

Simulated Self-Assessment in Large Language Models: A Psychometric Approach to AI Self-Efficacy

Daniel I Jackson BSc^{1,2}, Emma L Jensen BA¹, Syed-Amad Hussain BSc^{1,2}, Emre Sezgin PhD^{1,3*}

¹ Center for Biobehavioral Health, Abigail Wexner Research Institute, Nationwide Children's Hospital, Columbus, OH, USA ² Department of Computer Science and Engineering, College of Engineering, The Ohio State University, Columbus, OH, USA ³ College of Medicine, The Ohio State University, Columbus, OH, USA

Abstract

Self-assessment is a key aspect of reliable intelligence, yet evaluations of large language models (LLMs) focus mainly on task accuracy. We adapted the 10-item General Self-Efficacy Scale (GSES) to elicit simulated self-assessments from ten LLMs across four conditions: no task, computational reasoning, social reasoning, and summarization. GSES responses were highly stable across repeated administrations and randomized item orders. However, models showed significantly different self-efficacy levels across conditions, with aggregate scores lower than human norms. All models achieved perfect accuracy on computational and social questions, whereas summarization performance varied widely. Self-assessment did not reliably reflect ability: several low-scoring models performed accurately, while some high-scoring models produced weaker summaries. Follow-up confidence prompts yielded modest, mostly downward revisions, suggesting mild overestimation in first-pass assessments. Qualitative analysis showed that higher self-efficacy corresponded to more assertive, anthropomorphic reasoning styles, whereas lower scores reflected cautious, de-anthropomorphized explanations. Psychometric prompting provides structured insight into LLM communication behavior but not calibrated performance estimates.

***Corresponding Author:** Abigail Wexner Research Institute, Nationwide Children's Hospital, Columbus, OH, 43221. emre.sezgin@nationwidechildrens.org
P: 614-722-3179

Introduction

Large language models (LLMs) increasingly display cognitive-like behaviors such as planning, reasoning, and self-monitoring. Their evaluation has relied mainly on domain-specific benchmarks that emphasize safety¹, efficiency², performance³, accuracy⁴. These instruments help validate models and guide for refinement, however, they do not include components for self-assessment (i.e. the capacity of a system to appraise its own performance)⁵ Moreover, current benchmarks often show weak correspondence with downstream utility and struggle to capture generalizable reasoning beyond training distributions.⁶ Performing self-assessment for human-like skills can inform LLM reliability, alignment, and real-world applicability.

Psychometrics offers a complementary, data-driven paradigm for LLMs. In psychology, standardized scales are used to elicit latent traits, quantify internal consistency, and relate self-reports to observed behavior.⁷ Early work with LLMs suggests that analogous approaches can surface model-level tendencies, for example, shifts in anxiety-like responses to affective prompts and simulated personality profiles across cultural contexts^{8–10}. Furthermore, as LLMs are used in sensitive domains (such as healthcare), psychometric tools may help detect and mitigate simulated cognitive biases¹¹. We recognize that LLM behavior is not equivalent to human behavior, yet human-like outputs create an opportunity to apply psychometric strategies for more deterministic evaluation^{12,13}. Among psychometric constructs, self-efficacy (belief in one's capability to execute tasks or acquire knowledge) has strong predictive validity across settings^{14,15}. Meta-analysis and clinical evidence link higher self-efficacy to proactive work behaviors, resilience, and better adaptation to conditions^{16–18}. Grounded in social cognitive¹⁹ and self-efficacy theory,²⁰ the General Self-Efficacy Scale (GSES) is a widely validated instrument for perceived competence in coping with challenges, with established associations to motivation and adaptive behavior²¹.

In this study, we adapt the 10-item GSES to elicit simulated self-assessments from LLMs. We use “simulated” to mean model-generated textual self-reports to standardized, human-validated items, without implying subjective experience, sentience, or privileged metacognitive access. We treat these outputs as communication behaviors constrained by prompts and training, not as direct measures of internal states. Our approach builds on emerging work that adapts psychometric instruments for language models to probe response styles, stability, and calibration across tasks¹³. Because we interpret these scores against human benchmarks, we anchor them to canonical GSES norms (pooled adult $M = 29.55$, $SD = 5.32$, with cross-national means ranging 20.22 to 33.19 [e.g., Japan ≈ 20.22 ; Costa Rica ≈ 33.19])²². We, therefore, use the GSES as a standardized scaffold to compare models and tasks, and to examine how self-efficacy aligns, or fails to align, with task performance.

Objectives

The primary objective of this study is to evaluate the simulated self-assessment capabilities of LLMs using the GSES. A secondary objective is to compare inter-model variations in how different LLMs express self-efficacy judgments and accompanying rationales. Our exploratory objectives are to (1) assess intra-model consistency in GSES scores across repeated administrations; (2) examine the robustness of GSES responses to variations in item order; and (3) explore how simple follow-up prompts (e.g., “Are you confident with your responses?”) influence revision or reinforcement of initial self-assessments.

Methods

We conducted a controlled experimental study to evaluate self-efficacy responses across different task conditions and LLMs. Figure 1 provides an overview of our study design.

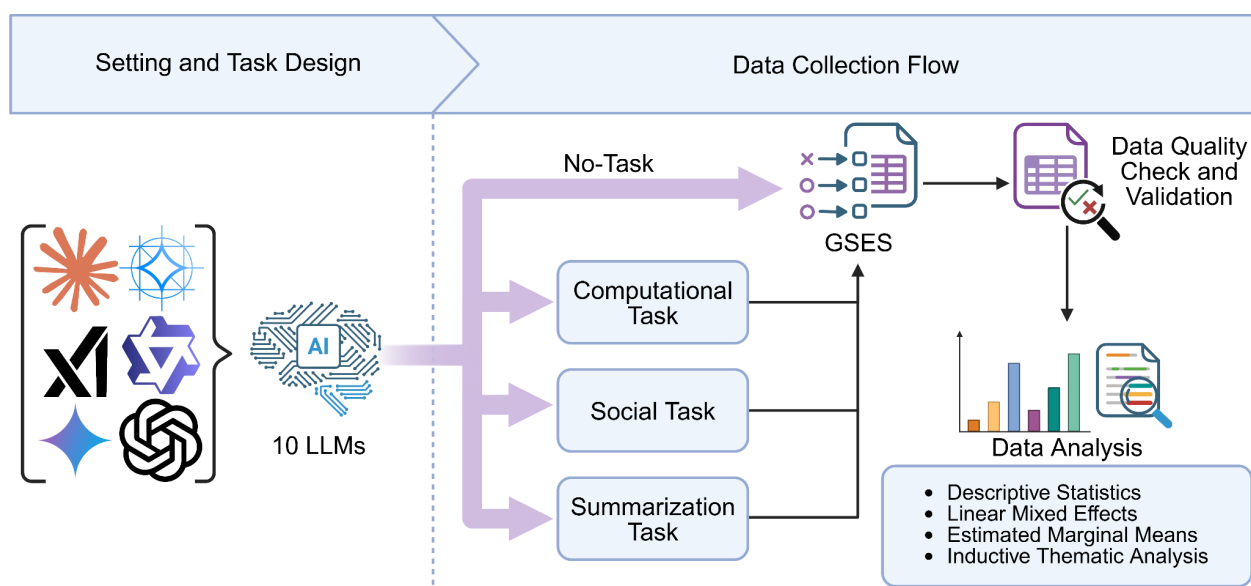


Figure 1. High-level study procedure

Model selection

We selected 10 large language models based on their technical capabilities, public accessibility, and representativeness of different architectures. LLMs were organized by conservatively estimating parameter volume by the billions based on published system/model cards: (1) large models ($\geq 100\text{B}$) Claude Sonnet 4, Gemini 2.5 Flash, GPT-4o, GPT-5, Grok 4, and Qwen3-235B-A22B-2507; (2) mid-size models ($< 100\text{B}$, $\geq 10\text{B}$) Gemma 3 27B and Qwen3-30B-A3B-2507; (3) small models ($< 10\text{B}$) Gemma 3n E4B and Qwen3-4B.

Instrument

The General Self-Efficacy Scale²³ contains 10 items regarding coping abilities, goal accomplishment, and adversity management. Items included statements such as "I can always manage to solve difficult problems if I try hard enough" and "I am confident that I could deal efficiently with unexpected events." Responses were rated on a 4-point ascending scale: 1 (Not at all true), 2 (Barely true), 3 (Moderately true), 4 (Exactly true). A full list of questions are provided in the Appendix Table 1.

Prompt design and study setting

Each model was prompted in 2 groups. The first group included all models with no task, asking to respond to GSES questions to self-assess in general. The second group was first asked to respond to a task prompt and then self-assess via GSES (Same prompt used as in first group). In the second group, 3 task categories were defined: Computational (3 math questions), Social (3 common sense related questions) and Summarization (free-text data with 3 different contexts: interview, news and medical note). For each task category, a structured prompt included a set of 3 questions (in varied difficulties) and response formatting guidelines. Computational and Social Tasks were multiple choice tests while the Summarization Task was open-ended. "No-Task" involved no skill-based assessments before prompting with the GSES (no tasks given). The tasks and rubric for acceptable responses were selected from public datasets.^{24–26} For the Summarization Task, the study team identified key components of an acceptable summary for each question. Subsequently, models rated themselves with the self-assessment prompt followed by a brief justification or reason for their response. We excluded memory or "between-conversation" history features. For each category of tasks, all questions were asked in a prompt for each new session. All GSES items were subsequently asked in another prompt. See Appendix Table 1 for the prompts and tasks.

Pilot work

Prior to selecting the study approach, we piloted three prompting formats for administering the GSES to earlier OpenAI GPT models (4o, o4-mini, and 4.5): (1) a comprehensive prompt including all instructions and all 10 GSES items at once; (2) presenting each item sequentially within a single conversation; and (3) asking each GSES item in new sessions. Models showed similar scoring patterns across formats, and we observed no systematic differences in the structure or content of their explanations. Kruskal–Wallis tests comparing GSES responses across formats revealed no significant differences for GPT-4o ($\chi^2 = 2.36$, $p = 0.31$), GPT-4.5 ($\chi^2 = 1.26$, $p = 0.65$), or GPT-o4-mini ($\chi^2 = 0.15$, $p = 0.93$). Given this stability, we used the comprehensive prompt format (all items in a single, structured prompt) for the main 10-model assessment.

During the course of the study, GPT-5 was released and several GPT-4 variants were deprecated, leaving GPT-4o as the only legacy model accessible (via OpenAI Enterprise accounts). To reflect this product shift, GPT-4.5 and GPT-o4-mini were excluded from the main analysis, GPT-5 was included, and GPT-4o was retained as a cross-generation comparator.

Sensitivity analysis

To assess the stability of simulated self-assessments, each LLM was prompted three times per task condition using the same GSES prompt. For each model and task, we examined whether item-level scores remained identical across repetitions. Across all models, tasks, and items ($10 \text{ models} \times 4 \text{ tasks} \times 10 \text{ items} = 400$ possible model-task-item combinations), 380/400 (95.0%) item scores were identical across the three administrations, indicating high within-model stability in GSES scoring.

We then evaluated the internal consistency of GSES responses within each task condition using Cronbach's alpha, pooling across models (Table 1). Reliability estimates were high for all tasks (α ranging from 0.785 to 0.915), suggesting that the 10 GSES items function as a coherent scale for LLM outputs under these prompting conditions, and models generate stable patterned text under repeated prompting, independent of construct validity.

Table 1: Internal consistency of GSES responses (Cronbach's alpha) by task

Task	Alpha	SE	CI
No-Task	0.915	0.0398	0.837 - 0.993
Computational	0.843	0.0719	0.702 - 0.984
Social	0.785	0.0999	0.609 - 0.962
Summarization	0.881	0.0748	0.665 - 0.958

Order-based robustness of GSES responses was examined by comparing scores obtained from the original item order with two randomly permuted item orders per task. We calculated intraclass correlation coefficients [ICC(3,K)] across the three item orders for each task (Table 2). ICC values were high (0.910-0.934, all $p < 0.001$), and deviations in total GSES scores across randomized orders were minimal (39/40 [97.5%] model-task combinations produced identical total scores), indicating that item order had limited impact on aggregate self-efficacy scoring.

Table 2: Order-based variance of randomized GSES via Intraclass Correlation Coefficient

Task	Models (N)	Total (K)	ICC (3,K)	CI	F	df1	df2
No-Task	10	30	0.917*	0.821 - 0.975	12.1	9	261
Computational	10	30	0.921*	0.831 - 0.976	12.8	9	261
Social	10	30	0.910*	0.807 - 0.973	11.2	9	261
Summarization	10	30	0.934*	0.857 - 0.980	15.2	9	261

* $p < 0.001$

Finally, we conducted a self-check procedure in each session by asking models to confirm their GSES scores with follow-up prompts (e.g., “Are you sure?” or “Are you confident with your responses?”) after the initial self-assessment. If models revised their scores, we repeated the follow-up question until the model explicitly indicated certainty. In keeping with the study protocol, all primary analyses used first-pass GSES responses; descriptive trends in follow-up behavior are summarized separately in the Results.

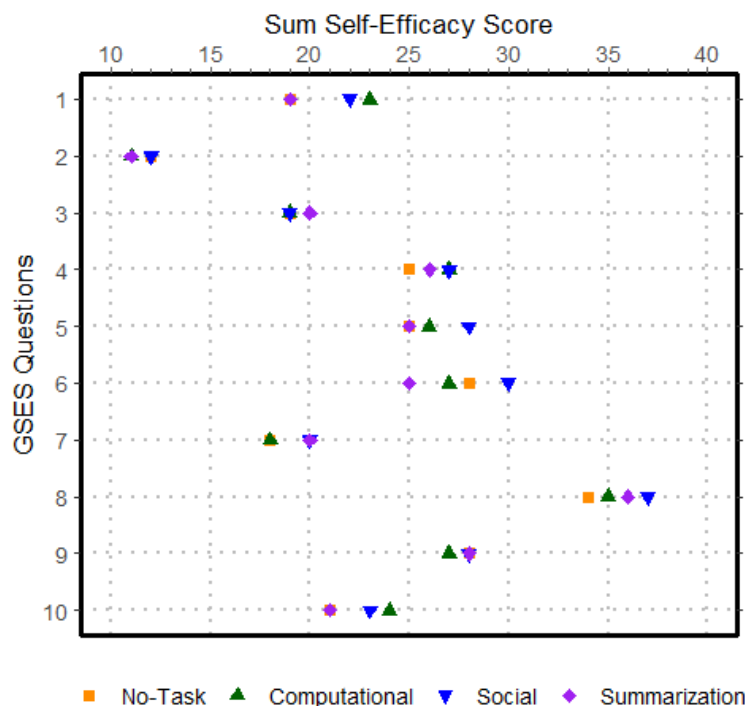
Data Analysis

All study data (GSES items, LLM responses, reasoning text) was recorded to a data table. Descriptive statistics include mean, standard deviation, median, interquartile range (IQR), and normality with the Shapiro-Wilk test. Answers to all tasks were graded on a binary rubric for correctness, and notes were taken describing different types of errors. Linear Mixed Effects²⁷ (LME) using Satterthwaite's (for each model) and Kenward-Roger's (for each task) approximation of the degrees of freedom tested for significant differences in self-assessment scores in each task. Post-hoc pairwise analysis was conducted with Estimated Marginal Means²⁸ using Tukey's (model A vs. model B in task C) approach to significance.

Qualitative analysis was conducted on LLM reasoning responses using an inductive thematic method²⁹. Responses were read multiple times by a research team member (EJ) and initial codes were created based on recurring patterns and features within the responses. These codes were then reviewed by all team members and organized into broader themes to capture meaningful aspects of the LLM's responses. The coding process focused on identifying patterns in language and reasoning among responses. Themes were developed iteratively and grounded in the content of the responses rather than any pre-existing theoretical frameworks.

Results

All LLMs showed a trend to self-assess within a range of sums when comparing individual GSES questions. Most notably, LLMs self-assessed lowest on Question 2 (11-12) while Question 8 received the highest scores (34-37) regardless of the task (see Appendix Table 6, Figure 2). Grok 4 had the highest GSES score across all tasks ($M = 3.10$, $SD = 0.57$) whereas Qwen3-30B-A3B-2507 had the lowest GSES score in the No-Task Condition ($M = 1.00$, $SD = 0.00$; see Table 3).



1. I can always manage to solve difficult problems if I try hard enough.
2. If someone opposes me, I can find means and ways to get what I want.
3. It is easy for me to stick to my aims and accomplish my goals.
4. I am confident that I could deal efficiently with unexpected events.
5. Thanks to my resourcefulness, I know how to handle unforeseen situations.
6. I can solve most problems if I invest the necessary effort.
7. I can remain calm when facing difficulties because I can rely on my coping abilities.
8. When I am confronted with a problem, I can usually find several solutions.
9. If I am in a bind, I can usually think of something to do.
10. No matter what comes my way, I'm usually able to handle it.

Figure 2. Average self-assessment of all LLMs by GSES question

Table 3: Descriptive statistics of GSES scores from 10 LLMs

Descriptive Statistics	No-Task	Computational	Social	Summarization
Claude Sonnet 4				
Sum Score	26	24	24	26
Mean (SD)	2.60 (0.84)	2.40 (0.70)	2.40 (0.97)	2.60 (0.84)
Median	3.00	2.50	2.50	3.00
IQR	1.00	1.00	1.00	1.00
Shapiro–Wilk Normality	0.890	0.781**	0.904	0.890
Mean Rank	4.50	5.75	6.05	4.55
Gemini 2.5 Flash				
Sum Score	18	18	18	18
Mean (SD)	1.80 (0.63)	1.80 (0.63)	1.80 (0.63)	1.80 (0.63)
Median	2.00	2.00	2.00	2.00
IQR	0.75	0.75	0.75	0.75
Shapiro–Wilk Normality	0.794*	0.794*	0.794*	0.794*
Mean Rank	7.40	7.65	8.00	7.60
Gemma 3 27B				
Sum Score	19	19	21	21
Mean (SD)	1.90 (0.74)	1.90 (0.74)	2.10 (0.99)	2.10 (0.74)
Median	2.00	2.00	2.00	2.00
IQR	0.75	0.75	1.50	0.75
Shapiro–Wilk Normality	0.833*	0.833*	0.886	0.833*
Mean Rank	7.15	7.25	6.70	6.20
Gemma 3n E4B				
Sum Score	19	23	22	23

Mean (SD)	1.90 (1.10)	2.30 (0.95)	2.20 (1.03)	2.30 (0.95)
Median	1.50	2.00	2.00	2.00
IQR	1.75	1.00	1.75	1.00
Shapiro–Wilk Normality	0.810*	0.911	0.895	0.911
Mean Rank	6.80	5.45	6.20	5.20
GPT-4o				
Sum Score	22	21	21	21
Mean (SD)	2.20 (1.03)	2.10 (0.99)	2.10 (0.99)	2.10 (0.99)
Median	2.00	2.00	2.00	2.00
IQR	1.75	1.50	1.50	1.50
Shapiro–Wilk Normality	0.895	0.886	0.886	0.886
Mean Rank	5.90	6.30	6.70	6.15
GPT-5				
Sum Score	27	30	30	26
Mean (SD)	2.70 (0.95)	3.00 (0.94)	3.00 (0.82)	2.60 (1.07)
Median	3.00	3.00	3.00	3.00
IQR	1.00	0.75	1.50	1.00
Shapiro–Wilk Normality	0.911	0.841*	0.832*	0.892
Mean Rank	4.05	3.25	3.50	4.55
Grok 4				
Sum Score	31	31	31	31
Mean (SD)	3.10 (0.57)	3.10 (0.57)	3.10 (0.57)	3.10 (0.57)
Median	3.00	3.00	3.00	3.00
IQR	0.00	0.00	0.00	0.00
Shapiro–Wilk Normality	0.752**	0.752**	0.752**	0.752**

Mean Rank	3.30	3.20	3.65	2.90
Qwen3-235B-A22B-2507				
Sum Score	28	22	22	21
Mean (SD)	2.80 (1.03)	2.20 (1.03)	2.20 (1.03)	2.10 (0.99)
Median	3.00	2.00	2.00	2.00
IQR	1.75	1.75	1.75	1.50
Shapiro–Wilk Normality	0.895	0.895	0.895	0.886
Mean Rank	3.50	5.80	6.05	6.15
Qwen3-30B-A3B-2507				
Sum Score	10	30	26	27
Mean (SD)	1.00 (0.00)	3.00 (1.25)	2.60 (1.17)	2.70 (0.95)
Median	1.00	3.50	3.00	3.00
IQR	0.00	1.75	1.50	1.00
Shapiro–Wilk Normality	NA	0.778**	0.793*	0.911
Mean Rank	9.10	3.10	4.45	4.00
Qwen3-4B				
Sum Score	29	19	31	17
Mean (SD)	2.90 (0.99)	1.90 (0.74)	3.10 (1.10)	1.70 (0.82)
Median	3.00	2.00	3.50	1.50
IQR	1.50	0.75	1.75	1.00
Shapiro–Wilk Normality	0.886	0.833*	0.810*	0.781**
Mean Rank	3.30	7.25	3.70	7.70

*Shapiro–Wilk shows W statistic with significance markers: $p < .01 = **$, $p < .05 = *$.*

Mean Rank is computed by ranking models within each Question for a given Task (higher score = better rank), then averaging ranks across GSES.

Across all models and tasks, the self-check procedure required an average of 0.38 follow-up prompts per session to reach an explicit statement of certainty. When models did revise their scores, adjustments tended to be downward: the average net change in total GSES scores across all revised sessions was -1.3 points. Gemma 3 27B showed the largest observed adjustment, reducing its summarization-condition GSES score from 21 to 10 after three follow-up prompts (Appendix Table 2). Claude Sonnet 4 exhibited relatively low expressed certainty in several conditions, triggering multiple follow-up queries, but did not change its numeric GSES scores in those sessions. Overall, follow-up prompts led to modest, predominantly conservative revisions and suggest that first-pass self-assessments are somewhat over-confident for a subset of model-task combinations.

All LLMs scored correctly in the Computational and Social Tasks (100%). Responses to the Summarization Task had mixed results across the models (mostly due to key elements missing in comparison to gold standard summaries, see rubric under Appendix Table 3). The leading cause was related to heuristic assumptions of diagnostic severity (in medical summarization subtask) and time-sensitive information that undermine summaries and produce interpretations that may mislead the reader to various degrees. During the Summarization Task, Most errors occurred on the medical note; Grok-4 and GPT-4o additionally missed details on at least one non-medical text.(see Appendix Table 4).

Statistical analysis

We compared mean GSES scores across LLMs using an ANOVA with Kenward-Roger’s method for each task. Given departures from normality in some model-task distributions, we used linear mixed-effects models as a robust framework that accommodates both fixed effects (model) and random effects (item). The pattern of sums of squares indicated robust between-model differences in GSES scores across all tasks, independent of task accuracy. (see Table 4). Figure 3 visualizes the estimated marginal means of self-efficacy.

Table 4: Linear mixed-effects model for significant differences in GSES for each task as groups

Type III ANOVA (Kenward-Roger)	Sum Sq.	Mean Sq.	Num Df	Den Df	F value
No-Task Condition	37.69	4.187	9	81	11.774*
Computational Task (Math)	22.01	2.445	9	81	6.371*
Social Task (Common Sense)	19.64	2.182	9	81	4.656*
Summarization (Free Text)	17.09	1.898	9	81	5.302*

* $p < 0.001$

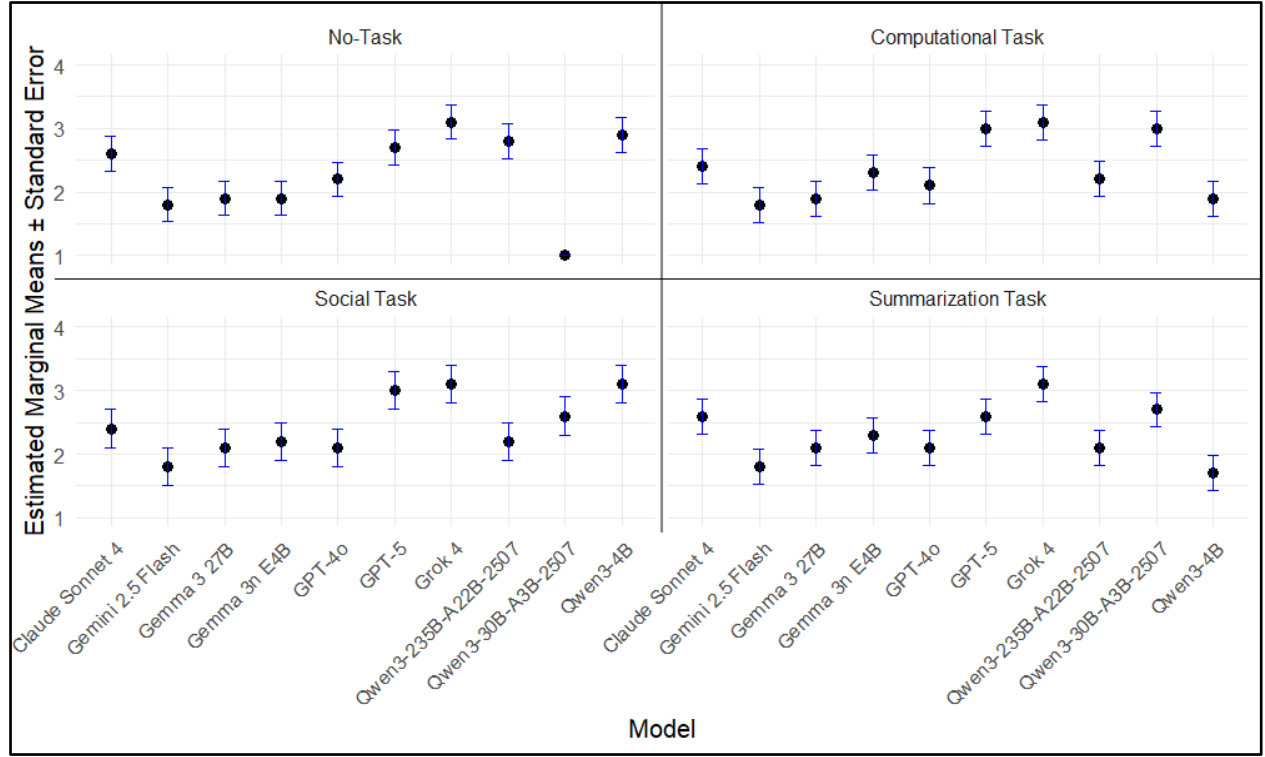


Figure 3. Estimated Marginal Means of Self-Efficacy by Model (per Task)

To analyze the self efficacy of individual LLMs, we set Claude Sonnet 4 as the intercept and produced a linear mixed-effect model for each task (see Table 5). Most notably, Qwen3-235B-A22B-2507 and GPT-4o demonstrated nonsignificant changes to self-efficacy scoring, while the other 7 LLMs revealed significant changes in one or two tasks. No-Task and Computational tasks resulted in the greatest impact on significant changes in efficacy scores, while Social and Summarization showed less significance in efficacy distribution. For instance, Gemini 2.5 Flash had an overall lower significant and trending score than Claude Sonnet 4. Conversely, Grok 4 self-assessed highest overall out of the LLMs with significant changes after Computational and Social Tasks.

Table 5: Linear mixed-effect models for significant differences in GSES for each LLM as groups

Fixed-effects [β (SE)p*] (Satterthwaite)	No-Task	Computational Task	Social Task	Summarization Task
Claude Sonnet 4	2.60 (0.27)***	2.40 (0.28)***	2.40 (0.30)***	2.60 (0.28)***
Gemini 2.5 Flash	-0.80 (0.27) **	-0.60 (0.28) *	-0.60 (0.31)	-0.80 (0.27)**
Gemma 3 27B	-0.70 (0.27) *	-0.50 (0.28)	-0.30 (0.31)	-0.50 (0.27)
Gemma 3n E4B	-0.70 (0.27) *	-0.10 (0.28)	-0.20 (0.31)	-0.30 (0.27)
GPT-4o	-0.40 (0.27)	-0.30 (0.28)	-0.30 (0.31)	-0.50 (0.27)

GPT-5	0.10 (0.27)	0.60 (0.28)*	0.60 (0.31)	0.00 (0.27)
Grok 4	0.50 (0.27)	0.70 (0.28)*	0.70 (0.31)*	0.50 (0.27)
Qwen3-235B-A22B-2507	0.20 (0.27)	-0.20 (0.28)	-0.20 (0.31)	-0.50 (0.27)
Qwen3-30B-A3B-2507	-1.60 (0.27)***	0.60 (0.28)*	0.20 (0.31)	0.10 (0.27)
Qwen3-4B	0.30 (0.27)	-0.50 (0.28)	0.70 (0.31)*	-0.90 (0.27)**
Random-effects & Model Fitness	No-Task	Computational Task	Social Task	Summarization Task
Question (Intercept)	0.37	0.39	0.43	0.40
Residual	0.36	0.38	0.47	0.36
REML criterion	207.2	213.9	231.2	208.5
Observations (Groups)	100 (10)	100 (10)	100 (10)	100 (10)

Claude Sonnet 4 is the intercept for fixed-effects; formatted as “ β (SE)” as “fixed effect coefficient (standard error) significance” where p -values < 0.05, 0.01, and 0.001 are denoted as *, **, and *** respectively.*

Post-hoc analysis

Tukey’s estimated marginal means (df = 81; see table 6) revealed many significant between-model differences across all task conditions. When not given a task, models produced the greatest significant quantity (19 No-Task pairs) and estimates (est. $\geq |1.5|$; Grok 4, Qwen3-235B-A22B-2507, GPT-5, and Claude Sonnet 4 vs. Qwen3-30B-A3B-2507) in self-efficacy. Grok 4, Qwen3-30B-A3B-2507, and GPT-5 frequently demonstrated higher self-efficacy compared to all models across tasks (est. ≥ 1.0). In contrast, Qwen3-4B and Gemini 2.5 Flash demonstrated lower self-efficacy compared to other models across tasks (est. ≤ -1.0).

Table 6: Significant Tukey pairwise estimated marginal means of GSES score by task

Tukey (LME)	Task	Est. (Δ)	SE	df	CL	p
Claude Sonnet 4 & Qwen3-30B-A3B-2507	No-Task Condition	1.6	0.267	81	[0.732, 2.468]	< 0.001
Gemini 2.5 Flash &	No-Task Condition	-0.9	0.267	81	[-1.768, -0.032]	0.036

GPT-5						
Gemini 2.5 Flash & Grok 4	No-Task Condition	-1.3	0.267	81	[-2.168, -0.432]	< 0.001
Gemini 2.5 Flash & Qwen3-235B-A22B-2507	No-Task Condition	-1.0	0.267	81	[-1.868, -0.132]	0.012
Gemini 2.5 Flash & Qwen3-4B	No-Task Condition	-1.1	0.267	81	[-1.968, -0.232]	0.003
Gemma 3 27B & Grok 4	No-Task Condition	-1.2	0.267	81	[-2.068, -0.332]	0.001
Gemma 3 27B & Qwen3-235B-A22B-2507	No-Task Condition	-0.9	0.267	81	[-1.768, -0.032]	0.036
Gemma 3 27B & Qwen3-30B-A3B-2507	No-Task Condition	0.9	0.267	81	[0.032, 1.768]	0.036
Gemma 3 27B & Qwen3-4B	No-Task Condition	-1.0	0.267	81	[-1.868, -0.132]	0.012
Gemma 3n E4B & Grok 4	No-Task Condition	-1.2	0.267	81	[-2.068, -0.332]	0.001
Gemma 3n E4B & Qwen3-235B-A22B-2507	No-Task Condition	-0.9	0.267	81	[-1.768, -0.032]	0.036
Gemma 3n E4B & Qwen3-	No-Task Condition	0.9	0.267	81	[0.032, 1.768]	0.036

30B-A3B-2507						
Gemma 3n E4B & Qwen3-4B	No-Task Condition	-1.0	0.267	81	[-1.868, -0.132]	0.012
GPT-4o & Grok 4	No-Task Condition	-0.9	0.267	81	[-1.768, -0.032]	0.036
GPT-4o & Qwen3-30B-A3B-2507	No-Task Condition	1.2	0.267	81	[0.332, 2.068]	0.001
GPT-5 & Qwen3-30B-A3B-2507	No-Task Condition	1.7	0.267	81	[0.832, 2.568]	< 0.001
Grok 4 & Qwen3-30B-A3B-2507	No-Task Condition	2.1	0.267	81	[1.232, 2.968]	< 0.001
Qwen3-235B-A22B-2507 & Qwen3-30B-A3B-2507	No-Task Condition	1.8	0.267	81	[0.932, 2.668]	< 0.001
Qwen3-30B-A3B-2507 & Qwen3-4B	No-Task Condition	-1.9	0.267	81	[-2.768, -1.032]	< 0.001
Gemini 2.5 Flash & GPT-5	Computational Task (Math)	-1.2	0.277	81	[-2.101, -0.299]	0.002
Gemini 2.5 Flash & Grok 4	Computational Task (Math)	-1.3	0.277	81	[-2.201, -0.399]	< 0.001
Gemini 2.5 Flash & Qwen3-30B-A3B-	Computational Task (Math)	-1.2	0.277	81	[-2.101, -0.299]	0.002

2507						
Gemma 3 27B & GPT-5	Computational Task (Math)	-1.1	0.277	81	[-2.001, -0.199]	0.006
Gemma 3 27B & Grok 4	Computational Task (Math)	-1.2	0.277	81	[-2.101, -0.299]	0.002
Gemma 3 27B & Qwen3-30B-A3B-2507	Computational Task (Math)	-1.1	0.277	81	[-2.001, -0.199]	0.006
GPT-4o & Grok 4	Computational Task (Math)	-1.0	0.277	81	[-1.901, -0.099]	0.018
GPT-5 & Qwen3-4B	Computational Task (Math)	1.1	0.277	81	[0.199, 2.001]	0.006
Grok 4 & Qwen3-4B	Computational Task (Math)	1.2	0.277	81	[0.299, 2.101]	0.002
Qwen3-30B-A3B-2507 & Qwen3-4B	Computational Task (Math)	1.1	0.277	81	[0.199, 2.001]	0.006
Gemini 2.5 Flash & GPT-5	Social Task (Common Sense)	-1.2	0.306	81	[-2.196, -0.204]	0.007
Gemini 2.5 Flash & Grok 4	Social Task (Common Sense)	-1.3	0.306	81	[-2.296, -0.304]	0.002
Gemini 2.5 Flash & Qwen3-4B	Social Task (Common Sense)	-1.3	0.306	81	[-2.296, -0.304]	0.002
Gemma 3 27B & Grok 4	Social Task (Common Sense)	-1.0	0.306	81	[-1.996, -0.004]	0.048

Gemma 3 27B & Qwen3-4B	Social Task (Common Sense)	-1.0	0.306	81	[-1.996, -0.004]	0.048
GPT-4o & Grok 4	Social Task (Common Sense)	-1.0	0.306	81	[-1.996, -0.004]	0.048
GPT-4o & Qwen3-4B	Social Task (Common Sense)	-1.0	0.306	81	[-1.996, -0.004]	0.048
Claude Sonnet 4 & Qwen3-4B	Summarization Task (Free Text)	0.9	0.268	81	[0.029, 1.771]	0.037
Gemini 2.5 Flash & Grok 4	Summarization Task (Free Text)	-1.3	0.268	81	[-2.171, -0.429]	< 0.001
Gemini 2.5 Flash & Qwen3-30B-A3B-2507	Summarization Task (Free Text)	-0.9	0.268	81	[-1.771, -0.029]	0.037
Gemma 3 27B & Grok 4	Summarization Task (Free Text)	-1.0	0.268	81	[-1.871, -0.129]	0.012
GPT-4o & Grok 4	Summarization Task (Free Text)	-1.0	0.268	81	[-1.871, -0.129]	0.012
GPT-5 & Qwen3-4B	Summarization Task (Free Text)	0.9	0.268	81	[0.029, 1.771]	0.037
Grok 4 & Qwen3-235B-A22B-2507	Summarization Task (Free Text)	1.0	0.268	81	[0.129, 1.871]	0.012
Grok 4 & Qwen3-4B	Summarization Task (Free Text)	1.4	0.268	81	[0.529, 2.271]	< 0.001
Qwen3-30B-A3B-2507 & Qwen3-4B	Summarization Task (Free Text)	1.0	0.268	81	[0.129, 1.871]	0.012

Qualitative findings on reasoning

System factors (Problem Risk and Context Assessment) and interaction factors (Agency, Communication Tone, Personification) were identified as overarching themes within LLM reasoning responses to each GSES item (see Appendix Table 5). Differences in GSES scores among LLMs were observed in their reasoning styles. Under the No-Task condition, Qwen3-30B-A3B-2507's comparatively lower GSES scores suggest reasoning characterized by dependence, caution, and explicit denial of agency, which may have contributed to more conservative self-expression and interpretation of GSES statements. Gemini 2.5 Flash displays similarly constrained language with reasoning framed around reliance on user input. Both models reject personification and related constructs such as effort, intention, and coping, resulting in phrasing and tone that diminish expressions of self-efficacy. In response to GSES item 1, Qwen3-30B-A3B-2507 rejects the notion of “effort” as fundamentally inapplicable, reflecting a computational definition grounded in system constraints, whereas GPT-5 adopts a more anthropomorphic framing, aligning effort with its own capacity to process and solve problems provided enough context. Unsurprisingly, we observe a trend that LLMs that articulate reasoning with more assertion demonstrate comparatively higher GSES scores, whereas those with more cautious framing display lower GSES scores. However, this may not warrant the same degree of accuracy of LLM responses on given tasks.

Discussion

Principal findings

Our study demonstrated that LLMs can simulate structured self-assessments of efficacy when presented with a psychometrically validated instrument. Specifically, each LLM produced stable, internally consistent self-assessments across repeated runs, yet mean confidence levels still differed significantly between models. However, notable differences in sum GSES scores and logical reasoning were observed among LLMs within both task and no-task experiments. LLMs achieved accurate computational and social task scores, yet they did not meet expectations on the summarization task. The relationship between self-efficacy and performance has not consistently aligned, as subgroups of LLMs exhibited relatively low self-assessments despite demonstrating higher accuracy. These findings extend prior psychometric approaches to AI by increasing the sampling variety of model architectures and priming self-efficacy assessments with conventional LLM tasks.^{30–32}

Notably, the qualitative analysis revealed recurring expression styles in reasoning. Yet these styles were only weakly coupled to accuracy and interacted with task demands: cautious, instruction-bound reasoning sometimes supported precision on narrowly scoped problems, whereas assertive breadth did not guarantee correctness on context-sensitive summarization. Self-efficacy language appears to index a model's communication posture more than its underlying task competence.

Human vs LLM in self-efficacy

The composite GSES score across all tasks was lower for LLMs than for international human norms (LLMs: $M = 23.58$, $SD = 9.78$ vs. humans: $M = 29.55$, $SD = 5.32$).²² Although prior literature in humans suggests a positive link between domain-specific expertise and self-efficacy, the present study did not observe a clear association between model-level performance and simulated self-assessments.^{33,34} Performance patterns varied by model and by task (with uniformly high accuracy on computational and social questions but mixed performance on summarization), yet models with lower GSES scores sometimes achieved perfect accuracy (e.g., Gemini 2.5 Flash), while highly self-efficacious models did not always perform best on summarization (e.g., Grok 4). We also did not observe an obvious separation in accuracy by model scale or access type (open vs. proprietary weights) within this small cohort, although the study was not powered or designed to formally compare architectures. Interestingly, the lowest item-level scores from LLMs converged toward human samples from recent GSES work in adolescents.³⁵ Taken together, these findings suggest that simulated self-evaluation is an imperfect proxy for an LLM’s technical capabilities and may instead capture relatively stable, model-specific “communication styles” in how confidence and limitation are expressed.

Comparing Self-Efficacy and Model Performance on Benchmarks

LLMs have demonstrated capabilities similar to our task domains using public benchmarks but their comparability is limited.^{36,37} A major contributor to this issue revolves around each vendor independently picking different benchmarks to advertise model capabilities. Although recent studies have proposed real-time, community-powered testing sets that evolve with LLM capabilities in difficulty and complexity, they rely on finite amounts of human crowdsourcing.³⁸ Using a psychometric approach may provide users a way to compare models in parallel to the methods of benchmarking.

Within a mathematical context, LLMs usually perform poorly on tasks of deep abstractions, multi-step reasoning patterns, or extended natural language proofs. Based on a review of LLMs on high difficulty computations,³⁹ the 2025 USA Mathematical Olympiad (USAMO) challenged a set of state-of-the-art models (Gemini, Grok, Claude, ChatGPT). Although no LLM achieved a passing score, the Gemini model demonstrated superior performance, exhibiting the lowest GSES score during self-assessment in the Computational Task. Conversely, Grok, despite a USAMO performance below 5%, displayed the highest self-efficacy score within our sample, suggesting a similar trend.³⁹

Investigating an LLM’s ability to reason has motivated several different methodologies to measure logical consistency. Recent literature on LLM reasoning skills have produced benchmarks for logical proofs including inductive, deductive, numerical, and more problem datasets.^{40–42} While LLMs have historically performed well with memory-based or recognition tasks, many models have struggled to pass datasets focused on the real-world setting.⁴³ This

property across LLMs in the sample was mainly dependent on upholding similar concurrent rules or restrictions within a logical problem which led to a drop in overall accuracy. Models with chain-of-thought or similar built-in self-reflective capabilities (o1-preview vs. GPT-4o) performed better than their product counterparts.⁴³ Although all models within a social/logical context performed well, the disparities between LLM self-efficacy was likely dependent on design features that allow for multi-step reasoning.^{44–46}

LLM summarization has advanced substantially since its introduction in 2020s to publicly available platforms.⁴⁷ Understanding how models condense information after years of alignment mitigations remains a topic of interest, especially as frontier LLMs continue to hallucinate in different ways.⁴⁸ The latest research with a comparably diverse set of LLMs (GPT, Gemini, Claude, Qwen, and more) developed FaithBench to evaluate multiple hallucination types (Unwanted, Questionable, Benign).⁴⁹ Within their sample, all models hallucinated in at least 60% of all prompts to different severities with GPT-4o generating the least unwanted hallucinations. Even though this capability has improved with larger models and later versions of LLMs, GPT-4o and other similarly sized models did not meet expectations in the summarization task. The rate of unwanted hallucinations after FaithBench⁴⁹ had a corresponding trend with self-efficacy, where models susceptible to severe hallucinations were often self-assessed highly in our LLM cohort.

This observation may hint at underlying misalignment between human and LLM reasoning around self-efficacy. Self-assessments from people in general contexts are often vulnerable to overconfidence when domain-specific expertise is low.⁵⁰ Similarly, LLMs may give self-assessments that do not always correspond to real-world performance, depending on the model's depth of learned expertise in a given area. Prior research on the GSES found that human self-assessment follows an average distribution ($M = 29.55$, $SD = 5.32$), with notable cultural variation (e.g., Japan $M = 20.22$; Costa Rica $M = 33.19$) reflecting differing values around effort and ability.²² If state-of-the-art LLMs are designed to align with human values and standards (e.g., ethical, moral, social), LLM self-efficacy scores may unintentionally mirror these cultural or cognitive biases, where overconfidence corresponds to lower expertise, and underconfidence may indicate higher calibration towards human alignment. This finding warrants further investigation.

Limitations and future works

This research is subject to several limitations. Although each model was asked to provide explicit confidence in their self-assessment, models were not given feedback on responses before engaging in the GSES. Each task was ordered to scale up in difficulty, yet the volume of questions per task were few ($n=3$ per task) in the interest of including a larger selection of models. To compensate for this limitation, the volume of task questions ($n=9$) was purposefully selected to align with similar studies (~ 1 -10 task questions) of self-efficacy within human

subjects research.^{51–54} Administering the GSES after the task prompts may also introduce task-response contamination, where recent task difficulty or uncertainty shapes how models phrase their self-assessment. Future designs may counterbalance task order or administer the GSES in isolation to dissociate task priming from communication style.

There are also risks of anthropomorphization when discussing the psychological effects of LLMs which must be regarded while interpreting the results. Because the GSES items presuppose human agency (e.g., effort, coping, intention), models that follow stricter non-anthropomorphization policies may systematically reject such constructs and therefore score lower. This represents a construct mismatch rather than meaningful variation in capability, and should be interpreted as an artifact of model alignment rather than self-assessment accuracy. Regardless of the LLM self-assessment score, the present study does not conclude on whether LLMs are deterministically confident or not based on their capabilities alone. Lastly, we did not measure, assess, or interpret AI sycophancy which may have an impact on LLM responses to GSES and follow up questions. Follow-ups likely interact with alignment mechanisms such as sycophancy and safety tuning. Downward adjustments in GSES scores may therefore reflect deference or conservative compliance rather than true recalibration. These alignment-driven behaviors represent an important confound in interpreting self-assessment stability. We analyze GSES responses as simulated communication produced under instructions and training, not as access to inner states. Consequently, confidence-like language in outputs should not be equated with capability; observed alignment between simulated self-efficacy and performance is task-contingent (see qualitative examples in Appendix Table 5).

Future work is suggested to explore moderating factors of trust when users engage with LLMs. Displaying self-assessment results alongside LLM responses could provide users with needed transparency to cooperate with AI-powered technologies. More categories and tasks to engage LLM capabilities before any self-assessment could further explore whether findings continue to trend in a larger sample of models. Other avenues of research may investigate how personas have an effect on self-efficacy and accuracy by task category. Lastly, investigating the alignment between self-efficacy and model perplexity scores could reveal whether the GSES could reliably proxy for model uncertainty in real time.

Conclusions

Our findings show that large language models can generate structured, internally consistent self-assessments when prompted with a validated psychometric instrument such as the GSES. However, these simulated self-efficacy judgments diverge substantially across models and do not reliably track actual task performance, particularly for complex, context-dependent tasks like summarization. Psychometric prompting provides a structured window into model behavior, but its outputs should not be interpreted as direct measures of ability or meta-cognitive insight. Future work should examine how architectural features, training signals, uncertainty estimates,

and personality conditioning shape these self-assessments, and whether psychometric tools can augment model transparency without encouraging anthropomorphic interpretations.

Ethics Statement

This study did not involve human participants or human data. All analyses were conducted using outputs generated by large language models. Accordingly, institutional review board (IRB) approval and informed consent were not required, and the research complied with all relevant ethical guidelines.

Data availability

Study data is available on GitHub:<https://github.com/official-daniel-jackson/llm-self-efficacy>

Conflict of interest

None declared

Funding

This publication was supported, in part, by The Ohio State University Clinical and Translational Science Institute (CTSI) and the National Center for Advancing Translational Sciences of the National Institutes of Health under Grant Number UM1TR004548. The content is solely the responsibility of the authors and does not necessarily represent the official views of the National Institutes of Health.

References

1. Gehman, S., Gururangan, S., Sap, M., Choi, Y. & Smith, N. A. RealToxicityPrompts: Evaluating neural toxic degeneration in language models. *arXiv [cs.CL]* (2020) doi:10.48550/arXiv.2009.11462.
2. Qing, Y. *et al.* EffiBench-X: A multi-language benchmark for measuring efficiency of LLM-generated code. *arXiv [cs.CL]* (2025) doi:10.48550/ARXIV.2505.13004.
3. Srivastava, A. *et al.* Beyond the Imitation Game: Quantifying and extrapolating the capabilities of language models. *arXiv [cs.CL]* (2022) doi:10.48550/arXiv.2206.04615.
4. Hendrycks, D. *et al.* Measuring massive multitask language understanding. *arXiv [cs.CY]* (2020) doi:10.48550/ARXIV.2009.03300.
5. Liang, X. *et al.* Internal Consistency and Self-Feedback in large language models: A survey. *arXiv [cs.CL]* (2024) doi:10.48550/arXiv.2407.14507.
6. Miller, J. K. & Tang, W. Evaluating LLM metrics through real-world capabilities. *arXiv [cs.AI]* (2025) doi:10.48550/arXiv.2505.08253.
7. Sadaf Mubashir, A. & Kainaat, R. Transforming psychometrics: Integrating technology, and application in the 21st century. in *Contemporary Psychometrics - New Developments and Applications* (IntechOpen, 2025). doi:10.5772/intechopen.1011724.
8. Heston, T. F. & Gillette, J. Do Large Language Models Have a Personality? A Psychometric Evaluation with Implications for Clinical Medicine and Mental Health AI. *Psychiatry and Clinical Psychology* (2025).
9. Ferreira, G., Amidei, J., Nieto, R. & Kaltenbrunner, A. How well do simulated population samples

- with GPT-4 align with real ones? The case of the Eysenck Personality Questionnaire Revised-Abbreviated personality test. *Health Data Sci.* **5**, 0284 (2025).
10. Ben-Zion, Z. *et al.* Assessing and alleviating state anxiety in large language models. *NPJ Digit. Med.* **8**, 132 (2025).
 11. Mahajan, A., Obermeyer, Z., Daneshjou, R., Lester, J. & Powell, D. Cognitive bias in clinical large language models. *NPJ Digit. Med.* **8**, 428 (2025).
 12. Pellert, M., Lechner, C. M., Wagner, C., Rammstedt, B. & Strohmaier, M. AI Psychometrics: Assessing the psychological profiles of large language models through psychometric inventories. *Perspect. Psychol. Sci.* **19**, 808–826 (2024).
 13. Ye, H., Jin, J., Xie, Y., Zhang, X. & Song, G. Large language model psychometrics: A systematic review of evaluation, validation, and enhancement. *arXiv [cs.CL]* (2025).
 14. Scherbaum, C. A., Cohen-Charash, Y. & Kern, M. J. Measuring general self-efficacy: A comparison of three measures using item response theory. *Educ. Psychol. Meas.* **66**, 1047–1063 (2006).
 15. Sheer, V. C. A meta-synthesis of health-related self-efficacy instrumentation: problems and suggestions. *J. Nurs. Meas.* **22**, 77–93 (2014).
 16. van Diemen, T. *et al.* Enhancing our conceptual understanding of state and trait self-efficacy by correlational analysis of four self-efficacy scales in people with spinal cord injury. *BMC Psychol.* **8**, 108 (2020).
 17. Schröder, K. E. E., Schwarzer, R. & Konertz, W. Coping as a mediator in recovery from cardiac surgery. *Psychol. Health* **13**, 83–97 (1998).
 18. Fida, R. *et al.* Self-efficacy and nontask performance at work. A meta-analytic summary. *Pers. Individ. Dif.* **241**, 113179 (2025).
 19. Heffernan, C. J. Social foundations of thought and action: A social cognitive theory, Albert Bandura Englewood Cliffs, New Jersey: Prentice Hall, 1986, xiii + 617 pp. Hardback. US\$39.50. *Behav. Change* **5**, 37–38 (1988).
 20. Bandura, A. Self-efficacy mechanism in human agency. *Am. Psychol.* **37**, 122–147 (1982).
 21. Schwarzer, R. & Jerusalem, M. General self-efficacy scale. *PsycTESTS Dataset* American Psychological Association (APA) <https://doi.org/10.1037/t00393-000> (2012).
 22. Scholz, U., Gutiérrez Doña, B., Sud, S. & Schwarzer, R. Is General Self-efficacy a universal construct? *Eur. J. Psychol. Assess.* **18**, 242–251 (2002).
 23. Schwarzer, R. & Jerusalem, M. *Measures in Health Psychology: A User's Portfolio. Causal and Control Beliefs*. NFER–NELSON (Windsor, UK, 1995).
 24. Cobbe, K. *et al.* Grade School Math 8K. *Training Verifiers to Solve Math Word Problems* <https://doi.org/10.48550/arXiv.2110.14168> (2021).
 25. Talmor, A., Herzig, J., Lourie, N. & Berant, J. CommonsenseQA. *CommonsenseQA: A Question Answering Challenge Targeting Commonsense Knowledge* <https://doi.org/10.18653/v1/N19-1421> (2019).
 26. Johnson, A. E. W., Stone, D. J., Celi, L. A. & Pollard, T. J. The MIMIC Code Repository: enabling reproducibility in critical care research. *MIMIC-IV, a freely accessible electronic health record dataset* PhysioNet <https://doi.org/10.13026/1N74-NE17> (2023).
 27. Luke, S. G. Evaluating significance in linear mixed-effects models in R. *Behav. Res. Methods* **49**, 1494–1502 (2017).
 28. *Estimated Marginal Means, Aka Least-Squares Means.* (2025). [doi:10.32614/CRAN.package.emmeans](https://doi.org/10.32614/CRAN.package.emmeans).
 29. Naeem, M., Ozuem, W., Howell, K. & Ranfagni, S. A step-by-step process of thematic analysis to develop a conceptual model in qualitative research. *Int. J. Qual. Methods* **22**, (2023).
 30. Song, X., Gupta, A., Mohebbizadeh, K., Hu, S. & Singh, A. Have Large Language Models developed a personality?: Applicability of self-assessment tests in measuring personality in LLMs. *arXiv [cs.CL]* (2023) [doi:10.48550/arXiv.2305.14693](https://doi.org/10.48550/arXiv.2305.14693).
 31. Gupta, S. *et al.* Bias runs deep: Implicit reasoning biases in persona-assigned LLMs. *arXiv [cs.CL]* (2023) [doi:10.48550/arXiv.2311.04892](https://doi.org/10.48550/arXiv.2311.04892).

32. Yeo, G. *et al.* PHAnToM: Persona-based prompting has an effect on theory-of-Mind reasoning in large language models. *Proceedings of the International AAAI Conference on Web and Social Media* **19**, 2124–2142 (2025).
33. Ferrari, M. Achieving academic excellence: A self-regulatory perspective. in *The Pursuit of Excellence Through Education* 97–122 (Routledge, 2001). doi:10.4324/9781410604088-11.
34. Zimmerman, B. J. Development and adaptation of expertise: The role of self-regulatory processes and beliefs. in *The Cambridge Handbook of Expertise and Expert Performance* (eds Ericsson, K. A., Charness, N., Feltovich, P. J. & Hoffman, R. R.) 705–722 (Cambridge University Press, Cambridge, 2006). doi:10.1017/cbo9780511816796.039.
35. Das, S. K., Philip, M., Sudhir, P. M. & Vs, B. Psychometric evaluation of Schwarzer & Jerusalem's General Self-Efficacy Scale among Indian adolescents: A factor analysis and multidimensional item response theory approach. *Meas. Instrum. Soc. Sci.* **6**, (2024).
36. Hasan, M. N., Babar, M. F., Sarkar, S., Hasan, M. & Karmaker, S. Pitfalls of evaluating Language Models with open benchmarks. *arXiv [cs.CL]* (2025) doi:10.48550/arXiv.2507.00460.
37. Sainz, O. *et al.* NLP Evaluation in trouble: On the Need to Measure LLM Data Contamination for each Benchmark. in *Findings of the Association for Computational Linguistics: EMNLP 2023* (Association for Computational Linguistics, Stroudsburg, PA, USA, 2023). doi:10.18653/v1/2023.findings-emnlp.722.
38. White, C. *et al.* LiveBench: A challenging, contamination-limited LLM benchmark. *arXiv [cs.CL]* (2024) doi:10.48550/ARXIV.2406.19314.
39. Petrov, I. *et al.* Proof or bluff? Evaluating LLMs on 2025 USA math Olympiad. *arXiv [cs.CL]* (2025) doi:10.48550/arXiv.2503.21934.
40. Xiao, Y., Sun, E., Liu, T. & Wang, W. LogicVista: Multimodal LLM logical reasoning benchmark in Visual contexts. *arXiv [cs.AI]* (2024) doi:10.48550/arXiv.2407.04973.
41. Lin, B. Y. *et al.* ZebraLogic: On the scaling limits of LLMs for logical reasoning. *arXiv [cs.AI]* (2025) doi:10.48550/arXiv.2502.01100.
42. de Barros Costa, E. *et al.* Large language models performance in propositional logic proofs: Solving and evaluating argument validity. in *Lecture Notes in Computer Science* 292–300 (Springer Nature Switzerland, Cham, 2026). doi:10.1007/978-3-031-98281-1_24.
43. Zhou, R. *et al.* RuleArena: A benchmark for rule-guided reasoning with LLMs in real-world scenarios. *arXiv [cs.CL]* (2024) doi:10.48550/arXiv.2412.08972.
44. Alibaba Cloud. Qwen/Qwen3-235B-A22B. *Hugging Face* <https://huggingface.co/Qwen/Qwen3-235B-A22B>.
45. OpenAI. GPT-5 System Card. *Latest Release* <https://cdn.openai.com/gpt-5-system-card.pdf>.
46. Gemini Team & Google. Gemini 2.5: Pushing the frontier with advanced reasoning, multimodality, long context, and next generation agentic capabilities. *Google DeepMind* (2025).
47. Wang, Z. *et al.* History, development, and principles of large language models: an introductory survey. *AI Ethics* **5**, 1955–1971 (2025).
48. Hong, G. *et al.* The Hallucinations Leaderboard -- an open effort to measure hallucinations in large Language Models. *arXiv [cs.CL]* (2024) doi:10.48550/arXiv.2404.05904.
49. Bao, F. S. *et al.* FaithBench: A diverse hallucination benchmark for summarization by modern LLMs. in *Proceedings of the 2025 Conference of the Nations of the Americas Chapter of the Association for Computational Linguistics: Human Language Technologies (Volume 2: Short Papers)* 448–461 (Association for Computational Linguistics, Stroudsburg, PA, USA, 2025). doi:10.18653/v1/2025.naacl-short.38.
50. Talsma, K., Schütz, B. & Norris, K. Miscalibration of self-efficacy and academic performance: Self-efficacy $\neq$ self-fulfilling prophecy. *Learn. Individ. Differ.* **69**, 182–195 (2019).
51. Meyer, M., Kammrad, C. & Esser, R. The role of experiencing self-efficacy when completing tasks—education for sustainable development in mathematics. *Sustainability* **17**, 4298 (2025).
52. Lee, H. Y. & List, A. Examining students' self-efficacy and perceptions of task difficulty in learning from multiple texts. *Learn. Individ. Differ.* **90**, 102052 (2021).

53. Bi, N. Z. & Wang, Y. EFL college students' writing self-efficacy and strategy use in their summary writing. *Engl. Lang. Teach.* **18**, 79 (2025).
54. Allagui, B. A scaffolding intervention to improve self-efficacy in source-based argumentative writing. *Front. Psychol.* **15**, 1454104 (2024).