



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

Fungerar stora "grundmodeller" för robotar faktiskt bättre? Ett noggrant svar — ja, i blygsam grad, och de flesta studier kan inte avgöra det

Toyota Research Institute tränade "stora beteendemodeller" — robotpolicyer som förtränats på ~1,700 timmar av varierade manipulationsdata — och testade dem mot policyer för en enda uppgift som tränats från grunden med ovanlig noggrannhet: blinda, randomiserade försök med stora stickprov (~1,800 i verkligheten, 47,000+ i simulering) och riktig statistik. Efter finjustering för varje uppgift presterade de stora modellerna bättre i genomsnitt, behövde ungefär 3–5× mindre uppgiftsspecifika data och var robustare när förhållandena förändrades; resultaten steg jämnt med mer förträningssdata. Men utan finjustering slog de inte konsekvent modeller för en enda uppgift, flera effekter var så små att de stora stickproven krävdes för att alls upptäcka dem, och ett vardagligt val av datanormalisering spelade större roll än arkitekturen. Det är ett återhållsamt stöd för riktningen mot grundmodeller för robotar — inte en robot för allmänna ändamål, inte en zero-shot-generalist och inte ett "emergent språng" — samt en skarp varning om att stora delar av robotiken kanske mäter brus.

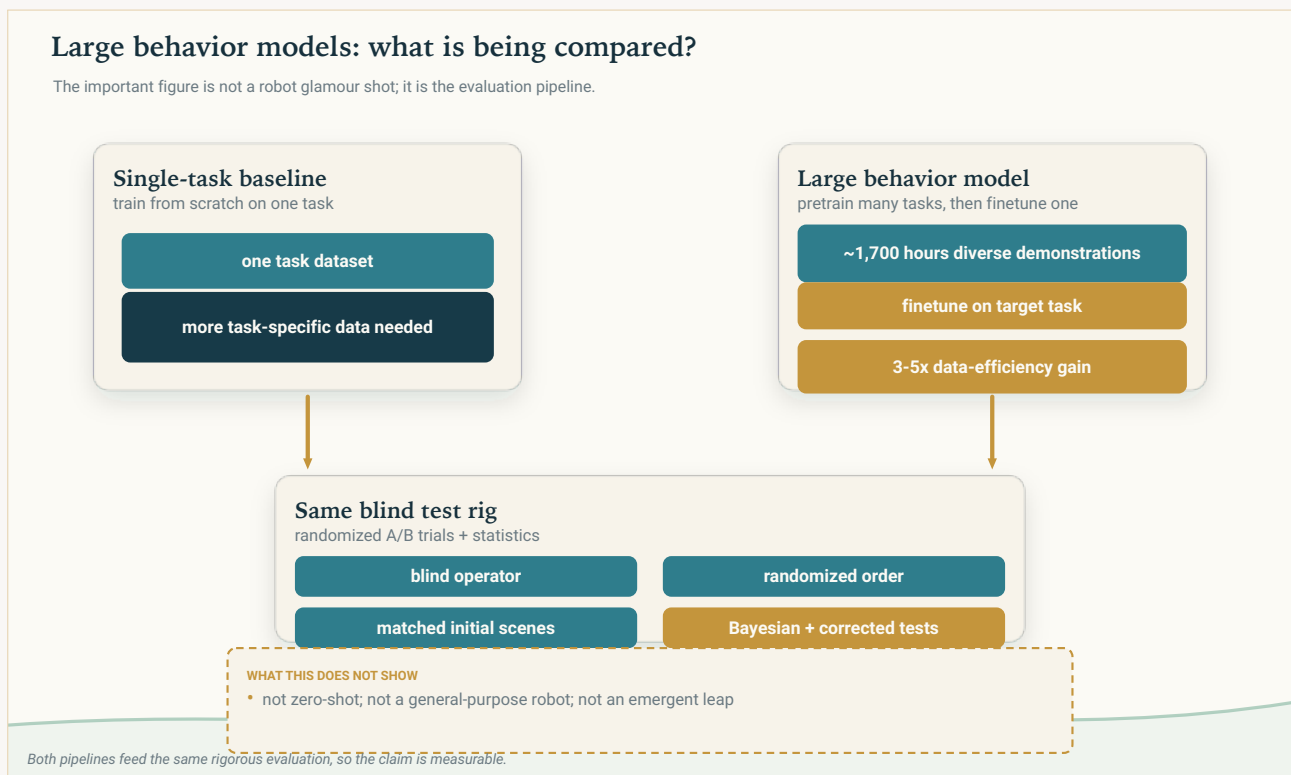
Written by Lucio Vaglio · Cross-checked by Laura Nesso · Edited by Michele Renda · Figures by Laura Nesso
· AI-assisted, human-reviewed

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

En robotikartikel vars verkliga ämne är ärlighet

Robotiken befinner sig mitt i en rusning efter ”grundmodeller”. Idén, lånad från språk- och bild-AI, är lockande: i stället för att träna en robot för en uppgift i taget tränar man en **”stor beteendemodell” (LBM)** på en enorm och varierad mängd demonstrationer och får ett system som har bred förmåga och snabbt kan anpassas. Entusiasmen — och investeringarna — är enorma. Rubriken skriver sig själv: *robotjärnor för allmänna ändamål är här*.

Den här artikeln från Toyota Research Institute är intressant just för att den vägrar skriva den rubriken. Dess verkliga bidrag är inte en mer spektakulär robot. Det är en grundlig granskning av en bedrägligt enkel fråga — *fungerar metoden med stora modeller faktiskt bättre, och hur skulle vi ens kunna veta det?* — besvarad med en statistisk noggrannhet som enligt författarna själva är ovanlig på området. Resultatet är ett genuint användbart ”ja, men” och en varning om att stora delar av robotiken kanske mäter brus.



Båda metoderna går vidare till samma blinda, randomiserade utvärdering. Artikelns påstående är inte att grundmodeller för robotar är zero-shot-generalister, utan att förträning kan förbättra dataeffektivitet och robusthet när mätningen görs omsorgsfullt.

Original The Clean Paper diagram CC BY 4.0

Vad författarna gjorde

De byggde LBM:er i en specifik, konkret bemärkelse: **diffusionsbaserade visuomotoriska policyer** (en Diffusion Transformer som läser kamerabilder, en kort språklig instruktion och robotens egna ledpositioner och matar ut korta serier av motorkommandon vid 10 Hz). Dessa **förtränades på ungefär 1,700 timmar** av robotdemonstrationer — fler än 500 olika uppgifter som samlats in internt, plus offentliga dataset — och **finjusterades** sedan för enskilda uppgifter. Jämförelsen görs genomgående med en **policy för en enda uppgift som tränats från grunden** på data från just den uppgiften.

Artikelns kärna är dock *utvärderingen*, som författarna behandlar som huvudresultatet. För att undvika att lura sig själva använde de:

- **Blinda, randomiserade A/B-tester** i verkligheten — personen som körde roboten visste inte vilken policy som testades, och ordningen var randomiserad.
- **Kontrollerade, upprepningsbara startförhållanden** — operatörerna passade in scenen mot en bildöverlagring före varje försök.
- **Ett stort antal försök** — 50 körningar i verkligheten per uppgift, per policy och per villkor; 200 per uppgift i simulering. Totalt: omkring **1,800 blinda körningar i verkligheten och fler än 47,000 simulerade körningar**.
- **Korrekt statistik** — bayesianska skattningar av sannolikheten för framgång och parvisa hypotesprövningar med korrigering för multipla jämförelser, i stället för att bedöma överlappande felstaplar på ett ungefär. De gjorde till och med en kvalitetssäkringsomgång för en fjärdedel av de försök som bedömts av människor för att mäta bedömningsfelet.

[video pending]

Jämförelse sida vid sida av modeller som dukar ett frukostbord: (vänster) baslinjen för en enda uppgift och (höger) LBM. Båda videorna spelas upp i 1x hastighet. Det här är en utvärderad uppgift, inte ett bevis på autonomi för allmänna ändamål.

Det här maskineriet är själva poängen. Hela artikeln argumenterar för att man utan det inte kan skilja en verklig förbättring från tur.

Vad de fann

Finjusterade stora modeller slår modeller för en enda uppgift som tränats från grunden — i genomsnitt. Sammantaget över uppgifterna överträffade en LBM som hade förtränats och sedan finjusterats för en uppgift på ett tillförlitligt sätt en policy som tränats från grunden för samma uppgift, både i simulering och i verkligheten, och skillnaden var statistiskt signifikant. För enskilda uppgifter var den finjusterade LBM:en statistiskt sett lika bra som eller bättre än modellen som tränats från grunden i nästan samtliga fall (3/3 uppgifter i verkligheten, 15/16 simulerade uppgifter).

Den största och tydligaste vinsten är dataeffektivitet. En finjusterad LBM nådde likvärdiga resultat med en modell som tränats från grunden med ungefär **3–5× mindre uppgiftsspecifika data**. I en

uppgift i verkligheten (att duka ett frukostbord) slog en LBM som finjusterats på bara **15%** av demonstrationerna en policy som tränats från grunden på **100%** av dem.

Förträning hjälper mest när förhållandena förändras. När testmiljön medvetet avvek från träningsförhållandena ("distributionsförskjutning") *ökade* den finjusterade LBM:ens försprång. I en uppsättning simuleringar slog den statistiskt modellen som tränats från grunden i 3 av 16 uppgifter under normala förhållanden, men i 10 av 16 vid distributionsförskjutning. Eftersom verkliga driftsmiljöer alltid glider bort från träningsförhållandena kan denna robusthet hävdas vara det praktiskt viktigaste resultatet.

Mer förträningsdata hjälpte, på ett jämnt sätt. Resultaten steg stadigt när mer förträningsdata lades till — utan **något plötsligt språng eller "emergent" kliv** vid de skalor som testades. Användbart, förutsägbart och odramatiskt.

Men berättelsen om en generalist utan finjustering höll inte. En förtränad LBM som användes *zero-shot* — utan uppgiftsspecifik finjustering — slog *inte* konsekvent policyer för en enda uppgift. Ett enda nätverk kunde utföra många uppgifter samtidigt, men drömmen om att "bara instruera den" fick inget stöd här; författarna tillskriver detta delvis den begränsade språkkodarens skörhet.

Och vinsterna var små nog för att vara lätta att missa — eller att framställa på falska grunder. Många av effekterna blev synliga först med de större stickproven än vanligt och de noggranna testerna. Författarna säger rent ut att det, med tanke på effekternas storlek och bruset, **finns en betydande risk att många robotikartiklar mäter statistiskt brus**. De fann också att ett vardagligt val — hur data *normaliseras* — påverkade resultaten mer än arkitekturförändringar, och att ett normaliseringsfel i förträningen upptäcktes först *efter* att utvärderingarna hade slutförts.

Vad detta sannolikt betyder

Den försvarbara tolkningen: storskalig förträning på varierade robotdata är en verklig och värdefull beståndsdel — den minskar mängden data som krävs för varje ny uppgift och gör policyerna stabilare när världen inte motsvarar träningen. Det ger ett genuint stöd åt den riktning som området satsar på. Men vinsterna är *måttliga och villkorade* (de visar sig främst efter finjustering och är tydligast sammantaget och under påfrestning), inte ankomsten av en färdig, allmän robot som bara kan tas i bruk.

Den tillsammans och viktigare innebörden är metodologisk. Artikeln fungerar i praktiken som en måttstock: den visar hur mycket evidens som faktiskt krävs för att göra ett trovärdigt påstående om en robotpolicy och antyder att mycket av den publicerade entusiasmen vilar på för lite. Det är en korrigeringsområde som behöver mer än ännu en modell.

Vad detta *inte* bevisar

- Det är **inte en robot för allmänna ändamål**. Vinsterna har visats för en specifik arkitektur (diffusionspolicyer) som finjusterats för varje uppgift, i kontrollerade miljöer och med fjärrstyrda demonstrationer — inte för en autonom robot som utför godtyckliga nya uppgifter på kommando.
- Det bekräftar **inte** användning med **zero-shot**. Utan finjustering slog den stora modellen inte konsekvent baslinjerna för en enda uppgift.
- Det är **inte** belägg för ett **”emergent språng”**. Uppskalningen förbättrade resultaten jämnt; här finns ingen diskontinuitet som stöder berättelser om att ”och sedan blev den plötsligt kapabel”.
- Siffrorna är **relativa och bundna till laboratoriet**. De absoluta framgångsnivåerna justerades avsiktligt mot ~50% för att göra jämförelserna känsliga; de mäter inte tillförlitlighet i verkligheten, och arbetet gäller en arkitektur från ett laboratorium.
- Det avgör **inte varför** någon policy lyckas eller misslyckas, och flera specifika uppgifter där den stora modellen presterade *sämre* redovisas men förklaras inte.
- Det säger **ingenting om säkerhet, autonomi eller driftsättning** utanför utvärderingsriggen.

Hur starka är beläggen?

För de centrala jämförande påståendena — att finjusterade LBM:er sammantaget slår baslinjer som tränats från grunden, behöver flera gånger mindre data och är robustare vid distributionsförskjutning — är beläggen starka *och* ovanligt välkontrollerade: blinda, randomiserade, med stora stickprov, statistiskt prövade och med en kvalitetssäkringsomgång av bedömningen. Det här är det sällsynta fall där metodiken är tillräckligt gedigen för att huvudslutsatserna ska kunna tas i stort sett för vad de är.

De ärliga förbehållen är sådana som författarna själva tar upp. Deras felstaplar fångar slumpmässigheten i *utvärderingen* men **inte** slumpmässigheten i *träningen* — tränar man samma modell två gånger kan man få policyer som skiljer sig på ett meningsfullt sätt, och den variationen ingår inte i statistiken. Uppgifterna i verkligheten hade 50 försök vardera, tillräckligt för att upptäcka medelstora effekter men med risk att missa små. Språkstyrningen använde en begränsad kodare, så påståenden om att ”bara säga åt roboten vad den ska göra” kan utfalla annorlunda för större system. Därtill kommer det rättframma avslöjandet av ett normaliseringsfel i efterhand. Inget av detta kullkastar huvudresultaten, men det är precis sådant som artikeln hävdar att området vanligtvis sopar åt sidan.

En anmärkning om källan, i samma anda: den här förklaringen bygger på författarnas preprint. Vi kunde inte hämta den tidskriftspublicerade versionen och har därför inte kontrollerat om något ändrades mellan preprinten och den publicerade texten.

Varför det spelar roll

”Grundmodeller för robotar” är en fras som gjord för överdrivna påståenden, och en studie som denna är lätt att feltolka i båda riktningarna — som ett triumferande *”det fungerar!”* eller ett avfär-

dande ”*det är överhajpat.*” Den korrekta slutsatsen är mer användbar än båda: förträning på varierade data ger verkliga, mätbara men måttliga fördelar — främst mindre data per uppgift och större robusthet — och utvecklingsvägen förbättras förutsägbart med skalan.

Det djupare skälet till att artikeln spelar roll är att den riktar sin noggrannhet mot det egna forskningsfältet. Genom att visa att de verkliga effekterna är små nog att försvinna i en slarvig utvärdering, och att ett alldagligt val som datanormalisering kan väga tyngre än en sinnrik ny arkitektur, argumenterar den för att stora delar av framstegen inom robotinlärning behöver stabilare mätning innan de kan anses trovärdiga. En artikel som använder sin trovärdighet till att bevaka skillnaden mellan ett resultat och en önskan gör något ovanligare och mer värdefullt än att toppa en rankingslista.

Ren sammanfattning

Forskare vid Toyota Research Institute tränade ”stora beteendemodeller” — diffusionsbaserade robotpolicyer som förtränats på ~1,700 timmar av varierade manipulationsdata — och testade dem mot policyer för en enda uppgift som tränats från grunden med ett ovanligt rigoröst protokoll: blint, randomiserat och med stora stickprov ($\approx 1,800$ försök i verkligheten och 47,000+ simulerade försök), med riktig statistik. Efter finjustering för varje uppgift presterade de stora modellerna tillförlitligt bättre sammantaget, behövde ungefär $3-5\times$ **mindre uppgiftsspecifika data** och var **robustare när förhållandena förändrades**, samtidigt som resultaten förbättrades jämnt när mängden förträningsdata ökade. Men när de användes **utan finjustering** slog de *inte* konsekvent modeller för en enda uppgift, flera effekter var så små att bara de stora stickproven avslöjade dem, och ett vardagligt val av datanormalisering spelade större roll än arkitekturen. Det är ett gediget, återhållsamt stöd för riktningen mot grundmodeller för robotar — **inte** en robot för allmänna ändamål, inte en zero-shot-generalist och inte ett ”emergent språng” — samt en skarp varning om att stora delar av robotiken kanske mäter brus.

Kontroll utan skitsnack

Vad artikeln visar: Med en rigorös, blind utvärdering med tillräcklig statistisk styrka ($\approx 1,800$ körningar i verkligheten + 47,000+ simulerade körningar) presterar diffusionspolicyer (LBM:er) som förtränats på flera uppgifter och sedan finjusterats sammantaget bättre än policyer för en enda uppgift som tränats från grunden, når likvärdiga resultat med $\sim 3-5\times$ mindre uppgiftsspecifika data och är robustare vid distributionsförskjutning; resultaten skalar jämnt med mängden förträningsdata.

Vad som är rimligt men inte bevisat: Att dessa fördelar kan överföras till mycket större modeller för syn, språk och handling (deras språkkodare var liten); att den jämna skalningen fortsätter bortom det testade dataintervallet.

Vad den inte visar: En allmän robot eller zero-shot-robot (ingen finjustering $\rightarrow$ ingen konsekvent fördel); något ”emergent” språng i förmåga; förklaringar till specifika misslyckanden på uppgiftsnivå;

tillförlitlighet i verkligheten i absoluta termer (framgångsnivåerna justerades till nära 50% för känslighet); något om säkerhet eller autonom driftsättning.

Huvudsakliga begränsningar: Statistiken fångar slumpmässighet i utvärderingen men inte mellan träningskörningar; 50 försök i verkligheten per uppgift kan missa små effekter; en arkitektur och ett laboratorium; begränsad språkkodare; ett datanormaliseringsfel upptäcktes efter utvärderingarna; analysen bygger på preprinten (den publicerade versionen har inte kontrollerats).

Hur stort förtroende bör en vanlig läsare ha? Högt förtroende för att multitask-förträning plus finjustering ger verkliga, måttliga fördelar — särskilt i dataeffektivitet och robusthet — och att dessa mättes ovanligt noggrant. Högt förtroende för att detta **inte** är en allmän robot eller zero-shot-robot och **inte** ett emergent språng. Medelhögt förtroende för hur långt vinsterna kan skalas till större modeller. Och något att ta på allvar: författarnas egen varning om att effekterna på området är så små att studier med otillräcklig statistisk styrka kan rapportera brus. Lämplig hållning: återhållsam optimism om metoden och sund skepticism mot resultat inom robot-AI som saknar denna typ av statistiskt stöd.

Källor

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/>