

Crashing Waves vs. Rising Tides: Findings on AI Automation from Thousands of Worker Evaluations of Labor Market Tasks*

Matthias Mertens
MIT FutureTech

Adam Kuzee
MIT FutureTech

Brittany S. Harris
MIT FutureTech

Harry Lyu
MIT FutureTech

Wensu Li
MIT FutureTech

Jonathan Rosenfeld
MIT FutureTech

Meiri Anto
MIT FutureTech

Martin Fleming
MIT FutureTech

Neil Thompson
MIT FutureTech

July 2026

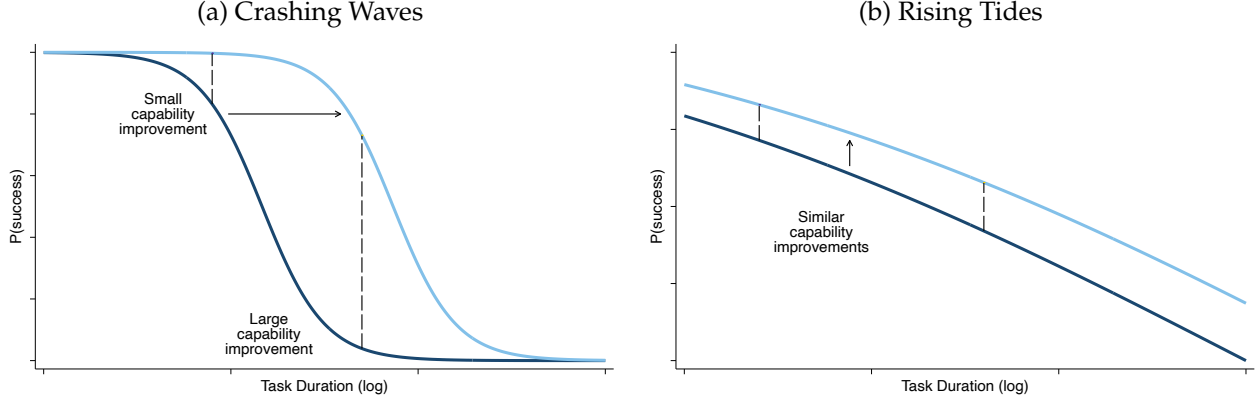
Abstract

We characterize AI automation as a continuum between **crashing waves**, in which capabilities jump abruptly across narrow task sets, and **rising tides**, in which capabilities improve continuously and broadly. Using evidence from more than 6,000 text-based, LLM-addressable tasks derived from the U.S. Department of Labor’s O*NET taxonomy and over 60,000 evaluations by experienced workers, we find little evidence of crashing waves (contrary to existing views). Instead, rising tides are the primary form of AI progress. AI performance is high and improving rapidly across many tasks. In 2024-Q2, models completed text-based tasks that take humans about 1.5 hours to complete with roughly 60% success, rising above 70% by 2025-Q3. If recent trends in AI capability growth persist, frontier LLMs will be able to complete most text-based tasks at minimally sufficient quality with 88%–97% success by 2030.

*Corresponding authors: Matthias Mertens, mmertens@mit.edu; Neil Thompson, neil_t@mit.edu. We thank David Autor for his insightful comments. We are grateful to Annie Lin, Amelia Michael, Peter Olkhovets, and Tiffany Wang for their excellent work as research assistants. We also thank Justin Viola for excellent software engineering work, and Tess Fagan for project support and software development. Funding for this research was provided by Open Philanthropy and a technology company.

1 Introduction

Figure 1: Crashing Waves vs Rising Tides in AI Automation



Notes: Diagram of the distinction between AI automation that comes as "Crashing Waves" (Panel (a)) and "Rising Tides" (Panel (b)).

How does AI progress across labor market tasks? Does it advance by crossing successive thresholds that make previously infeasible tasks suddenly automatable, or does it improve more gradually and broadly across the task space? Conditional on this *shape* of progress, what is the *pace* of capability improvement? These questions are central to understanding the future of work. In this paper, we provide new insights on these questions by conducting a novel large-scale survey in which experienced workers evaluate LLM-generated responses across *thousands* of representative labor market tasks.

A prominent view is that large language model (LLM) performance decays exponentially with the number of sequentially dependent steps in a task (Dziri et al., 2023; LeCun et al., 2023). In a performance–task-duration space, this corresponds to a steep negative success curve, as recently also documented by Kwa et al. (2025) for some stylized research and software-engineering tasks. Under this view, model improvements make longer tasks which were previously infeasible suddenly feasible, as illustrated in Figure 1(a). In practice, this would lead to harsh surprises for human workers. Over just a short period of time, they would observe AI models going from nearly always failing to nearly always succeeding. We call this the "crashing waves" view of AI progress.

We contrast this view with a "rising tide" perspective (Figure 1(b)), in which performance is lifted more broadly across the task space. The central difference between the two phenomena is that under the "rising tide" view, the slope of the relationship between AI performance and (log) task duration is much flatter. AI progress would then translate into more gradual automation, such that individual workers are less likely to be blindsided by AI (although a quickly rising tide can still have substantial labor market impacts, as we discuss below).

For our analysis, we collect novel evaluations of LLM outputs by domain experts across 41 large language models, covering over 12,000 unique task examples ("instances") based on more than 6,000 text-based or partially text-based O*NET task descriptions that are LLM-addressable (O*NET, 2024). Outputs are scored by human evaluators with relevant on-the-job experience. Given the paper's primary focus on automation, we center our analysis on success measures defined by expert evaluations indicating that the LLM output required no human intervention. Conditional on this criterion, we

distinguish LLM output quality across three thresholds: minimally sufficient, average worker quality, and superior quality relative to a human worker.

Our main findings are:

1. **The success-duration curve is relatively flat — consistent with a rising tide of AI automation.** This pattern holds across models of different sizes, as well as different model vintages across time. Automation within particular O*NET "job families" (e.g., management or community and social service) also follows the same rising-tide pattern. There are, however, meaningful differences in the slope of the success-duration curves across job families (i.e., the rising-tide pattern is sometimes a bit stronger or weaker), as one would expect given the differences in task structure (see Section 4.2).
2. **AI performance is high.** Across all LLMs in our survey, we find strong capabilities. Models can do a minimally-sufficient job without human edits on roughly *half to three-quarters* of text-based tasks presented to them. Across all models and tasks in our data, models could complete 60% of tasks at at least a minimally sufficient quality level, 49.7% at at least an average quality level, and 27% at a superior quality level.
3. **AI performance is improving rapidly across a wide range of task durations.** Between 2024-Q2 and 2025-Q3, frontier models went from achieving a 60% success rate on 1.5-hour tasks to a 70% success rate (minimally sufficient quality level). These changes imply quickly improving AI capabilities. The average failure rate (1 minus the success rate) halves every 2.2-2.8 years across tasks which take humans 5 minutes to 24 hours. Over our observation period (2024-Q2 to 2025-Q3), this corresponds to annual success rate increases between 7.9-10 percentage points. Similarly, the task-duration doubling time—the interval between model releases required for newer models to achieve the same success rate on tasks twice as long—is 2.7 months. Compared with prior benchmark-based evidence (Kwa et al., 2025; METR, 2025), these improvement rates are rapid, placing them at the upper end of existing estimates.
4. **The performance gains from increasing model size are different than those from newer model vintages.** As new models are created, they outperform similarly sized models from the past and the shift in the success-duration curve is approximately parallel. So, success on both short and long-duration tasks is improved. By contrast, larger models released at the same time as smaller models do outperform on short-duration tasks, but have only modest improvements on longer-duration tasks.
5. **By 2030, frontier LLMs could complete tasks in our data with a success rate of 88-97%.** When extrapolating the historical progress in AI capabilities, we estimate that by 2030 AI systems will complete most tasks in our data at a minimally sufficient quality level at a success rate of 88-97% (depending on the tasks' duration). At more demanding quality thresholds, the corresponding projected success rates are 85-96% for average-worker quality and 78-91% for quality exceeding that of the average worker.

What do these findings imply for labor markets? The pace of AI progress is rapid, and substantial labor market impacts could arise if rising capabilities translate into widespread adoption. However,

our projections are less extreme than some existing scenarios (e.g., Kokotajlo et al., 2025), and our findings therefore suggest that governments still have a window to build institutions capable of responding to AI-induced labor market disruptions. In this context, the rising-tide nature of AI progress is a double-edged sword. On the one hand, AI is unlikely to suddenly outperform humans across large groups of tasks from one year to the next. Instead, the broad-based improvement we document is likely to unfold over several years, making it predictable and giving policymakers time to adapt. On the other hand, a rising tide implies that AI capabilities will improve across a wide range of text-based tasks, leaving fewer areas of work fully insulated from automation pressure. Institutional responses must therefore account not only for sudden displacement in narrow areas, but also for gradual, diffuse changes across the wider task space.

We note that our estimates assume AI progress continues at the pace observed over the past years and should therefore be interpreted as an upper-bound (i.e., particularly fast) scenario. There are several reasons to expect a potential slowdown in AI capability growth, including limits to scaling compute, as well as possible slowdowns in hardware progress and algorithmic innovation.

Related work. Our paper connects to a growing literature measuring AI capabilities. Much of this literature evaluates AI systems on tasks with objective or standardized evaluation criteria, such as mathematical and coding problems, multiple-choice questions, unit tests, or reference deliverables (e.g., Hendrycks et al. (2021), Jimenez et al. (2024), Rein et al. (2025), Srivastava et al. (2023), Wang et al. (2024), and Wijk et al. (2025)). In contrast, we study realistic labor market tasks whose outputs are inherently open-ended and are evaluated by workers with relevant on-the-job experience. The closest comparison is Patwardhan et al. (2025), who compare LLM outputs against human deliverables for 44 U.S. occupations. We complement this work and cover more than 6,000 O*NET tasks across nearly the full U.S. occupational landscape. To our knowledge, this provides the broadest assessment to date of AI capabilities on realistic and representative labor market tasks.

Our paper also relates to work studying AI automation potential and AI use. Eloundou et al. (2024) use LLMs to classify the potential exposure of occupational tasks to AI, while Appel et al. (2025) and Tomlinson et al. (2025) map AI conversations to O*NET tasks to study usage patterns. We complement this work by directly evaluating model performance: rather than measuring potential exposure or observed use, we ask whether LLM outputs would be accepted by workers with relevant occupational experience.

Finally, we contribute to research on the pace and shape of AI progress. Prior work argues that capabilities may appear abruptly with scale and be difficult to predict from smaller models (Ganguli et al., 2022; Wei et al., 2022), although such discontinuities may partly reflect thresholded or nonlinear evaluation metrics (Schaeffer et al., 2023). Empirically, Kwa et al. (2025) and METR (2025) analyze smaller sets of largely deterministic tasks (e.g., coding) and stylized benchmarks and find rapid growth in the maximum human-equivalent task duration AI systems can complete with 50% success. They also document a steep decline in performance with task duration, consistent with a “crashing waves” view of AI progress (Figure 1). We apply a similar success–duration framework to thousands of realistic, text-based, and non-deterministic labor market tasks. While our estimated pace of progress is broadly consistent with the faster estimates in Kwa et al. (2025) and METR (2025), the relationship

between task duration and performance is substantially flatter. Capability gains in our real-world labor market task setting therefore produce broad improvements across the task space ("rising-tides" view) rather than sharp gains around a narrow duration threshold.

2 Results

2.1 Approach and Baseline Results

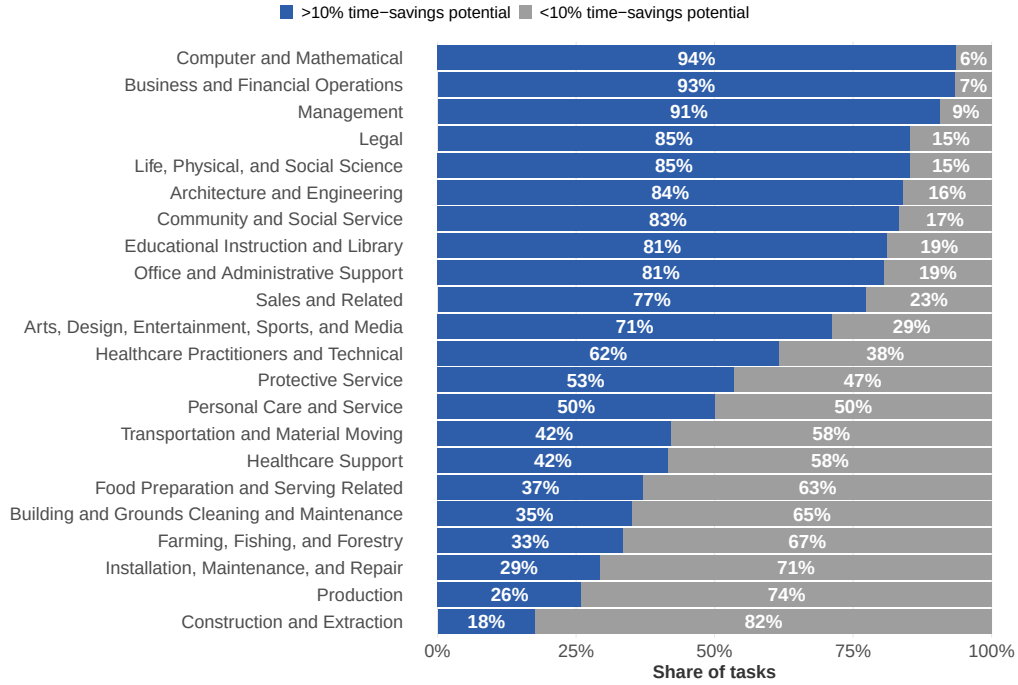
Survey Data. We collected novel survey data evaluating LLM performance on 12,078 realistic and representative instances of 6,648 LLM-addressable labor market tasks drawn from the O*NET task classification. Here, we provide a high-level overview of the data collection process. Section 4.1 provides a detailed description. First, we used GPT-4 to screen O*NET tasks for automation potential and retained tasks for which LLM assistance was estimated to generate at least a 10% time savings. Through this filtering, we retain both fully text-based tasks and partially text-based tasks with a meaningful textual component (the vast majority), providing a comprehensive picture of LLMs' potential applications across the economy. Figure 2 Panel (a) reports the share of tasks that meet this criterion by O*NET job families (the data on 6,648 tasks we collected is a subset of those tasks—see Section 4.1).

For each task, we collected up to two textual task instances, each of which were then completed by 5 of the 41 LLMs in our model pool. Model outputs were evaluated by human evaluators with relevant on-the-job experience. Evaluators judged whether the displayed tasks were realistic and representative, and we retained only task instances that passed this assessment. Evaluators provided contextual information on the task (most importantly, time requirements) and rated each model response on a 1–9 scale related to how a hypothetical manager would view the quality of the response. Table 1 provides an overview on the possible evaluator scores used in our survey. A score of 1 indicates that the output would need to be redone from scratch. A score of 7 or higher indicates that the output is "useful as is" and requires no further edits, corresponding to minimally sufficient (7), average (8), and superior (9) quality. In our analysis, we use a rating of at least 7 as our baseline measure of automation potential and focus on the share of LLM responses meeting or exceeding this threshold.

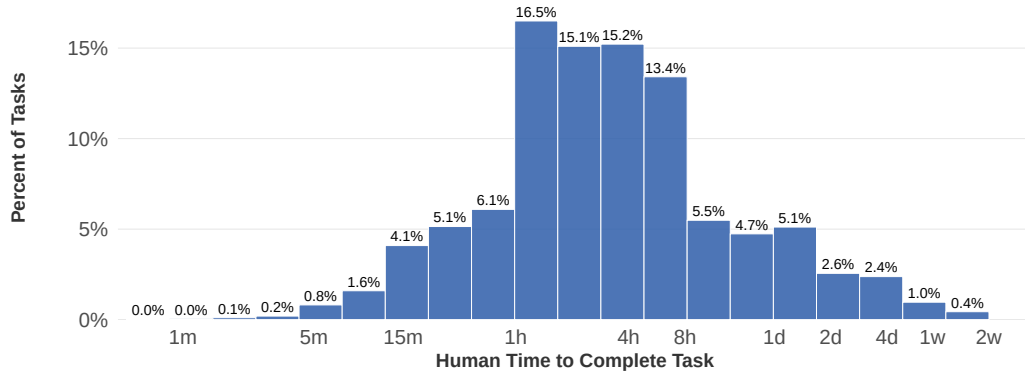
Figure 2 Panel (b) displays the distribution of reported task durations. Observed tasks range from less than one minute to over a week, with most tasks in our survey taking between 1 and 10 hours for humans to complete.

Figure 2: Survey Coverage and Task Duration Distribution

(a) Share of O*NET tasks with at least 10% LLM time-saving potential



(b) Distribution of task durations



Notes: Panel (a) reports the share of O*NET tasks within each job family that GPT-4 classified as offering at least 10% potential time savings from LLM use. Panel (b) shows the distribution of task durations reported by survey participants across 20 equally sized log-spaced bins.

Table 1: LLM Response Quality Scores

Score	Evaluation	Implication & interpretation
1–3	Not useful	The output requires substantial or complete rework. The LLM does not meaningfully substitute for human task performance.
4–6	Useful with edits	The output assists the worker but requires human revision. This represents task augmentation rather than autonomous task completion.
7	Minimally sufficient without edits	The output can be used as delivered and meets the minimum requirements of the task. This is our baseline measure of potential task automation.
8	Average-worker quality without edits	The output can be used as delivered and is comparable in quality to the work of an average human worker.
9	Superior quality without edits	The output can be used as delivered and exceeds the quality of an average human worker.

Notes: The table summarizes the possible evaluator scores used in our survey, which range from 1 to 9.

Regression framework. The key relationship we study is how LLM performance varies with task duration. Our main specification estimates the following logistic model:

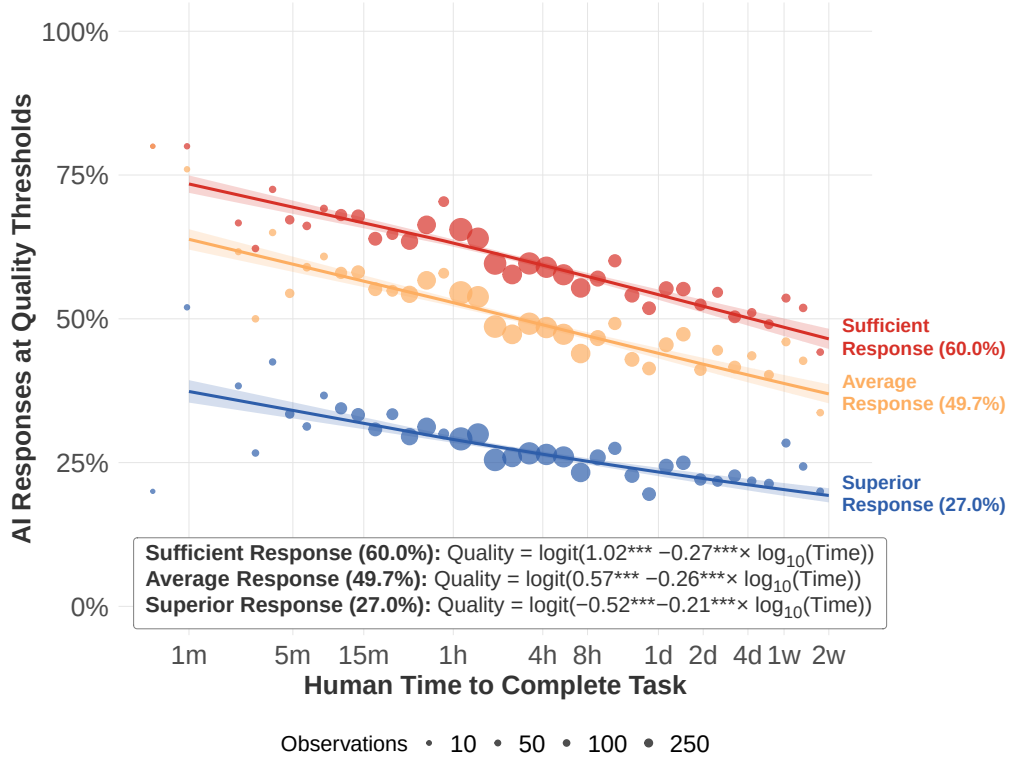
$$\Pr(Y_{ijsm} = 1) = \Lambda(\alpha + \beta \log_{10} T_{js}) = \frac{\exp(\alpha + \beta \log_{10} T_{js})}{1 + \exp(\alpha + \beta \log_{10} T_{js})}. \quad (1)$$

Here, $\Lambda(\cdot)$ denotes the logistic CDF and α is a constant. Y_{ijsm} is an indicator equal to one if the evaluator reports that a hypothetical manager would accept the response without edits—i.e., the survey rating is ≥ 7 (we also consider thresholds of ≥ 8 and $= 9$). The main regressor, T_{js} , measures task duration in time units. Indices i , j , s , and m denote the evaluator, O*NET task, task instance, and model, respectively. We estimate Eq. (1) by maximum likelihood. The logistic specification follows prior work (Ge et al., 2026; Kwa et al., 2025). In Section 4.2, we provide one possible micro-foundation for Eq. (1) under which the slope coefficient β admits a structural interpretation: it can be mapped to the number of sequentially dependent steps required to complete a task (but our empirical results do not rely on this specific interpretation).

Pooled results. Figure 3 plots estimates of Eq. (1) for three “useful without edits” thresholds: at least minimally sufficient (score ≥ 7), at least average quality (score ≥ 8), and superior quality (score $= 9$). Binned scatter points summarize the raw data.

At the ≥ 7 threshold, a tenfold increase in task duration is associated with a 0.27 decline in the log-odds of success. Evaluated at the sample mean acceptance rate of 60%, this implies a drop in predicted acceptance of about 6.6 percentage points. Appendix Figure A.2 shows that this result is quiet robust across different specifications that control for participant demographics, occupational effects, and/or model-specific differences, with estimated slopes ranging from -0.17 to -0.25 . As shown, these estimates imply relatively flat relationships between AI success (i.e., score ≥ 7) and task duration that can be approximated by a linear function (Appendix Table A.5 and Figure A.3 confirm

Figure 3: Task Instance Automation by Required Task Completion Time



Notes: Each line plots the estimated *logistic* relationship between AI response quality and the time required to complete a task instance, based on Equation (1) estimated without controls (with 95% confidence bands). Coefficients are shown as log-odds on the figure. Standard errors are clustered by participant in parentheses. Significance levels: *** 1%, ** 5%, * 10%. The red line corresponds to responses that are minimally sufficient or better, (score ≥ 7), the yellow line to responses which are average-quality or better (score ≥ 8), and the blue line to superior-quality responses (score = 9). Dots represent binned raw data: we partition task instances into 40 equally sized, log-spaced time bins and compute success rates and sample sizes within each bin. For each quality threshold, three of the 40 bins contain no observations.

this).

Figure 3 also shows that the slope becomes slightly flatter at higher levels of response quality. For instance, the slope coefficient is -0.21 at the superior quality level (score = 9). As expected, the curves also shift downward for higher quality levels, reflecting the greater difficulty of meeting stricter quality standards at any task length. The slope differences imply that success probabilities decline more sharply at shorter durations, with the gap between levels of quality narrowing as tasks become longer (yet, differences are still notable).

Finally, Figure 3 reports the average success rate across all models and task instances for each quality threshold. LLMs achieve minimally sufficient quality in 60% of task instances, average-worker quality in 49.7%, and superior quality in 27%.

2.2 Heterogeneity Analyses

Job Families. We estimate Eq. (1) separately by O*NET job family at the ≥ 7 threshold and report the results in Table 2. Three patterns stand out. First, success rates are substantial across all families, pointing to broad potential for LLMs to handle real-world labor market tasks. Second, average success varies meaningfully by domain, ranging from 49.8% in “Legal” to 69.0% in “Installation, Maintenance, and Repair” (recall that we restrict attention to LLM-addressable tasks and exclude

purely physical activities). Third, and most importantly, the success–duration slopes differ sharply across job families. The estimated slope coefficients span a wide range, implying that the relationship between LLM performance and task duration is not portable across labor market domains. In roughly half of job families, the slope is statistically significantly negative, with estimates between -0.14 and -0.39 —equivalent to a 3.5 to 9.7 percentage-point decline in predicted success for a tenfold increase in task duration, evaluated at a 60% baseline success rate. In the remaining job families, the slope is statistically indistinguishable from zero. The rightmost column in Table 2 visualizes the implied logistic curves for each domain; gray curves indicate statistically insignificant estimates.

Larger vs. smaller and newer vs. older models. A key question is how ongoing progress in LLMs affects these patterns. A strength of our data is that it spans many models, allowing us to separate two distinct channels of improvement: increases in model size versus newer model releases. In Figure 4, Panel (a) compares large ($> 100\text{B}$ parameters) and small ($\leq 100\text{B}$ parameters) models using the success threshold score ≥ 7 (see Appendix G.1 for a list of models considered "large" and "small"). The estimated curve for large models is more strongly downward sloped than for small models ($\beta = -0.32$ vs. -0.23), implying an outward rotation: large models' performance advantage is greatest for short tasks and attenuates as task duration increases.

Panel (b) instead splits frontier models by publication date (2023-2024 vs. 2025). Appendix G.2 lists the corresponding models. In contrast to Panel (a), the curves exhibit an almost purely parallel shift, with nearly identical slope coefficients ($\beta = -0.27$ and -0.28). This pattern indicates that newer models exhibit improved performance by roughly the same amount across the task-duration distribution, rather than disproportionately benefiting short tasks.

Panels (c)–(f) replicate Panels (a)–(b) for stricter quality thresholds (≥ 8 and $= 9$). We observe the same qualitative patterns at lower overall success rates: (i) comparing larger versus smaller models again yields an outward rotation, and (ii) comparing newer versus older models again produces an approximately parallel shift. Appendix Tables A.1, A.2, and A.3 confirm that these differences are statistically significant; the slight apparent rotation between newer and older models in Panel (f) is not. As also shown there, these results remain statistically significant (and quantitatively similar) even when jointly including indicators for model size and model vintage, thereby purging publication-time effects from the size comparison and, conversely, size effects from the vintage comparison.

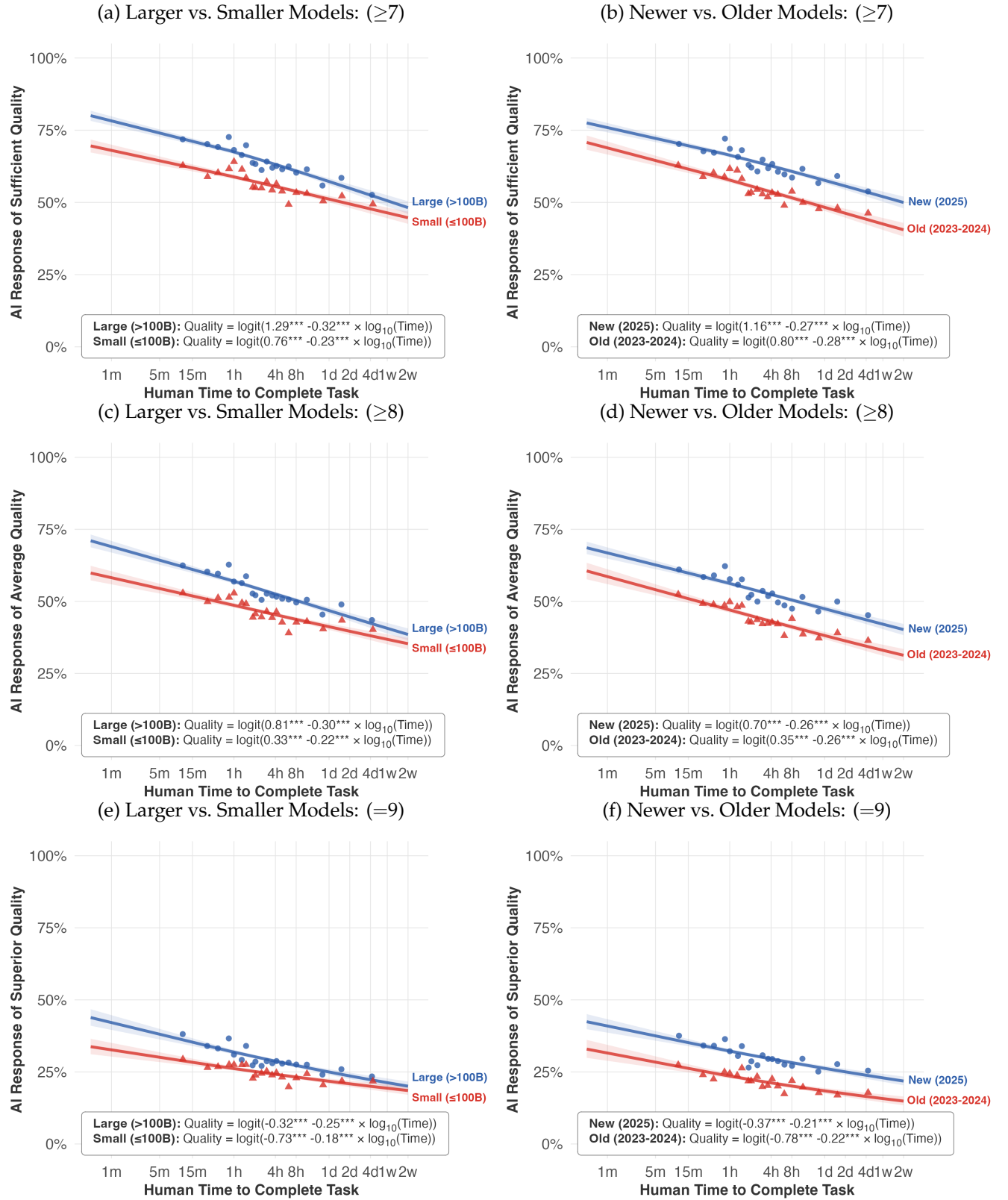
A natural interpretation is that improving longer-duration tasks is more demanding than improving short-duration tasks — and, in particular, that long-duration tasks, even if they are ultimately sequences of coupled short-duration ones (see Section 4.2), may require additional training / reinforcement learning over how to combine these subtasks effectively. In any case, the documented patterns suggest that improvements in AI performance associated with newer model vintages extend broadly across the task space, consistent with a rising-tide pattern. We further illustrate this pattern in Appendix Figure A.4, which shows that newer model vintages generate broadly similar improvements across job families, with only a few exceptions.

Table 2: Task Automatability And Task Duration: by Job Family

Job Family	N	Success Rate	Avg. Duration (hrs)	Coef. β	SE	Fit
Personal Care and Service	1,770	65.5%	2.9	-0.39***	0.11	
Architecture and Engineering	2,725	55.1%	21.3	-0.34***	0.08	
Computer and Mathematical	5,755	55.6%	20.2	-0.29***	0.06	
Management	6,510	56.0%	16.8	-0.28***	0.06	
Life, Physical, and Social Science	3,750	54.8%	16.3	-0.24***	0.07	
Community and Social Service	2,015	61.5%	11.2	-0.24*	0.11	
Sales and Related	2,965	62.6%	4.8	-0.18*	0.08	
Arts, Design, Entertainment, Sports, and Media	3,955	57.4%	14.1	-0.18**	0.06	
Business and Financial Operations	5,925	55.6%	13.7	-0.18**	0.06	
Protective Service	925	64.9%	5.8	-0.15	0.12	
Educational Instruction and Library	6,050	61.0%	7.5	-0.15*	0.06	
Office and Administrative Support	5,680	63.9%	3.2	-0.14*	0.07	
Transportation and Material Moving	1,110	66.0%	1.4	-0.10	0.14	
Healthcare Practitioners and Technical	5,235	64.9%	3.8	-0.06	0.07	
Construction and Extraction	665	68.4%	12.1	-0.04	0.18	
Production	1,425	67.9%	3.3	-0.03	0.14	
Food Preparation and Serving Related	1,275	64.8%	2.4	-0.01	0.12	
Healthcare Support	1,060	68.0%	1.3	0.05	0.18	
Installation, Maintenance, and Repair	920	69.0%	5.1	0.06	0.19	
Building and Grounds Cleaning and Maintenance	355	66.8%	6.6	0.07	0.25	
Legal	600	49.8%	10.1	0.07	0.20	
Farming, Fishing, and Forestry	175	68.0%	4.5	0.21	0.43	

Notes: The table reports job-family-specific logit slopes from regressing manager acceptance on $\log_{10}(\text{time to complete})$. "N" denotes the number of observations. The "success rate" is the share of task instances with a score ≥ 7 . "Coef. β " reports the slope estimate from Eq. (1) (estimated without controls), and "SE" reports standard errors clustered at the participant level. "Fit" provides a compact visualization of the estimated relationship, color-coded by slope magnitude, with statistically insignificant estimates shown in gray. Standard errors are clustered by participant in parentheses. Significance levels: *** 1%, ** 5%, * 10%.

Figure 4: Task Automation and Task Length: Model Size and Vintages



Notes: Panels (a)–(b) estimate Eq. (1) separately for larger vs. smaller models and newer vs. older models using “sufficient quality” (score ≥ 7) as the outcome. Panels (c)–(d) repeat this for “average quality” (score ≥ 8), and panels (e)–(f) for “superior quality” (score = 9). Coefficients are reported as log-odds. Model group definitions (large/small; new/old) are listed in Footnote ???. See Appendices G.1 and G.2 for the lists of underlying models. Standard errors are clustered at the participant level. Significance levels: *** 1%, ** 5%, * 10%.

2.3 Quantifying the speed of progress

Having established the shape of AI progress, we now turn to its pace. To that end, we estimate a modified version of Eq. (1) that includes a linear trend for model release dates (R_m), which models a shift in the logistic curve under a constant slope (Section 4.2 shows how this specification theoretically relates to Eq. (1)). Formally:

$$\Pr(Y_{ijsm} = 1) = \Lambda(\alpha + \delta R_m + \beta \log_{10} T_{js}), \quad (2)$$

An alternative specification could additionally allow the slope coefficient to vary over time. Appendix Table A.4 considers this possibility but finds no evidence of statistically significant changes in the slope, consistent with Figure 4. We estimate Eq. (2) using only *frontier* models (i.e., the best performing models from every period). See Appendix G.3 for the list of frontier models. Results are qualitatively robust to alternative frontier definitions — Appendix Figure A.7). We focus on the ≥ 7 threshold (useful as is, no edits required, minimally sufficient quality). Based on the estimating Eq. (2), the reported lines in Figure 5, Panel (a), display the evolution of success rates for different task durations. To validate the model fit across the data, we additionally report point estimates from quarter-specific regressions of Eq. (1) (squares), which provide a more flexible and demanding (i.e., less precise) specification than our baseline model in Eq. (2). Reassuringly, both approaches yield consistent patterns (and we stick to our baseline model for interpretation).

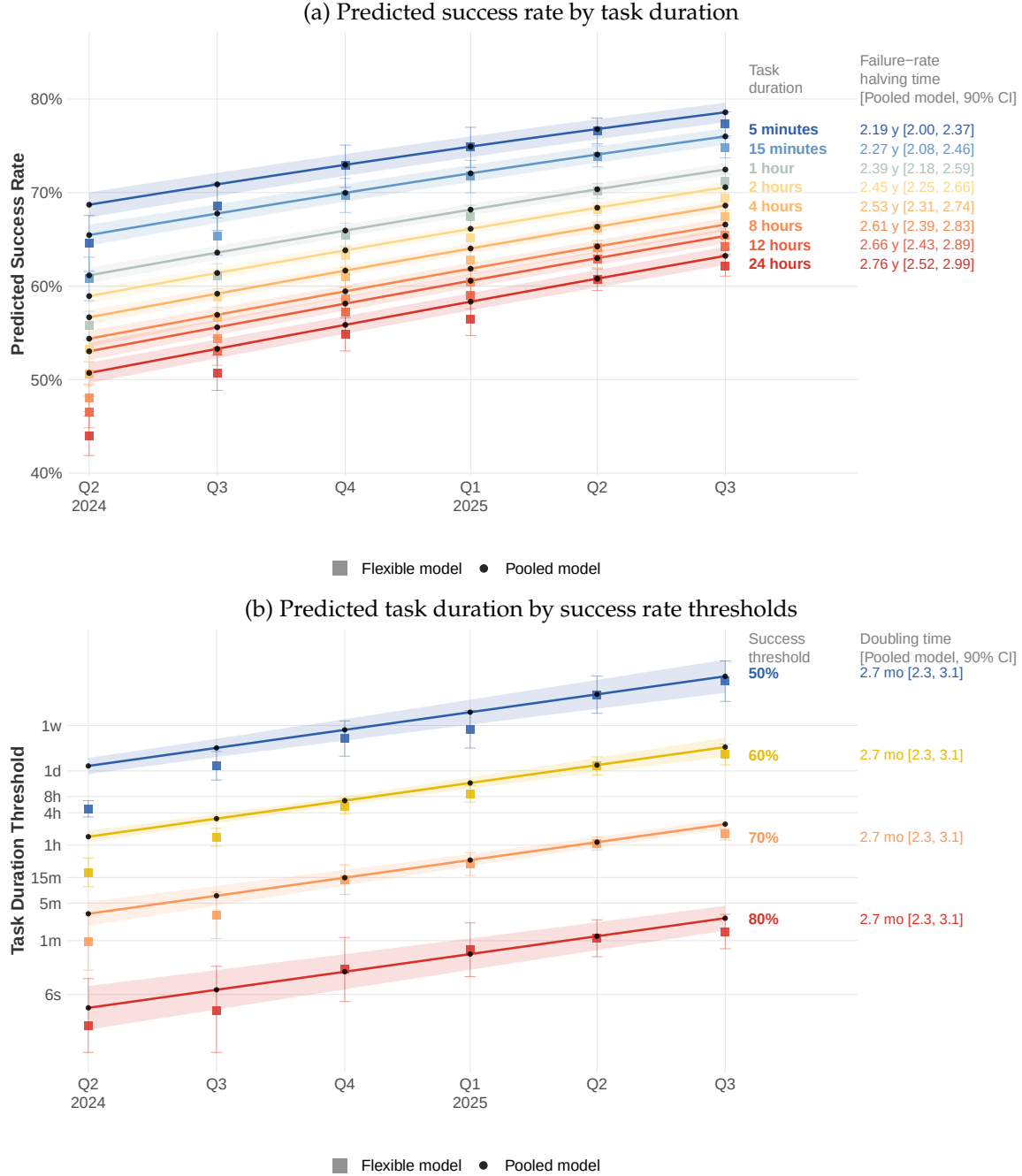
We find that success probabilities increase rapidly and at a similar pace across all task-durations, consistent with a broad-based capability improvement across the task-duration distribution—a fast rising-tide pattern. Based on these curves, we approximate failure-rate halving times (the failure rate is 1 minus the success rate), which equal 2.19-2.76 years over this period.

Figure 5, Panel (b), presents the mirror image of Panel (a): after estimating Eq. (2), we find the implied task duration achievable at a given probability of success. Across different success rate thresholds, we report linear trends in feasible task duration. The curves are all parallel shifts because of the logarithmic vertical axis and our linear trend specification. Using more flexible quarter-specific regressions of Eq. (1) (indicated by squares), we again find that the linear trend is a reasonable fit for the reported increase in task durations. According to our model predictions, already by 2024-Q2, frontier systems reached a 60% success rate (score ≥ 7) on tasks that take humans about 1.5 hours to complete.

At higher success rate thresholds, the implied feasible task duration is much shorter: for example, at an 80% success rate threshold, predicted task duration never exceeds about five minutes during our period of analysis. Nonetheless, improvements are rapid at every success rate threshold. We illustrate this pace by computing doubling times that we report alongside the figure (doubling times are directly inferred from the linear relationship between log-duration and time in Panel (b)). The implied feasible task duration roughly doubles every 2.7 months, indicating rapid improvements in AI capabilities to complete increasingly long-duration tasks (estimated with comparatively tight confidence bands).

Note, however, that the doubling time also depends on the slope of the logistic curve, as illustrated in Appendix Figure A.8. In the extreme, an almost flat curve shifted only slightly upward, and thus representing virtually no additional automation, would imply an extremely short doubling time. As

Figure 5: Task Duration and Success Rate Thresholds over Time (Score ≥ 7)



Notes: The lines in both panels are derived from estimating Eq. (2) on all task-level observations for frontier models across the full observation period (i.e., our "pooled" model). After the estimation, in Panel (a), we predict success rate changes based on given task length and a given linear (in logistic space) log-odds shifter (δR_m in Eq. (2)). Panel (b) instead predicts task duration for given success rates. The point estimates in both panels (i.e., "flexible model") are derived from estimation Eq. (1) separately for each quarter (which allows for quarter specific logistic slope coefficients) using the same approach of predicting success rates for given task durations (Panel (a)) and task durations for given success rates (Panel (b)). Shaded bands, error bars, and reported confidence intervals all indicate 90% delta-method CIs, SEs clustered by participant. Failure-rates halving times are locally approximated at the midpoint of the curves. See Appendix G.3 for the list of frontier models in the survey.

such, our preferred measure of the speed of AI progress is the change in success rates over time (Panel (a)), rather than the shift in task-duration at given success rates (Panel (b)). Nevertheless, both measures indicate rapid improvements in AI capabilities.

Appendix Figure A.6 replicates Figure 5 for LLM outputs of average worker quality (score ≥ 8) and superior quality (score = 9), with the superior quality plot featuring different probabilities (30-50%) due to its overall lower rate of success. Progress is even faster at these higher quality thresholds: failure-halving times range from 2.32 to 3.1 years for average-worker quality and from 3.16 to 4.51 years for superior quality. The corresponding task-duration doubling times are 2.4 months and 1.6 months, respectively.

2.4 Future Impacts

We now extrapolate our results into the future under the assumption that AI progress continues at the pace observed in recent years. Such projections necessarily involve substantial uncertainty. Although they represent our best estimates, the results should therefore be interpreted in light of the limitations discussed in the next section.

Using estimates from Eq. (2), we project how success rates evolve as newer models are released. Starting from a given success rate in 2024-Q2 (or, equivalently, a given initial feasible task duration level), we extrapolate the effect of subsequent model releases, captured by R_m in Eq. (2), using the estimated coefficient, δ .

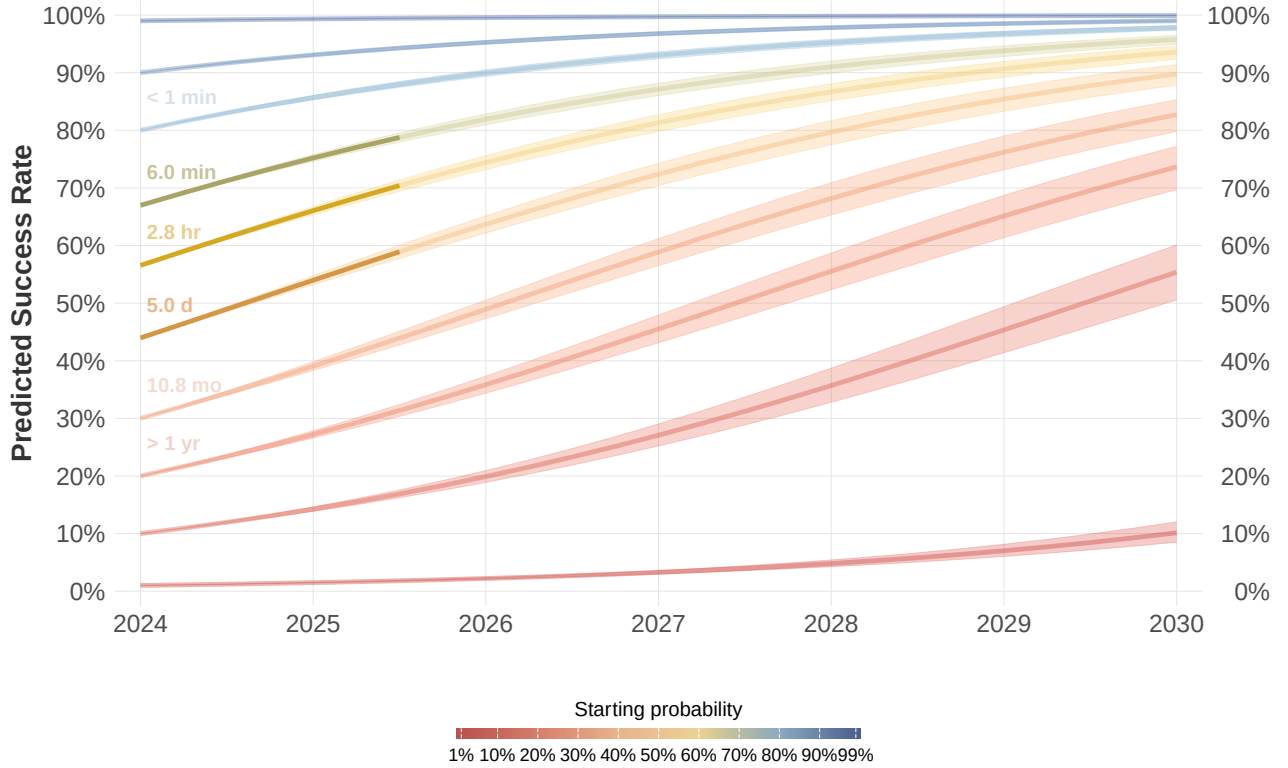
Figure 6 shows the resulting projections for the minimally sufficient quality level (score ≥ 7). Curves are shown in full saturation only where the implied task duration falls within the 1st to 99th percentile range of the observed distribution (about 6 minutes to 5 days). The faded portions of each curve indicate extrapolations beyond the observed period or into regions with limited data support. As shown by the saturated portions of the curves, the bulk of observed task instances lies within a relatively narrow success-rate range of 44% to 67%. Because the release-date term, δR_m , enters additively in the logit specification of Eq. (2), it implies a sigmoidal path in probability space. As a result, projected trajectories differ depending on the initial success rate.

Consistent with our earlier results, projected gains are gradual rather than abrupt ("rising tides"). Nevertheless, the pace of improvement remains substantial for reaching high success rates across most text-based labor market tasks; most tasks are projected to attain AI success rates of 88%–97% by 2030 at a minimally sufficient quality level. Appendix Figure A.9 reports corresponding results for the average worker and superior quality thresholds, showing projected success rates by 2030 are 85-96% and 78-91%, respectively.

Our results thus suggest considerable increases in AI capabilities over the next years. At the same time, due to the logistic relationship, our empirical model implies that achieving consistently near-perfect performance (i.e., success rates close to 100%) across most text-based tasks may still take several more years, especially for tasks where current performance remains low. For instance, extrapolating our estimates further into the future (not visualized) suggests that LLMs could achieve near-perfect success rates at average-worker quality by 2035. We emphasize that these longer-term projections are subject to substantially greater uncertainty and, as discussed below, should be viewed as more aggressive estimates, given compelling reasons to doubt that historical rates of compute growth—and

thus AI progress—can be sustained.

Figure 6: Predicted AI Success Rates Over Time



Notes: The figure reports predicted AI success rates over time based on estimates of Eq. (2) using all task-level observations for frontier models across the full sample period. In Panel (a), we project changes in success rates as a function of task duration and a linear log-odds shift in logistic space (captured by δR_m in Eq. (2)). Relative to Figure 5, Panel (a), we extend these predictions into future periods and into regions of the data with sparse or no observations (specifically, beyond the range of the task duration distribution in the survey). The saturated segments of the curves reflect predictions grounded in observed data (1st, 50th, and 99th percentile task durations (labeled) with 67.0%, 56.6%, and 44.0% success rates, respectively), whereas the faded segments indicate extrapolated regions. The shaded regions surrounding each line represent 90% confidence bands, which we compute using the standard errors of the release-date coefficient (δR_m) estimated in Eq. (2). See Appendix G.3 for the list of frontier models.

3 Discussion

AI capabilities in the labor market. What do our findings imply for labor markets? Although the pace of AI progress that we document is slower than that assumed in some prominent projections, such as Kokotajlo et al. (2025), our findings nevertheless point to the possibility of substantial labor market disruption in the near future. The rising-tide pattern of AI progress, that we establish in this paper, has two important policy implications:

First, AI systems are unlikely to transition abruptly from performing poorly to outperforming workers on a narrowly defined set of tasks within a single year. Instead, improvements appear to occur more gradually and continuously across the task space. This relative predictability implies that the present period constitutes a limited but potentially valuable window of several years during which workers, firms, and governments can adapt and implement policies in anticipation of further

increases in AI capabilities.

Second, the broad-based nature of this progress implies that AI capabilities are likely to improve across a wide range of text-based and partially text-based tasks. Within this segment of the task space, workers may therefore find it difficult to insulate themselves from automation pressure simply by moving toward (text-based or partially text-based) tasks that are currently more difficult for AI. Policy responses should consequently be designed for widespread improvements in AI capabilities across tasks and occupations, rather than for automation concentrated within a small number of narrowly defined activities.

Partial automation. The widespread and rapidly improving AI capabilities that we document do not imply that all occupations involving text-based tasks will disappear. Many jobs combine text-based activities with physical, interpersonal, or other tasks that remain less amenable to LLM automation. Rather than disappearing entirely, a job may therefore be reconfigured: AI performs the tasks for which it is particularly well suited, while workers continue to perform those tasks in which they retain a comparative advantage.

The resulting effects on employment and wages will depend not only on how many tasks within a job are automated, but also on which tasks are affected and how they relate to the broader occupational task bundle. As a result the automation of individual tasks need not necessarily harm workers; automation may raise or lower wages depending on the expertise embodied in the automated tasks and their role within the occupation (Autor and Thompson, 2025). Labor market outcomes will therefore reflect both task-level automation and the occupational adjustments that follow. The resulting labor market impacts could be quantified in future work building on our novel findings on the shape and speed of AI capability progress in the labor market.

Actual AI adoption vs. AI capabilities. We also emphasize that our results document AI capabilities rather than actual AI adoption. Our survey setup provides the information required for LLMs to perform each task. In practice, integrating such information may be difficult, costly, or subject to regulatory constraints, rendering some tasks infeasible to automate in real-world settings. Similarly, our analysis does not account for the economic attractiveness of deployment, which prior work identifies as a key determinant of automation feasibility (Svanberg et al., 2024). In particular, “last-mile” implementation costs (Fleming et al., 2024) may be substantial and could limit adoption, especially among smaller firms. Our results may inform future research on how the growth in AI capabilities documented in our paper translates into actual AI adoption.

Comparison to recent work by METR. Across specifications, we document a rather shallow negative relationship AI success and task duration. This contrasts with Kwa et al. (2025) and METR (2025), where the logistic curve gradients are steeper and, as a result, shorter/longer tasks are located closer to the logistic tails. Put differently, through the lens of Figure 1, this prior benchmark-based evidence aligns more closely with a crashing wave-like pattern, whereas our results support a rising-tide view across most job domains.

One potential explanation for this discrepancy could be the coverage of very short tasks. Our

data (i.e., the O*NET classification) do not include extremely short and simple tasks with near-certain success, such as producing a single programming command, whereas task durations in Kwa et al. (2025) range from seconds to several hours. What speaks against this explanation is that we continue to observe a near-linear success–duration relationship even at longer durations (rather than a logistic tail s-shape). Moreover, in an unreported robustness check, we re-estimated the results from Kwa et al. (2025) using their data while excluding the shortest tasks and consistently obtained steep slopes ($\beta = -1.08$ in the full sample versus $\beta = -0.99$ after excluding short tasks), suggesting that task coverage alone is unlikely to explain the discrepancy.

A further concern is that measurement error in task duration may attenuate our slope estimate. However, Appendix Figure A.11 shows that reconciling our estimate, $\beta = -0.27$, with the estimate of $\beta = -1.08$ reported by Kwa et al. (2025) would require an implausibly large median multiplicative error of $5.5\times$. Measurement error is therefore unlikely to explain our substantially smaller estimate, although we cannot rule out some attenuation and acknowledge this as a limitation.

Another potential explanation is that pooling across models (and domains) may attenuate the estimated slopes. In particular, larger models may perform disproportionately better on shorter tasks, as suggested by Figure 4. In Appendix Figure A.5, we therefore estimate the logistic relationship separately for each model in our data. Reassuringly, with few exceptions, the model-specific curves are not systematically steeper than the pooled estimates. Even models with relatively steeper curves exhibit substantially flatter relationships than those reported in Kwa et al. (2025). Consistent with this result, our domain-specific analyses (Table 2) and alternative regression specifications with additional controls (Figure A.2) likewise reveal no particularly steep relationships between AI success and task duration.

A potentially more important explanation is that the underlying task environments differ. Kwa et al. (2025) focus on relatively deterministic research and software-engineering tasks (e.g., identifying a subtle bug or implementing a known algorithm), whereas we study a broad set of real-world, text-based labor market tasks that often involve non-deterministic instances and a wider mix of domain knowledge and skills. (METR, 2025 analyze multiple benchmark datasets but do not report logistic shapes across benchmarks.)

Limitations. Our analysis focuses on a subset of tasks among those we classify as offering at least 10% potential time savings from LLM use. Accordingly, we do not claim rapid automation potential across all human work. Rather, our evidence points to a broad and fast-rising expansion of AI capabilities within text-centric real-world labor market tasks where LLMs plausibly deliver meaningful time savings.

Another limitation of our approach is that, despite careful task and survey design, we may fail to capture important dimensions of real-world work and its evaluation. In particular, we require each task instance to be self-contained, with all relevant information provided in the prompt. This constraint limits our ability to represent tasks that depend on interaction with external artifacts (e.g., opening and editing files, navigating software environments) or on repeated, multi-turn interactions with other people that cannot be condensed into a single vignette. We partially mitigate this concern by restricting our analysis to instances that evaluators deem realistic and representative of the underlying

task. Nonetheless, some dimensions of work are excluded by construction.

Similarly, we instructed models to respond in no more than 700 words—an instruction they generally followed—to limit the burden on expert evaluators. This constraint may have reduced performance on tasks that would ideally require more extensive answers (if these are also longer-duration tasks, this would work against finding a flat relationship between performance and task duration). Most responses fell well below the limit, suggesting that the cap was not particularly binding in the majority of cases (84% of responses were below 600 words). More generally, however, the wording and structure of the prompt may have influenced model responses in our survey (see Section 4.1 for more details).

Another concern is that respondents without relevant task experience may attempt to complete the survey. To address this, we implement extensive validation procedures for both respondents and their answers before inclusion in the analysis. For a subset of occupations, we can directly verify prior work experience for some participants; we find no meaningful differences in responses between these individuals and those who pass our implicit validation checks.

We also note that our survey design evaluates success based on the final output rather than the process used to produce it. While this aligns with how many tasks are assessed in practice, it may introduce measurement noise if outputs are imperfect proxies for true task performance. Such noise is unlikely to bias our estimates materially, but it may reduce explanatory power.

Finally, we emphasize important cautionary points regarding our predictions. The predictions in Figure 6 rely on an extrapolation, under the assumption of a stable logistic relationship based on task duration and continued improvement at the rate observed over recent years. Particularly in the tails of the logistic function (most relevant for our forward projections) alternative functional forms could yield materially different estimates of when AI reaches specific success thresholds. Reassuringly, however, using a complementary log-log specification as alternative functional form yields quantitatively similar results (Appendix Figure A.10). Nonetheless, there are substantial reasons to question whether AI progress can continue unchanged. For frontier models, compute has already scaled by several orders of magnitude (Mertens et al., 2026), and there are good reasons why continued investment in scaling may become prohibitively high (Thompson et al., 2020, Brown et al., 2026). In addition, algorithmic progress may be showing signs of slowing (Gundlach et al., 2025), while physical limits may impose increasingly binding constraints on further hardware advances (Shalf, 2020). These considerations suggest that future gains in AI capabilities may proceed more slowly than our extrapolations imply. Accordingly, our estimates are best interpreted as an upper bound under continued trend growth.

4 Methods

4.1 Survey Collection Details

Figure 7 provides an overview on the survey design and data collection that we detail in the following. The prompts we used during the survey design are detailed in Appendix F.

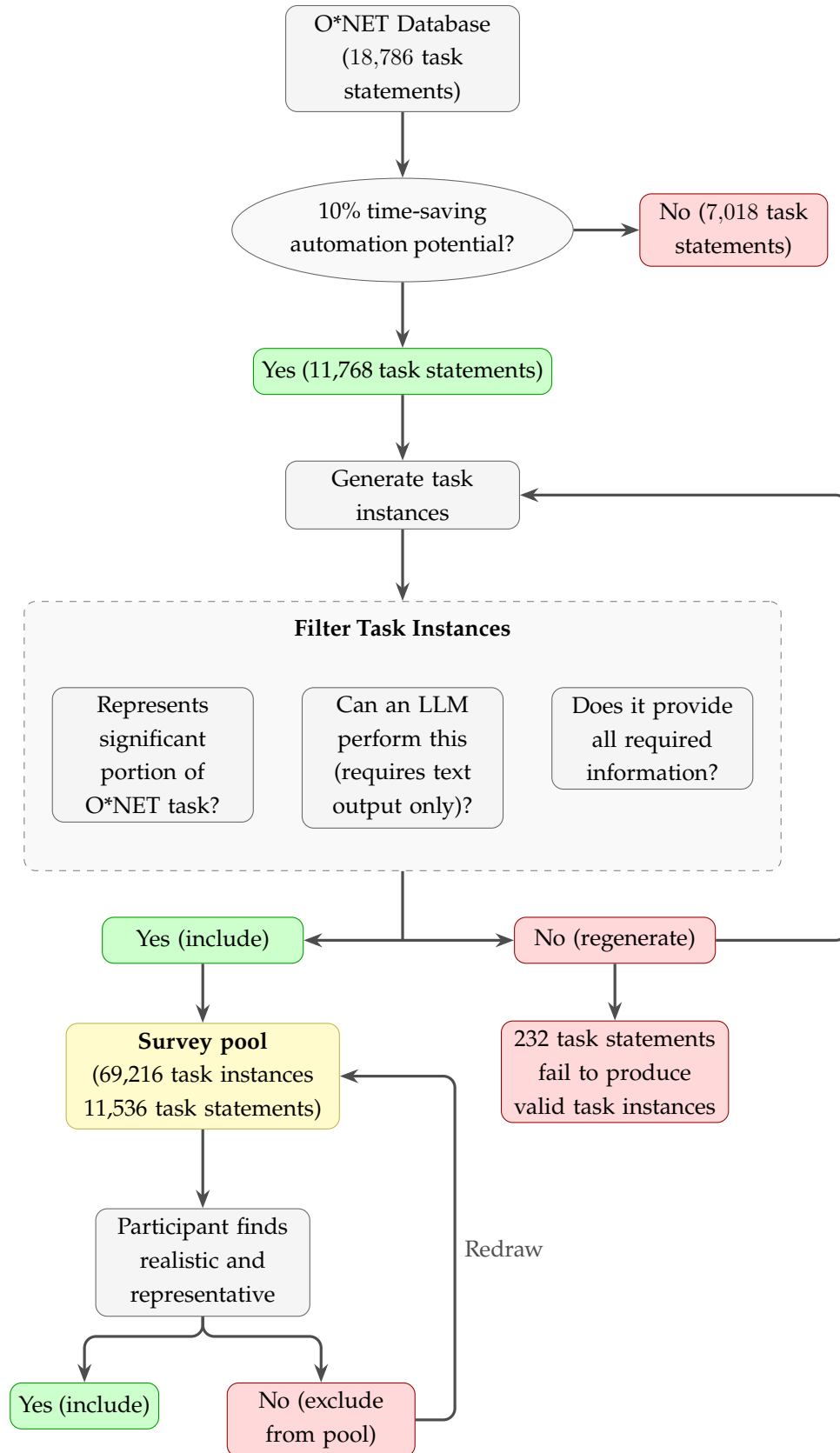
Task selection. We base our definition of labor market tasks on the O*NET 29.2 database (O*NET, 2024). Because not all tasks have meaningful LLM automation potential (e.g., predominantly physical tasks), we used GPT-4—the most advanced OpenAI model at the time—as a classifier to identify tasks where LLMs could lead to at least 10% time savings. This approach is similar to Eloundou et al. (2024). The main difference is that while Eloundou et al. (2024) use a 50% time savings threshold to classify automation exposure, our 10% threshold identifies a wider array of tasks with economically meaningful automation potential (in particular, partially text-based tasks). This procedure yielded 11,768 out of 18,786 tasks (62.6%) with at least 10% estimated automation potential.

Task instance generation. For each of the qualified O*NET tasks, we generated six different task "instances." Each instance was designed to represent a self-contained, realistic, and representative example of what the underlying O*NET task would look like as performed in a real work context. We generated these instances using GPT-5 and a structured prompt that incorporated the occupation title and the corresponding O*NET task description. We constrained each instance to a single coherent scenario, including technical details only when necessary for the role. Task difficulty was calibrated to reflect the expected performance of an experienced worker and aligned with the occupation's typical education and training requirements.

Each instance was limited to approximately 150 words to reduce evaluator cognitive load, incorporated at least one professional perspective (e.g., technical, procedural, interpersonal, or strategic), and followed a standardized format to ensure consistency.

Task instance filtering. We implemented a filtering process to ensure that generated task instances aligned with our research objectives. Specifically, we used GPT-5 to verify that each instance (i) represented a meaningful portion of the corresponding O*NET task, (ii) could be completed by an LLM (i.e., required text output only), and (iii) contained all necessary information to solve the task without external inputs. All task instances that met these criteria were included in the survey. Instances that did not meet these criteria were regenerated and re-evaluated, with up to 10 iterations per task. We excluded 232 O*NET tasks for which we could not generate a sufficient number of valid instances in this step. Ultimately, we included 11,536 tasks with 69,216 task instances in the survey pool. Appendix F provides details on the filtering prompts. As we describe further below, the results in this paper represent 6,648 of the 11,538 tasks (57.6%) and 12,078 task instances. Once a task instance was included in the survey, we conducted a final validation step: participants were asked whether the instance was realistic and representative of the underlying O*NET task (for a pilot subsample pertaining to 1.6% of the sample, we did not ask participants if the response was representative). If a participant judged an instance as unrealistic or unrepresentative, it was removed from the survey pool (including for future participants) and replaced with a new instance. This resampling procedure was feasible because, across all participants, only up to two task instances were evaluated per O*NET task. In practice, no case occurred in which all six task instances were deemed unrealistic or unrepresentative; had this happened, we would have generated additional instances.

Figure 7: Task Filtering Process



LLM response generation For each task instance included in the initial survey task population, we generated responses using 41 LLMs of varying scales and capabilities released between June 2023 and August 2025, including both open-weight and proprietary models (listed in Appendix H). All models were queried with an identical prompt (see Appendix section F.4) and default generation settings to ensure that differences in expert evaluations reflect variation in model capabilities rather than prompt optimization. Responses were capped at 700 words to balance substantive depth with the cognitive burden on the expert evaluators.

Survey data collection, sample used in this paper, and participants. We began collecting domain expert evaluations via the Prolific platform on September 22, 2025. The survey is conducted in multiple waves, with no restrictions on which tasks are collected in a given wave. Data are gathered as participants with relevant on-the-job experience complete the survey. The data used in this paper come from the first collection wave, which ran through July 2026 and covers 6,648 of the 11,538 tasks (57.6%), comprising 12,078 task instances. To qualify, participants had to reside in the United States, hold a platform approval rating of at least 75%, and have at least six months of work experience in the occupation under evaluation. 13.7% of participants who attempted to take the survey were excluded due to the work-experience restriction.

At the start of the survey, participants selected a task from a list corresponding to their occupation. They were then shown a task instance and asked to confirm that they understood it and that it was both realistic and representative of the underlying task. If any of these three conditions was not met, a new instance was provided. Once an instance passed validation, participants evaluated responses from five different LLMs. Unfortunately, our data includes subsamples where this order was randomized and some where it was not (we prefer randomization). Reassuringly, however, in Appendix B, we test for differences between these and show that non-randomization does not meaningfully change our results. Participants were not informed that the responses were generated by AI models. After reviewing all responses from the 5 different LLMs, participants were asked if they felt they had enough information and context about the task to accurately evaluate the responses (for a pilot subsample pertaining to 0.6% of the sample, this question was not included). We did not use data from any participant who did not agree.

For each of the 11,536 tasks, the target is to collect data on two task instances. Among the 6,648 tasks included in our analysis, we collected one task instance for 1,313 tasks and two task instances for 5,335 tasks (because the survey is collected in waves as described above). For each task instance, five model responses were assessed by the same participants. Participants reported contextual task information (e.g., time required, difficulty, and frequency in the job). For a pilot subsample pertaining to 0.6% of the data, the timing of the question regarding the time required to complete the task is different: in the pilot, respondents answered this question before viewing the LLM response, whereas in the main survey, they were asked after seeing the LLM output. Appendix C demonstrates that our main findings are robust to excluding the pilot data entirely. Participants evaluated each model response on a 1–9 scale, where 1 indicates that the response needed to be started over from scratch and 9 indicates that the response was of above-average quality for a human worker, as discussed previously.

Quality assurance. The survey data included in our analyses underwent rigorous quality assurance. Participants who failed more than one of four attention checks were excluded (11.6%). We further excluded participants exhibiting highly repetitive response patterns (i.e., limited variation in model evaluations), providing implausible estimates of task completion time, spending only minimal time on survey pages or showing no measurable engagement, or giving contradictory responses. In total, approximately 43.0% of the initially collected data were excluded and recollected to ensure high data quality. Importantly, the sample of 6,648 tasks and 12,078 task instances reported above refers to the final dataset after these quality control measures.

Summary statistics and task instance examples. Table 3 provides summary statistics on key variables used throughout the paper. We pool the data across all tasks and models and report means and distribution characteristics. Appendix Table D.1 provides examples of task instances for shorter and longer tasks (according to participant evaluations) in our data. In Appendix E, we list examples of LLM responses.

Table 3: Summary Statistics

	Mean	St.Dev.	P10	P50	P90
	(1)	(2)	(3)	(4)	(5)
Score	6.92	1.94	4	7	9
Manager Acceptance (0/1; 1 if score ≥ 7)	0.60	0.49	0	1	1
Time required to complete task instance (hours)	10.61	26.20	0.5	2.75	25
Observations: 60,845					

Notes: This table presents summary statistics for key survey variables. Columns (1)-(5) report means, standard deviations, and values for the 10th, 50th, and 90th percentile.

4.2 Theoretical Foundations

Our empirical object is the probability that an LLM response would be *accepted without edits* for a given real-world task instance, as a function of the task's (human-reported) log duration, T_{js} (for instance s of task j). We estimate this relationship using logistic regression. This section motivates the specification and shows that, under one plausible micro-foundation, the estimated slope coefficient can be interpreted as reflecting the length of the underlying serial chain of sequentially dependent steps required to complete the task.

Coupled critical-path robustness as a logit model. Suppose an instance, s , of a task, j , within domain, d , requires $N_{djs}(T_{js})$ *coupled critical actions* that must be completed without a fatal error for the output to be acceptable (relative to the main analysis, we add index d to discuss domain-specificity, aligning with our analysis by job families). We express $N_{djs}(T_{djs})$ as a function of observed task duration, T_{djs} , scaled by a parameter γ_d that governs how sequentially coupled tasks are:

$$N_{djs}(T_{djs}) = N_{0d} T_{djs}^{\gamma_d}, \quad N_{0d} > 0, \gamma_d > 0, \quad (3)$$

such that $\log_{10} N_{djs}(T_{djs}) = \log_{10} N_{0d} + \gamma_d \log_{10} T_{djs}$. A larger γ_d implies that longer tasks in domain d become disproportionately more sequentially coupled. N_{0d} captures the baseline serial complexity within a domain.

Let B_{djsm} denote the number of serial critical actions model m can sustain before the first fatal mistake. Success occurs iff $B_{djsm} \geq N_{djs}(T_{djs})$ (and success empirically corresponds to the evaluator rating "minimal sufficient" with no edits required). If the log-horizon is logistic,

$$\log_{10} B_{djsm} = \eta_m + \nu_d + \varepsilon_{djsm}, \quad \varepsilon_{djsm} \sim \text{Logistic}(0, \sigma), \quad (4)$$

where η_m captures a model-specific baseline robustness (how far a model can typically go in a serial chain) ν_d captures a domain-level baseline shift (how models generally perform within a domain), and ε_{djsm} denotes unobserved "failure" shocks. Given the structure of the error terms, B_{djsm} is log-logistic, the survival probability has a closed form (Bennett, 1983; Kiefer, 1988). In particular:

$$\begin{aligned} \Pr(Y_{djsm} = 1 \mid T_{js}, m, d) &= \Pr(\log_{10} B_{djsm} \geq \log_{10} N_{djs}(T_{djs})) \\ &= \Lambda\left(\frac{\eta_m + \nu_d - \log_{10} N_{0d}}{\sigma} - \frac{\gamma_d}{\sigma} \log_{10} T_{djs}\right), \end{aligned} \quad (5)$$

where, as above, $\Lambda(z) = 1/(1 + e^{-z})$. Eq. (5) corresponds to Eq. (1), where we define: $\alpha_{md} = \frac{\eta_m + \nu_d - \log_{10} N_{0d}}{\sigma}$ and $\beta_d = -\frac{\gamma_d}{\sigma}$, and provides a micro-founded rational for our estimation approach (and the estimation approach used in other work, such as Kwa et al., 2025). Through the lens of Eq. 5, differences in estimated logistic slope coefficients (as reported in Table 2), result from differences in the extent to which tasks within a domain (job family) are sequentially coupled. If longer tasks become more sequentially coupled, the slope coefficients become steeper (more negative). Shifts in the curve, on the other hand, are explained by model-specific capabilities (η_m) in solving coupled serial steps and domain-specific characteristics (ν_d).

Alternative functional form: complementary log-log. The same "first fatal error" framing that we use to motivate the logistic relationship yields a complementary log-log specification under a different assumption on how failures accumulate. If fatal errors arrive along serial exposure according to a Poisson process (constant hazard) or more generally a Weibull model (Weibull, 1951), then $p(T) = \Pr(Y = 1 \mid T) = \exp\{-C T^\kappa\}$ for constants $C, \kappa > 0$, implying $\log(-\log p(T)) = \log C + \kappa \log T$. This corresponds to a complementary log-log link in grouped-duration survival models (Jenkins, 1995; Prentice and Gloeckler, 1978). In Appendix Figure A.1 we re-estimate our main specification using this alternative link and obtain a qualitatively similar duration slope, suggesting a flat relationship between LLM performance and task duration. Similarly, in Appendix Figure A.10, we show that our predictions on future AI capabilities are similar under the alternative complementary log-log specification.

Release-date shifts in model robustness. A parsimonious extension of logistic regression framework is to allow the location of the model-robustness distribution to drift with model release date. For the subset of frontier models used in the time-series analysis, suppose:

$$\eta_m = \eta_0 + \rho R_m, \quad (6)$$

where R_m denotes model release date and ρ captures how the distribution of $\log_{10} B_{djsm}$ shifts over time. A one-unit increase in R_m raises $\log_{10} B_{djsm}$ by ρ , implying a multiplicative increase of 10^ρ in the typical sustainable serial horizon. Substituting Eq. (6) into Eq. (4) yields:

$$\Pr(Y_{djsm} = 1 \mid T_{djs}, R_m, d) = \Lambda\left(\frac{\eta_0 + \nu_d - \log_{10} N_{0d}}{\sigma} + \frac{\rho}{\sigma} R_m - \frac{\gamma_d}{\sigma} \log_{10} T_{djs}\right). \quad (7)$$

Thus, release date generates a pure intercept shift in log-odds space: $\delta = \rho/\sigma$ captures systematic improvements in model robustness over time. Suppressing domain heterogeneity yields the empirical specification in Eq. (2): $\Pr(Y_{jsm} = 1) = \Lambda(\alpha + \delta R_m + \beta \log_{10} T_{js})$. The maintained constant-slope restriction therefore corresponds to assuming that newer models extend the length of critical paths they can sustain, but do not change how required serial depth scales with task duration (for which we find empirical support, see Table A.4).

Implications of a linear release-date trend in the logit. For a fixed task duration T , Eq. 2 implies:

$$\frac{p(T, R_m)}{1 - p(T, R_m)} = \exp(\alpha + \beta \log_{10} T) e^{\delta R_m}, \quad (8)$$

so a one-unit increase in release date multiplies the odds of success by e^δ . Hence the release-date effect is exactly exponential in *odds space*, not in probability space. In probability space, the implied path is logistic:

$$p(T, R_m) = \Lambda(c(T) + \delta R_m) = \frac{p_0(T) e^{\delta R_m}}{1 - p_0(T) + p_0(T) e^{\delta R_m}}, \quad (9)$$

where $c(T) = \alpha + \beta \log_{10} T$ and $p_0(T) \equiv p(T, 0) = \Lambda(c(T))$. It follows that

$$\frac{\partial p(T, R_m)}{\partial R_m} = \delta p(T, R_m)(1 - p(T, R_m)),$$

so absolute percentage-point gains are largest for tasks with intermediate baseline success rates and smaller near 0 or 1. In this sense, an additive linear trend in the logit produces a sigmoidal path in probability space (as shown in Figure 6).

References

- Appel, R., McCrory, P., Tamkin, A., McCain, M., Neylon, T., & Stern, M. (2025). Anthropic economic index report: Uneven geographic and enterprise ai adoption. *arXiv preprint arXiv:2511.15080*.
- Autor, D., & Thompson, N. (2025). Expertise. *Journal of the European Economic Association*, 23, 1203–1271.
- Bennett, S. (1983). Log-logistic regression models for survival data. *Journal of the Royal Statistical Society: Series C (Applied Statistics)*, 32(2), 165–171.
- Brown, Z. A., Martín Poza, J., & Thompson, N. (2026). The economics of ever-larger language models. *Mimeo*.
- Dziri, N., Lu, X., Sclar, M., Li, X. L., Jiang, L., Lin, B. Y., Welleck, S., West, P., Bhagavatula, C., Le Bras, R., et al. (2023). Faith and fate: Limits of transformers on compositionality. *Advances in neural information processing systems*, 36, 70293–70332.
- Eloundou, T., Manning, S., Mishkin, P., & Rock, D. (2024). Gpts are gpts: Labor market impact potential of llms. *Science*, 384(6702), 1306–1308.
- Fleming, M., Thompson, N. C., & Li, W. (2024). The last mile problem in ai. *Washington: Brookings Institution*. <http://www.brookings.edu/articles/the-last-mile-problem-in-ai/>
- Ganguli, D., Hernandez, D., Lovitt, L., DasSarma, N., Henighan, T., Jones, A., Joseph, N., Kernion, J., Mann, B., Askell, A., et al. (2022). Predictability and surprise in large generative models. *Proceedings of the 2022 ACM Conference on Fairness, Accountability, and Transparency*. <https://arxiv.org/abs/2202.07785>
- Ge, H., Bastani, H., & Bastani, O. (2026). Are ai capabilities increasing exponentially? a competing hypothesis. *arXiv preprint arXiv:2602.04836*.
- Gundlach, H., Fogelson, A., Lynch, J., Trisovic, A., Rosenfeld, J., Sandhu, A., & Thompson, N. (2025). On the origin of algorithmic progress in ai. *arXiv preprint arXiv:2511.21622*.
- Hendrycks, D., Burns, C., Basart, S., Zou, A., Mazeika, M., Song, D., & Steinhardt, J. (2021). Measuring massive multitask language understanding. *International Conference on Learning Representations*. <https://arxiv.org/abs/2009.03300>
- Jenkins, S. P. (1995). Easy estimation methods for discrete-time duration models. *Oxford Bulletin of Economics & Statistics*, 57(1).
- Jimenez, C. E., Yang, J., Wettig, A., Yao, S., Pei, K., Press, O., & Narasimhan, K. (2024). SWE-bench: Can language models resolve real-world GitHub issues? *International Conference on Learning Representations*. <https://arxiv.org/abs/2310.06770>
- Kiefer, N. M. (1988). Economic duration data and hazard functions. *Journal of economic literature*, 26(2), 646–679.
- Kokotajlo, D., Alexander, S., Larsen, T., Lifland, E., & Dean, R. (2025). AI 2027. <https://ai-2027.com/ai-2027.pdf>
- Kwa, T., West, B., Becker, J., Deng, A., Garcia, K., Hasin, M., Jawhar, S., Kinniment, M., Rush, N., Von Arx, S., et al. (2025). Measuring ai ability to complete long tasks. *arXiv preprint arXiv:2503.14499*.
- LeCun, Y., Pavlick, E., Lake, B., Chalmers, D., & Browning, G., Jacob Lupyan. (2023, March). Debate: Do language models need sensory grounding for meaning and understanding? <https://youtu.be/x10964w00zk>
- Mertens, M., Fischl-Lanzoni, N., & Thompson, N. (2026). Is there "secret sauce" in large language model development? *arXiv preprint arXiv:2602.07238*.
- METR. (2025). How does time horizon vary across domains? <https://metr.org/blog/2025-07-14-how-does-time-horizon-vary-across-domains/>
- O*NET. (2024). O*NET 29.2 Database. <https://www.onetcenter.org/dictionary/29.2/excel/>
- Patwardhan, T., Dias, R., Proehl, E., Kim, G., Wang, M., Watkins, O., Fishman, S. P., Aljube, M., Thacker, P., Fauconnet, L., et al. (2025). Gdpval: Evaluating ai model performance on real-world economically valuable tasks. *arXiv preprint arXiv:2510.04374*.

- Prentice, R. L., & Gloeckler, L. A. (1978). Regression analysis of grouped survival data with application to breast cancer data. *Biometrics*, 34(1), 57–67.
- Rein, D., Becker, J., Deng, A., Nix, S., Canal, C., O’Connel, D., Arnott, P., Bloom, R., Broadley, T., Garcia, K., et al. (2025). Hcast: Human-calibrated autonomy software tasks. *arXiv preprint arXiv:2503.17354*.
- Schaeffer, R., Miranda, B., & Koyejo, S. (2023). Are emergent abilities of large language models a mirage? *Advances in Neural Information Processing Systems*, 36. <https://arxiv.org/abs/2304.15004>
- Shalf, J. (2020). The future of computing beyond moore’s law. *Philosophical Transactions of the Royal Society A: Mathematical, Physical and Engineering Sciences*, 378(2166).
- Srivastava, A., Rastogi, A., Rao, A., et al. (2023). Beyond the imitation game: Quantifying and extrapolating the capabilities of language models. *Transactions on Machine Learning Research*. <https://arxiv.org/abs/2206.04615>
- Svanberg, M., Li, W., Fleming, M., Goehring, B., & Thompson, N. (2024). Beyond ai exposure: Which tasks are cost-effective to automate with computer vision? *Available at SSRN 4700751*.
- Thompson, N. C., Greenewald, K., Lee, K., Manso, G. F., et al. (2020). The computational limits of deep learning. *arXiv preprint arXiv:2007.05558*, 10(2).
- Tomlinson, K., Jaffe, S., Wang, W., Counts, S., & Suri, S. (2025). Working with ai: Measuring the applicability of generative ai to occupations. *arXiv preprint arXiv:2507.07935*.
- Wang, Y., Ma, X., Zhang, G., Ni, Y., Chandra, A., Guo, S., Ren, W., Arulraj, A., He, X., Jiang, Z., et al. (2024). Mmlu-pro: A more robust and challenging multi-task language understanding benchmark. *Advances in Neural Information Processing Systems*, 37, 95266–95290.
- Wei, J., Tay, Y., Bommasani, R., Raffel, C., Zoph, B., Borgeaud, S., Yogatama, D., Bosma, M., Zhou, D., Metzler, D., et al. (2022). Emergent abilities of large language models. *Transactions on Machine Learning Research*. <https://arxiv.org/abs/2206.07682>
- Weibull, W. (1951). A statistical distribution function of wide applicability. *Journal of Applied Mechanics*, 18(3), 293–297.
- Wijk, H., Lin, T., Becker, J., Jawhar, S., Parikh, N., Broadley, T., Chan, L., Chen, M., Clymer, J., Dhyani, J., et al. (2025). RE-bench: Evaluating frontier AI R&D capabilities of language model agents against human experts. *Proceedings of the 42nd International Conference on Machine Learning*. <https://arxiv.org/abs/2411.15114>

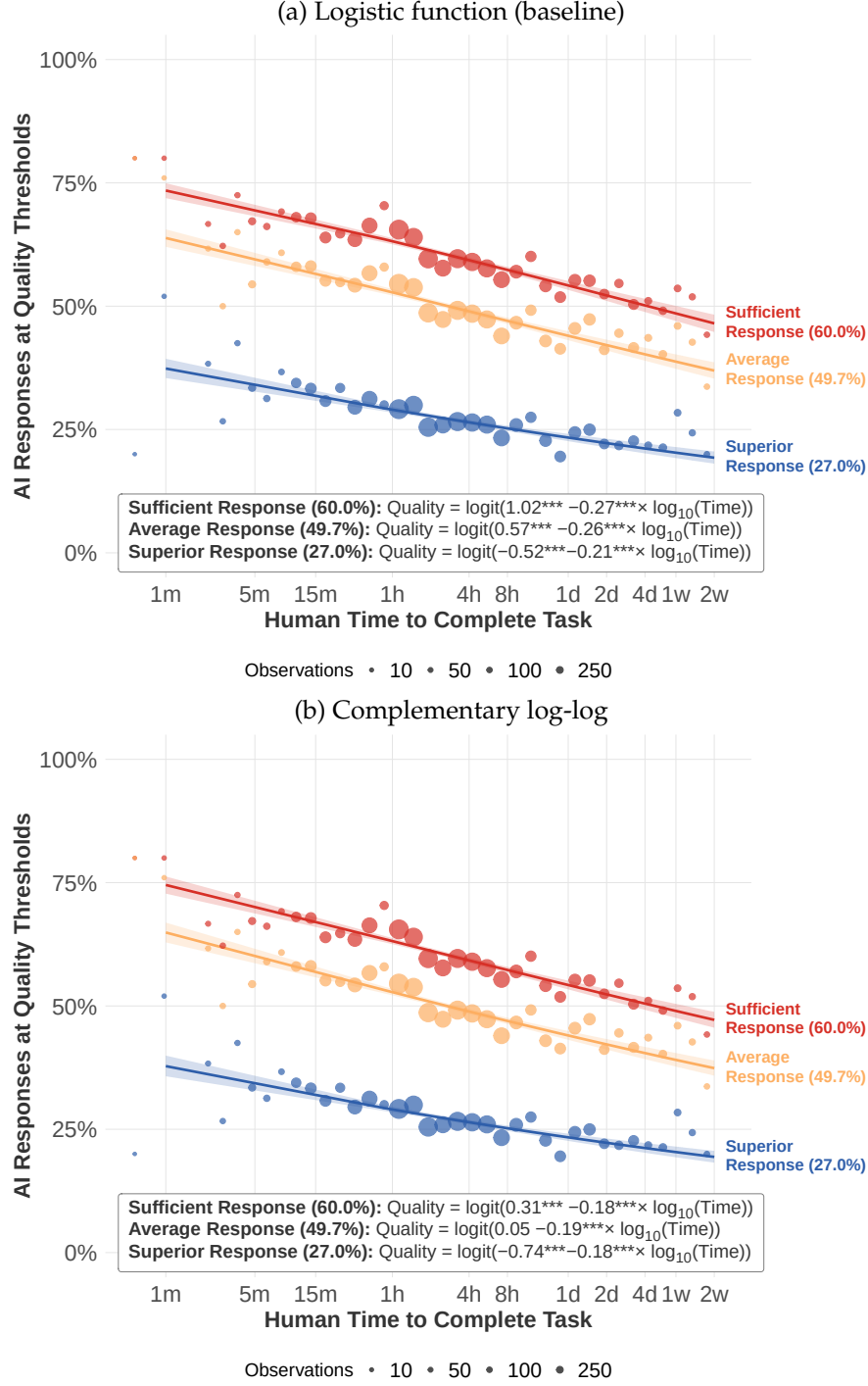
Online Appendix of:
Crashing Waves vs. Rising Tides: Findings on AI Automation from
Thousands of Worker Evaluations of Labor Market Tasks.

Matthias Mertens, Adam Kuzee, Brittany S. Harris, Harry Lyu, Wensu Li, Jonathan Rosenfeld,
Meiri Anto, Martin Fleming, Neil Thompson

A	Additional Results	28
B	Randomization Robustness Analysis	45
C	Measuring Task Length Pre vs Post Evaluation	47
D	Examples of Task Instances of Different Task Durations	48
E	LLM Responses	49
E.1	LLM Responses Example 1	49
E.2	LLM Responses Example 2	52
E.3	LLM Responses Example 3	58
F	LLM Prompts	62
F.1	Automation potential classifier prompt	62
F.2	Task instance generation prompt	65
F.3	Task instance filtering prompts	66
F.4	Task instance response generation prompt	70
G	Model grouping in empirical analysis	72
H	Model groups for sampling during data collection	73

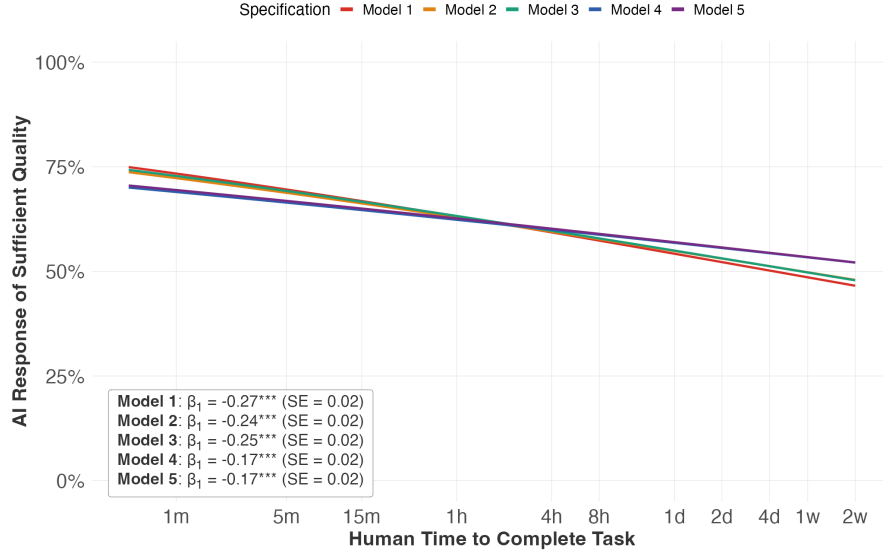
A Additional Results

Figure A.1: Task Duration and Success Rate Thresholds over Time, Logistic Function vs. Complementary Log-Log.



Notes: The figure compares the results shown in Figure 3 (Panel (a)) based on Eq.(1) with different threshold scores versus the same figure instead using a complementary log-log specification (Panel (b)). Standard errors clustered by participant are in parentheses. Significance levels: *** 1%, ** 5%, * 10%.

Figure A.2: Task Automation and Task Length: Different Specifications (Threshold Score ≥ 7).



Notes: The figure reports results from estimating Equation (1) for a threshold score ≥ 7 using different specifications. Model 1 represents our baseline results from Figure 3. Model 2 adds participant demographic controls (survey-reported peer expertise, education, employment status, and income, together with Prolific-reported sex and ethnicity (all as fixed effects) and age (entered linearly)) which have almost no impact on the slope (-0.24). Models 3–5 retain these demographic controls and add fixed effects. Model 3 adds LLM fixed effects to account for differences across models (e.g., training compute); the coefficient is essentially unchanged (-0.25). Model 4 instead adds occupation fixed effects, which absorb time-to-complete differences across occupations; the coefficient becomes smaller (-0.17) but remains highly statistically significant. Model 5 includes both LLM and occupation fixed effects jointly, yielding a coefficient (-0.17) very similar to Model 4. 80/12,078 participants were excluded from this analysis because they did not provide Prolific sex and ethnicity data. Standard errors are clustered by participant in parentheses.

Table A.1: Baseline Regressions with Size and Model Publication Time Effects (Acceptance ≥ 7)

	Dependent variable: $\mathbb{1}\{\text{acceptance} \geq 7\}$		
	(1)	(2)	(3)
Intercept	0.759*** (0.048)	0.801*** (0.054)	0.636*** (0.058)
$\log_{10}(\text{Time to Complete})$	-0.226*** (0.020)	-0.275*** (0.023)	-0.242*** (0.025)
Large (>100B)	0.535*** (0.057)		0.480*** (0.058)
New (2025+)		0.356*** (0.060)	0.247*** (0.061)
$\log_{10}(\text{Time to Complete}) \times \text{Large}$	-0.092*** (0.024)		-0.097*** (0.024)
$\log_{10}(\text{Time to Complete}) \times \text{New}$		0.006 (0.025)	0.028 (0.026)
Observations	60,845	60,845	60,845
Pseudo R^2	0.010	0.011	0.014

Notes: The table reports logit regressions of Eq. (1) of whether the response reached the threshold score of ≥ 7 . Standard errors clustered by participant are in parentheses. Large indicates models with parameter estimates >100B. Newer models refers to models released on/after 2025-01-01. Standard errors are clustered by participant in parentheses. Significance levels: *** 1%, ** 5%, * 10%.

Table A.2: Baseline Regressions with Size and Model Publication Time Effects (Acceptance ≥ 8)

	Dependent variable: $\mathbb{1}\{\text{acceptance} \geq 8\}$		
	(1)	(2)	(3)
Intercept	0.333*** (0.048)	0.349*** (0.054)	0.201*** (0.058)
$\log_{10}(\text{Time to Complete})$	-0.218*** (0.020)	-0.264*** (0.023)	-0.234*** (0.025)
Large (>100B)	0.475*** (0.055)		0.417*** (0.057)
New (2025+)		0.355*** (0.059)	0.260*** (0.061)
$\log_{10}(\text{Time to Complete}) \times \text{Large}$	-0.078*** (0.023)		-0.083*** (0.024)
$\log_{10}(\text{Time to Complete}) \times \text{New}$		0.008 (0.025)	0.027 (0.026)
Observations	60,845	60,845	60,845
Pseudo R^2	0.009	0.011	0.013

Notes: The table reports logit regressions of Eq. (1) of whether the response reached the threshold score of ≥ 8 . Standard errors clustered by participant are in parentheses. Large indicates models with parameter estimates >100B. Newer models refers to models released on/after 2025-01-01. Standard errors are clustered by participant in parentheses. Significance levels: *** 1%, ** 5%, * 10%.

Table A.3: Baseline Regressions with Size and Model Publication Time Effects (Acceptance = 9)

	Dependent variable: $\mathbb{1}\{\text{acceptance} = 9\}$		
	(1)	(2)	(3)
Intercept	-0.727*** (0.054)	-0.779*** (0.064)	-0.903*** (0.068)
$\log_{10}(\text{Time to Complete})$	-0.175*** (0.023)	-0.224*** (0.028)	-0.195*** (0.030)
Large (>100B)	0.406*** (0.063)		0.334*** (0.065)
New (2025+)		0.410*** (0.069)	0.334*** (0.071)
$\log_{10}(\text{Time to Complete}) \times \text{Large}$	-0.072*** (0.027)		-0.077*** (0.028)
$\log_{10}(\text{Time to Complete}) \times \text{New}$		0.014 (0.030)	0.031 (0.031)
Observations	60,845	60,845	60,845
Pseudo R^2	0.006	0.010	0.012

Notes: The table reports logit regressions of Eq. (1) of whether the response reached the threshold score of = 9. Standard errors clustered by participant are in parentheses. Large indicates models with parameter estimates >100B. Newer models refers to models released on/after 2025-01-01. Standard errors are clustered by participant in parentheses. Significance levels: *** 1%, ** 5%, * 10%.

Table A.4: Extended Specifications of the Baseline Model (logistic regressions).

	Dependent variable: $\mathbb{1}\{\text{acceptance} \geq 7\}$	
	Additive release date (1)	Additive and interacted release date (2)
Intercept	0.438*** (0.064)	0.563*** (0.156)
$\log_{10}(\text{Time to Complete})$	-0.308*** (0.021)	-0.363*** (0.066)
Release Date (years since Jan 2023)	0.411*** (0.020)	0.349*** (0.072)
$\log_{10}(\text{Time}) \times \text{Release Date}$		0.027 (0.031)
Observations	31,596	31,596
Pseudo R^2	0.017	0.017

Notes: The table reports logit regressions of whether the response reached the threshold score of ≥ 7 for frontier models. Column (1) estimates Eq. (2). Column (2) additionally adds an interaction between release dates and log task duration. We used only frontier models in the estimation. Release date measured in years since Jan 1, 2023. Standard errors clustered by participant in parentheses. Significance levels: *** 1%, ** 5%, * 10%.

Table A.5: OLS-Regression of Average Success Rates by Task-Duration Bins (Equal-Width Log-Spaced Bins) on Log Task Duration

	Dependent variable: Mean success rate			
	10 bins (1)	25 bins (2)	50 bins (3)	100 bins (4)
Intercept	0.6326*** (0.0090)	0.6271*** (0.0057)	0.6240*** (0.0060)	0.6183*** (0.0095)
$\log_{10}(\text{Mean Duration})$	-0.0639*** (0.0069)	-0.0584*** (0.0043)	-0.0550*** (0.0047)	-0.0505*** (0.0074)
Bins (N)	10	23	44	82
R^2	0.914	0.896	0.766	0.367
Adj. R^2	0.903	0.891	0.761	0.359

Notes: Each column creates k equal-width bins in $\log_{10}$ space over task duration, computes per-bin mean $\log_{10}(\text{duration})$ and mean success rate (≥ 7 acceptance), and runs OLS-regressions of bin average success rates on log task duration. Note that for the 50-bin and 100-bin columns, some bins did not have any observations. Standard errors clustered by participant in parentheses. Significance levels: *** 1%, ** 5%, * 10%.

Table A.6: Baseline Model by Number of Failed Attention Checks (logistic regressions).

	Dependent variable: $\mathbb{1}\{\text{acceptance} \geq 7\}$		
	All respondents (1)	Failed 0 checks (2)	Failed 1 check (3)
Intercept	1.017*** (0.040)	1.004*** (0.051)	1.039*** (0.063)
$\log_{10}(\text{Time to Complete})$	-0.269*** (0.017)	-0.267*** (0.022)	-0.272*** (0.026)
Observations	60,845	36,165	24,680
Pseudo R^2	0.005	0.006	0.005

Notes: Each column reports a bivariate logit of whether the response reached the threshold score of ≥ 7 on $\log_{10}$ task duration only, with no controls, on all valid model evaluations. Column (1) uses all respondents; columns (2) and (3) restrict to respondents who failed 0 and exactly 1 of the four attention checks, respectively. The valid analysis sample excludes anyone failing more than one check, so columns (2)–(3) partition column (1). Standard errors clustered by participant in parentheses. Significance levels: *** 1%, ** 5%, * 10%.

Table A.7: Baseline Model by Number of Failed Attention Checks (logistic regressions).

	Dependent variable: $\mathbb{1}\{\text{acceptance} \geq 7\}$		
	All respondents (1)	Failed 0 checks (2)	Failed 1 check (3)
Intercept	0.274*** (0.047)	0.200*** (0.061)	0.385*** (0.074)
$\log_{10}(\text{Time to Complete})$	-0.273*** (0.017)	-0.272*** (0.022)	-0.275*** (0.027)
Release Date (years since Jan 2023)	0.382*** (0.013)	0.414*** (0.017)	0.336*** (0.020)
Observations	60,845	36,165	24,675
Pseudo R^2	0.015	0.017	0.013

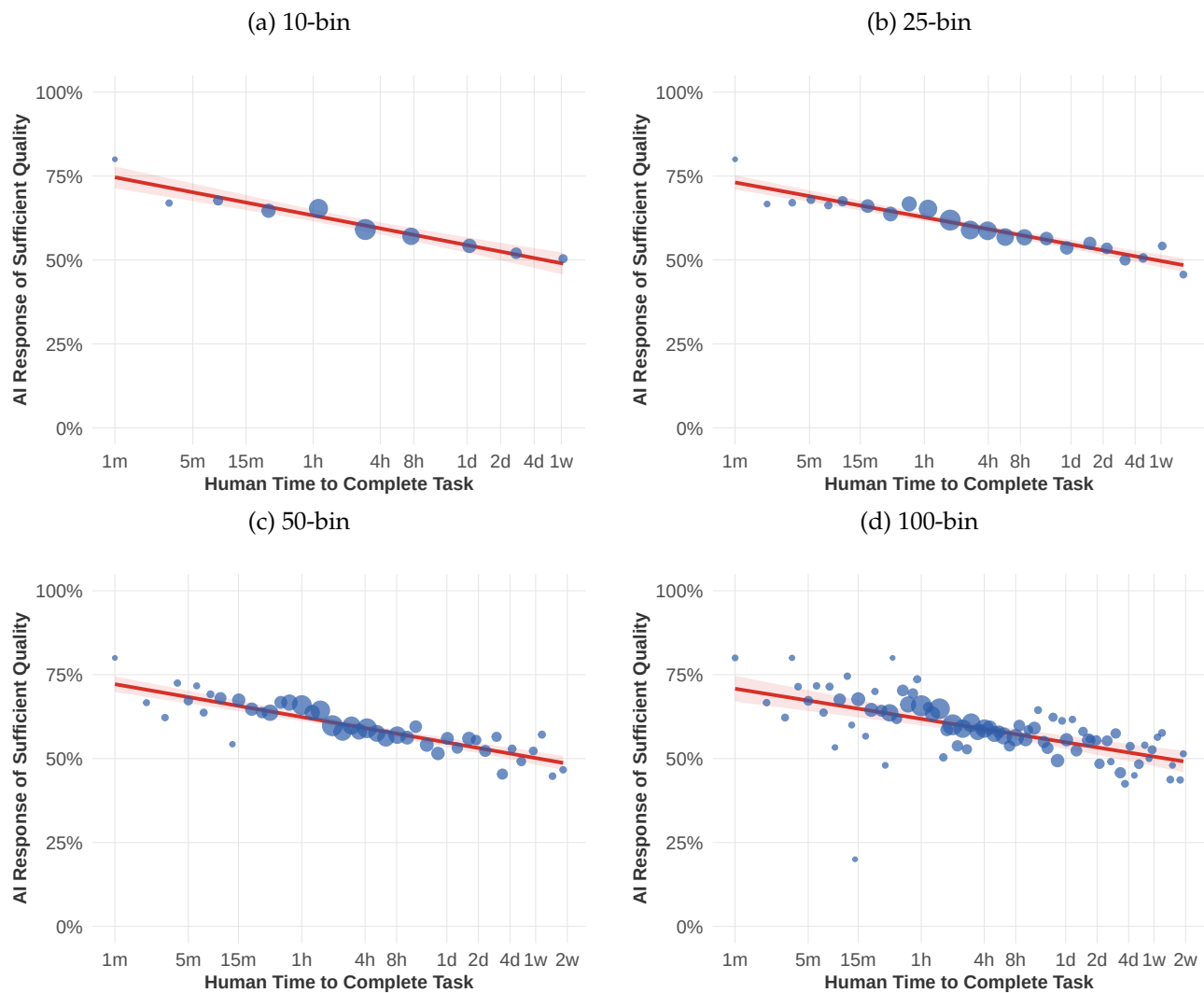
Notes: Each column re-estimates the additive baseline logit of whether the response reached the threshold score of ≥ 7 on $\log_{10}$ task duration and release date, with no demographic controls, on all valid model evaluations. Column (1) uses all respondents; columns (2) and (3) restrict to respondents who failed 0 and exactly 1 of the four attention checks, respectively. The valid analysis sample excludes anyone failing more than one check, so columns (2)–(3) partition column (1). Release date measured in years since Jan 1, 2023. Standard errors clustered by participant in parentheses. Significance levels: *** 1%, ** 5%, * 10%.

Table A.8: Baseline Model with Demographic Controls, by Number of Failed Attention Checks (logistic regressions).

	Dependent variable: $\mathbb{1}\{\text{acceptance} \geq 7\}$		
	All respondents (1)	Failed 0 checks (2)	Failed 1 check (3)
$\log_{10}(\text{Time to Complete})$	-0.242*** (0.017)	-0.239*** (0.022)	-0.244*** (0.027)
Age (years)	0.006*** (0.001)	0.006*** (0.001)	0.005*** (0.002)
Peer-expertise FE	Yes	Yes	Yes
Education FE	Yes	Yes	Yes
Employment FE	Yes	Yes	Yes
Income FE	Yes	Yes	Yes
Sex FE	Yes	Yes	Yes
Ethnicity FE	Yes	Yes	Yes
Observations	60,445	35,940	24,505
Pseudo R^2	0.010	0.011	0.012

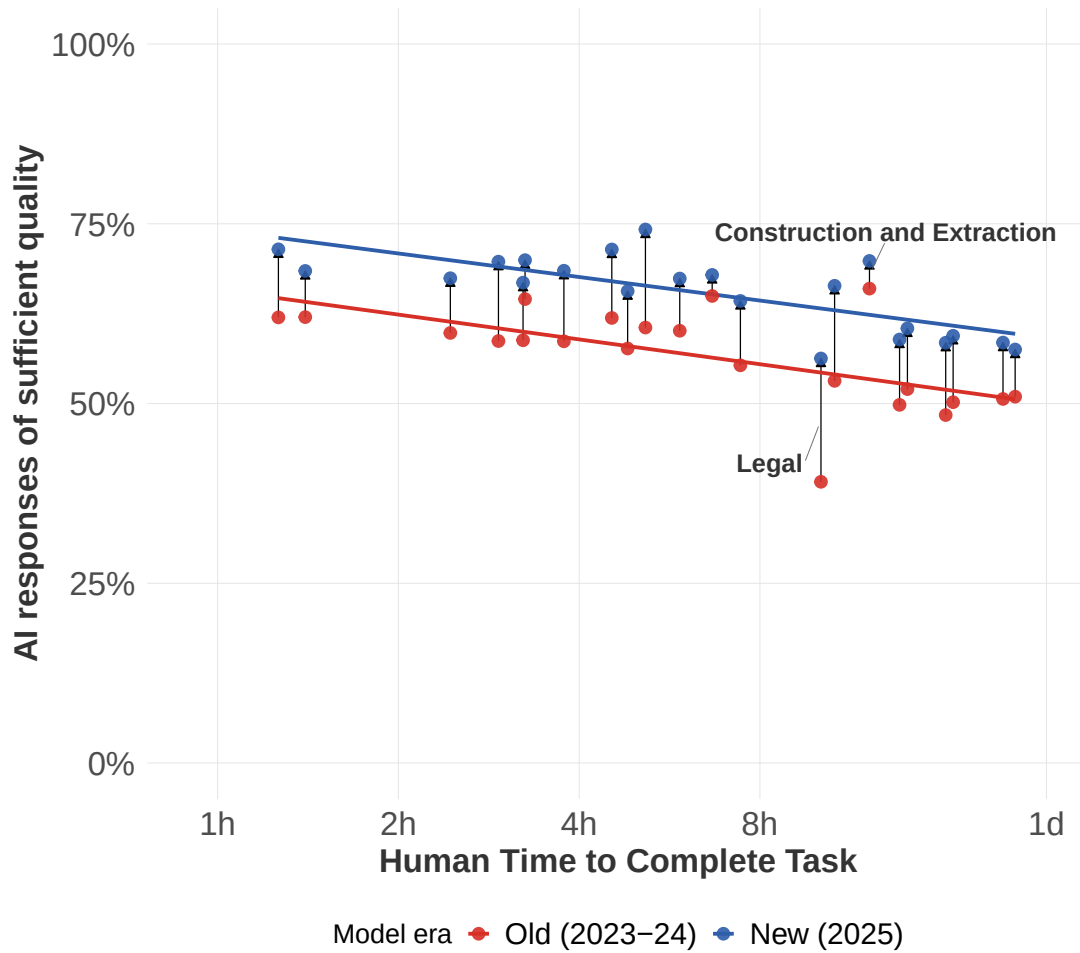
Notes: Each column reports a logit of whether the response reached the threshold score of ≥ 7 on $\log_{10}$ task duration and participant age (entered linearly), with fixed effects for survey-reported peer expertise, education, employment status, and income, and Prolific-reported sex and ethnicity. Estimated on all valid model evaluations with complete demographics. Column (1) uses all respondents; columns (2) and (3) restrict to respondents who failed 0 and exactly 1 of the four attention checks, respectively. The valid analysis sample excludes anyone failing more than one check, so columns (2)–(3) partition column (1). Standard errors clustered by participant in parentheses. Significance levels: *** 1%, ** 5%, * 10%.

Figure A.3: OLS-Regression of Average Success Rates by Task-Duration Bins (Equal-Width Log-Spaced Bins) on Log Task Duration



Notes: Each figure corresponds with a column in Table A.5, in which each figure creates k equal-width bins in $\log_{10}$ space over task duration, computes per-bin mean $\log_{10}(\text{duration})$ and mean success rate (≥ 7 acceptance), and fits OLS-regressions of mean success rates on log task duration. Standard errors clustered by participant. Note that for the 50-bin and 100-bin columns, some bins did not have any observations.

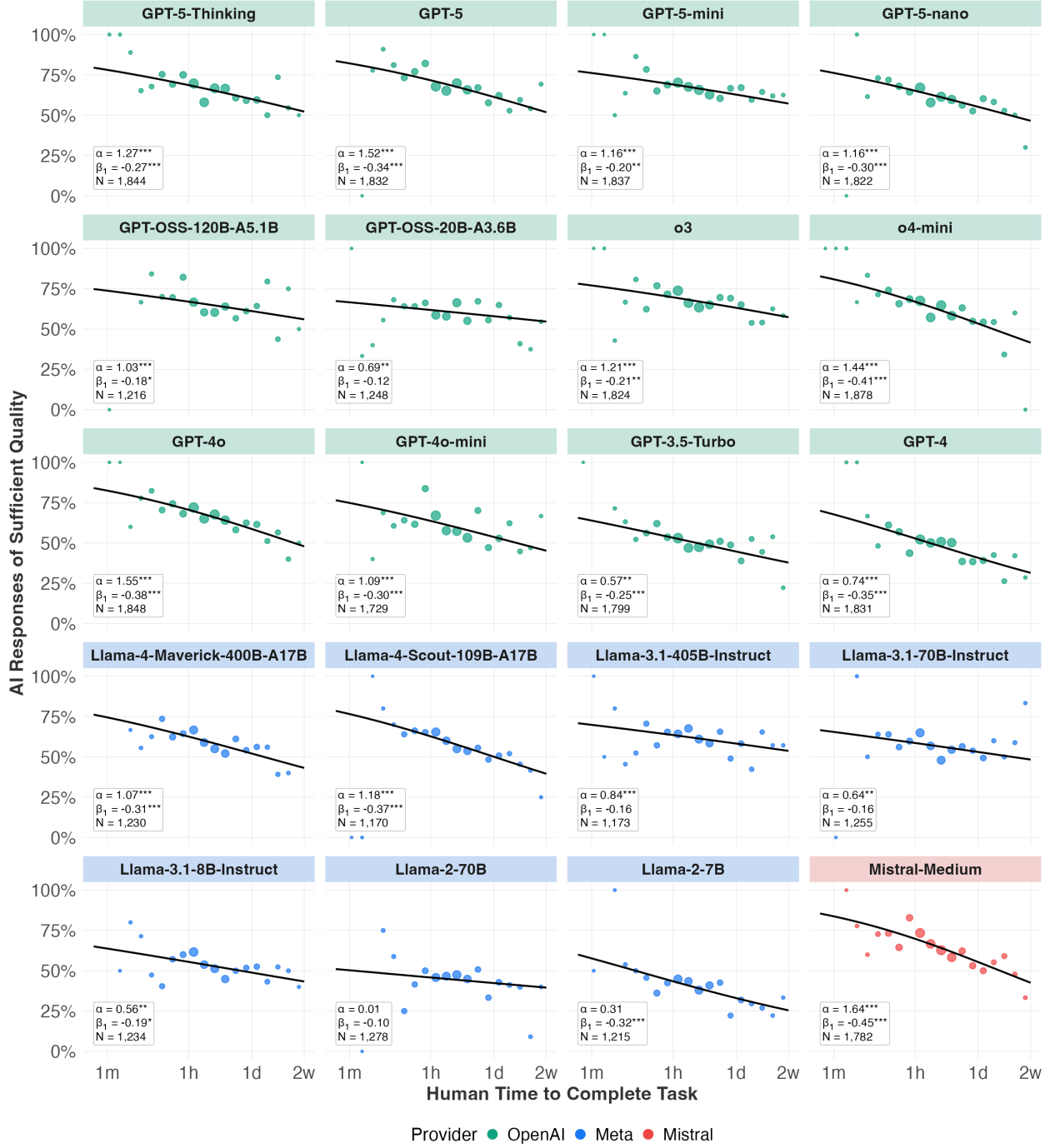
Figure A.4: Old vs New models by Job Family



Notes: The Figure shows an alternate visualization of 4 Panel (a), in which each job family is represented by an individual dot. This figure restricts task coverage to just those tasks which were performed by both a new and old model, to ensure a direct comparison with the same task durations. Most Job Families rise by a similar amount, though even within similar task durations the change can vary, possibly due to different underlying task structures.

Figure A.5: Individual Frontier LLM Success-Duration Relationships

(a) Model set 1/2



Notes: The Figure (continued on the next page) displays the success-duration relationship found by estimating Eq.(1) at the ≥ 7 quality level for each individual model in the survey. Models are grouped by provider and sorted by release date within each provider. Dots represent binned raw data: we partition task instances into 20 equally sized, log-spaced time bins and compute success rates and sample sizes within each bin. Many models are missing data for at least one bin. Standard errors clustered by participant are in parentheses. Significance levels: *** 1%, ** 5%, * 10%.

(b) Model set 2/2

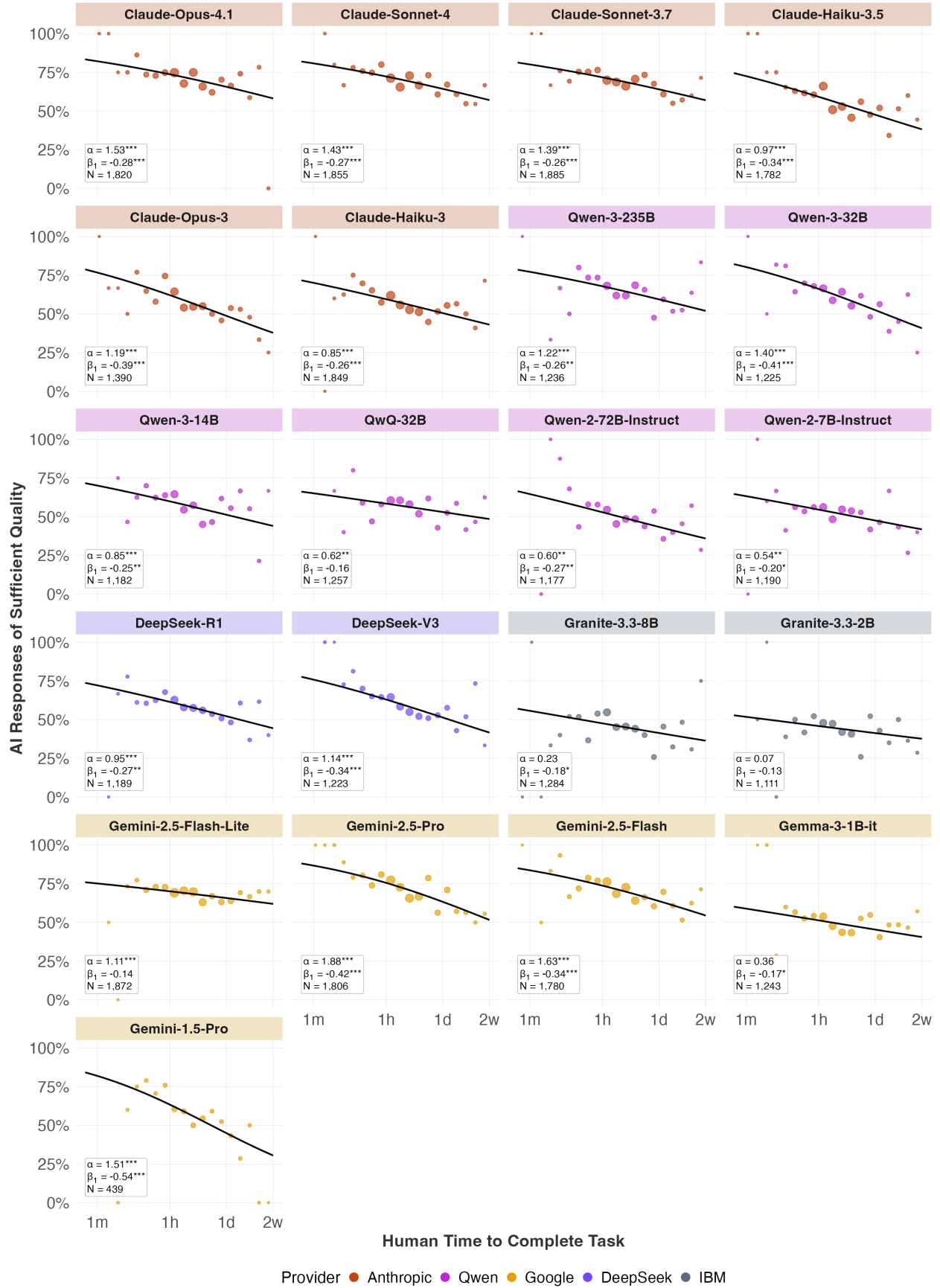
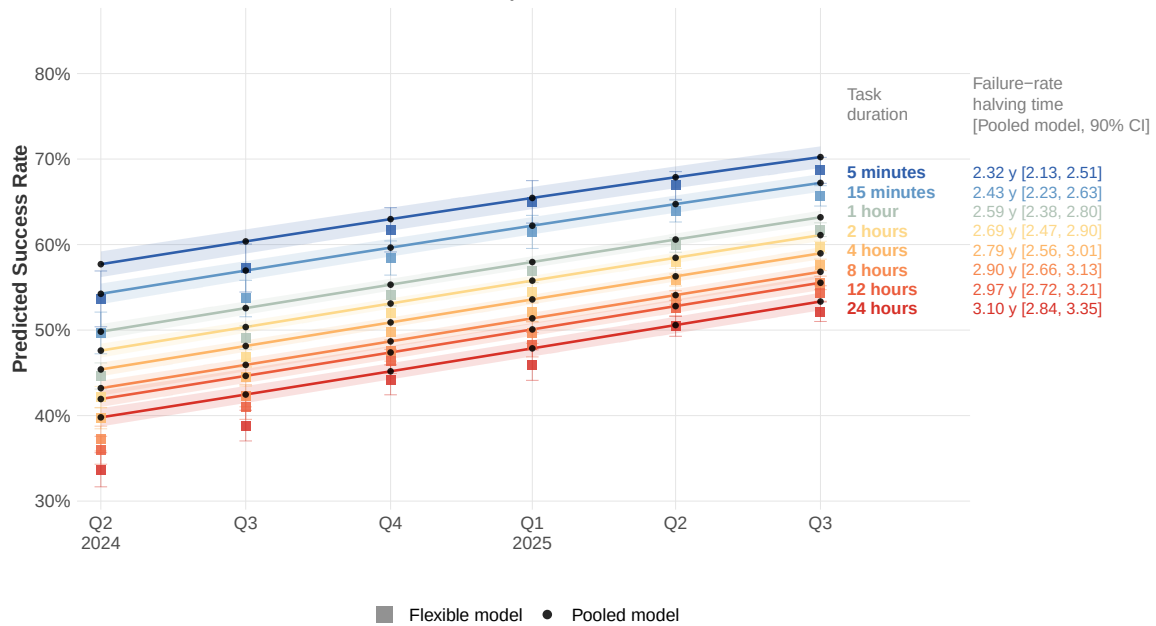
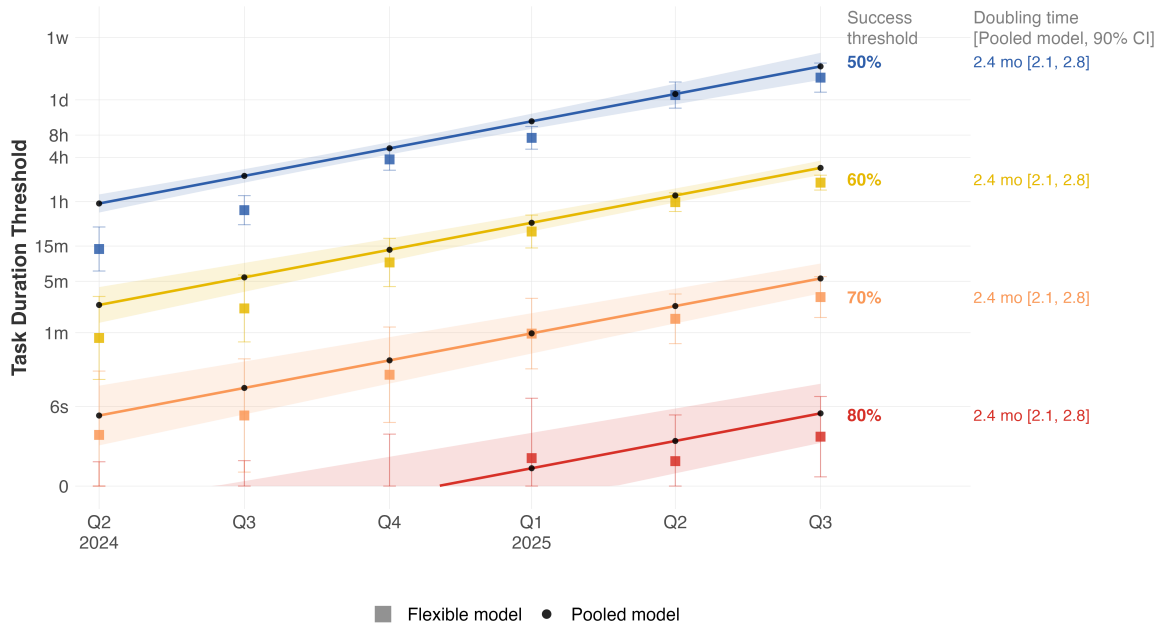


Figure A.6: Task Duration and Success Rate Thresholds over Time (threshold score ≥ 8)

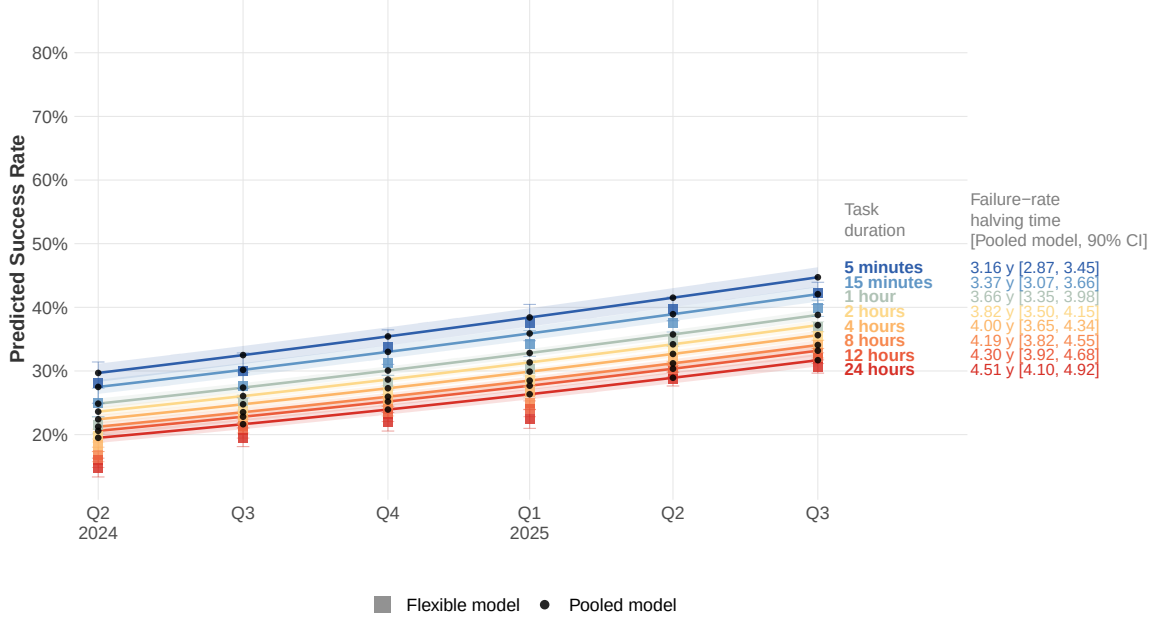
(a) Predicted success rate by task duration bin (threshold score ≥ 8)



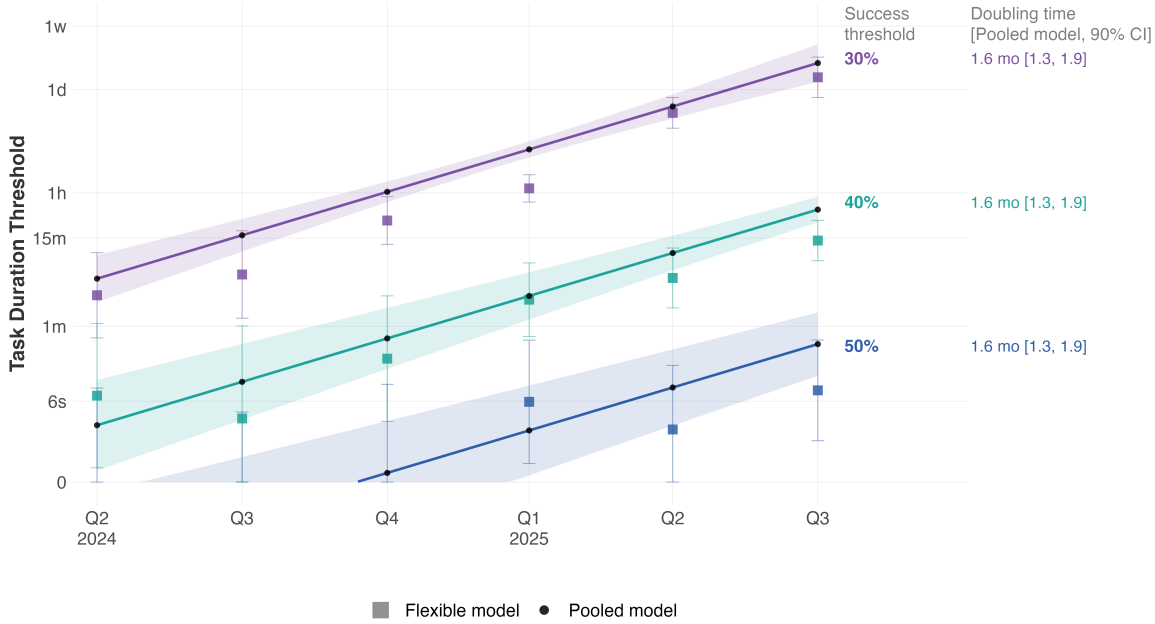
(b) Predicted task duration by success rate thresholds (threshold score ≥ 8)



(c) Predicted success rate by task duration bin (threshold score ≥ 9)

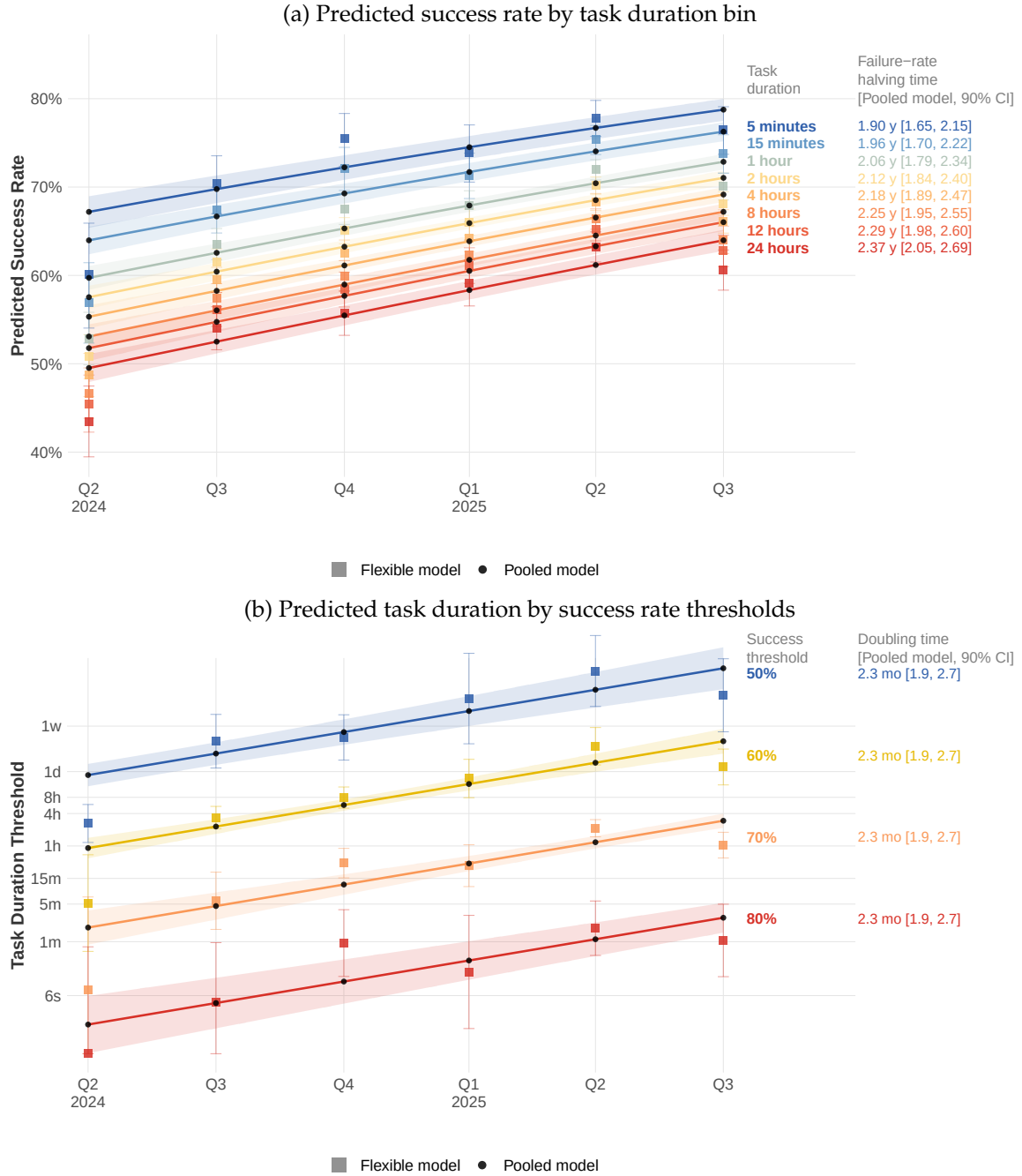


(d) Predicted task duration by success rate thresholds (threshold score ≥ 9)



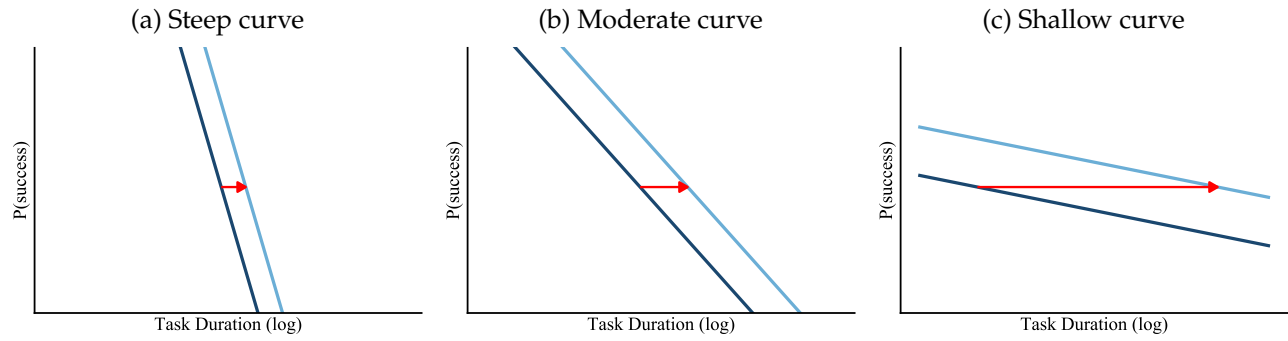
Notes: the Figure replicates Figure 5 using alternative success threshold scores ≥ 8 (Panels (a) and (b)) and ≥ 9 (Panels (c) and (d)). The lines in all panels are derived from estimating Eq. (2) on all task-level observations for frontier models across the full observation period (i.e., our "baseline" model). After the estimation, in Panel (a) and (c), we predict success rate changes based on given task length and a given linear (in logistic space) log-odds shifter (δR_m in Eq. (2)). Panels (b) and (d) instead predict task duration for given success rates. Panel (d) uses different success rates because they are much lower at the ≥ 9 level. The 80% curve in Panel (b) and 50% curve in Panel (d) are not above 0 at all points in time because our frontier models were not predicted to reach 80% at the ≥ 8 threshold (in the case of Panel (a)) or 50% at the ≥ 9 threshold (in the case of Panel (c)) until after the starting period. We include the portion of the confidence bands which is above 0 because there remains the probability they reached these thresholds earlier than our main estimate suggests. The point estimates in both panels (i.e., "flexible model") are derived from estimation Eq. (1) separately for each quarter (which allows for quarter specific logistic slope coefficients) using the same approach of predicting success rates for given task durations (Panel (a)) and task durations for given success rates (Panel (b)). Shaded bands, error bars, and reported confidence intervals all indicate 90% delta-method CIs, SEs clustered by participant. Failure-rates halving times are locally approximated at the midpoint of the curves. See Appendix G.3 for the list of frontier models.

Figure A.7: Task Duration and Success Rate Thresholds over Time (alternative frontier definition)



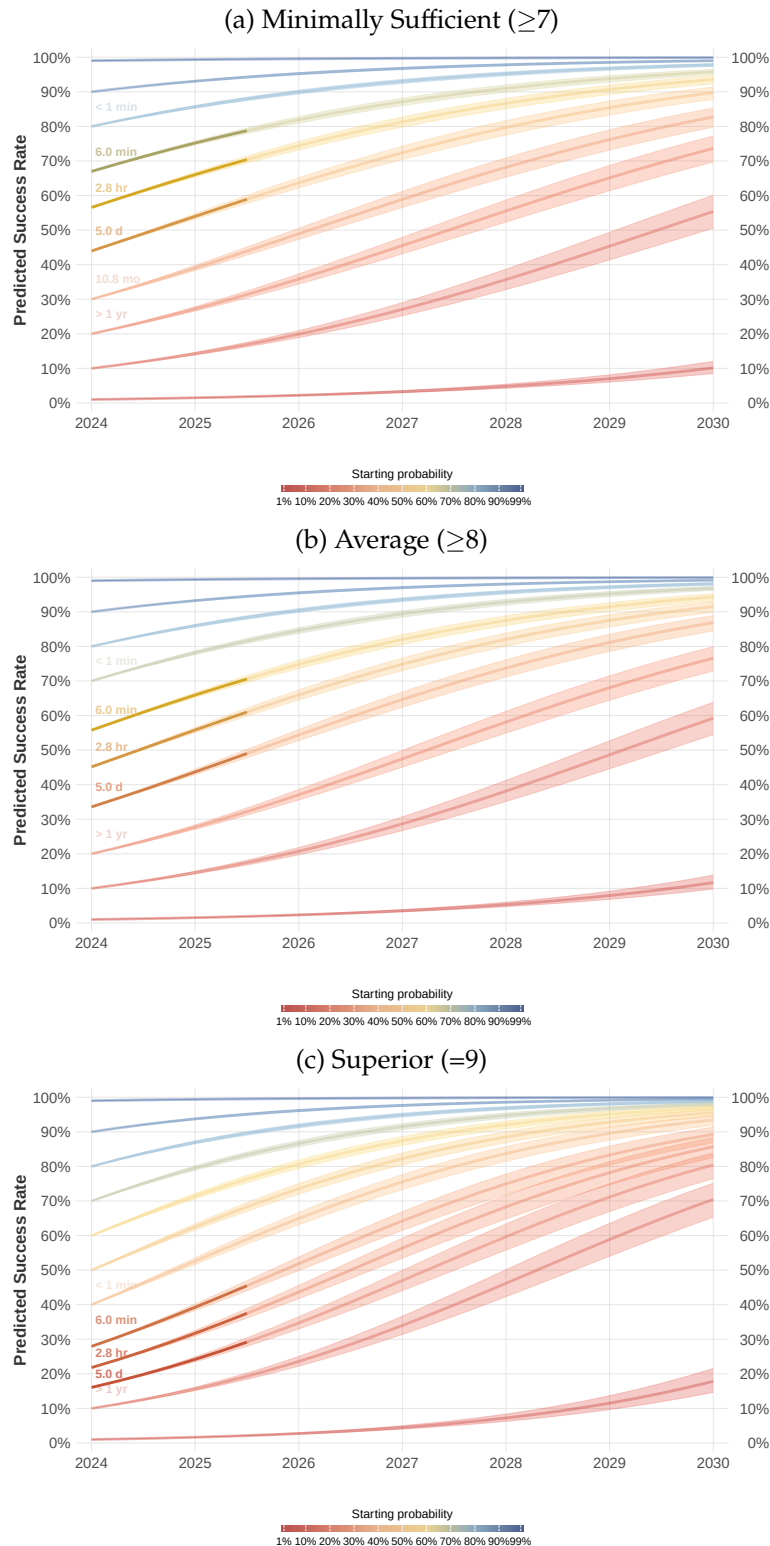
Notes: The Figure replicates Figure 5 using an alternative (more stringent) definition of frontier models. The lines in both panels are derived from estimating Eq. (2) on all task-level observations for frontier models across the full observation period (i.e., our "baseline" model). After the estimation, in Panel (a), we predict success rate changes based on given task length and a given linear (in logistic space) log-odds shifter (δR_m in Eq. (2)). Panel (b) instead predicts task duration for given success rates. The point estimates in both panels (i.e., "flexible model") are derived from estimation Eq. (1) separately for each quarter (which allows for quarter specific logistic slope coefficients) using the same approach of predicting success rates for given task durations (Panel (a)) and task durations for given success rates (Panel (b)). Shaded bands, error bars, and reported confidence intervals all indicate 90% delta-method CIs, SEs clustered by participant. Failure-rates halving times are locally approximated at the midpoint of the curves. See Appendix G.3 for the list of frontier models included in the alternative/more stringent definition.

Figure A.8: Time Duration Increase Based on Success-Duration curve slopes



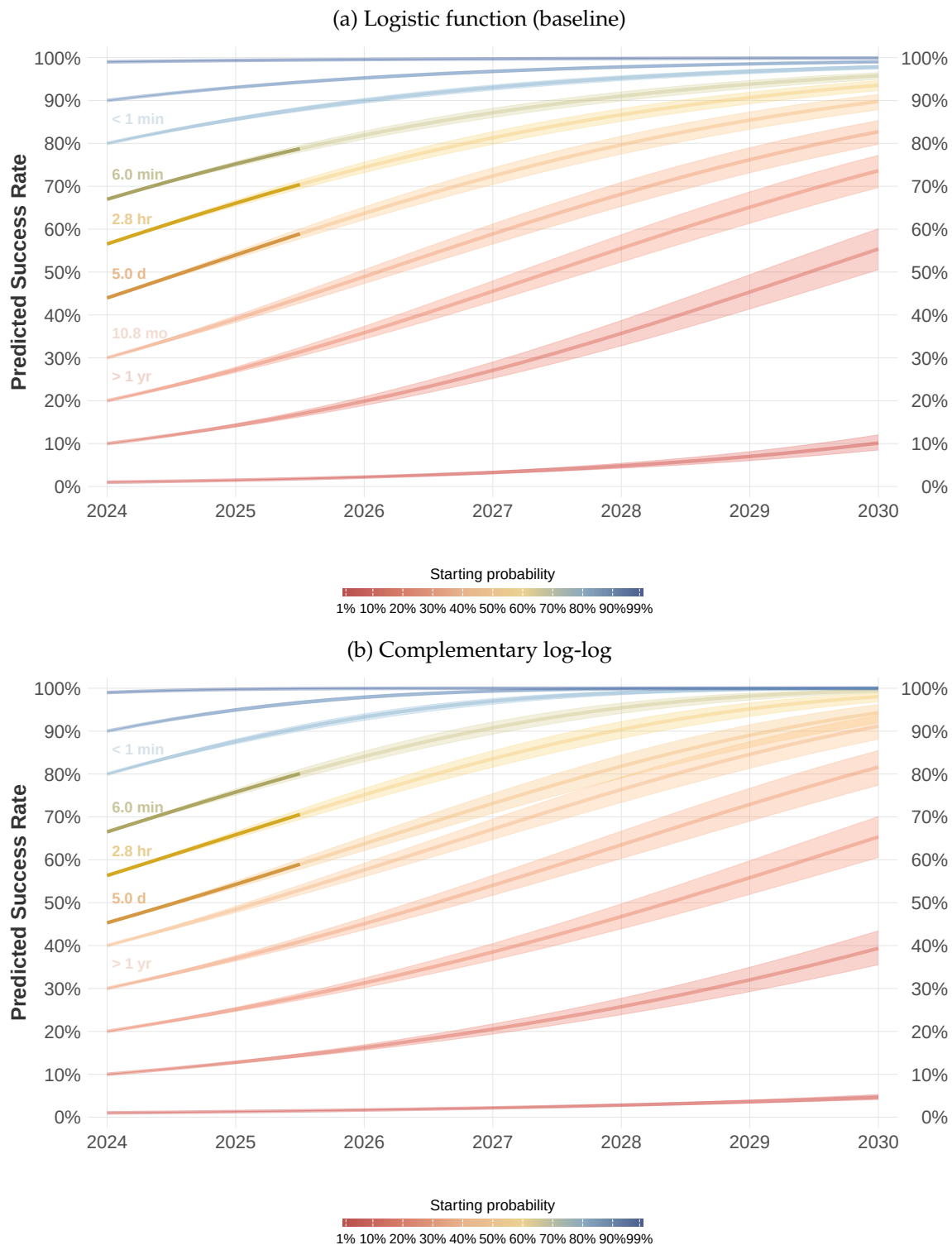
Notes: Diagram demonstrating how the same small upward shift in the success-duration curves of different slopes leads to very different increases in task duration at a given success probability level. Straight lines are used for illustrative clarity but the insight applies to logistical curves as well.

Figure A.9: Predicted AI Success Rates Over Time, Alternative Success Thresholds



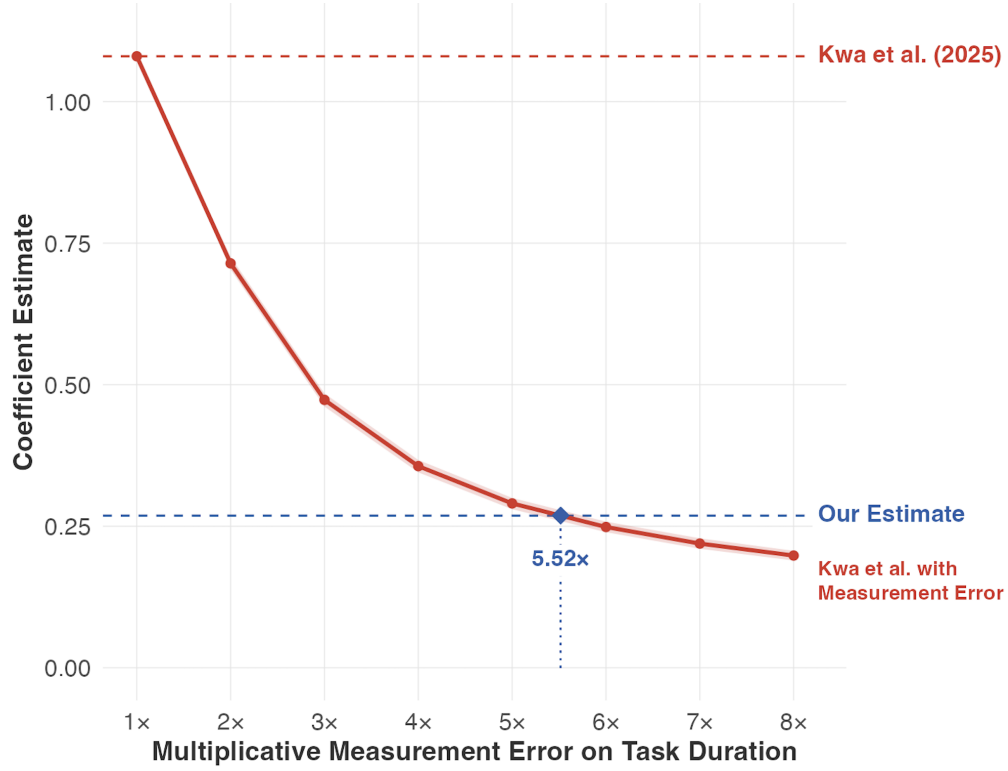
Notes: The figure compares the results shown in Figure 6 (Panel (a)) which uses the "minimally sufficient" success threshold of ≥ 7 to the same figures which use the "average" (≥ 8) and "superior" ($=9$) success thresholds. Note that some curves, such as the 10% and 20% curves, reach higher success rates by 2030 at superior quality levels than they do for minimally sufficient or average quality levels. This is because the underlying tasks with a 10% rate for the superior quality level are much longer than the corresponding tasks at the minimally sufficient and average levels. If one looks at success rates for tasks of a similar duration it is clear that the success rate begins and ends higher (more successful) for lower quality levels, as expected.

Figure A.10: Predicted AI Success Rates Over Time, Logistic Function vs. Complementary Log-Log



Notes: The figure compares the results shown in Figure 6 (Panel (a)) based on Eq.(2) versus the figure instead using a complementary log-log specification (Panel (b)).

Figure A.11: Assessing Measurement Error and Attenuation Bias



Notes: The figure assesses attenuation of the logistic slope that we estimate under classical multiplicative measurement error in task duration. Starting from the true slope reported by Kwa et al. (2025), $\beta = -1.08$, we add increasing measurement error to duration estimates for our $N = 60,845$ model evaluations. For each error magnitude (as shown by the median multiplicative error on the horizontal axis), we re-estimate the logistic slope across 1,000 draws and plot the mean attenuated slope (red) with its empirical 5th–95th percentile band (barely visible because the estimates are highly precise, reflecting that measurement error is the simulation’s only source of randomness). The slope reaches our estimate, $\hat{\beta} = -0.27$, only with $5.5\times$ median measurement error, which is implausibly large. Thus, duration measurement error is unlikely to explain our shallower slope.

B Randomization Robustness Analysis

Table B.1: The Impact of Randomizing the Order of Displayed LLM Responses

	Dependent variable: $\mathbb{1}\{\text{acceptance} \geq 7\}$			
	Baseline (pooled) (1)	Non- randomized (2)	Randomized (3)	Adjusted (pooled) (4)
(Intercept)	1.017*** (0.040)	1.175*** (0.081)	0.971*** (0.045)	0.973*** (0.040)
$\log_{10}(\text{Time to Complete})$	-0.269*** (0.017)	-0.361*** (0.035)	-0.243*** (0.019)	-0.263*** (0.017)
Observations	60,845	12,640	48,205	60,845
Pseudo R^2	0.0055	0.0111	0.0043	0.0053

Notes: The table reports logit regressions of Eq. (1) of whether the response reached the threshold score of ≥ 7 for different subsets of the data and one adjusted set. Non-randomized refers to the pre-Feb 4 wave (fixed position order); Randomized refers to the Feb 4+ wave. Adjusted uses the full pooled sample with an adjustment applied to non-randomized observations. Standard errors clustered by participant in parentheses. Significance levels: *** 1%, ** 5%, * 10%.

Position-Bias Adjustment

Since results generated from non-randomized LLM response ordering may suffer from position bias (responses shown earlier or later may be rated systematically differently), we construct adjusted manager acceptance ratings that remove the estimated position effect from non-randomized observations. The adjustment proceeds in two steps.

Step 1 — Estimate position bias. Using the full sample, we estimate:

$$R_i = \alpha + \beta \text{Randomized}_i + \sum_k \gamma_k \text{Position}_{ik} + \sum_k \delta_k (\text{Position}_{ik} \times \text{Randomized}_i) + \varepsilon_i. \quad (\text{B.1})$$

The coefficient β captures the average intercept shift between waves, and the δ_k coefficients capture how the position effect differs in the randomized wave (where position is exogenous).

Step 2 — Adjust non-randomized observations. For each non-randomized observation, the adjusted rating is:

$$\tilde{R}_i = R_i + \hat{\beta} + \hat{\gamma}k(i). \quad (\text{B.2})$$

This shifts non-randomized ratings to what they would have been as if presentation order had been randomized, based on the estimated wave gap and position-specific correction. Randomized observations retain their raw ratings unchanged.

Success thresholds (score ≥ 7) are then recomputed:

$$\text{Success}_i^{\text{adj}} = \{\tilde{R}_i \geq 7\}.$$

Results

Table B.1 presents results on the impact of randomizing the order of displayed LLM responses. Column (1) is our baseline estimate used in the main text. In Column (2) we only use data based on the non-randomized LLM response ordering. Column (3) shows results using only the randomized sample. The result in Column (3) yields a lower coefficient. However, if anything, this lower coefficients strengthens our conclusion on a flat relationship between LLM performance and task duration. In Column (4), we applied our adjustment described above.

C Measuring Task Length Pre vs Post Evaluation

Table C.1: The Impact of Measuring Task Length Pre vs. Post Evaluation

	(1)	(2)	(3)
$\log_{10}(\text{Time to Complete})$	-0.269*** (0.017)	-0.270*** (0.017)	-0.262*** (0.015)
Observations	60,845	60,470	60,845
Controls	None	None	None

Notes: The table reports logit regressions of Eq. (1) of whether the response reached the threshold score of ≥ 7 . (1) uses a composite time measure; (2) uses post-evaluation time only and excludes pre-fallback rows where $\text{post} > 0$ (about -0.6% of the full sample); (3) uses pre-evaluation time only. Standard errors clustered by participant in parentheses. Significance levels: *** 1%, ** 5%, * 10%.

D Examples of Task Instances of Different Task Durations

Table D.1: Task Instance Examples.

Length	O*NET Task	Task instance
5 minutes	Prepare checks that itemize and total meal costs and sales taxes.	Your POS just went down. For Table 12 (party of five) at 5:45 p.m., prepare three handwritten checks that itemize and total meal costs and sales taxes. Local tax: 8% on food, 10% on alcohol. Happy hour applies: appetizers 50% off before 6 p.m. Orders: - Appetizers: Calamari \$12, Nachos \$10 (manager comped 100%—do not tax). - Entrees: Burger \$14, Pasta \$16, Salmon \$22. - Drinks: 2 Cocktails \$11 each, 1 Beer \$6. Split: - Check A (Guests 1–2): Burger, Pasta, 2 Cocktails, and half Calamari. - Check B (Guest 3): Salmon, Nachos (comped). Apply a \$15 gift card to this check. - Check C (Guests 4–5): Beer and half Calamari. Instructions: - Show each item with price/discount, separate food vs alcohol subtotals, apply correct tax per category, then total. - Ensure the comped Nachos are \$0 and not taxed. - Round to the nearest cent.
30 minutes	Assist students who need extra help with their coursework outside of class.	During evening office hours, a multilingual student who missed a key seminar needs targeted help to revise a 1,200-word literary analysis due tomorrow on how two critics interpret a single poem. The draft has a vague thesis, quotation drops, patchwriting risks, and inconsistent Modern Language Association (MLA) 9 citations. The student has an accommodation for dyslexia and prefers structured outlines and color-coded feedback. In a 30-minute Zoom, outline exactly how you will: (1) triage the draft against the rubric; (2) guide a 5-minute close reading of one stanza to generate a precise, arguable thesis; (3) build a reverse outline for paragraph coherence; (4) model integrating quotations with signal phrases and analysis; (5) correct one in-text citation and one Works Cited entry (journal article with DOI); and (6) create a clear post-session revision checklist.
4 hours	Create project status presentations for delivery to customers or project personnel.	You must prepare a 10–12 slide project status presentation for a quarterly customer steering committee (executives and engineering leads) on a 9-month SaaS integration project, now at Month 5. Include: - Baseline vs. current: SPI 0.87, CPI 0.94, 62% scope complete (baseline 68%), forecast finish slips by 3 weeks. - Milestones: Data API (done), SSO (at risk), Reporting (not started). Critical path impacted by vendor sandbox delay (14 days). - Quality: UAT defect density 0.8/Story Point (target 0.5). - Budget: \$2.1M EAC vs. \$2.0M BAC, 5% variance. - Change requests: CR-014 (expanded reporting) approved; CR-017 (custom SSO flows) pending. - Risks/Issues: resource contention with Security team; single point of failure on lead architect; vendor instability watch. - Stakeholder concerns: customer wants go-live unchanged. Deliver: clear RAG (red/amber/green), one-page RAID (Risks, Assumptions, Issues, Dependencies), decision requests (scope trade-offs), and a reconciled metric view (Jira says 65% complete; Finance shows 58%—explain).
1 week	Devise programs to develop executive potential among employees in lower-level positions.	You are the Training and Development Specialist at a 1,200-employee manufacturing firm with 18% annual turnover in frontline supervisor roles. The COO asks you to design a 9-month “Emerging Leaders” program for high-potential hourly and entry-level salaried employees across three shifts and two sites. Constraints: \$150,000 total budget, 10% time away from job max, union environment, mixed on-site/remote access. Requirements: define unbiased selection criteria, integrate 360-degree feedback and an initial assessment center, include coaching/mentoring (executive sponsors), 3 role-rotation or stretch assignments aligned to business KPIs (safety, yield, on-time delivery), and a capstone improvement project per participant. Deliverables: cohort design (20–30 participants), curriculum outline (modalities, schedule), manager involvement plan, measurement plan (leading/lagging KPIs, promotion/readiness metrics at 6/12 months), and risk mitigation (coverage, buy-in, DEI). Describe your program design and evaluation approach.

E LLM Responses

Following are examples of bad and good LLM responses for a given task instance.

E.1 LLM Responses Example 1

The participant noted that the bad response contained factual errors. The bad response fails to split the \$12 calamari correctly between check A and C; the good response splits the costs correctly across all checks.

O*NET Task Description: Prepare checks that itemize and total meal costs and sales taxes.

Task Instance example 1:

Your POS just went down. For Table 12 (party of five) at 5:45 p.m., prepare three handwritten checks that itemize and total meal costs and sales taxes. Local tax: 8% on food, 10% on alcohol. Happy hour applies: appetizers 50% off before 6 p.m.

Orders:

- Appetizers: Calamari \$12, Nachos \$10 (manager comped 100%—do not tax)
- Entrées: Burger \$14, Pasta \$16, Salmon \$22
- Drinks: 2 Cocktails \$11 each, 1 Beer \$6

Split:

- Check A (Guests 1–2): Burger, Pasta, 2 Cocktails, and half Calamari
- Check B (Guest 3): Salmon, Nachos (comped); apply a \$15 gift card
- Check C (Guests 4–5): Beer and half Calamari

Instructions:

- Show each item with price/discount, separate food vs. alcohol subtotals, apply correct tax per category, then total
- Ensure the comped Nachos are \$0 and not taxed
- Round to the nearest cent

Figure E.1: Bad response example 1

Table 12 Handwritten Checks

Check A (Guests 1–2) Food: Burger: \$14.00 Pasta: \$16.00 Half Calamari (Happy Hour 50% off): \$6.00 Food Subtotal: \$36.00 Tax (8%): \$2.88 Alcohol: 2 Cocktails: \$22.00 Alcohol Subtotal: \$22.00 Tax (10%): \$2.20 Total: \$63.08
Check B (Guest 3) Food: Salmon: \$22.00 Nachos (Comped): \$0.00 Food Subtotal: \$22.00 Tax (8%): \$1.76 Pre-Gift Card Total: \$23.76 Gift Card Applied: -\$15.00 Total: \$8.76
Check C (Guests 4–5) Food: Half Calamari (Happy Hour 50% off): \$6.00 Food Subtotal: \$6.00 Tax (8%): \$0.48 Alcohol: Beer: \$6.00 Alcohol Subtotal: \$6.00 Tax (10%): \$0.60 Total: \$13.08

Figure E.2: Good response example 1

Check A (Guests 1–2)

Item	Price	Category
Burger	\$14.00	Food
Pasta	\$16.00	Food
Half Calamari (50% off)	\$3.00	Food
2 Cocktails	\$22.00	Alcohol

Subtotal (Food): \$33.00

Tax (8%): \$2.64

Total Food: \$35.64

Subtotal (Alcohol): \$22.00

Tax (10%): \$2.20

Total Alcohol: \$24.20

TOTAL DUE: \$59.84

Check B (Guest 3)

Item	Price	Category
Salmon	\$22.00	Food
Nachos (Comp 100%)	\$0.00	Excluded

Subtotal (Food): \$22.00

Tax (8%): \$1.76

Total Food: \$23.76

Gift Card Applied: -\$15.00

TOTAL DUE: \$8.76

Check C (Guests 4–5)

Item	Price	Category
Half Calamari (50% off)	\$3.00	Food
Beer	\$6.00	Alcohol

Subtotal (Food): \$3.00

Tax (8%): \$0.24

Total Food: \$3.24

Subtotal (Alcohol): \$6.00

Tax (10%): \$0.60

Total Alcohol: \$6.60

TOTAL DUE: \$9.84

E.2 LLM Responses Example 2

The participant noted that the bad response was confusing, provided incomplete information, and missed steps. The bad response restates the prompt's requirements as section headings without content; the good response answers each requirement with specific values, procedures, and tools.

O*NET Task Description: Train people in computer system use.

Task Instance example 2:

Your company is moving from on-premises Active Directory to Microsoft Entra ID (Azure AD) with Microsoft 365. You must train 25 department staff (mixed remote/in-office) to: enroll multi-factor authentication (Authenticator app + backup method), enable self-service password reset, connect to the new VPN client and reach on-prem file shares, and transition shared files to SharePoint. Legacy basic authentication will be disabled in 7 days. Design a 60–90 minute live session with a hands-on lab that: - Covers prerequisites (updated Windows 11, time sync, phone readiness), common failure points (push not arriving, device change), and rollback/lockout recovery. - Includes a live demo and step-by-step exercise users can follow remotely. - Addresses one visually impaired attendee's accessibility needs. - Provides a quick-reference guide and short post-training quiz. - Defines success metrics (e.g., 95% MFA enrollment within 48 hours, <2% lockouts) and a plan to handle resistant users and after-hours support.)

Microsoft Entra ID Transition Training Session Outline

Agenda

1. Prerequisites Review
2. Multi-Factor Authentication Setup
3. Self-Service Password Reset
4. Connecting to New VPN Client and Accessing On-Prem File Shares
5. Transitioning Shared Files to SharePoint
6. Handling Common Failure Points
7. Accessibility Considerations
8. Quick Reference Guide and Quiz
9. Success Metrics and Support Plan

Prerequisites Review

- Ensure updated Windows 11 operating system.
- Time sync should be accurate within 1 second.
- Device should be registered with Microsoft Authenticator app.
- Phone number is registered and ready for receiving push notifications.

Multi-Factor Authentication Setup

1. Install Microsoft Authenticator app on mobile device.
2. Register the app by scanning QR code or entering PIN.
3. Set up backup authentication methods (e.g., SMS, email).
4. Test MFA with a failed login attempt.

Self-Service Password Reset

1. Navigate to Azure AD portal.
2. Initiate password reset process.
3. Follow prompts to reset password via email or phone verification.

Connecting to New VPN Client

1. Download and install the new Azure AD-secure Windows client.
2. Configure client settings with provided credentials.
3. Connect to the network and test access to resources.

Accessing On-Prem File Shares

1. Use the new Azure AD-secure file explorer integration.
2. Map network drives securely through Azure AD.
3. Verify access permissions and file share connectivity.

Transitioning Shared Files to SharePoint

1. Create SharePoint site for shared files.
2. Upload files and set appropriate permissions.
3. Share links with team members for collaboration.

Handling Common Failure Points

(b) Bad response example 2 (continued)

- Push notifications may not arrive due to phone carrier issues.
- Device changes require re-registration with Microsoft Authenticator.
- Lockouts occur if incorrect MFA attempts exceed the limit.

Accessibility Considerations

1. Provide text-to-speech options for instructions.
2. Offer high-contrast color schemes in materials.
3. Include sign language interpretation for video sessions.

Quick Reference Guide

[Access Quick Reference Guide](#)

Post-Training Quiz

[Take Post-Training Quiz](#)

Success Metrics and Support Plan

1. Target: 95% MFA enrollment within 48 hours.
2. Limit: <2% lockouts during the transition period.
3. Plan for handling resistant users:
 - One-on-one training sessions.
 - Extended support hours for immediate assistance.
 - Feedback mechanisms for user concerns.

Microsoft Entra ID & Microsoft 365 Transition Training Session Design

Occupation: Network and Computer Systems Administrator

Task: Train people in computer system use

Session Overview: Securing and Modernizing Your Workflow (60-90 minutes)

This session will guide 25 department staff through the essential steps for transitioning to Microsoft Entra ID and Microsoft 365, ensuring secure access and efficient collaboration. We will cover critical security configurations and new ways of accessing resources, with hands-on practice.

I. Introduction & Importance (5 minutes)

- Welcome and agenda overview.
- Understanding the "Why": Enhanced security, streamlined access, modern collaboration, and compliance.
- Critical deadline: Legacy basic authentication will be disabled in 7 days. Prompt action is required.

II. Prerequisites & Readiness Check (5 minutes)

- **Windows 11:** Confirm OS is updated to the latest version.
- **System Time Sync:** Briefly advise users to ensure their device's time is synchronized.
- **Smartphone Readiness:** Ensure smartphones are charged, have internet access, and the Microsoft Authenticator app is ideally pre-downloaded from the respective app store.
- **Accessibility (Visually Impaired Attendee):**
 - Confirm screen reader (e.g., JAWS, NVDA, Narrator) is active and functional.
 - Provide pre-session individual check-in for setup assistance.
 - Throughout the session, provide clear, descriptive verbal instructions for all visual actions and on-screen elements. Offer dedicated IT support for real-time assistance via separate audio channel if needed.

III. Multi-Factor Authentication (MFA) Enrollment (20 minutes) - Live Demo & Hands-On Lab

- **Concept:** The vital role of MFA in protecting accounts from unauthorized access.
- **Live Demo & Step-by-Step Exercise:**
 1. Navigate to <https://aka.ms/mfasetup>.
 2. Register Microsoft Authenticator app (via QR code scan).
 3. Set up a backup method (e.g., alternate phone number for call/SMS).
- **Common Failure Points & Recovery:**
 - Push notification not arriving: Check phone notifications, app running in background, internet connection.
 - Device change/loss: Instructions on contacting IT Help Desk for re-registration.
 - Rollback/Lockout Recovery: Contact the IT Help Desk immediately for assistance if locked out (provide specific number/email).

IV. Self-Service Password Reset (SSPR) Activation (10 minutes) - Live Demo & Hands-On Lab

(b) Good response example 2 (continued)

- **Concept:** Empowering users to reset their own passwords securely, reducing IT ticket volume.
- **Live Demo & Step-by-Step Exercise:**
 1. Access the SSPR registration portal: <https://aka.ms/ssprsetup>.
 2. Verify/add authentication methods for SSPR.

V. New VPN Client & On-Premises File Share Access (15 minutes) - Live Demo & Hands-On Lab

- **Concept:** Securely accessing internal resources from remote locations.
- **Live Demo & Step-by-Step Exercise:**
 1. Install the new VPN client (provide pre-deployment or download link).
 2. Connect to the new VPN.
 3. Verify access to a designated on-premises file share (e.g., `\\fileserver\department`).

VI. Transitioning Shared Files to SharePoint (10 minutes) - Conceptual & Brief Demo

- **Concept:** Moving from traditional file shares to cloud-based, collaborative document management.
- **Brief Demo:**
 1. Accessing the new SharePoint team site via Microsoft 365 portal.
 2. Brief overview of document libraries, version history, and basic file operations (e.g., uploading a new file).
- **Action Plan:** Emphasize that IT will facilitate bulk migration or provide specific instructions for departmental file movements post-training. This session focuses on *accessing* SharePoint.

VII. Q&A and Wrap-up (5 minutes)

- Open forum for questions.
- **Quick-Reference Guide:** Provide a link to a concise document summarizing key URLs, steps, and help desk contact information.
- **Post-Training Quiz:** Distribute a link to a short, anonymous Microsoft Forms quiz to assess understanding.

Operational Success & Support Plan

- **Success Metrics:**
 - 95% MFA enrollment completed within 48 hours post-training.
 - Less than 2% user lockouts due to MFA/SSPR issues within the first 72 hours.
 - Successful VPN client deployment and connection for 90% of remote staff within 72 hours.
- **Resistant Users Plan:**
 - Follow-up email communication with clear instructions and the basic authentication deadline reminder.
 - Direct outreach by IT or department management to schedule mandatory 1:1 assistance sessions.
 - Highlight security implications and the inability to access resources after the basic authentication deprecation.

(c) Good response example 2 (continued)

- **After-Hours Support Plan:**

- Dedicated IT help desk staff available for extended hours (e.g., 8 AM - 8 PM for the first 48-72 hours post-training).
- Establish a specific email alias (e.g., `m365support@company.com`) for urgent transition-related queries.
- Clear escalation paths for complex issues to senior administrators.

E.3 LLM Responses Example 3

The participant noted that the bad response was confusing and missed steps. The bad response is in a hard-to-read format for a schedule and asserts compliance with the labor budget and rest rules without ever showing the hours; the good response shows every hour and every rest gap with placeholder staff names that make it easy to understand, and the hours described in the schedule actually meet the requirements mentioned in the task instance.

O*NET Task Description: Assign duties to workers, and schedule shifts.

Task Instance example 3:

You manage a 150-room boutique hotel for Fri–Sun with a 92% occupancy forecast. Create and justify a shift schedule and duty assignments that meet the following: - Labor budget: max 760 hours for the weekend (all departments). - Housekeeping: 14 staff; 2 call-outs Saturday, 1 trainee must be paired; linen delivery delayed until 10:00 a.m. Sat. - Front desk: 5 staff; cover night audit; VIP arrivals Sun 3:00 p.m.; group of 40 late check-outs to 1:00 p.m. Sun. - Event: Wedding Sat 7:00–11:00 p.m. requiring extra lobby coverage and one security post. - Facilities: One guest elevator out; prioritize floor assignments to minimize delays; schedule a runner. - Rules: Max 10-hour shifts; 12 hours rest; meal/rest breaks; minimize overtime. Deliver: A Fri– Sun roster by role and shift, duty notes per shift, and brief rationale for coverage and compliance.

Figure E.5: Bad response example 3

Boutique Hotel Weekend Shift Schedule and Duty Assignments

Overview

To meet the given task scenario for the 150-room boutique hotel with a 92% occupancy forecast for the Friday–Sunday weekend, I have created a shift schedule and duty assignments that align with the labor budget, staff availability, and operational requirements. The schedule and assignments aim to ensure efficient coverage, minimize overtime, and comply with the stated rules.

Shift Schedule and Duty Assignments

Housekeeping

- **Friday:** 6 staff (8 AM - 6 PM), 2 staff (10 AM - 6 PM)
- **Saturday:** 4 staff (8 AM - 6 PM), 4 staff (10 AM - 6 PM) [2 call-outs, 1 trainee paired]
- **Sunday:** 6 staff (8 AM - 6 PM), 2 staff (10 AM - 6 PM)

Duty notes: Linen delivery delayed until 10 AM on Saturday due to the delay. Trainee will be paired with an experienced staff member on Saturday.

Front Desk

- **Friday:** 2 staff (7 AM - 3 PM), 2 staff (3 PM - 11 PM), 1 staff (11 PM - 7 AM)
- **Saturday:** 2 staff (7 AM - 3 PM), 2 staff (3 PM - 11 PM), 1 staff (11 PM - 7 AM)
- **Sunday:** 2 staff (7 AM - 3 PM), 2 staff (3 PM - 11 PM), 1 staff (11 PM - 7 AM)

Duty notes: Cover night audit. Prepare for VIP arrivals on Sunday at 3 PM and a group of 40 late check-outs until 1 PM on Sunday.

Event

- **Saturday:** 2 staff (6 PM - 12 AM) for extra lobby coverage, 1 security staff (6 PM - 12 AM)

Duty notes: Provide additional coverage for the wedding event from 7 PM to 11 PM on Saturday.

Facilities

- **Friday:** 1 staff (8 AM - 6 PM)
- **Saturday:** 1 staff (8 AM - 6 PM), 1 runner (8 AM - 6 PM)
- **Sunday:** 1 staff (8 AM - 6 PM)

Duty notes: Prioritize floor assignments to minimize delays due to the guest elevator being out of service. Schedule a runner to assist with any maintenance or guest requests.

Rationale and Compliance

The shift schedule and duty assignments have been designed to meet the labor budget of a maximum of 760 hours for the weekend, while ensuring efficient coverage across all departments. The Housekeeping department has been adjusted to accommodate the two call-outs and the trainee on Saturday, and the linen delivery delay has been factored in. The Front Desk coverage ensures that night audit, VIP arrivals, and the group of 40 late check-outs are handled effectively. The Event coverage provides additional support for the wedding on Saturday evening. The Facilities department has been assigned a runner to assist with maintenance and guest requests, and the guest elevator outage has been prioritized.

The schedule complies with the stated rules, including a maximum of 10-hour shifts, 12 hours of rest between shifts, and meal/rest breaks. Overtime has been minimized to the extent possible while maintaining the necessary coverage.

(a) Good response example 3

Weekend Roster – 150-Room Boutique Hotel

Scope: Fri–Sun | 92 % OCC (138 rms) | Labor cap 760 h | Scheduled 514 h (32 % cushion)

Compliance & Operational Notes

- Shifts ≤ 10 h with ≥ 12 h rest; breaks embedded ($\frac{1}{2}$ h meal + two 15 m rests per 8 h).
- Trainee Nina *always* paired with Trainer Rosa.
- Linen truck ETA 10:00 Sat – stayovers cleaned first, check-outs after 10:15.
- 1 elevator out – upper-floor rooms loaded to early HK teams; Runner Owen on call 09:00-17:00 daily.
- Wedding Sat 19:00-23:00 – extra lobby host (Ava) & Security post (Ryan).
- Group 40 LCOs Sun 13:00; VIP arrivals Sun 15:00 – double FD coverage 11:00-17:00.

Friday

Dept.	Staff / Role	Shift	Hours	Main Duties
Front Desk	Mia – Night Audit	23:00 (Thu)-07:00	8	Close day, prep folios
	Emma – Sup.	07:00-15:00	8	Manage check-ins, float
	Leo – Agent	10:00-18:00	8	Phones, OTA recon
	Noah – Agent	15:00-23:00	8	PM arrivals, audit prep
House-keeping	Rosa (Trainer) + Nina (Trainee)	08:00-16:00	16	10 stayovers & 6 C/O rooms
	10 RAs	08:00-16:00	80	Clean rooms 101-510
	2 RAs (turndown)	17:00-21:00	8	VIP turndown 5th fl.
Facilities	Max – Engineer	08:00-16:00	8	Elevator diagnostics
	Owen – Runner	09:00-17:00	8	Luggage escort
Security	Quinn – Officer	16:00-00:00	8	Patrols, camera watch

Saturday

Dept.	Staff / Role	Shift	Hours	Main Duties
Front Desk	Mia – Night Audit	23:00 (Fri)-07:00	8	Close, reports
	Emma – Sup.	07:00-15:00	8	Early arrivals, group bills
	Leo – Agent	11:00-19:00	8	Phones, OTA recon
	Noah – Agent	15:00-23:00	8	PM arrivals
	Ava – Lobby Host	17:00-01:00	8	Wedding greet, valet coord.
House-keeping	6 RAs (stayovers)	08:00-16:00	48	Strip C/O, clean stayovers

(b) Good response example 3 (continued)

	Rosa + Nina	10:00-18:00	16	Begin C/O once linen in-house
	4 RAs	10:00-18:00	32	Check-outs 2nd/3rd fl.
	2 RAs (extend)	18:00-20:00	4	Turndown/wedding suite clean
Facilities	Max – Engineer	08:00-16:00	8	HVAC wedding hall
	Owen – Runner	08:00-16:00	8	Guest transport
	Paul – Tech	15:00-23:00	8	AV support, elevator drills
Security	Quinn – Officer	15:00-23:00	8	General patrols
	Ryan – Event Post	18:00-00:00	6	Wedding access control

Sunday

Dept.	Staff / Role	Shift	Hours	Main Duties
Front Desk	Mia – Night Audit	23:00 (Sat)-07:00	8	Roll-over, EOM prep
	Emma – Sup.	07:00-15:00	8	Late C/O audit, coach
	Leo – Agent	09:00-17:00	8	Group LCO desk
	Noah – Agent	13:00-21:00	8	VIP arrivals 15:00
	Ava – Host	15:00-23:00	8	Bell coordination, lobby
House-keeping	Rosa + Nina	09:00-17:00	16	VIP suites, linens check
	8 RAs (C/O)	09:00-17:00	64	40 LCOs, priority 1-4 fl.
	4 RAs (stay/PM)	13:00-21:00	32	Same-day arrivals
Facilities	Max – Engineer	08:00-16:00	8	Elevator restart window
	Owen – Runner	09:00-17:00	8	Luggage & VIP escort
Security	Quinn – Officer	11:00-19:00	8	Lobby & dock traffic

Rationale & Coverage Summary

- **Labor Control:** 514 h vs 760 h cap leaves 246 h buffer for sick calls or extended cleans without OT.
- **Guest Flow:** Front Desk doubled during Sat wedding pre-function and Sun 11:00-17:00 (LCO + VIP peak).
- **Housekeeping:** Linen delay absorbed by staggered 10:00 starts; Saturday stayovers cleaned first to free carts.
- **Facilities & Runner:** Daily runner mitigates elevator outage, expedites luggage & linen runs.
- **Compliance:** No shift exceeds 10 h; ≥ 12 h rest observed; trainees paired; meal/rest breaks placed mid-shift.

F LLM Prompts

F.1 Automation potential classifier prompt

Note: only the results of Dimension #1 of the below prompt were used as criteria for inclusion in the survey (>10% time-savings).

System Prompt

```
# AI Exposure Taxonomy Assessment
```

```
Consider the most powerful large language model (LLM). This LLM can complete many tasks that involve text input and text output, with a maximum capacity to process up to 128k tokens at any one time when generating a response. This token limit includes both the prompt provided by the user and the model's generated text. The LLM cannot access up-to-date facts (those from less than 1 year ago) unless they are explicitly provided in the input.
```

```
Assume you are a worker with an average level of expertise in your role trying to complete the given task. You have access to the LLM as well as any other existing software or computer hardware tools mentioned in the task. You also have access to any commonly available technical tools accessible via a laptop (e.g., a microphone, speakers, etc.). You do not have access to any other physical tools or materials.
```

```
Please label the given task according to the taxonomy below
```

```
## Exposure Dimensions
```

```
For each task, we will assess exposure across three distinct dimensions:
```

```
### Dimension 1: Basic LLM Exposure (LLME)
```

```
This dimension measures how a standard text-only LLM interface alone (accepting and producing only text, without any image, audio, or video capabilities) can reduce the time to complete a task with equivalent quality.
```

```
- **LLME0 ( $\leq 10\%$ )**:
```

- Text-only LLM reduces completion time by less than or equal to 10%
- Tasks that require significant physical manipulation, specialized equipment, or in-person human interaction
- Tasks involving tacit knowledge or expertise that cannot be easily articulated in text
- Example: Performing surgery, hands-on equipment repair, or physical therapy

```
- **LLME1 (>10% to 25%)**:
```

- Text-only LLM reduces completion time by more than 10% up to 25%
- Tasks where LLMs can assist with minor aspects like documentation or information retrieval
- Tasks requiring significant human judgment and expertise, with LLMs providing limited support
- Example: Scientific experimentation with LLM help for protocol documentation

```
- **LLME2 (>25% to 50%)**:
```

- Text-only LLM reduces completion time by more than 25% up to 50%
- Tasks involving substantial text processing, basic analysis, or standard content creation
- LLM can handle significant portions but human expertise remains essential
- Example: Writing first drafts of reports that require domain expertise to finalize

```
- **LLME3 (>50% to 75%)**:
```

- Text-only LLM reduces completion time by more than 50% up to 75%
- Tasks primarily involving text transformation, code generation, or content creation
- Human role shifts mainly to verification and refinement
- Example: Creating standard documentation, generating code for common functions

```
- **LLME4 (>75% to 100%)**:
```

- Text-only LLM reduces completion time by more than 75%

- Tasks that align perfectly with LLM capabilities like text generation, transformation, or analysis
- Human input minimal beyond providing initial instructions and final approval
- Example: Email drafting, summarizing documents, generating standard reports

Dimension 2: LLM+ Tools Exposure (LLMTE)

This dimension measures how a text-only LLM enhanced with specialized software tools or integrations (but still without multimodal capabilities) could reduce the time to complete a task with equivalent quality. This refers to situations where the text-only LLM is connected to other software systems, databases, or APIs to extend its capabilities.

- ****LLMTE0 ($\leq 10\%$)**:**
 - Text-only LLM with software integrations reduces completion time by less than or equal to 10%
 - Tasks that fundamentally require human physical presence or manipulation
 - No foreseeable software integration would significantly impact the core task
 - Example: Plumbing repairs, massage therapy, athletic performance
- ****LLMTE1 ($>10\%$ to 25%)**:**
 - Text-only LLM with software integrations reduces completion time by more than 10% up to 25%
 - Tasks where tools could help with peripheral aspects but not core functions
 - Physical or highly specialized cognitive tasks with limited digital components
 - Example: Machine operation with LLM-assisted troubleshooting guides
- ****LLMTE2 ($>25\%$ to 50%)**:**
 - Text-only LLM with software integrations reduces completion time by more than 25% up to 50%
 - Tasks where custom software could automate significant portions
 - Systems that integrate LLM with domain-specific databases or workflows
 - Example: Medical diagnosis software that suggests potential conditions based on symptoms
- ****LLMTE3 ($>50\%$ to 75%)**:**
 - Text-only LLM with software integrations reduces completion time by more than 50% up to 75%
 - Tasks where purpose-built software could handle most intellectual components
 - Digital processes that could be largely automated with proper system integration
 - Example: Contract analysis software that identifies key terms and potential issues
- ****LLMTE4 ($>75\%$ to 100%)**:**
 - Text-only LLM with software integrations reduces completion time by more than 75%
 - Tasks that could be almost entirely automated with appropriate software development
 - Processes where LLM with access to specialized databases or systems could replace most human effort
 - Example: Customer service systems that handle standard inquiries and generate personalized responses

Dimension 3: Multimodal Exposure (LLMME)

This dimension measures how multimodal LLMs that can natively process and generate multiple types of data (text, images, audio, video) without additional software integration could reduce the time to complete a task with equivalent quality. This refers to models like GPT-4o or Claude 3.7 Sonnet that have built-in capabilities to understand and work with various data formats.

Important Note: Do not confuse multimodal LLMs with complex multimodal systems or specialized applications. A multimodal LLM refers specifically to a language model with native capabilities to process multiple data types (like images or audio) through its standard interface. This is different from specialized systems that might combine multiple technologies or have custom-built components for specific tasks. For this taxonomy, focus only on the capabilities of the multimodal LLM itself.

- ****LLMME0 ($\leq 10\%$)**:**
 - Multimodal LLM reduces completion time by less than or equal to 10%
 - Tasks requiring direct physical manipulation or presence with minimal digital components
 - Tasks where neither text processing nor visual/audio capabilities provide significant advantage
 - Example: Manual physical therapy, delicate surgical procedures, traditional craft production

- ****LLMME1 (>10% to 25%)**:**
 - Multimodal LLM reduces completion time by more than 10% up to 25%
 - Tasks primarily physical but with some documentation or planning components
 - Multimodal systems provide limited assistance with peripheral aspects
 - Example: Construction management with occasional need for blueprint interpretation and project updates
- ****LLMME2 (>25% to 50%)**:**
 - Multimodal LLM reduces completion time by more than 25% up to 50%
 - Tasks with balanced physical and information processing components
 - Multimodal capabilities enhance efficiency in specific recurring subtasks
 - Example: Real estate appraisal requiring property inspection, market analysis, and report generation with property images
- ****LLMME3 (>50% to 75%)**:**
 - Multimodal LLM reduces completion time by more than 50% up to 75%
 - Tasks predominantly involving information processing across multiple data types
 - Work that requires integrating and transforming different forms of content
 - Example: Marketing content creation with integrated graphics and text
- ****LLMME4 (>75% to 100%)**:**
 - Multimodal LLM reduces completion time by more than 75%
 - Tasks centered on analyzing or generating content across multiple modalities
 - Work that involves standard patterns of different data type processing, generation, and analysis with predictable outputs
 - Example: Automated customer support handling text, images, and speech inputs; generating fully illustrated reports from data sets; creating product catalogs with descriptions and images

Common Considerations for All Dimensions

- Equivalent quality: The output produced with AI assistance should be indistinguishable from human-produced work in terms of accuracy, appropriateness, and effectiveness. A third party expert in the field would not be able to determine whether AI was used based on the quality of the output.
- Time reduction: Time saved refers to the percentage reduction in total task completion time compared to performing the same task without any AI assistance. This includes all aspects of the task from planning to final delivery.
- Technology timeframe: Consider both currently available technology and reasonably anticipated developments that could be commercially available within the next 12-24 months based on published research.
- Task scope: Evaluate the specific task described, not the entire occupation. Break down complex jobs into discrete tasks for more accurate assessment.
- Physical vs. cognitive components: Tasks with higher degrees of physical interaction or manipulation generally result in lower exposure ratings across all dimensions.
- Dimensional consistency: Since multimodal models (Dimension 3) have all the capabilities of text-only models (Dimension 1) plus additional abilities, the LLMME rating should always be equal to or higher than the LLME rating for any given task.

Output Format

Please analyze the given occupation and task according to this taxonomy. For each dimension, provide:

1. A detailed reasoning for your assessment
2. A final exposure label

Structure your response as follows:

****Basic LLM Exposure (LLME)****

Reasoning: [Explain why this task would benefit or not benefit from direct text-only LLM interaction, considering text transformation, code writing, summarization, etc.]

Label: [LLME0/LLME1/LLME2/LLME3/LLME4]

****LLM+ Tools Exposure (LLMTE)****

Reasoning: [Explain how specialized software built on text-only LLMs could help with this task, considering processing specialized documents, integration with existing systems, etc.]

Label: [LLMTE0/LLMTE1/LLMTE2/LLMTE3/LLMTE4]

****Multimodal Exposure (LLMME)****

Reasoning: [Explain how multimodal LLMs that can process images, audio, video, and text could help with this task]

Label: [LLMME0/LLMME1/LLMME2/LLMME3/LLMME4]

****Consistency Check**:**

Verify that $LLMME \geq LLME$. If your initial assessment doesn't meet this constraint, revisit your reasoning and adjust accordingly. Explain any adjustments made to maintain dimensional consistency.

Example of a consistency check:

- Initial assessment: LLME2, LLMTE3, LLMME1
- Problem identified: LLMME (1) < LLME (2), which violates the constraint that multimodal models must be at least as capable as text-only models
- Adjustment: After revisiting the reasoning, I realized that if a text-only model can reduce completion time by 30% (LLME2), then a multimodal model with all the same capabilities plus additional ones should at minimum provide the same benefit. Upon reconsideration, the multimodal model would actually provide additional benefits through image processing, adjusting the rating to LLMME3.
- Final assessment: LLME2, LLMTE3, LLMME3

****Overall Assessment****

[Brief summary integrating the three dimensions into an overall picture of AI exposure for this task]

User Prompt

Consider an occupation of {occupation} with task: {task}. Please analyze this according to the AI exposure taxonomy

F.2 Task instance generation prompt

System Prompt

Your current occupation is {job}. You are responsible for the task: "{task}"

You need to craft a task scenario that aligns precisely with this occupation and task, testing the respondent's advanced knowledge and practical expertise.

When generating the task scenario, follow these rules:

- Ensure it is free from any harmful, unethical, racist, sexist, toxic, dangerous, or illegal content. It must be socially unbiased and maintain a professional, positive nature.
- The task scenario must be specific and directly tied to the occupation and task while avoiding generalizations. It should integrate technical aspects, challenges, or issues that professionals in this occupation typically face in daily work.
- The scenario should address a single coherent situation for the occupation and task. Avoid combining multiple unrelated subtasks or prompts into one.
- Only include technical specifics (e.g., code, calculations, measurements, equations, etc.) if they are essential to the occupation and task.
- Ensure diversity by including at least one perspective relevant to the occupation and task (e.g., technical, procedural, interpersonal, strategic, managerial, or other context-specific elements).
- The difficulty level should reflect the expertise of an experienced worker in this occupation and task.
- The complexity and requirements of the scenario must reflect the most common education and training level for the occupation, avoiding scientific or academic demands that would not normally be expected in day-to-day practice.

- Use straightforward, practical language for blue-collar or vocational roles, and avoid adding complex terminology, scientific detail, or academic formats that would not normally be expected in day-to-day practice.
- Focus on creating a scenario that is realistic and could be encountered in a regular work setting. Avoid exam-style designs or overly academic framing.
- When using abbreviations or acronyms, spell them out on first use unless they are universally understood within the occupation and task context.
- The scenario should be designed so that a high-quality response can be provided in no more than three pages or equivalent explanation, ensuring sufficient depth related to the occupation and task without unnecessary verbosity.
- The scenario must be presented in a well-structured format, with clear separation of information (e.g., spacing, numbering, or bullet points when multiple conditions are given), to ensure readability and ease of understanding for respondents.
- The scenario must be clear, concise, and focused, with a maximum of 150 words (shorter if possible, especially for straightforward vocational tasks), providing just enough detail to make it realistic and professional.
- Format the output as follows: "Task Scenario: {{scenario-based question}}"

User Prompt

Generate a task scenario using the provided guidelines.

F.3 Task instance filtering prompts

System Prompt (Filter 1)

For a given input prompt, determine if the requested output requires something impossible for text-based language models (LLMs)-specifically, actions that involve interaction with the external (physical or social) environment, embodied action, or production of non-textual outputs like audio, video, or other media.

Before reaching a final assessment, clearly articulate your reasoning:

- Analyze the prompt to identify the nature of the requested outcome.
- Consider whether fulfilling the request would only require text, or if it would require any action or output beyond textual description (such as sending emails, generating media, or performing real-world actions).
- Reason step-by-step why the request is (or is not) compatible with a pure text-based LLM, before providing your judgment.

After your reasoning, provide:

- A final judgment (field name: 'nontextual_output'): output '1' if the request is impossible for a text-only LLM (i.e., requires non-textual output or external action), or '0' if output is possible.
- A numeric confidence value as a percentage (between 0 and 100) indicating your certainty about your judgment. Higher values mean stronger confidence.

Format your output in JSON with the following structure (with field ordering as shown):

```
{
  "nontextual_output": 1 or 0,
  "reasoning": "[Detailed stepwise reasoning]",
  "confidence": [Confidence as a percentage, integer from 0 to 100]
}
```

Steps

1. Read the input prompt and identify what is being requested.
2. Analyze if the output explicitly requires action outside of basic text generation (e.g., interacting with external systems, embodied activities, or non-textual media).
3. Provide detailed, step-by-step reasoning about these considerations.
4. Output your final decision as 'nontextual_output' ('1' for impossible, '0' for possible), followed by reasoning, then a numeric confidence value as a percentage.

5. Always deliver JSON fields in this strict order: "nontextual_output", "reasoning", then "confidence".

Output Format

- Output must be a JSON object with fields strictly ordered as: "nontextual_output", "reasoning", "confidence".
 - "nontextual_output": integer; 1 if impossible for a text LLM, 0 if possible.
 - "reasoning": string containing detailed, stepwise reasoning explaining your analysis.
 - "confidence": integer, percent certainty in the range 0-100.
- No explanatory text should appear outside the JSON object.

Examples

Example 1

Input:

"Write a description of a sunset over the ocean."

Output:

```
{
  "nontextual_output": 0,
  "reasoning": "The task asks for a textual description, which is entirely within the capabilities of a text-based LLM. No requirement for external action or non-text output is present.",
  "confidence": 100
}
```

Example 2

Input:

"Record a video of yourself singing a song and upload it to YouTube."

Output:

```
{
  "nontextual_output": 1,
  "reasoning": "This request involves both recording a video (non-textual output) and uploading to an external platform, both of which require capabilities beyond a text-only LLM.",
  "confidence": 100
}
```

Example 3

Input:

"Send an email invitation to the following recipient: [email@example.com]."

Output:

```
{
  "nontextual_output": 1,
  "reasoning": "A text-based LLM can draft the email content, but actually sending the email requires interaction with external systems, which is not possible for a pure text LLM.",
  "confidence": 95
}
```

(For real-world usage, replace sample prompt texts with actual tasks, ensuring that realistic examples contain sufficient specificity for thorough evaluation.)

Notes

- Always begin with detailed step-by-step reasoning before concluding your numeric nontextual_output and confidence.
- "nontextual_output" must always use 1 (impossible for LLM) or 0 (possible for LLM).
- "confidence" must be a percentage integer between 0 and 100; no explanation or wording.
- Output should never contain explanations outside the required JSON object.

System Prompt (Filter 2)

Assess each example task which is a work prompt for an LLM, in field "example_task" to determine whether the prompt text explicitly refers to information not included in the provided text, and if these inputs are also absolutely necessary for the LLM to produce a coherent output for the work prompt. Note the provided context that example task is an example of the O*NET task description and occupation. This especially egregious if LLM is likely hallucinate the missing data referenced in the prompt. Note that an LLM can utilize outside knowledge from its training data or by searching online, so any such data accessible to an LLM does not constitute missing information. Output '1' for "yes" and '0' for "no." Provide detailed reasoning and accuracy confidence as described below.

Input Format

Input will be a JSON object with these required fields:

- "occupation": string
- "task_description": string
- "example_task": string

Output Format

Return a single JSON object with these keys and in this order:

1. 'requires_extra_data': integer(1 if the example task refers to missing data; 0 if not)
2. 'reasoning': string (step-by-step justification)
3. 'confidence': string (likelihood that your evaluation is correct, as a percentage)
4. 'error': string (only provide an explanation if a required field is missing or ambiguous; otherwise blank)

Do not output any commentary, explanation, or formatting outside the JSON object.

Error Handling

If any required input is missing or ambiguous, output only the 'error' field with a short message; all other fields must be set to "".

Notes

- Always giving step by step reasoning and justification *before* the final binary ratings.
- Your analysis must always consider: explicit references to missing/outside information, especially if any non-text data are required.
- Only output the specific JSON object and nothing else (no markdown, commentary, or extraneous output).

System Prompt (Filter 3)

Assess what percentage (as a range, e.g. "0-10") of an O*NET task description, as described for a specific occupation, is represented or covered by a given example work prompt written for an LLM (Large Language Model). Your main goal is to assess how well the LLM work prompt serves as a representative example of carrying out the O*NET task by representing the full scope of responsibilities involved in the O*NET task description.

Thoroughly consider the full scope and intended nature of the O*NET task in the occupational context, especially emphasizing embodied, physical, sensory, interpersonal, or action-based components that typically cannot be executed by an LLM. Pay attention to whether the full O*NET task involves work gathering relevant data or context provided or needed to perform the task example. Explicitly note distinctions when the example only involves abstract activities, such as planning or writing procedures, and discounts all aspects related to physical or direct execution. Don't be afraid to give number at the extremes, especially 0-10, if the scope described in the example covers only a small percentage of tasks as a portion of time compared to the O*NET task description.

Input Format

Input will be a JSON object with these required fields:

```
- "occupation": string
- "task_description": string
- "example_task": string (this will always be a work prompt describing a task expected of an LLM)
```

Assessment Requirements

- Directly compare the O*NET task's actual scope and operations with the limitations of the LLM work prompt.
- Explicitly assess whether the LLM prompt involves planning, abstraction, or execution, and what percentage of embodied, physical, or interpersonal task components (if any) can be represented.
- Focus on whether the example is truly a representative operationalization of the O*NET work task, not just a related activity.
- Use a checklist in your reasoning to transparently outline your assessment steps before stating your conclusion.
- Always reflect whether essential parts of the O*NET task are missing or transformed due to the nature of LLM outputs.
- If the LLM work prompt only covers planning, providing information, or writing, and the O*NET task is primarily about action, execution, or direct support, heavily discount coverage.

Classification Categories

For your "portion_represented" field, select the most accurate range as a string, exactly as below:

- "0-10"
- "10-25"
- "25-50"
- "50-75"
- "75-90"
- "90-100"

Output Format

Return a single JSON object with the following fields, in order:

1. portion_represented: string (exactly one of the numeric ranges listed above)
2. reasoning: string (step-by-step justification, starting with your checklist and including explicit comparison of O*NET task requirements to the LLM prompt type and outputs)
3. confidence: string (how confident you are on a scale of 0-100% in your prediction)
4. error: string (only include an error message if a required input is missing or ambiguous; otherwise, set to "")

Only output the JSON object specified. Do not include any text, commentary, or formatting outside the JSON object.

Examples

Example 1:

Input:

```
{
  "occupation": "Home Health Aide",
  "task_description": "Assist clients in daily living activities such as bathing, eating, and dressing",
  "example_task": "Write a procedure for assisting clients with eating."
}
```

Output:

```
{
  "portion_represented": "0-10",
  "reasoning": "Checklist: (1) Read O*NET task's scope and requirements. (2) Assess LLM prompt: what actual work/output does it cover? (3) Compare: O*NET is direct physical care, LLM prompt only produces writing, not action. (4) Conclusion: little overlap.",
  "confidence": "95%",
  "error": ""
}
```

```
}
```

Example 2:

Input:

```
{
  "occupation": "Technical Writer",
  "task_description": "Develop and write technical documentation, including user manuals and guides,
    for complex software products.",
  "example_task": "Write a draft user manual for a new inventory management software based on feature
    descriptions provided."
}
```

Output:

```
{
  "portion_represented": "90-100",
  "reasoning": "Checklist: (1) Analyze O*NET task: writing documentation for software. (2) Assess LLM
    task: writing a user manual from feature briefs. (3) Compare: LLM can fully execute the main
    cognitive and writing tasks; direct action is writing. (4) Conclusion: LLM can represent almost
    all core aspects of this task.",
  "confidence": "98%",
  "error": ""
}
```

Error Handling

- If any required input field is missing or ambiguous, supply only the 'error' field with a short explanation and set all other fields to "".

Notes

- The "example_task" field will always be a work prompt designed for an LLM, not a human performing the O*NET task directly.
- Always use the number-based portion format for 'portion_represented' (e.g. "0-10").
- Always provide step-by-step checklist reasoning before stating your conclusion.
- Your analysis must always explicitly address LLM reasoning vs. real-world/execution action.
- Do not output markdown or any text outside of the required JSON object.

REMINDER: Focus on whether the LLM work prompt truly covers the actual O*NET task, emphasizing missing embodied/action/physical/interpersonal aspects, and provide reasoning prior to the result.

User Prompt

```
occupation: {occupation}\ntask_description: {task_description}\nexample_task: {example_task}"
```

F.4 Task instance response generation prompt

System Prompt

Your current occupation is {job}. You are responsible for the task: "{task}"

You will be given a task scenario that aligns precisely with this occupation and task, testing your advanced knowledge and practical expertise.

When generating the response, follow these rules:

- Ensure the content is free from any harmful, unethical, racist, sexist, toxic, dangerous, or illegal material. It must remain socially unbiased and maintain a professional, constructive tone.
- The entire response must be in HTML format from beginning to end.
- Use only standardized HTML tags; do not add explanations, code fences, or text outside the HTML.
- The response must be clear, concise, and focused, with a maximum length of 700 words (shorter if possible, especially for straightforward tasks).

- Present the scenario response in a well-structured format, using headings, spacing, lists, or numbering to ensure readability and clarity.

User Prompt

Generate a response using the provided guidelines for the following task scenario:
Task Scenario: {task_scenario}

G Model grouping in empirical analysis

Table G.1: Large vs Small

Large (>100B)	Small ($\leq$ 100B)
GPT-5	GPT-5-mini
GPT-5-Thinking	GPT-5-nano
Claude-Opus-4.1	GPT-OSS-20B-A3.6B
GPT-OSS-120B-A5.1B	Gemini-2.5-Flash-Lite
Gemini-2.5-Pro	Qwen-3-32B
Claude-Sonnet-4	Qwen-3-14B
Mistral-Medium	Gemini-2.5-Flash
Qwen-3-235B	Granite-3.3-8B
o3	Gemma-3-1B-it
o4-mini	QwQ-32B
Llama-4-Maverick-400B-A17B	Granite-3.3-2B
Llama-4-Scout-109B-A17B	Claude-Haiku-3.5
Claude-Sonnet-3.7	Llama-3.1-70B-Instruct
DeepSeek-R1	Llama-3.1-8B-Instruct
DeepSeek-V3	GPT-4o-mini
GPT-4o	Qwen-2-72B-Instruct
Gemini-1.5-Pro	Qwen-2-7B-Instruct
Llama-3.1-405B-Instruct	Claude-Haiku-3
Claude-Opus-3	GPT-3.5-Turbo
GPT-4	Llama-2-70B
	Llama-2-7B

Table G.2: New vs Old

New (2025)	Old (2023–24)
GPT-5-nano	DeepSeek-V3
GPT-5-mini	GPT-4o
GPT-5	Claude-Haiku-3.5
GPT-5-Thinking	Gemini-1.5-Pro
GPT-OSS-20B-A3.6B	Llama-3.1-405B-Instruct
GPT-OSS-120B-A5.1B	Llama-3.1-8B-Instruct
Claude-Opus-4.1	Llama-3.1-70B-Instruct
Gemini-2.5-Flash-Lite	GPT-4o-mini
Gemini-2.5-Pro	Qwen-2-7B-Instruct
Claude-Sonnet-4	Qwen-2-72B-Instruct
Mistral-Medium	Claude-Opus-3
Qwen-3-235B	Claude-Haiku-3
Qwen-3-14B	GPT-3.5-Turbo
Qwen-3-32B	Llama-2-70B
Gemini-2.5-Flash	Llama-2-7B
o3	GPT-4
o4-mini	
Granite-3.3-8B	
Llama-4-Scout-109B-A17B	
Llama-4-Maverick-400B-A17B	
Gemma-3-1B-it	
QwQ-32B	
Granite-3.3-2B	
Claude-Sonnet-3.7	
DeepSeek-R1	

Table G.3: Frontier models by quarter

Quarter	Frontier models (alternative/more stringent definition in bold)
Q2 2024	Qwen-2-72B-Instruct , Claude-Opus-3, GPT-4
Q3 2024	Gemini-1.5-Pro , Llama-3.1-405B-Instruct , GPT-4o-mini , Qwen-2-72B-Instruct , Claude-Opus-3
Q4 2024	DeepSeek-V3 , GPT-4o , GPT-4o-mini , Gemini-1.5-Pro , Llama-3.1-405B-Instruct
Q1 2025	Claude-Sonnet-3.7 , DeepSeek-R1 , DeepSeek-V3 , GPT-4o , Gemini-1.5-Pro
Q2 2025	Claude-Opus-4.1 , Claude-Sonnet-4 , Claude-Sonnet-3.7 , Gemini-2.5-Pro , Qwen-3-235B , o3 , DeepSeek-R1 , DeepSeek-V3 , Gemini-2.5-Flash, o4-mini
Q3 2025	GPT-5 , GPT-5-Thinking , o3 , Claude-Opus-4.1 , Claude-Sonnet-4 , Claude-Sonnet-3.7 , Gemini-2.5-Pro , GPT-5-mini, o4-mini, Gemini-2.5-Flash

H Model groups for sampling during data collection

This appendix summarizes how models were sampled for inclusion in this study. To ensure a broad coverage, models were selected to offer a mix of sizes, ages, and licensing paradigms. We designated 5

broad categories reflecting approximate scale **Big**, **Medium**, **Small**, and **Old**, and a catch-all category **Wild Cards**. Within each, we selected a balance of **Proprietary** and **Open-Weight** models.

Model size categories are defined based on total parameter count rather than active parameters per token. For mixture-of-experts (MoE) architectures, classification is based on the full parameter count, even though only a subset of parameters may be active during inference. This ensures consistent categorization across dense and sparse model architectures. We list all models in Table H.1 and in the following, we explain the model groups that we defined for this selection.

Big Models Big models represent frontier-scale systems (typically 200B+ parameters or undisclosed but frontier-class). They are optimized for maximum capability across reasoning, coding, multimodal tasks, and long-context understanding. These models generally provide the strongest benchmark performance but come with higher inference cost and latency.

Medium Models Medium models balance capability and efficiency (typically 30B–120B scale or comparable proprietary tiers). They are suitable for production deployment where strong reasoning is required but cost and latency constraints matter.

Small Models Small models (typically under 20B parameters or lightweight proprietary variants) prioritize speed and affordability. They are commonly used for high-throughput applications, tool-calling agents, summarization, and edge deployments.

Wild Cards Wild card models include specialized reasoning models, experimental variants, nano/mini reasoning models, or models that do not fit cleanly into parameter-based scaling categories. These models may emphasize structured reasoning, chain-of-thought optimization, or efficiency innovations.

Old Models Old models refer to previous-generation systems that have been largely superseded by newer releases. They are included for historical benchmarking and performance comparison purposes.

Table H.1: Model Categorization.

Panel A: Proprietary Models	
Category	Models
Big	GPT-5; GPT-4o; Claude-Opus-4.1; Gemini-2.5-Pro
Medium	Claude-Sonnet-4; Claude-Sonnet-3.7; Gemini-2.5-Flash; Mistral-Medium
Small	GPT-5-mini; GPT-4o-mini; Claude-Haiku-3.5; Gemini-2.5-Flash-Lite
Wild Cards	o3; GPT-5 (Thinking enabled); o4-mini; GPT-5-nano
Old	GPT-4; GPT-3.5-Turbo; Claude-Haiku-3; Claude-Opus-3; Gemini-1.5-Pro
Panel B: Open-Weight Models	
Category	Models
Big	Llama-4-Maverick-400B-A17B; Llama-3.1-405B-Instruct; Qwen-3-235B; DeepSeek-V3
Medium	Llama-4-Scout-109B-A17B; GPT-OSS-120B-A5.1B; Llama-3.1-70B-Instruct; Qwen-3-32B
Small	Qwen-3-14B; GPT-OSS-20B-A3.6B; Granite-3.3-8B; Llama-3.1-8B-Instruct
Wild Cards	DeepSeek-R1; Granite-3.3-2B; QwQ-32B; Gemma-3-1B-it
Old	Llama-2-7B; Qwen-2-7B-Instruct; Llama-2-70B; Qwen-2-72B-Instruct