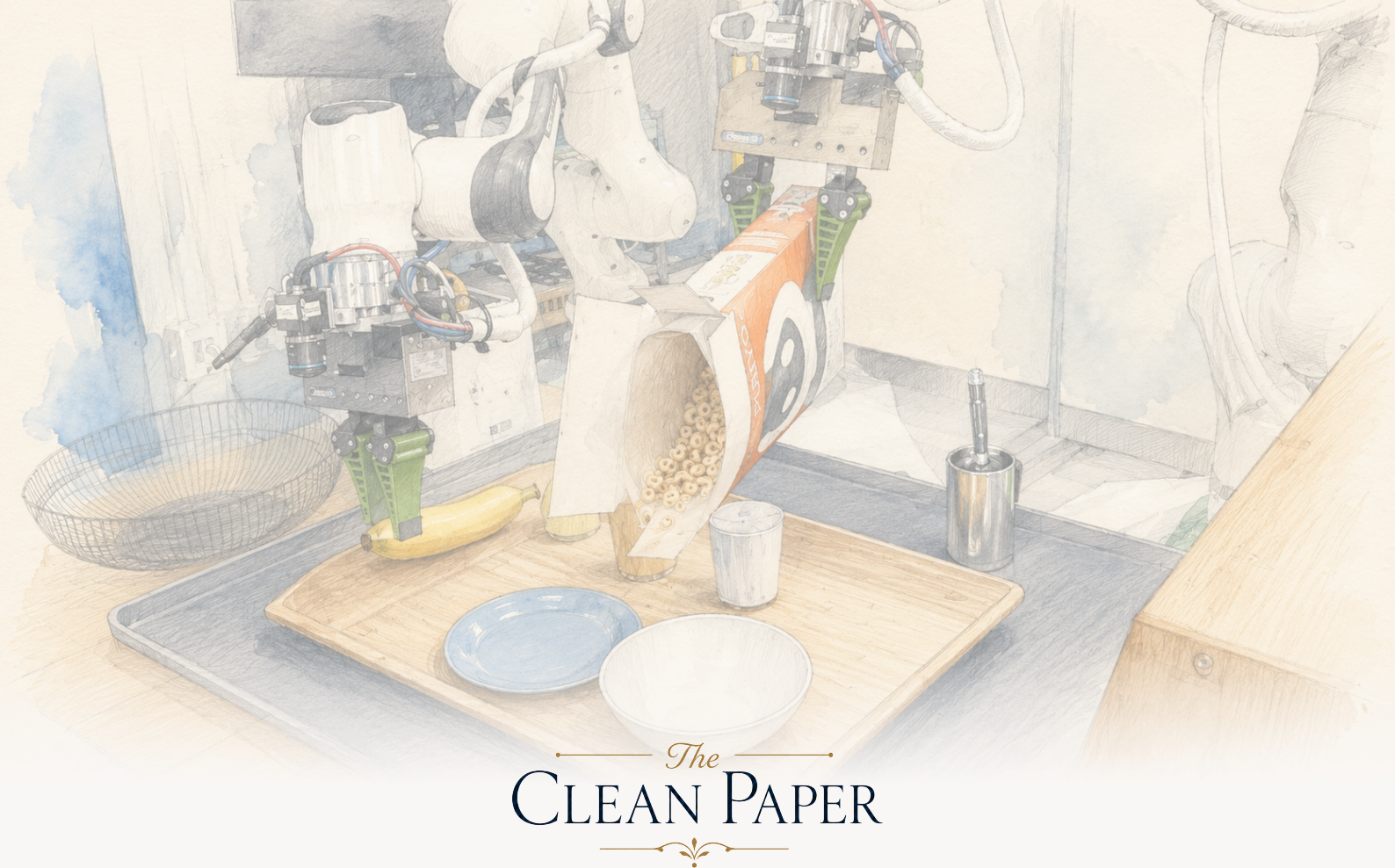




The CLEAN PAPER

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

大型机器人「基础模型」真的更好吗？一个谨慎的回答——答案是肯定的，但幅度有限，而大多数研究无法判断

丰田研究所训练了「大型行为模型」——在约 1,700 小时多样化操作数据上预训练的机器人策略——并以不寻常的严谨性将其与从零开始的单任务策略进行了对比：盲法、随机化、大样本量试验（约 1,800 次真实世界，47,000+ 次模拟）并采用正式统计检验。经过逐任务微调后，大模型平均表现更好，只需原来的约 1/3 至 1/5 任务特定数据，并且在条件变化时更鲁棒；性能随预训练数据量平稳上升。但在没有微调的情况下，它们并未一致地击败单任务模型，若干效应小到只有大样本量才能揭示，而且一个平凡的数据标准化选择比架构影响更大。这是对机器人基础模型方向的有节制支持——并非通用机器人，不是零样本万能者，也不是「涌现式的飞跃」——加上一条尖锐的警告：机器人学中的许多研究可能正在测量噪声。

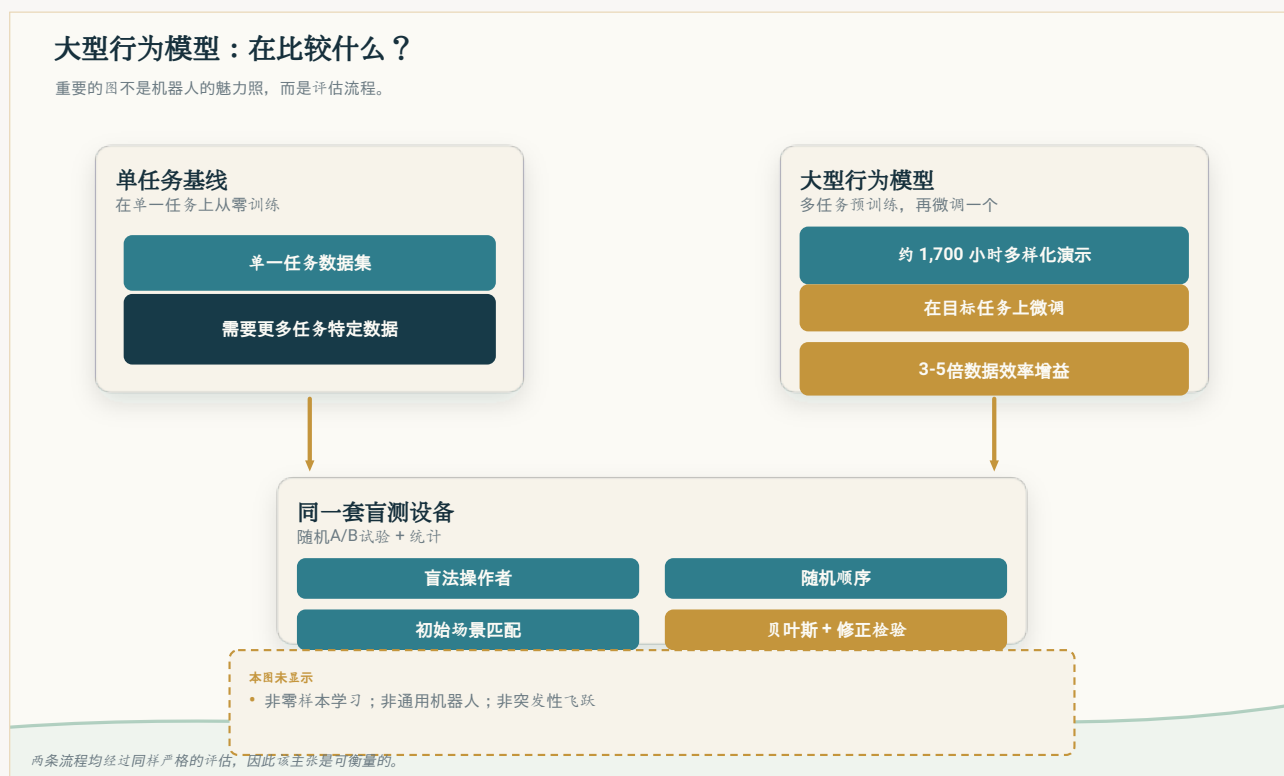
Written by Lucio Vaglio · Cross-checked by Laura Nesso · Edited by Michele Renda · Figures by Laura Nesso and Nova Vela · AI-assisted, human-reviewed

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

一篇真正主题是诚实的机器人论文

机器人技术正处在一股“基础模型”热潮中。这个借鉴自语言和图像 AI 的想法很有吸引力：与其每次只训练机器人完成一项任务，不如在大量、规模庞大且多样化的示范数据上训练一个“**大型行为模型**”（**LBM**），从而得到一个能力广泛、适应迅速的系统。热情和投入都十分巨大。头条标题几乎不言自明：通用机器人“大脑”来了。

这篇来自丰田研究院的论文之所以有意思，恰恰在于它拒绝写出这样的标题。它真正的贡献不是让机器人更炫，而是认真审视一个看似简单、实则容易误判的问题——大型模型路线真的更有效吗？我们又该如何知道？——并以作者自己所说的、在这个领域并不常见的统计严谨程度给出回答。结果是一个真正有用的“是的，但要看情况”，同时也提醒我们：机器人领域的许多研究，可能测到的只是噪声。



两种方法都进入同一个盲测、随机化评估流程。论文主张的并不是机器人基础模型能够零样本完成所有任务，而是：经过仔细测量后，预训练可以提升数据效率和稳健性。

Original The Clean Paper diagram CC BY 4.0

作者做了什么

他们以一种具体明确的方式构建了 LBM：**基于扩散的视觉运动策略**（一种扩散 Transformer，读取摄像头图像、简短的语言指令和机器人自身的关节位置，并以 10 Hz 输出短段电机指令）。这些模型在大约 **1,700 小时** 的机器人示范数据上进行了**预训练**——包括在内部收集的 500 多项不同任务，以及公开数据集——随后再针对单个任务进行**微调**。全文的比较对象始终是针对某一项任务、仅使用该任务数

据从头训练的单任务策略。

不过，这篇论文的核心其实是评估，作者将其视为主要成果。为了避免自欺，他们采用了：

- **在真实世界进行盲测、随机化的 A/B 测试**——操控机器人运行的人不知道当时测试的是哪种策略，测试顺序也是随机的。
- **受控且可重复的初始条件**——每次试验前，操作员都根据图像叠加提示，将场景调整到匹配状态。
- **大量试验次数**——每项任务、每种策略、每种条件在真实世界运行 50 次；在模拟环境中每项任务运行 200 次。总计约有 **1,800 次盲测真实世界运行，以及超过 47,000 次模拟运行**。
- **恰当的统计方法**——使用成功概率的贝叶斯估计，以及经过多重比较校正的成对假设检验，而不是仅凭误差条是否重叠来目测结果。他们甚至对四分之一由人工评分的试验进行质量保证复核，以测量评分误差。

[video pending]

模型摆放早餐餐桌的并排对比：（左）单任务基线模型，（右）LBM。两段视频都以 1 倍速播放。这只是一个经过评估的任务，并不能证明具备通用自主性。

这些流程才是重点。整篇论文都在论证：没有这些方法，你就无法区分真正的改进和偶然的好运。

他们发现了什么

微调后的大型模型平均胜过从头训练的单任务模型。在所有任务的汇总结果中，先经过预训练、再针对任务进行微调的 LBM，无论在模拟环境还是真实世界，都可靠地优于在同一任务上从头训练的策略，而且差异具有统计显著性。在单项任务上，微调后的 LBM 几乎每次都在统计上达到同等或更好表现（真实世界任务为 3/3，模拟任务为 15/16）。

最大、最清晰的优势是数据效率。微调后的 LBM 只需使用约为从头训练所需数据的 **1/3~1/5** 的任务特定数据，就能达到与从头训练相当的表现。在一项真实世界任务（摆放早餐餐桌）中，只用 **15%** 示范数据微调的 LBM，胜过使用全部 **100%** 示范数据从头训练的策略。

当条件发生变化时，预训练最有帮助。当测试环境被有意设置为偏离训练条件（“分布偏移”）时，微调后 LBM 的优势反而**扩大了**。在一组模拟实验中，正常条件下它在 16 项任务中有 3 项在统计上胜过从头训练的策略；而在分布偏移条件下，这一数字变成了 10/16。由于真实部署总会与训练条件有所偏离，这种稳健性或许是最具实际意义的发现。

更多预训练数据带来了平滑的提升。随着加入的预训练数据增多，性能稳步上升——在所测试的规模范围内，**没有突然跃升或“涌现式”飞跃**。这种提升有用、可预测，而且并不戏剧化。

但“不经微调的通才”这一说法没有成立。预训练后的 LBM 以零样本方式使用——不进行任务特定微调——并没有持续胜过单任务策略。一个网络确实可以同时完成许多任务，但“只要给它提示就行”的设想在这里没有得到验证；作者认为，部分原因在于他们使用的小型语言编码器不够稳健。

而且这些增益小到很容易被漏掉——甚至被伪造出来。许多效果只有在样本量大于通常水平、测试又足够谨慎时才显现。作者直截了当地指出，考虑到效果的大小和噪声，**许多机器人论文测到的可能是**

统计噪声，这一风险相当高。他们还发现，一个看似普通的选择——数据如何进行归一化——对结果的影响超过了架构变化；预训练中的一个归一化**错误**，直到评估结束后才被发现。

这大概意味着什么

最站得住脚的解读是：在多样化机器人数据上进行大规模预训练，是一个真实且值得采用的组成部分——它让新任务所需的数据更少，也让策略在现实不符合训练条件时更加稳健。这确实支持了该领域正在押注的方向。但这些收益是有限且有条件的（它们主要在微调之后出现，并且在汇总结果和压力条件下最为清晰），并不意味着一个即插即用的通用机器人已经到来。

更安静、也更重要的含义在方法论层面。这篇论文实际上是一把“标尺”：它展示了要对机器人策略提出可信主张，究竟需要多少证据，也暗示许多已发表论文中的乐观结论建立在过少的证据之上。这是该领域比另一款模型更需要的纠偏。

这并不能证明什么

- 它**不是通用机器人**。这些优势是在特定架构（扩散策略）上、针对每项任务进行微调，并在受控环境中利用遥操作示范数据得到的；它不是一个能按要求自主完成任意新工作的机器人。
- 它并没有验证**零样本**使用。没有微调时，这个大型模型并没有持续胜过单任务基线。
- 它并不是“**涌现式飞跃**”的证据。扩展规模带来了平滑的改进；这里没有任何不连续性可以支持“然后它突然获得了能力”之类的叙事。
- 这些数字是**相对的，而且局限于实验室**。绝对成功率被有意调到约 50%，以提高比较的敏感度；它们不是现实世界可靠性的衡量，而且这项工作只涉及一个实验室的一种架构。
- 它并没有确定任何策略成功或失败的原因，而且论文报告了几个大型模型表现**更差**的具体任务，却没有解释原因。
- 对评估装置之外的**安全性、自主性或部署**，它什么也没有说明。

证据有多强？

对于论文的核心比较性主张——微调后的 LBM 在汇总结果中胜过从头训练的基线、所需数据量少数倍，并且在分布偏移下更加稳健——证据既强，又经过异常严格的控制：盲测、随机化、大样本、统计检验，并对评分进行质量保证复核。这是少见的情况：其方法学严谨到足以让人近乎按字面接受它的主要结论。

作者自己提出了诚实的保留意见。他们的误差条反映了评估的随机性，**却没有**反映训练的随机性——同一个模型训练两次，可能得到明显不同的策略，而这种变化没有纳入统计结果。真实世界任务每项有 50 次试验，足以发现中等幅度的效果，却可能漏掉较小的效果。语言条件输入使用了一个小型编码器，因此对于“只要告诉机器人该做什么”的主张，更大型系统可能会给出不同结果。此外，还有一个在事后披露的归一化错误。这些问题都不足以推翻主要发现，但正是这类问题，论文认为该领域通常会将

其一扫而过。

还有一点关于资料来源的说明，也符合这篇论文的精神：本文解读基于作者的预印本。我们未能获取期刊发表版本，因此没有核对预印本与发表文本之间是否有任何变化。

为什么这很重要

“机器人基础模型”是一个很容易被过度解读的说法，而像这样的研究也很容易被误读成两个极端——胜利式的“它成功了！”，或者轻蔑式的“这只是炒作。”。准确的看法比这两种都更有用：在多样化数据上进行预训练，确实带来了可测量但幅度适中的收益——主要是每项任务需要更少数据，以及更强的稳健性——而且随着规模扩大，这条路径以可预测的方式变好。

它更深层的意义在于，这篇论文把严谨性用来审视自己的领域。它表明，真实效果小到足以在草率评估中消失，而数据归一化这种无聊的选择，影响甚至可能超过聪明的新架构；因此，许多机器人学习领域的进展，在被相信之前都需要更稳固的测量。一篇花费自身信誉去厘清结果与愿望之间差别的论文，比另一篇登上排行榜榜首的论文更罕见，也更有价值。

简明总结

丰田研究院的研究人员训练了“**大型行为模型**”——在约 1,700 小时多样化操作数据上预训练的基于扩散的机器人策略——并使用一套异常严格的流程，将它们与从头训练的单任务策略进行比较：盲测、随机化、大样本（约 1,800 次真实世界试验和 47,000 多次模拟试验），并采用真正的统计分析。经过逐任务微调后，大型模型在汇总结果中可靠地表现更好，所需的任务特定数据约为从头训练所需数据的 **1/3~1/5**，并且在条件发生变化时**更加稳健**；随着预训练数据增加，性能也平滑提升。但在**不进行微调**的情况下，它们并没有持续胜过单任务模型；有些效果小到只有大样本量才能显现，而一个普通的数据归一化选择所产生的影响超过了架构。它为机器人基础模型的发展方向提供了扎实而克制的支持——**不是通用机器人**，不是零样本通才，也不是“涌现式飞跃”——同时尖锐地提醒我们，机器人领域的大量研究可能测到的是噪声。

No-BS 检查

论文展示了什么：通过严谨、盲测且具备统计效力的评估（约 1,800 次真实世界运行 + 47,000 多次模拟运行），先进行多任务预训练、再进行微调的扩散策略（LBM）在汇总结果中胜过从头训练的单任务策略，只需约 **1/3~1/5** 的任务特定数据就能达到相当表现，并且在分布偏移下更加稳健；随着预训练数据增加，性能平滑提升。

合理但尚未证明的内容：这些收益是否会迁移到规模大得多的视觉-语言-动作模型（他们的语言编码器规模较小）；这种平滑扩展是否会在所测试数据范围之外继续存在。

它没有展示什么：通用或零样本机器人（不进行微调 → 没有持续优势）；任何“涌现式”能力跃升；对

具体任务层面失败的解释；现实世界中以绝对值衡量的可靠性（成功率被调到接近 50%，以提高敏感度）；任何关于安全性或自主部署的内容。

主要局限：统计结果反映了评估的随机性，却没有反映训练运行的随机性；每项任务 50 次真实世界试验可能漏掉较小的效果；只有一种架构和一个实验室；语言编码器规模适中；评估结束后才发现一个数据归一化错误；分析基于预印本（未核对发表版本）。

普通读者应有多大信心？对于“多任务预训练加微调能够带来真实但适中的收益”这一点，信心应当很高——尤其是数据效率和稳健性，而且这些收益经过了异常仔细的测量。对于这不是**通用**或**零样本**机器人、也不是“涌现式飞跃”，信心同样应当很高。对于这些收益能在多大程度上扩展到更大型号，信心为中等。还应认真看待作者自己的警告：该领域的效果小到样本效力不足的研究可能报告的只是噪声。恰当的态度是：对这一路线持克制的乐观，同时对缺乏这类统计支持的机器人 AI 结果保持健康的怀疑。

来源

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/>