



The CLEAN PAPER

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

Fungují velké „základní modely“ robotů opravdu lépe? Opatrná odpověď zní: do určité míry ano

Toyota Research Institute natrénovával „velké modely chování“ – řídicí politiky robotů předtrénované na zhruba 1 700 hodinách rozmanitých manipulačních dat – a v mimořádně přísných, zaslepených a randomizovaných zkouškách je porovnal s politikami pro jedinou úlohu trénovanými od nuly. Po doladění si velké modely vedly v průměru lépe, potřebovaly asi třikrát až pětkrát méně dat pro konkrétní úlohu a lépe odolávaly změnám podmínek. Bez doladění však jednoúlohové modely soustavně nepřekonávaly, některé účinky byly velmi malé a normalizace dat měla větší vliv než architektura. Výsledek podporuje tento směr, ale neprokazuje univerzálního robota, schopnost řešit nové úlohy bez ukázek ani náhlý „emergentní skok“.

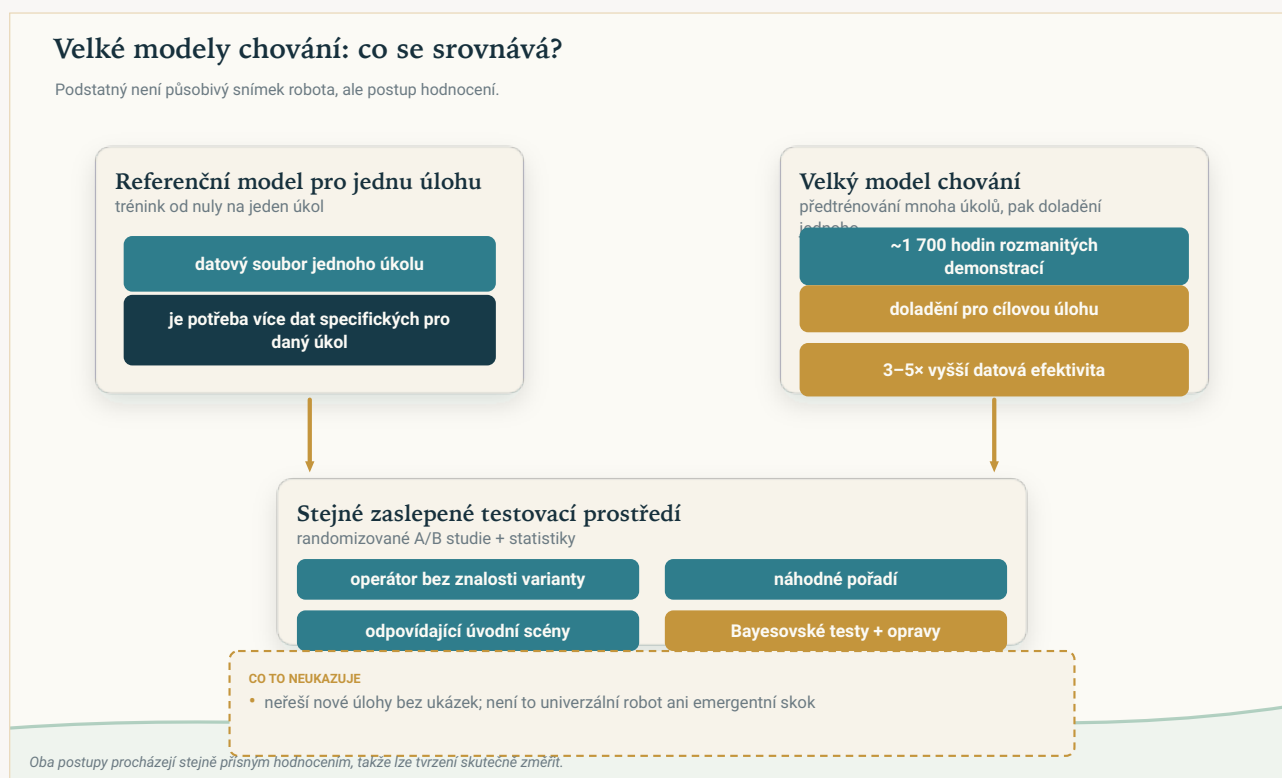
Written by Lucio Vaglio · Cross-checked by Laura Nesso · Edited by Michele Renda · Figures by Laura Nesso
· AI-assisted, human-reviewed

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

Článek o robotech, jejichž skutečným tématem je poctivost

Robotika zažívá horečku „základního modelu“. Díky jazykové a obrazové umělé inteligenci je tato myšlenka lákavá: místo toho, abyste robota učili každý úkol zvlášť, trénujte jeden velký **“velký model chování” (LBM)** na obrovské, rozmanité sadě ukázek a získejte systém s širokými možnostmi, který se rychle přizpůsobí novým úkolům. Nadšení – a investice – jsou obrovské. Titulek sám o sobě píše: *univerzální robotické mozky jsou zde*.

Práce Toyota Research Institute je zajímavá právě tím, že odmítá napsat takový titulek. Jeho hlavním přínosem není působivější robot, ale náročná analýza zdánlivě jednoduché otázky - *opravdu funguje přístup velkého modelu lépe a jak bychom to věděli?* - provedená s důrazem na statistiku, která je podle samotných autorů v oboru neobvyklá. Výsledkem je opravdu užitečné „ano, ale“ a varování, že velká část robotiky umí měřit hluk.



Oba přístupy procházejí stejným zaslepeným, randomizovaným hodnocením. Práce netvrdí, že robotické základní modely jsou univerzální systémy zero-shot, ale že pečlivě měřené předtrénování může zlepšit efektivitu a robustnost dat.

Original The Clean Paper diagram CC BY 4.0

Co autoři udělali

LBM vytvořili v konkrétním smyslu jako **difuzní vizuomotorické řídicí politiky**. Diffusion Transformer zpracovává obrazy z kamer, krátký jazykový pokyn a polohu kloubů robota a poté generuje

krátké série pohybových příkazů s frekvencí 10 Hz. Modely byly **předtrénovány na přibližně 1 700 hodinách** robotických ukázek – více než 500 různých úloh z interních i veřejných datových sad – a následně **doladěny** pro jednotlivé úkoly. Referenčním bodem byla vždy **jednouúlohová řídicí politika trénovaná od nuly** výhradně na datech daného úkolu.

Jádrem práce je však *hodnocení*, které autoři považují za hlavní výsledek. Aby se neoklamali, použili:

- **Zaslepené, randomizované A/B testy** ve fyzickém světě – operátor robota nevěděl, kterou řídicí politiku testuje, a pořadí bylo náhodné.
- **Kontrolované a opakovatelné počáteční podmínky** – před každým pokusem operátoři srovnali scénu s překrytým referenčním obrazem.
- **Velké počty pokusů** – 50 fyzických provedení na úkol, politiku a podmínku a 200 na úkol v simulaci. Celkem přibližně **1 800 zaslepených provedení ve fyzickém světě a více než 47 000 v simulaci**.
- **Odpovídající statistiku** – bayesovské odhady pravděpodobnosti úspěchu a párové testy hypotéz s korekcí na vícenásobná srovnání namísto posuzování překryvu chybových úseček pouhým okem. Navíc znovu zkontrolovali čtvrtinu pokusů hodnocených lidmi, aby změřili chybu bodování.

[video pending]

Porovnání modelů nastavení snídaňového stolu: vlevo základní jednouúkolový model, vpravo LBM. Obě videa se přehrávají normální rychlostí. Toto je jeden hodnocený úkol, nikoli obecný účelový důkaz autonomie.

Všechna tato opatření mají jasný smysl. Bez nich podle autorů nelze odlišit skutečné zlepšení od šťastné náhody.

Co objevili

Doladěné velké modely v průměru překonaly jednouúlohové modely trénované od nuly. Po sloučení výsledků napříč úkoly předtrénované a následně doladěné LBM spolehlivě překonaly řídicí politiky trénované od nuly na stejných datech – v simulaci i ve fyzickém světě – a rozdíl byl statisticky významný. Téměř u každého jednotlivého úkolu byl doladěný LBM statisticky stejně dobrý nebo lepší než model trénovaný od nuly: u 3 ze 3 fyzických a 15 ze 16 simulovaných úloh.

Největším a nejzřetelnějším přínosem je datová efektivita. Doladěné LBM dosáhly výkonu modelu trénovaného od nuly s přibližně **3–5krát menším množstvím dat pro daný úkol**. V jednom fyzickém úkolu – prostírání snídaňového stolu – LBM doladěný na pouhých **15 %** ukázek překonal řídicí politiku trénovanou od nuly na **100 %** dat.

Předtrénování pomáhalo nejvíce při změně podmínek. Když se testovací prostředí záměrně odchýlilo od tréninkových podmínek, tedy při *posunu distribuce*, výhoda doladěného LBM *vzrostla*. V jedné sadě simulací model statisticky překonal variantu trénovanou od nuly u 3 ze 16 úloh za běžných podmínek, ale u 10 ze 16 po posunu distribuce. Protože se skutečné nasazení od tréninkových podmínek vždy nějak liší, je tato odolnost patrně nejdůležitějším praktickým výsledkem.

Pomohlo více předtréninkových dat – hladce. Skóre se s množstvím dat plynule zvyšovalo, **žádný náhlý nárůst nebo „vznikající“ průlom** na testované škále. Užitečné, předvídatelné, žádné drama.

Příběh univerzálního zero-shot modelu však neobstál. LBM použitý *zero-shot* – bez ladění úkolu – trvale nepřekonal jednoúlohové modely. Jedna síť mohla plnit mnoho úkolů, ale sen o „prostě tomu dejte příkaz“ se zde nepotvrdil; autoři to částečně připisují křehkosti malého jazykového kodéru.

Přínosy byly natolik malé, že je bylo snadné přehlédnout – nebo zdánlivě vyrobit. Mnohé efekty se ukázaly až s většími než obvyklými vzorky a pečlivým testováním. Autoři přímo píší, že vzhledem k velikosti efektů a šumu **existuje značné riziko, že mnoho robotických studií měří statistický šum.** Zjistili také, že všední rozhodnutí o tom, jak data *normalizovat*, ovlivnilo výsledky více než změny architektury; **chybu** v normalizaci předtrénovacích dat navíc odhalili až *po* dokončení hodnocení.

Co to pravděpodobně znamená

Rozumná interpretace zní: předtrénování na velkém množství různorodých robotických dat je skutečně cenné – snižuje množství dat potřebných pro nový úkol a zvyšuje odolnost řídicích politik vůči rozdílům mezi tréninkem a skutečným světem. Výsledek podporuje směr, kterým se obor ubírá. Výhody jsou však *mírné a podmíněné*: projevují se hlavně po doladění, nejzřetelněji ve sloučených výsledcích a při změně podmínek. Neznamenají vznik univerzálního robota připraveného k použití.

Méně nápadný a důležitější význam je metodologický. Práce v podstatě nastavuje měřítko: ukazuje, kolik důkazů je zapotřebí, aby tvrzení o robotických řídicích politikách byla věrohodná, a naznačuje, že značná část publikovaného nadšení stojí na příliš malém množství dat. Obor takovou korekci potřebuje více než další model.

Co to nedokazuje

- Toto **není univerzální robot**. Výhody byly prokázány u konkrétní architektury – difuzních řídicích politik – samostatně doladěné pro jednotlivé úkoly, v kontrolovaných podmínkách a na teleoperovaných ukázkách. Nešlo o autonomního robota plnícího libovolné nové příkazy.
- Výsledky **nepodporují výkon zero-shot**. Bez doladění velký model jednoúlohové základní modely soustavně nepřekonával.
- Toto **není** důkaz „**emergentního skoku**“. S rostoucím rozsahem se výsledky zlepšovaly plynule; žádná diskontinuita nepodporuje příběh „a najednou to uměl“.
- Čísla jsou **relativní a laboratorní**. Absolutní úspěšnost byla záměrně nastavena poblíž 50 %, aby bylo srovnání citlivé; neměří spolehlivost ve skutečném světě a práce zkoumá jedinou architekturu z jedné laboratoře.
- Výzkum **neurčuje, proč** konkrétní řídicí politiky uspějí nebo selžou. Několik úloh, v nichž byl velký model *horší*, popisuje, ale nevysvětluje.
- Neříká **nic o bezpečnosti, autonomii ani nasazení** mimo testovací prostředí.

Jak silné jsou důkazy?

Důkazy pro hlavní srovnávací tvrzení – že doladěné LBM jako skupina překonávají modely trénované od nuly, vyžadují několikanásobně méně dat a lépe zvládají posuny distribuce – jsou silné a výjimečně dobře kontrolované: zaslepené, randomizované, rozsáhlé, statisticky testované a s kontrolou kvality bodování. Jde o vzácný příklad metodiky natolik robustní, že lze její hlavní závěry přijmout téměř doslova.

Omezení otevřeně uvádějí sami autoři. Intervaly chyb zahrnují náhodnost *hodnocení*, nikoli však náhodnost *tréninku* – dvě trénování téhož modelu mohou vytvořit výrazně odlišné řídicí politiky, což statistika nezachycuje. Padesát fyzických pokusů na úkol stačí odhalit středně velké efekty, ale malé mohou uniknout. Jazyková část používala skromný kodér, takže výsledky typu „řekni robotovi, co má dělat“ mohou u větších systémů vypadat jinak. Autoři také otevřeně přiznávají chybu normalizace. Žádný z těchto problémů nevyvrací hlavní výsledky, ale právě takové věci podle práce obor často zametá pod koberec.

Poznámka ke zdroji: diskuse je založena na předtisku autorů. Nepodařilo se nám získat publikovanou verzi v časopise, proto jsme nezkoumali změny mezi předtiskem a konečným textem.

Proč je to důležité

„Základní robotické modely“ jsou pojem jako stvořený pro nadsázku a tuto práci lze snadno vyložit chybně v obou směrech – vítězoslavně jako „*funguje to!*“ nebo odmítavě jako „*je to přeceňované*“. Přesnější pohled je užitečnější: předtrénování na různorodých datech přináší skutečné, měřitelné, ale mírné výhody – především menší potřebu dat pro konkrétní úkol a větší odolnost – a výsledky se s rozsahem předvídatelně zlepšují.

Hlubší význam spočívá v tom, že práce vnáší přísnost do vlastního oboru. Autoři ukazují, že skutečné efekty jsou tak malé, že při neopatrném hodnocení mizí, a že všední rozhodnutí o normalizaci dat může znamenat více než chytrá architektura. Tvrdí tak, že velká část pokroku v robotickém učení potřebuje robustnější měření, než mu budeme moci věřit. Studie, která pečlivostí hlídá hranici mezi výsledkem a přáním, dělá něco vzácnějšího a cennějšího než jen obsadit první místo v žebříčku.

Stručné shrnutí

Výzkumníci z Toyota Research Institute trénovali „velké modely chování“ – difuzní řídicí strategie robotů předtrénované na přibližně 1 700 hodinách různorodých manipulačních dat – a porovnávali je se strategiemi pro jediný úkol, trénovanými od nuly. Použili mimořádně přísný protokol: zaslepené, randomizované hodnocení na velkých vzorcích (asi 1 800 pokusů ve fyzickém světě a přes 47 000 v simulaci) se skutečnou statistickou analýzou. Po doladění na konkrétní úkol dosahovaly velké modely celkově spolehlivě lepších výsledků, vyžadovaly přibližně **3–5krát méně dat pro daný úkol** a byly

odolnější vůči změnám podmínek; výkon přitom plynule rostl s množstvím předtrénovacích dat. Při použití **bez doladění** však jednoúlohové modely soustavně nepřekonávaly, některé efekty byly patrné jen ve velkých vzorcích a na obyčejné normalizaci dat záleželo více než na architektuře. Jde o pevnou a vyváženou podporu směru základních robotických modelů – **nikoli** o univerzálního robota, univerzální model zero-shot ani „náhlý emergentní skok“ – a o důrazné varování, že velká část robotiky možná měří jen šum.

Kontrola bez příkras

Co práce ukazuje: V přísném zaslepeném hodnocení s dostatečnou statistickou silou – přibližně 1 800 fyzických a více než 47 000 simulovaných provedení – víceúlohové předtrénované a následně doladěné difuzní řídicí politiky LBM jako skupina překonaly jednoúlohové politiky trénované od nuly. Dosáhly srovnatelných výsledků s přibližně 3–5krát menším množstvím dat pro daný úkol, lépe zvládaly posun distribuce a jejich skóre plynule rostlo s množstvím předtrénovacích dat.

Co je pravděpodobné, ale neprokázané: Přenos výhod na mnohem větší vizuomotorické modely – použitý jazykový kodér byl malý – a pokračování plynulého škálování za hranici zkoumaného množství dat.

Co práce neukazuje: Univerzálního robota ani spolehlivý výkon zero-shot – bez doladění nebyla výhoda konzistentní; žádný emergentní skok ve schopnostech; vysvětlení selhání konkrétních úloh; absolutní spolehlivost ve skutečném světě, protože úspěšnost byla kvůli citlivému srovnání nastavena poblíž 50 %; ani nic o bezpečnosti či autonomním nasazení.

Hlavní omezení: Statistiky zahrnují náhodnost hodnocení, nikoli tréninku; 50 fyzických pokusů na úkol může přehlédnout malý efekt; jedna architektura a jedna laboratoř; skromný jazykový kodér; chyba normalizace dat zjištěná až po hodnocení; a analýza založená na preprintu, bez kontroly změn v publikované verzi.

Kolik důvěry by měl mít čtenář? Vysokou v závěr, že víceúlohové předtrénování spojené s doladěním přináší skutečné, mírné výhody – zejména vyšší datovou efektivitu a odolnost – a že byly změřeny mimořádně pečlivě. Stejně vysokou v to, že **nejde** o univerzálního robota, spolehlivý systém zero-shot ani náhlý emergentní zlom. Střední v přínosy dalšího škálování na větší modely. Vážně je třeba brát i varování autorů, že efekty v tomto oboru jsou tak malé, že studie s nedostatečnou statistickou silou mohou vykazovat jen šum. Správný postoj: vyvážený optimismus vůči přístupu a zdravá skepse k výsledkům robotické AI bez srovnatelně pevné statistiky.

Zdroje

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/>