



The CLEAN PAPER

WHAT THE PAPER SAYS · WHAT IT DOESN'T · WHY IT MATTERS

Je, “foundation models” kubwa za roboti kweli hufanya vizuri zaidi? Jibu la makini — ndiyo kwa kiasi, na tafiti nyingi haziwezi kutofautisha

Toyota Research Institute ilifundisha “large behavior models” — sera za roboti zilizopata pretraining kwa ~saa 1,700 za data mbalimbali za manipulation — na kuzilinganisha na sera za kazi moja za from-scratch kwa ukali usio wa kawaida: majaribio ya blind, randomized na sampuli kubwa (~1,800 halisi, 47,000+ simulation) pamoja na takwimu sahihi. Baada ya finetuning kwa kila kazi, modeli kubwa zilifanya vizuri zaidi kwa wastani, zilihitaji karibu mara 3–5 data kidogo maalum kwa kazi, na zilikuwa imara zaidi hali zilipobadilika; utendaji uliongezeka kwa ulaini na data zaidi ya pretraining. Lakini bila finetuning hazikushinda kwa uthabiti modeli za kazi moja, athari kadhaa zilikuwa ndogo kiasi cha kuhitaji sampuli kubwa ili zionekane, na chaguo la kawaida la data normalisation lilikuwa muhimu kuliko architecture. Ni msaada wa kiasi kwa mwelekeo wa robot foundation models — si roboti ya matumizi ya jumla, si zero-shot generalist, si “emergent leap” — pamoja na onyo kwamba robotiki nyingi huenda zinapima kelele.

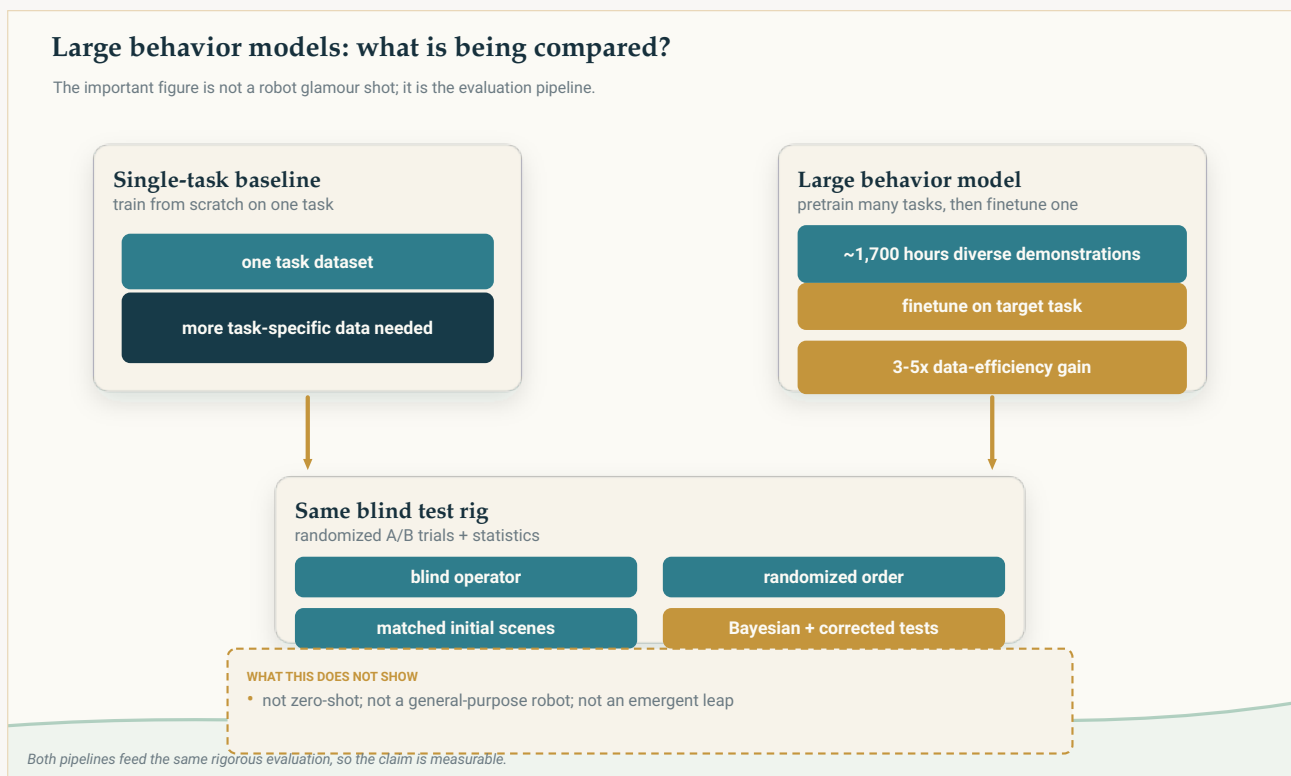
Written by Lucio Vaglio · Cross-checked by Laura Nesso · Edited by Michele Renda · Figures by Laura Nesso and Nova Vela · AI-assisted, human-reviewed

Paper: *A Careful Examination of Large Behavior Models for Multitask Dexterous Manipulation*, Toyota Research Institute Large Behavior Model Team — J. Barreiros, A. Beaulieu, et al.; senior authors incl. R. Ambrus, B. Burchfiel, S. Feng, H. Kress-Gazit (Cornell), R. Tedrake, Science Robotics (2026); preprint arXiv:2507.05331 · DOI: 10.1126/scirobotics.aea6201

Makala ya roboti ambayo mada yake halisi ni uaminifu

Roboti ziko katikati ya mbio za “foundation modeli.” Wazo, lililokopwa kutoka AI ya lugha na picha, linavutia: badala ya kufundisha roboti kazi moja baada ya nyingine, fundisha **“large behavior modeli” (LBM)** moja kubwa kwa mkusanyiko mkubwa na tofauti wa maonyesho, upate mfumo wenye uwezo mpana na unaoweza kuzoea haraka. Msisimko — na uwekezaji — ni mkubwa sana. Kichwa cha habari kinajiandika chenyewe: *akili za roboti za matumizi ya jumla zimefika.*

Makala hii, kutoka Toyota Research Institute, inavutia hasa kwa sababu inakataa kuandika kichwa hicho. Mchango wake halisi si roboti yenye kuvutia zaidi. Ni uchunguzi mgumu wa swali linaloonekana rahisi kwa udanganyifu — *je, mbinu ya modeli kubwa kweli hufanya kazi vizuri zaidi, na tungejuaje?* — ukipewa jibu kwa umakini wa takwimu ambao, kwa maelezo ya waandishi wenyewe, si wa kawaida katika taaluma hii. Matokeo ni “ndiyo, lakini” yenye manufaa kweli, na onyo kwamba sehemu kubwa ya robotiki inaweza kuwa inapima kelele.



Mbinu zote mbili zinaingia katika tathmini ileile iliyofichwa na ya mpangilio nasibu. Dai la makala si kwamba foundation modeli za roboti ni generalists za zero-shot, bali kwamba mafunzo ya awali inaweza kuboresha ufanisi wa data na uthabiti inapopimwa kwa makini.

Original The Clean Paper diagram CC BY 4.0

Walichofanya waandishi

Walijenga LBM kwa maana maalum na halisi: **sera za visuomotor zinazotegemea diffusion** (Diffusion Transformer inayosoma picha za kamera, agizo fupi la lugha, na nafasi za viungo vya roboti yenyewe, kisha kutoa milipuko mifupi ya amri za mwendo kwa 10 Hz). Zilipewa **mafunzo ya awali kwa takriban saa 1,700** za maonyesho ya roboti — zaidi ya kazi 500 tofauti zilizokusanywa ndani ya taasisi, pamoja na hifadhidata za umma — kisha zikafanyiwa **mafunzo maalum ya mwisho** kwa kazi binafsi. Ulinganisho wote ni dhidi ya **sera ya kazi moja iliyofundishwa kutoka mwanzo** kwa data ya kazi hiyo moja.

Lakini moyo wa makala ni *tathmini*, ambayo waandishi wanaitazama kama matokeo makuu. Ili wasi-jidanganye walitumia:

- **Upimaji A/B wa blind na randomised** katika ulimwengu halisi — mtu anayeendesha roboti haku-jua ni sera ipi ilikuwa inajaribiwa, na mpangilio uliwekwa nasibu.
- **Hali za mwanzo zilizodhibitiwa na zinazoweza kurudiwa** — waendeshaji walilinganisha eneo na overlay ya picha kabla ya kila jaribio.
- **Idadi kubwa ya majaribio** — rollouts 50 za ulimwengu halisi kwa kila kazi, sera na hali; 200 kwa kila kazi katika uigaji. Kwa jumla: takriban **rollouts 1,800 za blind katika ulimwengu halisi na zaidi ya 47,000 katika uigaji**.
- **Takwimu zinazofaa** — makadirio ya Bayesian ya uwezekano wa mafanikio, na majaribio ya hypothesis ya jozi yenye marekebisho ya multiple ulinganisho, badala ya kuangalia tu error bars zinazopishana. Hata walifanya QA kwenye robo ya majaribio yaliyopimwa na binadamu ili kupima kosa la scoring.

[video pending]

Ulinganisho wa modeli mbili zikifanya maandalizi ya meza ya kifungua kinywa: (kushoto) baseline ya kazi moja, na (kulia) LBM. Video zote mbili zinachezwa kwa kasi ya 1x. Hii ni kazi moja iliyotathminiwa, si ushahidi wa autonomy ya matumizi ya jumla.

Mfumo huu wa tathmini ndio hoja. Makala nzima ni hoja kwamba bila mfumo huu, huwezi kuto-fautisha uboreshaji halisi na bahati.

Walichogundua

Modeli kubwa zilizofanyiwa mafunzo maalum ya mwisho zilishinda modeli za kazi moja zilizofundishwa kutoka mwanzo — kwa wastani. Zikiunganishwa katika kazi mbalimbali, LBM iliyopata mafunzo ya awali na kisha mafunzo maalum ya mwisho kwa kazi fulani ilifanya vizuri zaidi kwa kuaminika kuliko sera iliyofundishwa kutoka mwanzo kwa kazi hiyo, katika uigaji na ulimwengu halisi, na tofauti ilikuwa na umuhimu wa takwimu. Kwa kazi binafsi, LBM iliyofanyiwa mafunzo maalum ya mwisho ilikuwa takwimu sawa au bora kuliko iliyofundishwa kutoka mwanzo karibu kila mara (3/3 kazi za ulimwengu halisi, 15/16 za uigaji).

Ushindi mkubwa na wazi zaidi ni ufanisi wa data. LBM iliyofanyiwa mafunzo maalum ya mwisho ilifikia utendaji sawa na iliyofundishwa kutoka mwanzo kwa kutumia karibu **mara 3–5 data kidogo maalum kwa kazi**. Katika kazi moja halisi (kuandaa meza ya kifungua kinywa), LBM iliyofanyiwa mafunzo maalum ya mwisho kwa **15%** tu ya maonyesho ilishinda sera ya iliyofundishwa kutoka mwanzo iliyofundishwa kwa **100%** ya maonyesho.

mafunzo ya awali husaidia zaidi hali zinapobadilika. Mazingira ya majaribio yalipobadilishwa kwa makusudi kutoka hali za mafunzo (“mabadiliko ya mgawanyo”), faida ya LBM iliyofanyiwa mafunzo maalum ya mwisho *iliongezeka*. Katika seti moja ya uigaji, ilishinda iliyofundishwa kutoka mwanzo kwa umuhimu wa takwimu kwenye kazi 3 kati ya 16 katika hali za kawaida lakini 10 kati ya 16 chini ya mabadiliko ya mgawanyo. Kwa kuwa matumizi halisi daima huondoka kidogo kutoka hali za mafunzo, uthabiti huu huenda ndio matokeo muhimu zaidi kwa matumizi.

Data zaidi ya mafunzo ya awali ilisaidia, hatua kwa hatua. Utendaji ulipanda kwa utaratibu walipoongeza data ya mafunzo ya awali — bila **kuruka ghafla au “emergent” leap** katika viwango walivyopima. Ni muhimu, kinachotabirika, na kisicho na drama.

Lakini simulizi la generalist bila mafunzo maalum ya mwisho halikushikilia. LBM iliyopata mafunzo ya awali ikitumiwa *zero-shot* — bila mafunzo maalum ya mwisho maalum kwa kazi — *haikushinda kwa uthabiti* sera za kazi moja. Mtandao mmoja uliweza kufanya kazi nyingi, lakini ndoto ya “mpe tu prompt” haikuthibitishwa hapa; waandishi wanahusisha sehemu ya tatizo na udhaifu wa encoder yao ndogo ya lugha.

Na faida zilikuwa ndogo kiasi cha kuwa rahisi kukosa — au kujidanganya nazo. Athari nyingi zilionekana tu kwa sampuli kubwa kuliko kawaida na majaribio makini. Waandishi wanasema wazi kwamba, kutokana na ukubwa wa athari na kelele, **kuna hatari kubwa kwamba makala nyingi za robotiki zinapima kelele za takwimu**. Pia waligundua kwamba chaguo la kawaida kabisa — namna data inavyofanyiwa *normalisation* — liliathiri matokeo zaidi kuliko mabadiliko ya architecture, na kwamba **bug** ya normalisation katika mafunzo ya awali iligunduliwa tu *baada* ya tathmini kukamilika.

Huenda hii inamaanisha nini

Tafsiri inayoweza kutetewa: mafunzo ya awali ya kiwango kikubwa kwa data mbalimbali za roboti ni kiungo halisi chenye thamani — inapunguza data inayohitajika kwa kila kazi mpya na kufanya sera ziwe imara zaidi wakati dunia hailiani na mafunzo. Hilo kwa kweli linaunga mkono mwelekeo ambao taaluma inawekeza ndani yake. Lakini faida ni *za wastani na zenye masharti* (zinaonekana hasa baada ya mafunzo maalum ya mwisho, na wazi zaidi zikiunganishwa na chini ya stress), si kuwasili kwa roboti general ya kuingiza tu na kutumia.

Maana tulivu lakini muhimu zaidi ni ya mbinu. Makala ni, kwa vitendo, fimbo ya kupimia: inaonyesha ni ushahidi kiasi gani unahitajika ili kutoa dai la kuaminika kuhusu sera ya roboti, na inadokeza kwamba msisimko mwingi uliochapishwa unategemea ushahidi mdogo sana. Taaluma inahitaji marekebisho hayo zaidi kuliko modeli nyingine.

Hili halithibitishi nini

- **Si roboti ya matumizi ya jumla.** Ushindi umeonyeshwa kwa architecture maalum (diffusion policies) iliyofanyiwa mafunzo maalum ya mwisho kwa kila kazi, katika mazingira yaliyodhibitiwa, kutoka maonyesho ya teleoperation — si roboti huru inayofanya kazi yoyote mpya kwa amri.
- **Haithibitishi matumizi ya zero-shot.** Bila mafunzo maalum ya mwisho, modeli kubwa haikushinda kwa uthabiti baseline za kazi moja.
- **Si ushahidi wa “emergent leap.”** Scaling iliboresha mambo kwa ulaini; hakuna discontinuity inayounga mkono simulizi ya “kisha ghafla ikawa na uwezo.”
- Namba ni **za kulinganisha na za maabara.** Viwango kamili vya mafanikio vilipangwa kwa makusudi kuwa karibu ~50% ili kuongeza sensitivity ya ulinganisho; si kipimo cha reliability ya dunia halisi, na kazi ni architecture moja kutoka maabara moja.
- **Haijibu kwa nini** sera fulani hufaulu au hushindwa, na kazi kadhaa mahususi ambapo modeli kubwa ilifanya *vibaya zaidi* zinaripotiwa lakini hazijaelezwa.
- Haisemi **chochote kuhusu usalama, autonomy au deployment** nje ya mfumo wa tathmini.

Ushahidi una nguvu kiasi gani?

Kwa madai yake makuu ya ulinganisho — LBM zilizofanyiwa mafunzo maalum ya mwisho hushinda baseline za iliyofundishwa kutoka mwanzo zikiunganishwa, zinahitaji data mara kadhaa kidogo, na ni imara zaidi chini ya mabadiliko ya mgawanyo — ushahidi ni wenye nguvu *na* umedhibitiwa vizuri isivyo kawaida: blind, randomised, sampuli kubwa, majaribio ya takwimu, pamoja na QA ya scoring. Hii ni kesi adimu ambapo methodology ni nzuri kiasi cha kuchukua hitimisho kuu karibu na thamani yake halisi.

Tahadhari za uaminifu ni zile ambazo waandishi wenyewe wanaibua. Error bars zao hushika randomness ya *tathmini* lakini **si** randomness ya *mafunzo* — fundisha modeli ileile mara mbili na unaweza kupata sera tofauti kwa maana, na tofauti hiyo haimo katika takwimu. Kazi za dunia halisi zilikuwa na majaribio 50 kila moja, ya kutosha kuona athari za wastani lakini yanaweza kukosa ndogo. lugha-conditioning ilitumia encoder ya kiasi, hivyo madai ya “mwambie tu roboti cha kufanya” yanaweza kubadilika kwa mifumo mikubwa. Na kuna ufichuzi wa wazi wa bug ya normalisation iliyogunduliwa baada ya tathmini. Hakuna kati ya haya kinachoangusha matokeo makuu, lakini ndio aina ya mambo ambayo makala inasema taaluma mara nyingi husukuma pembeni.

Dokezo la chanzo kwa roho hiyo: maelezo haya yanategemea preprint ya waandishi. Hatukuweza kupata toleo lililochapishwa katika jarida, kwa hiyo hatujakagua kama kuna mabadiliko kati ya preprint na maandishi yaliyochapishwa.

Kwa nini ni muhimu

“Robot foundation modeli” ni kauli iliyotengenezwa kwa ajili ya madai yaliyopitiliza, na utafiti kama huu ni rahisi kuusoma vibaya pande zote — kama ushindi wa “*inafanya kazi!*” au kejeli ya “*imehupewa.*” Tafsiri sahihi ni yenye manufaa zaidi kuliko zote: mafunzo ya awali kwa data mbalimbali hutoa faida halisi, zinazopimika lakini za wastani — hasa data kidogo kwa kila kazi na uthabiti zaidi — na njia hiyo inaboreka kwa namna inayotabirika kadiri scale inavyoongezeka.

Sababu ya kina zaidi ya umuhimu ni kwamba makala inaelekeza ukali wake kwa taaluma yake yenyewe. Kwa kuonyesha kwamba athari halisi ni ndogo kiasi cha kutoweka chini ya tathmini dhaifu, na kwamba chaguo la kawaida kama data normalisation linaweza kuwa muhimu kuliko architecture mpya yenye akili, inaonyesha kwamba maendeleo mengi ya robot learning yanahitaji vipimo imara zaidi kabla ya kuaminiwa. Makala inayotumia uaminifu wake kulinda tofauti kati ya matokeo na tamaa inafanya jambo adimu zaidi, na lenye thamani zaidi, kuliko kuongoza leaderboard.

Muhtasari safi

Watafiti wa Toyota Research Institute walifundisha **large behavior models (LBM)** — sera za roboti za diffusion zilizopata mafunzo ya awali kwa takribani saa 1,700 za data kutoka kazi nyingi za kuendesha vitu — kisha wakazilinganisha na sera za kazi moja zilizofundishwa kutoka mwanzo. Itifaki ilikuwa kali isivyo kawaida: tathmini zisizojua modeli iliyotumika, ugawaji wa nasibu, sampuli kubwa, takribani majaribio 1,800 katika dunia halisi na zaidi ya 47,000 katika uigaji, pamoja na uchanganuzi rasmi wa takwimu. Baada ya mafunzo maalum ya mwisho kwa kila kazi, LBM zilifanya vizuri zaidi kwa jumla, zilihitaji takribani **mara 3–5 data kidogo maalum kwa kazi**, na zilikuwa **imara zaidi mazingira yalipobadilika**; utendaji pia uliongezeka kwa ulaini kadiri data ya mafunzo ya awali ilivyoongezeka. Lakini bila mafunzo maalum ya mwisho hazikushinda kwa uthabiti modeli za kazi moja, baadhi ya athari zilikuwa ndogo kiasi kwamba sampuli kubwa tu ndizo zilizozionyesha, na namna ya kurekebisha data ilikuwa muhimu kuliko usanifu wa modeli. Huu ni ushahidi mzuri lakini wa kiasi kwa wazo la modeli za msingi za roboti — si roboti ya matumizi ya jumla, si mfumo wa zero-shot unaoweza kufanya kila kitu, na si kuruka kwa ghafla kwa uwezo.

Ukaguzi bila kupamba

Makala inaonyesha nini: Kwa tathmini kali, blind na yenye nguvu ya takwimu ($\approx 1,800$ rollouts za dunia halisi + 47,000+ za uigaji), diffusion policies zilizopata multitask mafunzo ya awali kisha mafunzo maalum ya mwisho (LBM) hushinda sera za kazi moja za iliyofundishwa kutoka mwanzo kwa jumla, hufikia utendaji sawa kwa $\sim$ mara 3–5 data kidogo maalum kwa kazi, na huwa imara zaidi chini ya mabadiliko ya mgawanyo; utendaji huongezeka kwa ulaini kadiri data ya mafunzo ya awali inavyoongezeka.

Kinachowezekana lakini hakijathibitishwa: Kwamba faida hizi zitahamia kwenye vision-lugha-

action modeli kubwa zaidi (encoder yao ya lugha ilikuwa ndogo); kwamba scaling laini itaendelea nje ya kiwango cha data kilichopimwa.

Isichoonyeshwa: Roboti ya matumizi ya jumla au zero-shot (hakuna mafunzo maalum ya mwisho → hakuna faida ya uthabiti); kuruka kwa ghafla kwa uwezo “emergent”; maelezo ya sababu za kushindwa kwa kazi maalum; reliability kamili ya dunia halisi (success rates zilipangwa karibu 50% kwa sensitivity); chochote kuhusu usalama au deployment huru.

Vikwazo vikuu: Takwimu hushika randomness ya tathmini lakini si ya training run; majaribio 50 ya dunia halisi kwa kila kazi yanaweza kukosa athari ndogo; architecture moja na maabara moja; encoder ya lugha ya kiasi; bug ya data normalisation iligunduliwa baada ya tathmini; uchambuzi unatokana na preprint (toleo lililochapishwa halikukaguliwa).

Msomaji wa kawaida anapaswa kuwa na uhakika kiasi gani? Uhakika mkubwa kwamba multitask mafunzo ya awali pamoja na mafunzo maalum ya mwisho hutoa faida halisi za wastani — hasa ufanisi wa data na uthabiti — na kwamba zilipimwa kwa umakini usio wa kawaida. Uhakika mkubwa kwamba hii **si** roboti ya matumizi ya jumla au zero-shot na **si** emergent leap. Uhakika wa kati kuhusu faida zitakua kiasi gani katika modeli kubwa zaidi. Na onyo la waandishi linastahili kuchukuliwa kwa uzito: athari katika taaluma ni ndogo kiasi kwamba tafiti zenye nguvu ndogo zinaweza kuripoti kelele. Msimamo unaofaa: matumaini yenye kipimo kuhusu mbinu, na shaka yenye afya kwa matokeo ya robot-AI yasiyo na msaada wa takwimu wa aina hii.

Vyanzo

arXiv:2507.05331

<https://arxiv.org/abs/2507.05331>

Peer-reviewed version (not retrieved) — Science Robotics, 10.1126/scirobotics.aea6201

<https://doi.org/10.1126/scirobotics.aea6201>

Project page

<https://toyotaresearchinstitute.github.io/lbm1/>