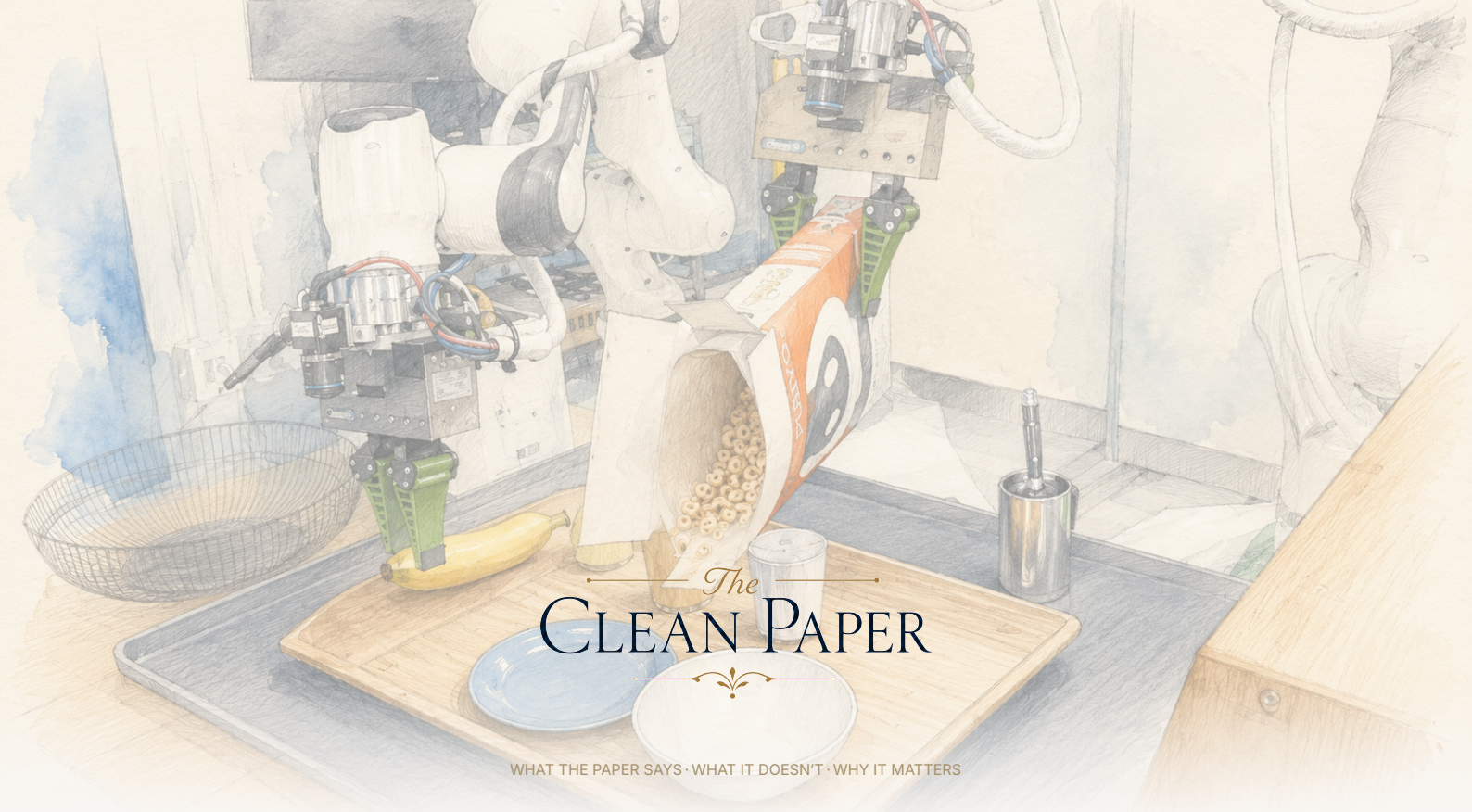




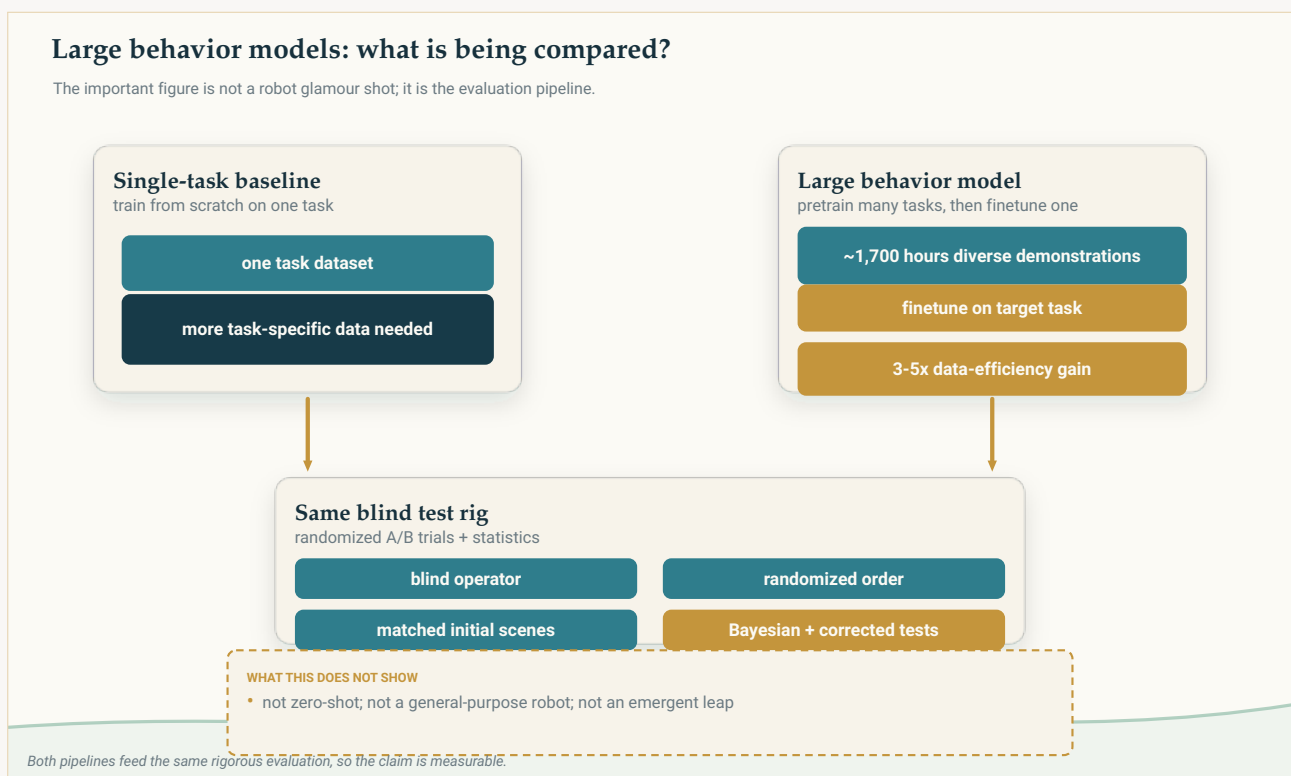
Λειτουργούν πράγματι καλύτερα τα μεγάλα «foundation models» για ρομπότ; Μια προσεκτική απάντηση — ναι, μετρίως, και οι περισσότερες μελέτες δεν μπορούν να το ξεχωρίσουν

To Toyota Research Institute εκπαιδευσε «large behavior models» — πολιτικές ρομπότ προεκπαιδευμένες σε περίπου 1.700 ώρες ποικίλων δεδομένων χειρισμού — και τα δοκίμασε απέναντι σε πολιτικές μίας εργασίας εκπαιδευμένες από το μηδέν, με ασυνήθιστη αυστηρότητα: τυφλές, τυχαιοποιημένες δοκιμές μεγάλου δείγματος (περίπου 1.800 στον πραγματικό κόσμο και πάνω από 47.000 σε προσομοίωση) με κανονική στατιστική ανάλυση. Μετά από fine-tuning ανά εργασία τα μεγάλα μοντέλα ήταν κατά μέσο όρο καλύτερα, χρειάζονται περίπου 3–5× λιγότερα δεδομένα ειδικά για την εργασία και ήταν πιο ανθεκτικά όταν άλλαζαν οι συνθήκες· η επίδοση αυξανόταν ομαλά με περισσότερα δεδομένα προεκπαίδευσης. Χωρίς fine-tuning όμως δεν ξεπερνούσαν συστηματικά τα μοντέλα μίας εργασίας, αρκετές επιδράσεις ήταν τόσο μικρές ώστε χρειάζονταν μεγάλα δείγματα για να φανούν, και μια πεζή επιλογή κανονικοποίησης δεδομένων επηρέαζε περισσότερο από την αρχιτεκτονική. Είναι μετρημένη υποστήριξη για την κατεύθυνση των foundation models στη ρομποτική — όχι ρομπότ γενικής χρήσης, όχι zero-shot γενικευτής, όχι «αναδυόμενο άλμα» — μαζί με μια αιχμηρή προειδοποίηση ότι μεγάλο μέρος της ρομποτικής ίσως μετρά θόρυβο.

Μια εργασία ρομποτικής με πραγματικό θέμα την εντιμότητα

Η ρομποτική βρίσκεται στη μέση ενός πυρετού για τα «foundation models». Η ιδέα, δανεισμένη από την τεχνητή νοημοσύνη γλώσσας και εικόνας, είναι δελεαστική: αντί να εκπαιδεύεις ένα ρομπότ για μία εργασία κάθε φορά, εκπαιδεύεις ένα μεγάλο **«large behavior model» (LBM)** σε ένα τεράστιο, ποικίλο σύνολο επιδείξεων και αποκτάς ένα σύστημα με ευρείες ικανότητες που προσαρμόζεται γρήγορα. Ο ενθουσιασμός — και οι επενδύσεις — είναι τεράστιες. Ο τίτλος γράφεται μόνος του: οι εγκέφαλοι ρομπότ γενικής χρήσης είναι εδώ.

Αυτή η εργασία, από το Toyota Research Institute, είναι ενδιαφέρουσα ακριβώς επειδή αρνείται να γράψει αυτόν τον τίτλο. Η πραγματική της συμβολή δεν είναι ένα πιο εντυπωσιακό ρομπότ. Είναι μια αυστηρή εξέταση μιας απατηλά απλής ερώτησης — λειτουργεί πράγματι καλύτερα η προσέγγιση του μεγάλου μοντέλου, και πώς θα το ξέραμε; — με επίπεδο στατιστικής φροντίδας που, σύμφωνα με τους ίδιους τους συγγραφείς, είναι ασυνήθιστο για τον κλάδο. Το αποτέλεσμα είναι ένα πραγματικά χρήσιμο «ναι, αλλά», μαζί με την προειδοποίηση ότι μεγάλο μέρος της ρομποτικής ίσως μετρά απλώς θόρυβο.



Και οι δύο προσεγγίσεις περνούν από την ίδια τυφλή, τυχαιοποιημένη αξιολόγηση. Ο ισχυρισμός της εργασίας δεν είναι ότι τα foundation models των ρομπότ είναι zero-shot γενικευτές, αλλά ότι η προεκπαίδευση μπορεί να βελτιώσει την αποδοτικότητα δεδομένων και την ανθεκτικότητα όταν μετριέται προσεκτικά.

Τι έκαναν οι συγγραφείς

Κατασκεύασαν LBM με μια συγκεκριμένη, χειροπιαστή έννοια: **οπτικοκινητικές πολιτικές βασισμένες σε diffusion** (έναν Diffusion Transformer που διαβάζει εικόνες κάμερας, μια σύντομη γλωσσική εντολή και τις θέσεις των αρθρώσεων του ίδιου του ρομπότ και παράγει μικρές ακολουθίες κινητικών εντολών στα 10 Hz). Αυτά **προεκπαιδεύτηκαν σε περίπου 1.700 ώρες** επιδείξεων ρομπότ — πάνω από 500 διαφορετικές εργασίες που συλλέχθηκαν εσωτερικά, μαζί με δημόσια σύνολα δεδομένων — και στη συνέχεια υποβλήθηκαν σε **fine-tuning** για μεμονωμένες εργασίες. Σε όλη την εργασία η σύγκριση γίνεται με μια **πολιτική μίας εργασίας εκπαιδευμένη από το μηδέν** στα δεδομένα αυτής της εργασίας.

Η καρδιά της εργασίας, όμως, είναι η αξιολόγηση, την οποία οι συγγραφείς αντιμετωπίζουν ως κύριο αποτέλεσμα. Για να μην εξαπατήσουν τους εαυτούς τους χρησιμοποίησαν:

- **Τυφλές, τυχαιοποιημένες δοκιμές A/B** στον πραγματικό κόσμο — ο άνθρωπος που χειριζόταν το ρομπότ δεν γνώριζε ποια πολιτική δοκιμαζόταν και η σειρά ήταν τυχαιοποιημένη.
- **Ελεγχόμενες, επαναλήψιμες αρχικές συνθήκες** — πριν από κάθε δοκιμή οι χειριστές ταίριαζαν τη σκηνή με μια υπερκείμενη εικόνα αναφοράς.
- **Μεγάλο αριθμό δοκιμών** — 50 πραγματικές εκτελέσεις ανά εργασία, πολιτική και συνθήκη· 200 ανά εργασία στην προσομοίωση. Συνολικά: περίπου **1.800 τυφλές εκτελέσεις στον πραγματικό κόσμο και πάνω από 47.000 εκτελέσεις σε προσομοίωση**.
- **Κανονική στατιστική** — μπειζιανές εκτιμήσεις της πιθανότητας επιτυχίας και ανά ζεύγη ελέγχους υποθέσεων με διορθώσεις για πολλαπλές συγκρίσεις, αντί για οπτική εκτίμηση επικαλυπτόμενων ράβδων σφάλματος. Έκαναν ακόμη και έλεγχο διασφάλισης ποιότητας στο ένα τέταρτο των δοκιμών που βαθμολογήθηκαν από ανθρώπους, για να μετρήσουν το σφάλμα βαθμολόγησης.

[video pending]

Σύγκριση δίπλα-δίπλα μοντέλων που στρώνουν τραπέζι για πρωινό: αριστερά η βασική πολιτική μίας εργασίας και δεξιά το LBM. Και τα δύο βίντεο παίζουν σε ταχύτητα 1x. Πρόκειται για μία αξιολογημένη εργασία, όχι για απόδειξη αυτονομίας γενικής χρήσης.

Αυτός ο μηχανισμός αξιολόγησης είναι το ουσιαστικό σημείο. Ολόκληρη η εργασία υποστηρίζει ότι χωρίς αυτόν δεν μπορείς να ξεχωρίσεις μια πραγματική βελτίωση από την τύχη.

Τι βρήκαν

Τα μεγάλα μοντέλα με fine-tuning ξεπέρασαν κατά μέσο όρο τα μοντέλα μίας εργασίας που εκπαιδεύτηκαν από το μηδέν. Συγκεντρωτικά σε όλες τις εργασίες, ένα LBM που είχε προεκπαιδευτεί και κατόπιν υποβλήθει σε fine-tuning για τη συγκεκριμένη εργασία ξεπερνούσε αξιόπιστα μια πολιτική εκπαιδευμένη από το μηδέν στην ίδια εργασία, τόσο στην προσομοίωση όσο και στον πραγματικό κόσμο, και η διαφορά ήταν στατιστικά σημαντική. Σε

μεμονωμένες εργασίες, το fine-tuned LBM ήταν στατιστικά εξίσου καλό ή καλύτερο σχεδόν σε κάθε περίπτωση (3/3 πραγματικές εργασίες, 15/16 εργασίες προσομοίωσης).

Το μεγαλύτερο και καθαρότερο κέρδος είναι η αποδοτικότητα δεδομένων. Ένα fine-tuned LBM έφτανε την επίδοση ενός μοντέλου από το μηδέν χρησιμοποιώντας περίπου **3–5× λιγότερα δεδομένα ειδικά για την εργασία**. Σε μία πραγματική εργασία (στρώσιμο τραπεζιού για πρωινό), ένα LBM με fine-tuning μόνο στο **15%** των επιδείξεων ξεπέρασε μια πολιτική από το μηδέν που είχε εκπαιδευτεί στο **100%**.

Η προεκπαίδευση βοηθά περισσότερο όταν αλλάζουν οι συνθήκες. Όταν το περιβάλλον δοκιμής μετατοπίζεται σκόπιμα από τις συνθήκες εκπαίδευσης («distribution shift»), το πλεονέκτημα του fine-tuned LBM μεγάλωνε. Σε ένα σύνολο προσομοιώσεων, ξεπερνούσε στατιστικά την πολιτική από το μηδέν σε 3 από 16 εργασίες υπό κανονικές συνθήκες αλλά σε 10 από 16 υπό distribution shift. Επειδή οι πραγματικές εφαρμογές πάντα απομακρύνονται κάπως από τις συνθήκες εκπαίδευσης, αυτή η ανθεκτικότητα είναι ίσως το πρακτικά σημαντικότερο εύρημα.

Περισσότερα δεδομένα προεκπαίδευσης βοηθούσαν, ομαλά. Η επίδοση αυξανόταν σταθερά καθώς προστίθεντο δεδομένα προεκπαίδευσης — χωρίς καμία ξαφνική εκτίναξη ή «αναδυόμενο» άλμα στις κλίμακες που δοκιμάστηκαν. Χρήσιμο, προβλέψιμο, καθόλου θεαματικό.

Αλλά η ιστορία του γενικευτή χωρίς fine-tuning δεν επιβεβαιώθηκε. Ένα προεκπαιδευμένο LBM σε χρήση zero-shot — χωρίς fine-tuning ειδικό για την εργασία — δεν ξεπερνούσε συστηματικά τις πολιτικές μίας εργασίας. Ένα μόνο δίκτυο μπορούσε να εκτελεί πολλές εργασίες, αλλά το όνειρο «απλώς πες του τι να κάνει» δεν επιβεβαιώθηκε εδώ· οι συγγραφείς αποδίδουν μέρος του προβλήματος στην ευθραυστότητα του μικρού γλωσσικού κωδικοποιητή τους.

Και τα κέρδη ήταν αρκετά μικρά ώστε εύκολα να χαθούν — ή να εμφανιστούν ψεύτικα. Πολλές επιδράσεις έγιναν ορατές μόνο με τα μεγαλύτερα από το συνηθισμένο δείγματα και τους προσεκτικούς ελέγχους. Οι συγγραφείς δηλώνουν καθαρά ότι, με δεδομένο το μέγεθος των επιδράσεων και τον θόρυβο, **υπάρχει σημαντικός κίνδυνος πολλές εργασίες ρομποτικής να μετρούν στατιστικό θόρυβο**. Βρήκαν επίσης ότι μια πεζή επιλογή — ο τρόπος κανονικοποίησης των δεδομένων — επηρέαζε τα αποτελέσματα περισσότερο από αλλαγές αρχιτεκτονικής, και ότι ένα σφάλμα κανονικοποίησης στην προεκπαίδευση αποκαλύφθηκε μόνο μετά την ολοκλήρωση των αξιολογήσεων.

Τι πιθανότατα σημαίνει αυτό

Η υπερασπίσιμη ανάγνωση: η προεκπαίδευση μεγάλης κλίμακας σε ποικίλα δεδομένα ρομπότ είναι πραγματικό και χρήσιμο συστατικό — μειώνει τα δεδομένα που χρειάζονται για κάθε νέα εργασία και κάνει τις πολιτικές πιο ανθεκτικές όταν ο κόσμος δεν ταιριάζει με την εκπαίδευση. Αυτό πράγματι στηρίζει την κατεύθυνση στην οποία ποντάρει ο κλάδος. Όμως τα κέρδη είναι μέτρια και υπό προϋποθέσεις (εμφανίζονται κυρίως μετά το fine-tuning και

φαίνονται καθαρότερα συγκεντρωτικά και υπό πίεση), όχι η άφιξη ενός έτοιμου γενικού ρομπότ.

Η πιο ήσυχη και σημαντικότερη σημασία είναι μεθοδολογική. Η εργασία λειτουργεί ουσιαστικά ως μέτρο σύγκρισης: δείχνει πόσα στοιχεία χρειάζονται πραγματικά για να γίνει ένας αξιόπιστος ισχυρισμός για μια πολιτική ρομπότ και υποδηλώνει ότι μεγάλο μέρος του δημοσιευμένου ενθουσιασμού στηρίζεται σε πολύ λιγότερα. Αυτή είναι διόρθωση που ο κλάδος χρειάζεται περισσότερο από ακόμη ένα μοντέλο.

Τι δεν αποδεικνύει

- **Δεν είναι ρομπότ γενικής χρήσης.** Τα κέρδη αποδεικνύονται για μια συγκεκριμένη αρχιτεκτονική (diffusion policies) με fine-tuning ανά εργασία, σε ελεγχόμενες συνθήκες, από τηλεχειριζόμενες επιδείξεις — όχι για ένα αυτόνομο ρομπότ που εκτελεί αυθαίρετες νέες δουλειές κατόπιν εντολής.
- **Δεν επικυρώνει τη χρήση zero-shot.** Χωρίς fine-tuning, το μεγάλο μοντέλο δεν ξεπερνούσε συστηματικά τις βασικές πολιτικές μίας εργασίας.
- **Δεν αποτελεί ένδειξη «αναδυόμενου άλματος».** Η κλιμάκωση βελτίωνε ομαλά τα πράγματα· δεν υπάρχει εδώ ασυνέχεια που να στηρίζει αφηγήσεις τύπου «και ξαφνικά απέκτησε ικανότητα».
- **Οι αριθμοί είναι σχετικοί και περιορισμένοι στο εργαστήριο.** Τα απόλυτα ποσοστά επιτυχίας ρυθμίστηκαν σκόπιμα κοντά στο ~50% ώστε οι συγκρίσεις να είναι ευαίσθητες· δεν αποτελούν μέτρο αξιοπιστίας στον πραγματικό κόσμο, και η δουλειά αφορά μία αρχιτεκτονική από ένα εργαστήριο.
- **Δεν εξηγεί γιατί** οποιαδήποτε πολιτική πετυχαίνει ή αποτυγχάνει, και αναφέρονται αρκετές συγκεκριμένες εργασίες όπου το μεγάλο μοντέλο τα πήγε χειρότερα χωρίς να εξηγείται το γιατί.
- **Δεν λέει τίποτα για ασφάλεια, αυτονομία ή ανάπτυξη** έξω από τη διάταξη αξιολόγησης.

Πόσο ισχυρά είναι τα στοιχεία;

Για τους κεντρικούς συγκριτικούς ισχυρισμούς — ότι τα fine-tuned LBM ξεπερνούν συγκεντρωτικά τις βασικές πολιτικές από το μηδέν, χρειάζονται αρκετές φορές λιγότερα δεδομένα και είναι πιο ανθεκτικά υπό distribution shift — τα στοιχεία είναι ισχυρά και ασυνήθιστα καλά ελεγχόμενα: τυφλά, τυχαιοποιημένα, με μεγάλα δείγματα, στατιστικά ελεγμένα και με έλεγχο ποιότητας στη βαθμολόγηση. Είναι η σπάνια περίπτωση όπου η μεθοδολογία είναι αρκετά καλή ώστε τα βασικά συμπεράσματα να μπορούν να ληφθούν σχεδόν τοις μετρητοίς.

Οι έντιμες επιφυλάξεις είναι εκείνες που θέτουν οι ίδιοι οι συγγραφείς. Οι ράβδοι σφάλματός τους περιλαμβάνουν την τυχαιότητα της αξιολόγησης αλλά **όχι** την τυχαιότητα της εκπαίδευσης — εκπαιδεύστε το ίδιο μοντέλο δύο φορές και μπορεί να πάρετε ουσιαστικά διαφορετική

πολιτική, και αυτή η διακύμανση δεν περιλαμβάνεται στη στατιστική. Οι πραγματικές εργασίες είχαν 50 δοκιμές η καθεμία, αρκετές για μέτριες επιδράσεις αλλά πιθανόν ανεπαρκείς για μικρές. Η γλωσσική προσαρμογή χρησιμοποιούσε έναν μέτριο κωδικοποιητή, οπότε οι ισχυρισμοί για το «απλώς πες στο ρομπότ τι να κάνει» μπορεί να διαφέρουν σε μεγαλύτερα συστήματα. Και υπάρχει η ειλικρινής αποκάλυψη ενός εκ των υστέρων σφάλματος κανονικοποίησης. Τίποτα από αυτά δεν καταρρίπτει τα κύρια ευρήματα, αλλά είναι ακριβώς το είδος λεπτομέρειας που η εργασία λέει ότι ο κλάδος συνήθως σπρώχνει στην άκρη.

Μια σημείωση για τις πηγές, στο ίδιο πνεύμα: αυτή η εξήγηση βασίζεται στο preprint των συγγραφέων. Δεν μπορέσαμε να ανακτήσουμε τη δημοσιευμένη έκδοση του περιοδικού, άρα δεν ελέγξαμε αν υπάρχουν αλλαγές μεταξύ preprint και δημοσιευμένου κειμένου.

Γιατί έχει σημασία

Το «robot foundation models» είναι φράση φτιαγμένη για υπερβολικούς ισχυρισμούς, και μια μελέτη σαν αυτή είναι εύκολο να παρερμηνευτεί προς οποιαδήποτε κατεύθυνση — ως θριαμβευτικό «δουλεύει!» ή ως απορριπτικό «είναι υπερβολικά διαφημισμένο». Η ακριβής ανάγνωση είναι χρησιμότερη και από τα δύο: η προεκπαίδευση σε ποικίλα δεδομένα δίνει πραγματικά, μετρήσιμα αλλά μέτρια οφέλη — κυρίως λιγότερα δεδομένα ανά εργασία και μεγαλύτερη ανθεκτικότητα — και η πορεία βελτιώνεται προβλέψιμα με την κλίμακα.

Ο βαθύτερος λόγος που έχει σημασία είναι ότι η εργασία στρέφει την αυστηρότητά της στον ίδιο της τον κλάδο. Δείχνοντας ότι οι πραγματικές επιδράσεις είναι αρκετά μικρές ώστε να εξαφανίζονται με πρόχειρη αξιολόγηση και ότι μια βαρετή επιλογή όπως η κανονικοποίηση δεδομένων μπορεί να μετρά περισσότερο από μια έξυπνη νέα αρχιτεκτονική, υποστηρίζει ότι μεγάλο μέρος της προόδου στη μάθηση ρομπότ χρειάζεται πιο γερή μέτρηση πριν γίνει πιστευτό. Μια εργασία που ξοδεύει την αξιοπιστία της αστυνομεύοντας τη διαφορά ανάμεσα σε αποτέλεσμα και ευχή κάνει κάτι σπανιότερο και πολυτιμότερο από το να βγει πρώτη σε ένα leaderboard.

Καθαρή σύνοψη

Ερευνητές στο Toyota Research Institute εκπαιδύσαν «large behavior models» — diffusion-based πολιτικές ρομπότ προεκπαιδευμένες σε ~1.700 ώρες ποικίλων δεδομένων χειρισμού — και τα δοκίμασαν απέναντι σε πολιτικές μίας εργασίας εκπαιδευμένες από το μηδέν με ασυνήθιστα αυστηρό πρωτόκολλο: τυφλό, τυχαιοποιημένο, μεγάλου δείγματος (~1.800 πραγματικές και 47.000+ προσομοιωμένες δοκιμές), με κανονική στατιστική. Μετά από fine-tuning ανά εργασία, τα μεγάλα μοντέλα ήταν αξιόπιστα καλύτερα συγκεντρωτικά, χρειάζονταν περίπου 3–5× λιγότερα δεδομένα ειδικά για την εργασία και ήταν πιο ανθεκτικά όταν άλλαζαν οι συνθήκες, ενώ η επίδοση βελτιωνόταν ομαλά καθώς αυξάνονταν τα δεδομένα προεκπαίδευσης. Χωρίς όμως **fine-tuning** δεν ξεπερνούσαν συστηματικά τα

μοντέλα μίας εργασίας, αρκετές επιδράσεις ήταν τόσο μικρές ώστε αποκαλύφθηκαν μόνο χάρη στα μεγάλα δείγματα, και μια πεζή επιλογή κανονικοποίησης δεδομένων επηρέαζε περισσότερο από την αρχιτεκτονική. Είναι στέρεη, μετρημένη υποστήριξη για την κατεύθυνση των foundation models στη ρομποτική — **όχι** ρομπότ γενικής χρήσης, **όχι** zero-shot γενικευτής και **όχι** «αναδυόμενο άλμα» — μαζί με μια αιχμηρή προειδοποίηση ότι μεγάλο μέρος της ρομποτικής ίσως μετρά θόρυβο.

Έλεγχος χωρίς υπερβολές

Τι δείχνει η εργασία: Με αυστηρή, τυφλή και στατιστικά επαρκή αξιολόγηση (≈ 1.800 πραγματικές + 47.000+ προσομοιωμένες εκτελέσεις), diffusion policies που προεκπαιδεύονται σε πολλές εργασίες και κατόπιν υποβάλλονται σε fine-tuning (LBM) ξεπερνούν συγκεντρωτικά πολιτικές μίας εργασίας από το μηδέν, φτάνουν ισοδύναμη επίδοση με $\sim 3\text{--}5\times$ λιγότερα δεδομένα ειδικά για την εργασία και είναι πιο ανθεκτικές υπό distribution shift· η επίδοση κλιμακώνεται ομαλά με τα δεδομένα προεκπαίδευσης.

Τι είναι εύλογο αλλά δεν έχει αποδειχθεί: Ότι αυτά τα οφέλη μεταφέρονται σε πολύ μεγαλύτερα vision-language-action μοντέλα (ο γλωσσικός κωδικοποιητής τους ήταν μικρός)· ότι η ομαλή κλιμάκωση συνεχίζεται πέρα από το εύρος δεδομένων που δοκιμάστηκε.

Τι δεν δείχνει: Ρομπότ γενικής χρήσης ή zero-shot (χωρίς fine-tuning $\rightarrow$ κανένα σταθερό πλεονέκτημα)· οποιοδήποτε «αναδυόμενο» άλμα ικανότητας· εξηγήσεις για συγκεκριμένες αποτυχίες ανά εργασία· απόλυτη αξιοπιστία στον πραγματικό κόσμο (τα ποσοστά επιτυχίας ρυθμίστηκαν κοντά στο 50% για ευαισθησία)· οτιδήποτε για ασφάλεια ή αυτόνομη ανάπτυξη.

Κύριοι περιορισμοί: Η στατιστική περιλαμβάνει την τυχαιότητα αξιολόγησης αλλά όχι τη διακύμανση μεταξύ ανεξάρτητων εκπαιδεύσεων· 50 πραγματικές δοκιμές ανά εργασία μπορεί να χάνουν μικρές επιδράσεις· μία αρχιτεκτονική και ένα εργαστήριο· μέτριος γλωσσικός κωδικοποιητής· βρέθηκε σφάλμα κανονικοποίησης δεδομένων μετά τις αξιολογήσεις· η ανάλυση βασίζεται στο preprint (δεν ελέγχθηκε η δημοσιευμένη έκδοση).

Πόση εμπιστοσύνη πρέπει να έχει ένας γενικός αναγνώστης; Υψηλή ότι η προεκπαίδευση πολλών εργασιών μαζί με fine-tuning δίνει πραγματικά, μέτρια οφέλη — ιδιαίτερα στην αποδοτικότητα δεδομένων και την ανθεκτικότητα — και ότι αυτά μετρήθηκαν ασυνήθιστα προσεκτικά. Υψηλή ότι αυτό **δεν** είναι ρομπότ γενικής χρήσης ή zero-shot και **δεν** είναι αναδυόμενο άλμα. Μέτρια για το πόσο μακριά κλιμακώνονται τα κέρδη σε μεγαλύτερα μοντέλα. Και αξίζει να ληφθεί σοβαρά η προειδοποίηση των ίδιων των συγγραφέων ότι οι επιδράσεις του κλάδου είναι αρκετά μικρές ώστε μελέτες χαμηλής στατιστικής ισχύος να μπορεί να αναφέρουν θόρυβο. Η κατάλληλη στάση: μετρημένη αισιοδοξία για την προσέγγιση και υγιής σκεπτικισμός απέναντι σε αποτελέσματα robot-AI χωρίς τέτοια στατιστική στήριξη.

Πηγές

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/>