



The CLEAN PAPER

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

Ali veliki robotski 'temeljni modeli' res delujejo bolje? Previden odgovor — nekoliko da, večina raziskav pa tega ne zna zanesljivo izmeriti

Toyota Research Institute je izuril 'large behavior models' — robotske politike, predhodno učene na približno 1.700 urah raznovrstnih manipulacij — in jih z nenavadno strogo metodologijo primerjal s politikami za posamezne naloge, učenimi od začetka: slepi, naključni, obsežni testi (~1.800 v resničnem svetu, več kot 47.000 v simulaciji) z ustrezno statistiko. Po dodatnem učenju za posamezno nalogo so veliki modeli v povprečju delovali bolje, potrebovali približno 3–5-krat manj podatkov za posamezno nalogo in bili robustnejši ob spremembi pogojev; uspešnost je gladko naraščala z več podatki za predhodno učenje. Brez dodatnega učenja pa niso dosledno prekašali modelov za eno nalogo, več učinkov je bilo tako majhnih, da so jih razkrili šele veliki vzorci, banalna odločitev o normalizaciji podatkov pa je vplivala bolj kot arhitektura. To je izmerjena podpora smeri robotskih temeljnih modelov — ne splošnonamenski robot, ne zero-shot generalist in ne 'emergentni skok' — ter ostro opozorilo, da velik del robotike morda meri šum.

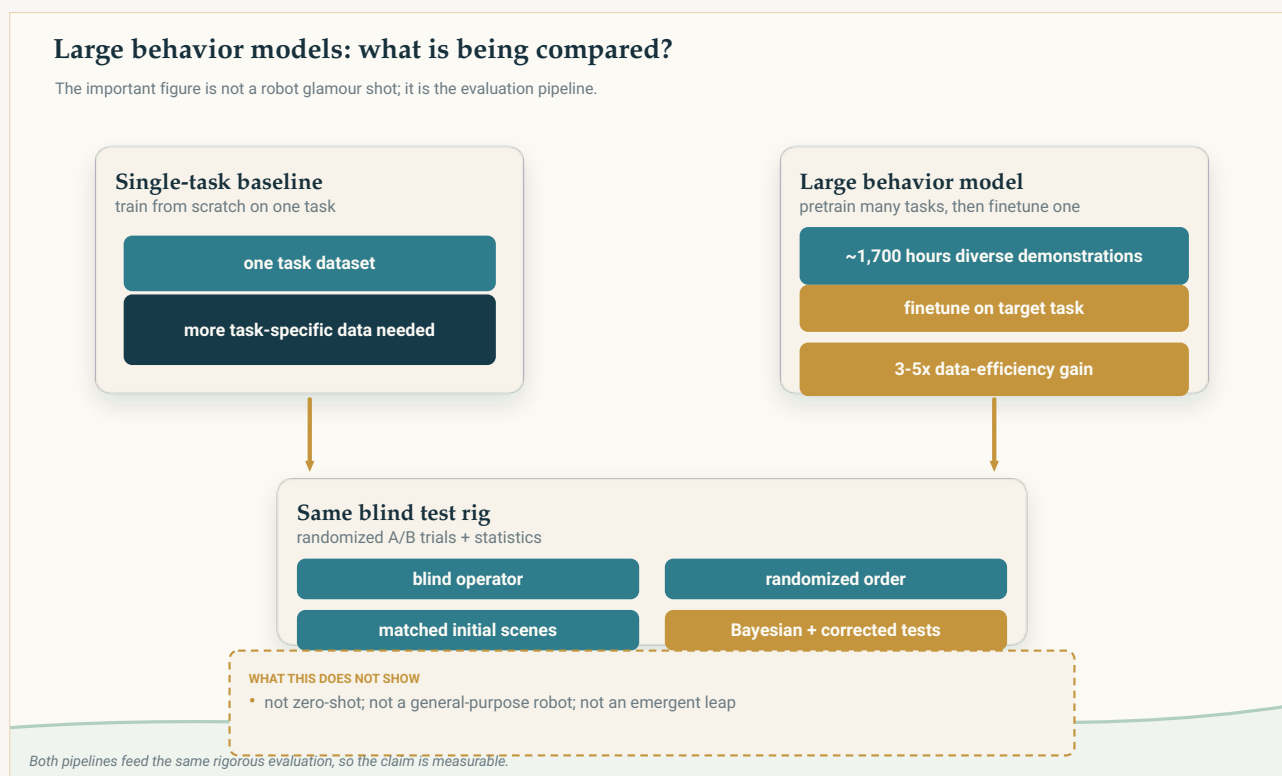
Written by Lucio Vaglio · Cross-checked by Laura Nesso · Edited by Michele Renda · Figures by Laura Nesso and Nova Vela · AI-assisted, human-reviewed

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

Članek o robotih, katerega prava tema je poštenost

Robotika je sredi vala »temeljnih modelov«. Zamisel, izposojena iz jezikovne in slikovne umetne inteligence, je mamljiva: namesto da bi robota učili vsako nalogo posebej, en velik »**large behavior model**« (LBM) naučimo na ogromni in raznoliki zbirki demonstracij ter dobimo sistem, ki zna veliko in se hitro prilagaja. Navdušenje — in naložbe — so ogromni. Naslov se skoraj napiše sam: *splošnonamenski robotski možgani so tu*.

Ta članek Toyota Research Institute je zanimiv prav zato, ker tega naslova noče napisati. Njegov glavni prispevek ni bolj bleščeč robot, ampak trd pogled na varljivo preprosto vprašanje — *ali pristop z velikim modelom res deluje bolje in kako bi to sploh zanesljivo vedeli?* — z ravno statistične skrbnosti, ki je po besedah avtorjev v področju nenavadna. Rezultat je resnično uporaben »da, vendar« in opozorilo, da velik del robotike morda meri šum.



Oba pristopa se končata v istem slepem, naključno razporejenem vrednotenju. Članek ne trdi, da so robotski temeljni modeli zero-shot generalisti, ampak da lahko predhodno učenje ob skrbnem merjenju izboljša podatkovno učinkovitost in robustnost.

Original The Clean Paper diagram CC BY 4.0

Kaj so naredili avtorji

LBM so zgradili v zelo konkretnem pomenu: kot **difuzijske vizualno-motorične politike** (Diffusion Transformer, ki bere slike kamer, kratko jezikovno navodilo in položaje robotovih sklepov ter pri

10 Hz izpisuje kratke nize motoričnih ukazov). Modele so **predhodno učili na približno 1.700 urah** robotskih demonstracij — več kot 500 različnih nalog, zbranih v laboratoriju, plus javni podatkovni nabori — nato pa jih **dodatno učili** za posamezne naloge. Primerjava je ves čas s **politiko za eno nalogo, učeno od začetka** samo na podatkih te naloge.

Toda srce članka je *vrednotenje*, ki ga avtorji obravnavajo skoraj kot glavni rezultat. Da se ne bi sami zavedli, so uporabili:

- **Slepo, naključno A/B-testiranje** v resničnem svetu — človek, ki je upravljal poskus, ni vedel, katera politika se testira, vrstni red pa je bil naključen.
- **Nadzorovane in ponovljive začetne pogoje** — izvajalci so prizor pred vsakim poskusom poravnali z grafičnim prekrivanjem referenčne slike.
- **Veliko število poskusov** — 50 izvedb v resničnem svetu na nalogo, politiko in pogoj; 200 na nalogo v simulaciji. Skupaj približno **1.800 slepih izvedb v resničnem svetu in več kot 47.000 simulacijskih izvedb**.
- **Ustrezno statistiko** — Bayesove ocene verjetnosti uspeha in parne hipotezne teste s popravki za večkratne primerjave, namesto vizualnega primerjanja prekrivajočih se intervalov napake. Izvedli so celo kontrolo kakovosti četrte poskusov, ki so jih ocenjevali ljudje, da bi izmerili napako ocenjevanja.

[video pending]

Vzporedna primerjava modelov pri pripravi mize za zajtrk: levo je osnovna politika za eno nalogo, desno LBM. Oba videa se predvajata z običajno hitrostjo 1×. To je ena ocenjena naloga, ne dokaz splošnonamenske avtonomije.

Ta metodologija je bistvo. Celoten članek je argument, da brez nje ni mogoče vedeti, ali je izboljšanje resnično ali le sreča.

Kaj so ugotovili

Veliki modeli po dodatnem učenju v povprečju prekašajo modele za eno nalogo, učene od začetka. Ko so rezultate združili čez naloge, je LBM, ki je bil predhodno učen in nato dodatno naučen za posamezno nalogo, zanesljivo prekašal politiko, učeno od začetka na isti nalogi, tako v simulaciji kot v resničnem svetu, razlika pa je bila statistično značilna. Pri posameznih nalogah je bil dodatno naučeni LBM statistično enako dober ali boljši skoraj povsod (3/3 nalog v resničnem svetu, 15/16 v simulaciji).

Največja in najjasnejša prednost je podatkovna učinkovitost. Dodatno naučeni LBM je dosegel uspešnost primerljivo s politiko, učeno od začetka, z okoli **3–5-krat manj podatki za posamezno nalogo**. Pri eni nalogi v resničnem svetu — pripravi mize za zajtrk — je LBM, dodatno naučen na samo **15 %** demonstracij, prekašal politiko, učeno od začetka na **100 %** podatkov.

Predhodno učenje najbolj pomaga, ko se pogoji spremenijo. Ko so testno okolje namenoma odmaknili od pogojev učenja (*distribution shift*), se je prednost dodatno naučenega LBM **povečala**. V

enem simulacijskem naboru je pod običajnimi pogoji statistično prekašal model od začetka pri 3 od 16 nalog, ob premiku porazdelitve pa pri 10 od 16. Ker se resnična okolja vedno nekoliko razlikujejo od učnih, je ta robustnost morda najbolj praktično pomembna ugotovitev.

Več podatkov za predhodno učenje je pomagalo — gladko. Uspešnost je enakomerno naraščala z dodatnimi podatki za predhodno učenje, **brez nenadnega skoka ali »emergentnega« preskoka** pri preizkušanih velikostih. Uporabno, predvidljivo, nedramatično.

Toda zgodba o generalistu brez dodatnega učenja ni držala. Predhodno naučeni LBM, uporabljen *zero-shot* — brez dodatnega učenja za nalogo — **ni** dosledno prekašal politik za eno nalogo. Enotna mreža je lahko izvajala veliko nalog, vendar se sanje »samo povej robotu, kaj naj naredi« tu niso potrdile; avtorji del tega pripisujejo omejitvam njihovega majhnega jezikovnega enkoderja.

In prednosti so bile dovolj majhne, da jih je zlahka spregledati — ali navidezno ustvariti. Mnogi učinki so postali vidni šele pri večjih od običajnih vzorcih in skrbnih statističnih testih. Avtorji izrecno napišejo, da ob takšni velikosti učinkov in šumu obstaja **resno tveganje, da številni članki v robotiki merijo statistični šum**. Ugotovili so tudi, da je banalna izbira — kako so podatki *normalizirani* — vplivala bolj kot spremembe arhitekture, napako v normalizaciji pri predhodnem učenju pa so odkrili šele *po* končanem vrednotenju.

Kaj to verjetno pomeni

Utemeljena razlaga: obsežno predhodno učenje na raznovrstnih robotskih podatkih je res koristna sestavina — zmanjša količino podatkov, potrebnih za novo nalogo, in naredi politike robustnejše, ko svet ne ustreza učnim pogojem. To res podpira smer, v katero področje vlaga. Vendar so prednosti *zmerne in pogojne* (večinoma se pokažejo po dodatnem učenju, najjasnejše pa so v agregatu in pod obremenitvijo), ne prihod robota, ki ga lahko preprosto uporabimo za karkoli.

Tišji in morda pomembnejši pomen je metodološki. Članek je pravzaprav merilna palica: pokaže, koliko dokazov je dejansko potrebnih za zaupanja vredno trditev o robotski politiki, in namigne, da velik del objavljenega navdušenja temelji na premalo meritvah. Področje takšen popravek potrebuje bolj kot še en model.

Česa to ne dokazuje

- To **ni splošnonamenski robot**. Prednosti so pokazane za konkretno arhitekturo (difuzijske politike), dodatno učeno za posamezne naloge, v nadzorovanih pogojih in iz teleoperiranih demonstracij — ne za avtonomnega robota, ki na ukaz opravlja poljubna nova dela.
- **Ne** potrjuje **zero-shot** uporabe. Brez dodatnega učenja veliki model ni dosledno prekašal osnovnih modelov za posamezne naloge.
- To **ni** dokaz **»emergentnega skoka«**. Skaliranje je izboljševalo rezultate gladko; tu ni diskontinuitete, ki bi podprla pripoved »in potem je nenadoma postal sposoben«.

- Številke so **relativne in laboratorijske**. Absolutne stopnje uspeha so bile namenoma nastavljene blizu ~50 %, da so bile primerjave občutljive; niso merilo zanesljivosti v resnični uporabi. Gre za eno arhitekturo iz enega laboratorija.
- **Ne** pojasni, **zakaj** določena politika uspe ali spodleti; poročani so tudi konkretni primeri, kjer je veliki model deloval slabše, brez razlage vzroka.
- O **varnosti, avtonomiji ali uporabi zunaj testne postavitve** ne pove ničesar.

Kako močni so dokazi?

Za osrednje primerjalne trditve — da dodatno naučeni LBM v agregatu prekaša osnovne modele od začetka, potrebuje večkrat manj podatkov in je robustnejši ob premiku porazdelitve — so dokazi močni *in* nenavadno dobro nadzorovani: slepo, naključno, z velikimi vzorci, statistično testirano in s kontrolo kakovosti ocenjevanja. To je redek primer, ko je metodologija dovolj dobra, da lahko glavne zaključke vzamemo skoraj dobesedno.

Poštene omejitve so tiste, ki jih izpostavijo avtorji sami. Njihovi intervali zajemajo naključnost *vrednotenja*, ne pa naključnosti *učenja* — isti model, naučen dvakrat, lahko da precej drugačno politiko, ta variabilnost pa ni vključena v statistiko. Naloge v resničnem svetu so imele po 50 poskusov, kar zadostuje za srednje velike učinke, ne nujno za majhne. Jezikovno pogojevanje je uporabljalo skromen enkoder, zato so lahko zaključki o »samo povej robotu« drugačni pri večjih sistemih. Poleg tega avtorji odkrito poročajo o pozneje odkriti napaki pri normalizaciji. Nič od tega ne podre glavnih rezultatov; prav to pa so podrobnosti, za katere članek trdi, da jih področje prepogosto pomete pod preprogo.

Opomba o viru v istem duhu: ta razlaga temelji na preprintu avtorjev. Objavljene revijalne različice nismo uspeli pridobiti, zato nismo preverili morebitnih sprememb med preprintom in objavljenim besedilom.

Zakaj je pomembno

»Robotski temeljni modeli« so izraz, narejen za pretiravanje, in takšno raziskavo je lahko napačno prebrati v obe smeri — kot zmagoslavni »*deluje!*« ali zaničevalni »*vse je hype*«. Natančnejša razlaga je uporabnejša od obeh: predhodno učenje na raznolikih podatkih prinese resnične, merljive, vendar zmerne koristi — predvsem manj podatkov na nalogo in več robustnosti — razvoj pa se s količino podatkov izboljšuje predvidljivo.

Globlji razlog za pomembnost članka je, da svojo strogost obrne proti lastnemu področju. Ko pokaže, da so resnični učinki dovolj majhni, da ob površnem vrednotenju izginejo, in da lahko dolgočasna odločitev, kot je normalizacija podatkov, pretehta pametno novo arhitekturo, argumentira, da velik del napredka v učenju robotov potrebuje trdnejše meritve, preden mu lahko verjamemo. Članek, ki svojo verodostojnost porabi za nadzor meje med rezultatom in željo, počne nekaj redkejšega in dragocenejšega od vzpona na vrh lestvice.

Kratek povzetek

Raziskovalci Toyota Research Institute so izurili »large behavior models« — difuzijske robotske politike, predhodno učene na približno 1.700 urah raznolikih podatkov o manipulaciji — in jih primerjali s politikami za posamezne naloge, učenimi od začetka, z nenavadno strogim protokolom: slepo, naključno, z velikimi vzorci (≈ 1.800 poskusov v resničnem svetu in več kot 47.000 v simulaciji) ter pravo statistiko. Po dodatnem učenju za posamezne naloge so veliki modeli v agregatu zanesljivo delovali bolje, potrebovali približno **3–5-krat manj podatkov za posamezno nalogo** in bili **robustnejši ob spremembi pogojev**, uspešnost pa je gladko naraščala z več podatki za predhodno učenje. Toda **brez dodatnega učenja** niso dosledno prekašali modelov za eno nalogo, več učinkov je bilo tako majhnih, da so jih pokazali šele veliki vzorci, banalna izbira normalizacije podatkov pa je vplivala bolj kot arhitektura. To je trdna, izmerjena podpora smeri robotskih temeljnih modelov — **ne** splošnonamenski robot, ne zero-shot generalist in ne »emergentni skok« — ter ostro opozorilo, da velik del robotike morda meri šum.

Preverjanje brez olepševanja

Kaj članek kaže: Pri strogem, slepem in statistično dovolj močnem vrednotenju (≈ 1.800 izvedb v resničnem svetu + več kot 47.000 v simulaciji) večnalogovno predhodno učene in nato dodatno naučene difuzijske politike (LBM) v agregatu prekašajo politike za eno nalogo, učene od začetka, dosežejo enakovredno uspešnost s približno 3–5-krat manj podatki za posamezno nalogo in so robustnejše ob premiku porazdelitve; uspešnost s količino predhodnih podatkov narašča gladko.

Kaj je verjetno, vendar ni dokazano: Da se bodo enake koristi prenesle na veliko večje modele vid–jezik–dejanje (njihov jezikovni enkoder je bil majhen); da se gladko skaliranje nadaljuje onkraj preizkušene obsega podatkov.

Česa ne kaže: Splošnonamenskega ali zero-shot robota (brez dodatnega učenja ni dosledne prednosti); kakršnegakoli »emergentnega« skoka sposobnosti; razlag za konkretne neuspehe pri posameznih nalogah; absolutne zanesljivosti v resničnem svetu (stopnje uspeha so bile zaradi občutljivosti testov nastavljene blizu 50 %); ničesar o varnosti ali avtonomni uporabi.

Glavne omejitve: Statistika zajema naključnost vrednotenja, ne variabilnosti med ponovnimi učenji; 50 resničnih poskusov na nalogo lahko spregleda majhne učinke; ena arhitektura in en laboratorij; skromen jezikovni enkoder; napaka v normalizaciji podatkov je bila odkrita po vrednotenju; analiza temelji na preprintu, objavljena različica ni bila preverjena.

Koliko zaupanja naj ima splošni bralec? Visoko, da večnalogovno predhodno učenje z dodatnim učenjem prinaša resnične, zmerne koristi — zlasti podatkovno učinkovitost in robustnost — in da so bile te izmerjene nenavadno skrbno. Visoko tudi, da to **ni** splošnonamenski ali zero-shot robot in **ni** emergentni skok. Srednje glede tega, kako daleč se bodo koristi skalirale pri večjih modelih. In opozorilo avtorjev je vredno vzeti resno: učinki so dovolj majhni, da lahko statistično šibke raziskave poročajo predvsem šum. Primerna drža: zmeren optimizem o pristopu in zdrava skepsa do rezultatov robotske AI brez takšne statistične podpore.

Viri

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/>