



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——适度地是的，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）并使用真实统计。经过逐任务微调后，大模型平均表现更好，需要大约少 3–5 倍的任务特定数据，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。

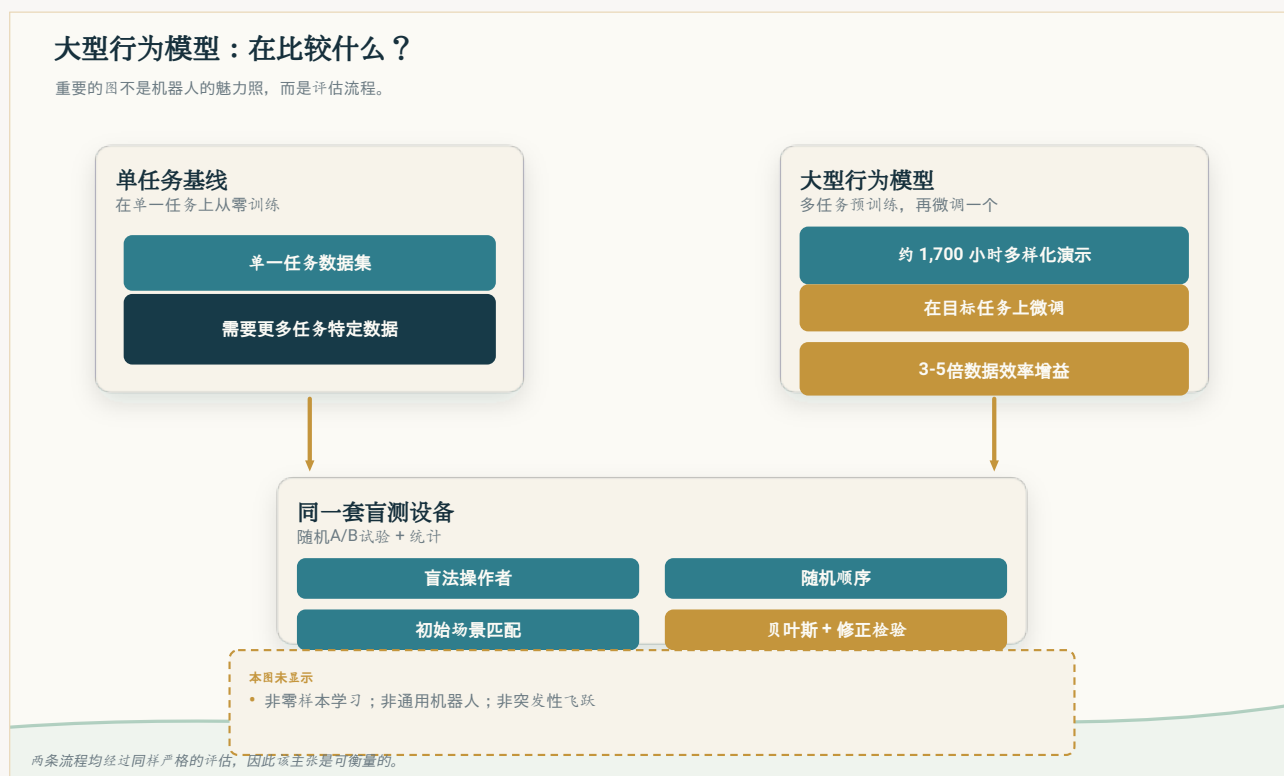
Written by Lucio Vaglio · Cross-checked by Laura Nesso · Edited by Michele Renda · Figures by Laura Nesso
· AI-assisted, human-reviewed

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

一篇关于机器人的论文，其真正主题是诚实

机器人学正处于「基础模型」热潮之中。这个想法借鉴自语言和视觉 AI，非常诱人：与其一次为一个任务训练一个机器人，不如在一个庞大、多样化的演示数据集上训练一个大型「**大型行为模型**」(LBM)，从而得到一个能力广泛且能快速适应的系统。热情——以及投资——是巨大的。标题不言自明：通用机器人来了。

这篇来自丰田研究所的论文之所以有趣，恰恰是因为它拒绝写出这样的标题。它真正的贡献不是一个更花哨的机器人。而是对一个看似简单的问题——大模型方法真的更好吗？我们又该怎么知道？——的一次严厉审视，并以在这个领域据作者自述不同寻常的统计严谨性作出了回答。其结果是「是的，但是」——一条有用的结论，以及一则警告：机器人学中的许多研究可能在测量噪声。



两种方法都进入相同的盲法随机化评估。论文的主张并非机器人基础模型是零样本万能者，而是预训练可以通过仔细的测量提高数据效率和鲁棒性。

Original The Clean Paper diagram CC BY 4.0

作者做了什么

他们以一种具体而明确的方式构建了 LBM：**基于扩散的视觉运动策略**（一种扩散 Transformer，读取相机图像、简短的语言指令和机器人自身的关节位置，并以 10 Hz 输出短促的运动指令）。这些策略在**约 1,700 小时**的机器人演示数据上进行了**预训练**——涵盖内部收集的 500 多个不同任务以及公共数据集——然后对单个任务进行**微调**。全文的比较对象始终是一个**从零开始训练的单一任务策略**。

然而，论文的核心是评估，作者将其视为主要成果。为了避免自欺欺人，他们使用了：

- **真实世界中的盲法、随机化 A/B 测试**——操作机器人的人不知道正在测试哪种策略，测试顺序也被随机化。
- **受控、可重复的初始条件**——操作员在每次试验前将场景与图像叠加层对齐。
- **大样本量**——每个任务、每种策略、每种条件下进行 50 次真实世界 rollout；模拟中每个任务进行 200 次。总计：约 **1,800 次盲法真实世界 rollout 和超过 47,000 次模拟 rollout**。
- **恰当的统计方法**——成功概率的贝叶斯估计，以及带有多重比较校正的配对假设检验，而非仅凭目测重叠的误差条。他们甚至对四分之一的人工评分试验进行了质量保证通行，以测量评分误差。

[video pending]

设置早餐桌的模型并排对比：（左）单任务基线，（右）LBM。两个视频均以 1 倍速度播放。这是一个被评估的任务，而非通用自主性的证明。

这套机制才是重点。整篇论文都在论证：没有它，你无法将真正的改进与运气区分开来。

他们发现了什么

经微调的大型模型平均而言优于从零开始的单任务模型。在所有任务上聚合，经过预训练然后微调的 LBM 在模拟和真实世界中均可靠地优于同一任务上从零开始训练的策略，且这种差异具有统计显著性。在单个任务上，微调 LBM 在几乎所有情况下均统计上达到同等或更优（真实世界任务 3/3，模拟任务 15/16）。

最大且最清晰的收益在于数据效率。一个微调 LBM 使用大约少 **3-5 倍** 的任务特定数据即可达到从零开始相当的性能。在一个真实世界任务（设置早餐桌）中，一个仅在 **15%** 演示数据上微调的 LBM 击败了 **100%** 数据上训练的单任务策略。

当条件发生变化时，预训练帮助最大。当测试环境发生变化（例如，不同的照明或物体位置）时，微调 LBM 比从零开始的策略保持了更多的性能。

收益随预训练数据量平稳增长。将预训练数据从几小时扩展到全部 1,700 小时，获得了持续而平滑的收益，没有饱和的迹象。没有突然的跳跃，没有「涌现」的门槛——只是一条持续向上的曲线。

一个单调的数据选择比其他任何因素都更重要。在撰写本文期间，他们发现训练数据中的一个标准化错误——一个影响输入编号方式的错误——正在损害从零开始的性能，使其看起来比应有的更差。修复后，两者之间的差距缩小了，尽管 LBM 仍然胜出。这一点之所以至关重要，不是因为结果被否定了，而是因为它凸显了该领域测量的脆弱性：如果一个数据预处理的错误可以在信号中造成真正模型差异大小的起伏，那么许多已发表的效果在没有这类严格基础设施支撑的情况下，很可能无法复现。

没有微调就没有优势

一个关键的否定性结果：**在没有微调的情况下，大模型并未一致地超越从零开始的模型。**一个仅在广泛数据集上预训练、没有特定任务微调的 LBM，在真实世界中的表现并不比一个对该任务从零开始训练的策略更好。这项工作中没有出现零样本万能者。LBM 是一种更好的起点，而不是一个更好的成品。

这没有证明什么

- 它**没有**展示一个通用机器人或零样本自主性。无微调 = 无优势。
- 它**没有**评估安全性、可靠性或部署就绪性。成功是二元的且经过调校，使得区分度最大化（使成功/失败接近 50/50），而非最大化绝对成功。
- 它**没有**将系统描述为能执行人类无法处理的任务。
- 它**没有**证明「涌现」或任何能力上的不连续跳跃。
- 它**没有**提供超出所研究设置范围的可泛化性证据——没有工业部署、没有家庭实验、没有长期自主性。

证据有多强？

对平均改进，证据是强有力的。真实世界和模拟中的聚合效应量在统计上显著且方向一致。严格的评估协议给机器人学领域带来了不寻常的可信度。

对单个任务，即使在 50 次试验中，效应也较小，一些置信区间重叠，且在某些案例中，从零开始的策略在数值上略胜一筹。论文对此是透明的；这与聚合收益的稳健性并不矛盾。

对数据效率，收益是清晰且可操作的。3-5 倍的减少是一个实质性的实用收获——少写数据——即使模型能力本身只是适度提高。

对可扩展性，证据是平滑且单调的，但仅基于一个模型的规模（从约 200 小时扩展到 1,700 小时预训练数据）。没有验证「越大越好」是否会在更大尺度上继续成立。

对标准化效应，教训是既谦卑又重要的：论文最大的效应量之一被追溯至一个数据错误，修复后该效应缩小了。这不是方法论上的失败——这正是那种在事后应被称赞的透明度——但它确实表明，该领域许多已发表的结果可能无法在这种程度的审查下幸存。

为什么重要

「机器人基础模型」是一个为夸大宣称而生的短语，而像这样的研究很容易在任一方向上被误读——作为「它奏效了！」的凯旋，或是「它被过度炒作了」的轻蔑。准确的认知比两者都更有用：在多样化数据集上进行预训练能带来真实、可测量但适度的收益——主要是每任务需要更少的数据，以及更强的鲁

棒性——而且这条路径随规模预测性地改善。

它更深层的原因在于这篇论文将其严谨性对准了它自己的领域。通过展示真实的效应量小到在不严谨的评估下会消失，并且像数据标准化这样一个乏味的选择可能压倒一个聪明的新架构，它提出了一个理由：机器人学习进展中的很大一部分需要更扎实的测量才能被相信。一篇花掉自己的可信度来警界线区分结果与期望的论文，正在做一些比在排行榜上登顶更罕见、也更有价值的事情。

清洁摘要

丰田研究所的研究人员训练了「大型行为模型」——基于扩散的机器人策略，在约 1,700 小时多样化操作数据上进行了预训练——并使用一套异常严格的协议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（约 1,800 次真实世界和 47,000+ 次模拟试验），并进行真实统计。经过逐任务微调后，大模型在聚合中可靠地表现更好，需要大约少 3-5 倍的任务特定数据，并且在条件变化时更鲁棒，性能随预训练数据的增加而平滑提高。但在没有微调的情况下，它们并未一致地击败单任务模型，若干效应量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坚实而有节制的支持——并非一个通用机器人，不是零样本万能者，也不是一次「涌现的飞跃」——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。

No-BS 检查

论文所展示的：通过一次严谨、盲法、统计有力的评估（约 1,800 次真实世界 + 47,000+ 次模拟 rollout），经过多任务预训练然后微调的扩散策略（LBM）在聚合中优于从零开始的单任务策略，使用大约少 3-5 倍的任务特定数据即可达到同等性能，并且在分布偏移下更鲁棒；性能随预训练数据平滑增长。

合理但未经证实的：这些收益可迁移到更大的视觉-语言-动作模型（他们的语言编码器很小）；平稳的增长趋势继续超出所测试的数据范围。

它没有展示的：通用或零样本机器人（无微调 → 无一致优势）；任何「涌现」式的能力跳跃；对特定任务层面失败的解释；绝对意义上的真实世界可靠性（成功率被调校至接近 50% 以获得灵敏度）；关于安全性或自主部署的任何内容。

主要局限：统计捕捉了评估随机性，但未捕捉训练运行的随机性；每个任务 50 次真实世界试验可能遗漏小效应；单一架构和单一实验室；适中的语言编码器；在评估后发现了一个数据标准化错误；分析基于预印本（未检查已发表版本）。

普通读者应该抱有多大信心？ 高度信心——多任务预训练加微调能带来真实、适度的收益——特别是在数据效率和鲁棒性方面——并且这些收益得到了异常仔细的测量。高度信心——这不是一个通用或零样本机器人，也不是一次涌现式的飞跃。中等信心——这些收益在多大规模上可扩展至更大的模型。值得认真对待：作者自己的警告，该领域的效应量已经小到统计力度不足的研究可能正在报告噪声。恰当的立场：对该方法持有节制的乐观，以及对缺乏这类统计基础的机器人 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/>