթոդաբանության. ավելի լավ է ոչինչ չհաղորդել, քան դրանք «լվանալ»:

9 Թեմաներ

Առաջին թողարկում. ցանկը նոր է սկսվում:

ԹԵՄԱ	ԿԱՐԳԱՎԻՃԱԿ	ԻՆՉԸ ԿՀԱՄԱՐՎԵՐ ՀԱՋՈՐԴ ԻՐԱԿԱՆ ԹԱՐՄԱՑՈՒՄ
Արդյո՞ք VLA-ի մասշտաբն է ապահովում լեզվով պայմանավորված կառավարումը	ԱՌԱՋԸՆԹԱՏ (PredVLA)	10 մլն-ից փոքր ռազմավարություն <i>իրական</i> սարքավորման վրա, բազմաառաջադրանքային, հրապարակված ելակետային մոդելի, ոչ թե ներքին մոդելի համեմատ
Հալման զգայում առանց շոշափողական սարքավորման	ԱՌԱՋԸՆԹԱՏ (VISTA, TacForcing)	Անկախ լաբորատորիա, որը վերարտադրում է տեսողական ձևափոխության զգայումն այլ բռնիչի վրա, կամ հրապարակված տեխնիկական բնութագրեր
Աշխարհ-գործողություն մոդելները որպես միասնական ռազմավարություն + սիմուլյատոր	ԱՌԱՋԸՆԹԱՏ (Riemann-1.0, CLAP, WALL-SS)	Ինչ-որ մեկը հրապարակված աշխարհի մոդելն օգտագործում է ռազմավարությունը բարելավելու համար, ոչ միայն տեսանյութ կանխատեսելու
Հայտարարված բաց թողարկումներն ընդդեմ իրական արտեֆակտների	ՎԻՃԵԼԻ (StreamPI՝ սպասվում է, Riemann՝ չստուգված, Muse՝ առանց քարտի)	Կշիռների հայտնվելը, և որ ինչ-որ մեկը դրանք գործարկի
Մարդանման ռոբոտների կապիտալն ընդդեմ ցուցադրված կարողության	ԱՌԱՋԸՆԹԱՏ (SoftBank/1X, չհաստատված)	1X-ի գործարքի հաստատումը կամ ձախողումը, մարդանման ռոբոտների ընկերություն, որը հրապարակում է մեթոդ, ոչ թե տեսանյութ
Մանիպուլյացիան գնվում է, ոչ թե կառուցվում	ԱՌԱՋԸՆԹԱՏ (Locus/Nexera)	Լոգիստիկայի երկրորդ հաստատված խաղացողը, որն այս եռամսյակում գնում է վերջնական գործիքի կամ ռազմավարությունների թիմ
Հենանիշների վավերականությունը ֆիզիկական ԱԲ-ում	ՆՈՐ (ավելորդության աուդիտ)	Խոշոր լաբորատորիա, որն ընդունում է հենանիշների կրճատված հավաքածու, կամ LIBERO-ն դուրս է գալիս հիմնական աղյուսակներից
Էժան բաց սարքավորումը որպես ուսուցման հենք	ՆՈՐ (Microduck, Poppy LQR)	Microduck-ի վաճառքի մեկնարկը, և երրորդ կողմը, որը դրա վրա ռազմավարություն է ուսուցանում
Գործարկման ուշացումը որպես կիրառման նեղ տեղ	ԱՌԱՋԸՆԹԱՏ (StreamPI, FlashVLA, MeanFlow)	Ուշացումը հաղորդվում է հաջողության մակարդակի կողքին՝ որպես կանոն
ԱՅԼ ՈՒՇԱԳՐԱՎ ՆՅՈՒԹԵՐ		
GaussVLA — Mamba հիմնական ցանց, որը մշակում է 3D գաուսյան թոքեններ: Ճարտարապետական դիտարկում՝ առանց գաուսյանները մեկուսացնող գնահատման: GaussianDream++ — աշխարհի վիճակի կոմպակտ գաուսյան թոքեններ VLA-ի ներսում: Համընկնում է CLAP-ի հետ, բայց առանց հրապարակման: Zero-WAM — համատեքստային նմանակում մարդու մեկ տեսանյութից: Հետաքրքիր պնդում, կող չկա, չվերարտադրված: R3 — RL հետուսուցում ազատ ձևի լեզվական դատողության համար, որն ուղղորդում է ցածր մակարդակի ռազմավարությունը: RA-VLA — որոնում՝ թեստի ժամանակ հարմարվելու համար: Որոնման կորպուսը նկարագրված չէ: One Policy, Many Embodiments — տեսախցիկակենտրոն գործողությունների երկրաչափություն՝ տարբեր մարմնավորումների միջև նախաուսուցման համար: Նույն տարածքը, ինչ CLAP-ը:		
LM-X — առաջխաղացման, իրադարձությունների և անորոշության գլուխներ VLA-ի վրա: Մեկնաբանելիության պնդում՝ առանց որոշումների համար նշանակալի գնահատման: Memory Anchors — շարունակական ուսուցում առանց ամբողջական վերաուսուցման: Առաջադրանքների փոքր խումբ: TemporalFlow-VLA — ժամանակային հոսքը ներկայացնում է որպես կատարման պատմություն: Մոտ է StreamPI-ին, ավելի թույլ գնահատմամբ: GRAFT — քեշավորված տեսողություն-լեզու նախաձեռն է ժանացնում է VLA-ի առցանց հարմարեցումը: Արժե վերադառնալ, եթե կողը հրապարակվի: FLARE — ձախողումների հայտնաբերում և վերականգնում՝ երկարաժամկետ VLA կատարման շուրջ: WALL-SS — հաջորդ մասշտաբի ավտոռեգրեսիա՝ աշխարհի մոդելի ավելի երկար փորձարկումների (rollouts) համար: 4DSynth — լեզվից կամ լուսանկարից՝ ֆիզիկայով խմբագրելի 4D տեսարաններ:		

Figure: Index — սմարթֆոններով խմբային (crowdsourced) տվյալների հավաքման խողովակաշար, խոստացված \$1 մլրդ, ոչինչ չի հրապարակվել:

Jetson Orin Nano 2 — 78 TOPS, 8 GB, 2× արագ գործարկում 40%-ով պակաս էներգիայով: Վաճառքը՝ 2027 թ. 1-ին կիսամյակում, ուստի դեռ ոչինչ չի փոխում:

Anthropic Model Hardware Standard — MCP-ի միջոցով դրայվերային պրիմիտիվներ՝ լաբորատոր մանիպուլյատորներ կառավարող գործակալների համար: Հետազոտական նախադիտում, ոչինչ չի հրապարակվել:

10 Բաց հարց

PredVLA-ի 4–7× առավելությունը կանխատեսողական կողավորման օգտին է վկայում, թե՛ այն բանի, որ Transformer-ները և գործողությունների փաթեթավորումը պարզապես չեն աշխատում մեկ միլիոն պարամետրից ցածր: Հոդվածը չի կարող տարբերել այս երկուսը. յուրաքանչյուր ելակետային մոդել նրա սեփական իրականացումն է մի ռեժիմում, որի համար ոչ ոք օպտիմալացում չի անում: Հարցը կլուծի կամ անկախ խմբի կողմից ճիշտ կարգավորված 1 մլն-ից փոքր Transformer-ը, կամ այն, որ PredVLA-ն պահպանի իր առավելությունը 10–100 մլն պարամետրի դեպքում, որտեղ Transformer ելակետային մոդելները հայտնի են որպես լավ աշխատող: Երկրորդը էժան է, և հեղինակները դա չեն արել:

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