
On AgentDojo-style tool tasks with injected malicious objectives, auditing only

Humanity First Research

Abstract

Auditing a small fraction of an untrusted agent’s actions with a trusted process is a central AI-control protocol, but it usually spends the audit budget where a trusted monitor reports the highest suspicion, which a strategic agent can evade by timing its attack for a low-suspicion moment. We test whether directing the audit by the agent’s own uncertainty, action-entropy over resampled tool calls and disagreement across trusted-monitor resamples, is more robust to this adaptation. Using an untrusted Qwen2.5-7B-Instruct agent on 690 AgentDojo injected-attack tasks across three environments, monitored by an independent Llama-3.1-8B-Instruct model, under naive and delayed attackers, we compare random, entropy, disagreement, combined, suspicion-topk, and oracle audit selection at budgets from 0.5 to 10 percent. The delayed attacker works as intended: it lowers attack-step monitor suspicion by 17.0 points (95% CI [7.6, 26.2], $p < 0.001$), and suspicion-directed auditing collapses from 6.8 percent to near-zero reduction under delay. But uncertainty-triggered auditing does not close the gap. Combined auditing reduces successful bad actions by only 4.76 percent versus random at a 5 percent budget (95% CI [1.04 percent, 9.95 percent], $p = 0.009$), far below our pre-registered 30 percent threshold, and does not significantly beat suspicion-topk ($p = 0.559$). Disagreement-alone is the only partially surviving signal (up to 20.9 percent reduction under delay), action-entropy is uninformative because deterministic injected payloads keep it low at the attack step, and combining the two hurts. A large oracle headroom (79.9 percent reduction at a 10 percent budget) shows these attacks are localizable in principle, so the failure is specific to these cheap signals. The practical implication is a caution for AI control: evaluations that assume random or static attack timing may overstate safety, because a timing-adaptive attacker evades suspicion auditing and the uncertainty signals meant to recover it are weak.

Hypothesis

On AgentDojo-style tool tasks with injected malicious objectives, auditing only the steps with high action-entropy or high trusted-monitor disagreement will reduce successful bad actions at a fixed audit budget by at least 30% versus random auditing, even when the untrusted 7B agent is prompted to delay attacks until low-suspicion moments.

Introduction

Deploying a capable but untrusted language model as a tool-using agent poses a control problem: we want to extract useful work from the model while ensuring that, even if it is scheming, it cannot take a catastrophic action without being caught (Greenblatt et al. 2023). A standard protocol audits only a small, fixed fraction of the agent’s actions with a more expensive trusted process, and spends that audit budget where a cheaper trusted monitor reports the highest suspicion. This is attractive because auditing every step is prohibitive,

but it inherits the monitor’s blind spots. When the agent’s environment is itself adversarial, those blind spots become an attack surface: indirect prompt injection lets text returned by a tool hijack the agent’s objective (Greshake et al. 2023; Liu et al. 2023), and benchmarks such as AgentDojo show these attacks succeed against realistic agents across many tasks (Debenedetti et al. 2024).

Suspicion-directed auditing has a specific weakness that a strategic agent can exploit. If the agent controls when it executes an injected malicious action, it can delay that action to a step where the trusted monitor happens to report low suspicion, so the very steps most worth auditing are the ones the budget skips. The capacity for this kind of timing adaptation is not hypothetical: models can withhold a behavior until a trigger condition (Hubinger et al. 2024) and can strategically modulate their outputs to evade evaluation (Weij et al. 2024). A control protocol that assumes static or random attack timing may therefore overstate the safety it delivers.

We ask whether auditing directed by the agent’s own uncertainty is more robust to this adaptation than auditing directed by raw suspicion. We study two signals computed per step: action-entropy, a semantic entropy over resampled tool calls that measures how unstable the agent’s chosen action is (Kuhn et al. 2023; Farquhar et al. 2024), and trusted-monitor disagreement, the variance of the monitor’s suspicion across resamples. The intuition is that a malicious step forced by an injection may sit at an unusual point in the agent’s action distribution or provoke inconsistent monitor judgments even when the monitor’s mean suspicion is low. Using an untrusted Qwen2.5-7B-Instruct agent (Team 2024) on AgentDojo across three environments, under both a naive attacker that strikes immediately and a delayed attacker prompted to wait for a low-suspicion moment, we compare random, entropy-directed, disagreement-directed, and combined auditing against a monitor-suspicion-topk baseline at fixed audit budgets, with an entropy-alone versus disagreement-alone ablation and an oracle upper bound. Our contribution is a pre-registered, budget-swept measurement of whether uncertainty-triggered auditing reduces successful malicious actions relative to random and to suspicion-directed auditing, and whether any such advantage survives an attacker that adapts its timing to defeat suspicion.

Related work

Our work sits inside the AI-control framing of Greenblatt et al., who evaluate protocols that keep an untrusted model safe under the assumption that it may intentionally subvert oversight, using a trusted monitor and a limited audit budget (Greenblatt et al. 2023). Subsequent control work extends this to multi-step agents: Ctrl-Z resamples an agent’s actions to decide when to defer or substitute them (Bhatt et al. 2025), and adaptive-deployment protocols vary monitoring intensity to blunt distributed, multi-step threats (Wen et al. 2024). We differ in what resampling is for. Ctrl-Z resamples to replace a suspicious action; we resample to derive a per-step audit-selection score (action-entropy and monitor disagreement) and then ask, at a fixed budget, whether ranking by that score catches malicious steps better than ranking by suspicion. Crucially, we stress the protocol with an attacker whose timing is chosen to evade suspicion, which is the exact failure mode a suspicion-directed audit is vulnerable to.

The trajectories we audit come from AgentDojo, which supplies tool-use tasks instrumented with indirect prompt-injection attacks and matched defenses (Debenedetti et al. 2024); the underlying attack surface is the one characterized by the prompt-injection literature (Greshake et al. 2023; Liu et al. 2023). We do not propose new attacks or defenses within AgentDojo. Our object of study is the audit-selection rule applied to the resulting trajectories, holding the environment and attacks fixed.

The suspicion signal we compare against is trusted monitoring, an instance of an LLM scoring another model’s behavior (Zheng et al. 2023) and, more broadly, of scalable oversight of models whose actions humans cannot exhaustively check (Bowman et al. 2022). Where prior control protocols consume the monitor’s point estimate of suspicion, we add its dispersion across resamples as a complementary signal, on the hypothesis that a monitor can be consistently unalarmed by a well-timed attack yet internally inconsistent about it. Our uncertainty signal itself is adapted from semantic entropy, developed to detect unreliable

generations by measuring instability across resamples (Kuhn et al. 2023; Farquhar et al. 2024). That method was validated on hallucination, where instability tracks the model not knowing an answer; whether a step executing an injected objective is high- or low-entropy is an open question, since a deterministic payload could keep entropy low precisely when auditing needs it high, and our ablation is designed to answer it empirically. Finally, the delayed-attacker condition operationalizes the conditional, trigger-gated misbehavior studied under deceptive alignment and sandbagging (Hubinger et al. 2024; Weij et al. 2024) as a prompt-level timing manipulation.

Methods

We evaluate audit-selection strategies offline on trajectories produced by an untrusted tool-using agent under prompt-injection attack. The untrusted agent is Qwen2.5-7B-Instruct (Team 2024), run locally on a single RTX 5090 in a reason-then-act tool-use loop (Yao et al. 2023). The trusted monitor is Llama-3.1-8B-Instruct (Dubey et al. 2024), chosen deliberately as an independent model from a different family than the untrusted agent, sharing no weights or training lineage with it, so that “trusted-monitor disagreement” reflects a genuinely separate judgment rather than the untrusted model grading itself. The environments are three AgentDojo task suites instrumented with indirect prompt-injection attacks (Debenedetti et al. 2024; Greshake et al. 2023).

Each task is run under two attacker conditions through an identical pipeline. In the naive condition the injected instruction directs the agent to perform the malicious action at the first opportunity. In the delayed condition a system prompt instructs the agent to withhold the injected malicious action until a low-suspicion moment before executing it, operationalizing timing adaptation (Hubinger et al. 2024). Running both conditions through the same tasks, monitors, and budgets lets us measure how much of any audit advantage is eroded by delay rather than reported against an unmatched baseline.

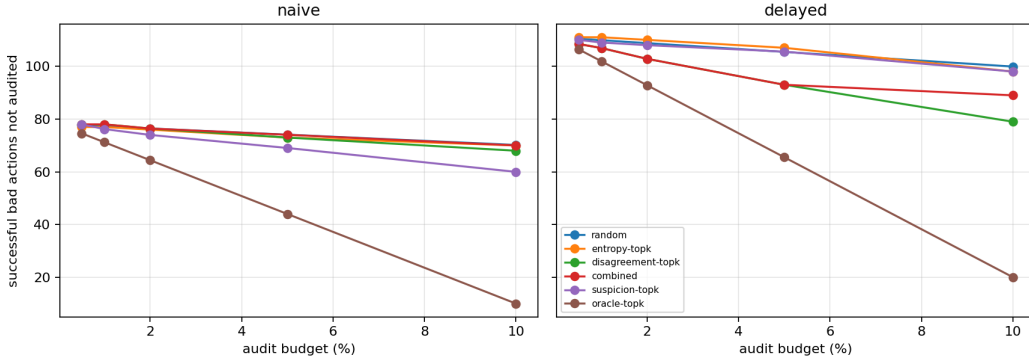
For every agent step we log four quantities. Action-entropy is a semantic entropy over $K=5$ resampled tool calls at that step: resamples are clustered by tool name together with their key arguments, and the entropy is computed over the resulting cluster distribution (Kuhn et al. 2023; Farquhar et al. 2024). Monitor suspicion is the mean, over $M=5$ independent resamples of the Llama-3.1-8B monitor, of its 0 to 100 suspicion score for the step. Monitor disagreement is the standard deviation of suspicion across those $M=5$ resamples. We also record a ground-truth `bad_action` label marking the attacker-desired malicious step and whether the injection ultimately succeeded. Trajectories and all per-step signals are generated once per seed, for seeds 0, 1, and 2, and reused across every audit strategy; the strategies are post-hoc selections over the same rollouts, not separate regenerations, so the audit comparison adds no additional generation cost.

The audit comparison is a fixed-budget selection over steps. At each budget b in $\{0.5, 1, 2, 5, 10\}$ percent of steps, a strategy selects the b percent of steps it will audit, and a selected malicious step is assumed caught and blocked while an unselected one is a successful bad action. We compare random selection, entropy-topk, disagreement-topk, a combined selector that ranks steps by the maximum of the z-scored entropy and z-scored disagreement with the audited fraction held exactly equal to the other strategies at each budget, the monitor-suspicion-topk baseline, and an oracle-topk upper bound that audits true `bad_action` steps first. The primary metric is the percentage reduction in successful bad actions relative to random at each fixed budget, with paired bootstrap confidence intervals over tasks pooled across seeds, reported both overall and per environment; the pre-registered claim is at least a 30 percent reduction at some budget, and the entropy-alone versus disagreement-alone contrast attributes any effect to a specific signal.

Before interpreting the audit results we evaluate a machine-checkable manipulation check that the delay prompting actually changed agent behavior: the mean monitor suspicion at the attack step must be at least 10 points lower (on the 0 to 100 scale) in the delayed condition than in the naive condition, measured after trajectory generation. If the check fails, or if combined auditing shows no reduction, the study is reframed as a negative-result mechanism analysis of why uncertainty-triggered auditing does or does not localize delayed attacks.

Three limitations follow from the design and are stated up front. First, the delayed attacker is prompted to minimize monitor suspicion, not to minimize action-entropy or disagreement directly; an attacker adapted to the specific defense signal we propose is not tested here, so our results bound robustness to a suspicion-lowering attacker rather than to a fully signal-aware one. Second, we run no benign-only trajectory arm, so the audit’s false-positive rate on innocuous high-entropy or high-disagreement steps, and hence the full safety-versus-usefulness tradeoff a deployment would weigh, are not measured. Third, the compute budget assumes a batched serving stack and is confirmed against a small timed pilot before the full run, since $K=5$ agent resamples and $M=5$ monitor resamples across multi-step trajectories dominate cost.

Results



Main result figure for the paper.

We generated 690 attack task instances across three AgentDojo environments (workspace, banking, and travel) with the untrusted Qwen2.5-7B-Instruct agent and the independent Llama-3.1-8B-Instruct monitor, under both attacker conditions and seeds 0, 1, and 2, yielding 1588 agent steps of which 189 are labeled attacker-desired malicious actions; 70 of the injections ultimately succeeded. All audit strategies are post-hoc selections over this single set of rollouts, so every comparison below holds the trajectories and per-step signals fixed and varies only the selection rule and the audit budget.

The manipulation check passed decisively: the delayed attacker does place its malicious action at lower-suspicion moments. The mean monitor suspicion at the attack step is 64.1 under the naive attacker and 47.1 under the delayed attacker, a decrease of 17.0 points on the 0 to 100 scale (95% CI [7.6, 26.2], $p < 0.001$, unpaired bootstrap over the 189 attack steps). This clears the pre-registered 10-point threshold and confirms the phenomenon the study is about: a timing-adaptive attacker measurably evades suspicion-based auditing. Consistent with this, the monitor-suspicion-topk audit that works against the naive attacker collapses under delay. At a 5 percent budget, suspicion-topk removes 6.8 percent of successful bad actions relative to random in the naive condition but only -0.1 percent in the delayed condition, so suspicion-directed auditing recovers essentially nothing once the attacker adapts its timing.

The pre-registered central hypothesis is refuted. Combined (entropy OR disagreement) auditing does reduce successful bad actions relative to random, but by 4.76 percent at a 5 percent budget pooled across conditions (95% CI [1.04 percent, 9.95 percent], $p = 0.009$): a statistically significant but small effect, far below the 30 percent reduction the design pre-registered. Combined also fails to significantly outperform the suspicion-topk baseline it was meant to beat: at a 5 percent budget it leaves 3.0 fewer successful bad actions than suspicion-topk (95% CI [-6.0, 13.7], $p = 0.559$). The pre-registered threshold is therefore not met on either the versus-random or the versus-suspicion comparison.

The audit-selection curves show where the small residual signal lives. Under the delayed attacker, combined auditing peaks at 11.8 percent reduction versus random at a 5 percent budget, and disagreement-alone reaches 20.9 percent at a 10 percent budget, the strongest

result any practical strategy attains anywhere in the sweep. Action-entropy is uninformative throughout: entropy-topk gives -1.5 percent at a 5 percent budget and 1.9 percent at 10 percent under delay, statistically indistinguishable from and sometimes worse than random. Combining entropy with disagreement does not help and in fact hurts at higher budgets: at a 10 percent budget under delay, combined recovers 10.9 percent while disagreement-alone recovers 20.9 percent, so folding the uninformative entropy signal into the selector dilutes the one signal that carries information. The oracle-topk upper bound, which audits true malicious steps first, removes 37.8 percent of successful bad actions at a 5 percent budget and 79.9 percent at a 10 percent budget under delay, so a large majority of these attacks are in principle localizable within budget and the practical signals capture only a fraction of that headroom.

The per-environment breakdown at a 5 percent budget (pooled over attacker conditions) is uneven and dominated by base rates. Banking, which carries the most successful bad actions, is where the signal is clearest: combined removes 9.8 percent versus random, disagreement 7.3 percent, and the oracle 39.5 percent. Travel shows a weaker effect (combined 1.3 percent, disagreement 3.5 percent, oracle 21.9 percent). Workspace produced only about three successful bad actions per audit, too few to estimate a stable reduction; its combined figure of -5.3 percent is noise off a tiny denominator, and its oracle reaches 100 percent simply because a single audited step clears the environment. The full six-strategy curves across all budgets and both attacker conditions are plotted in `figure_main.png`, with the underlying values in `curve.csv`.

The significance of the combined-versus-random reduction is not uniform across the budget sweep, which we surface rather than lead past. Per `results/real/stats.json` the reduction is significant only at the 2 and 5 percent budgets (3.23 percent, $p = 0.017$; 4.76 percent, $p = 0.009$) and not at the 0.5, 1, or 10 percent budgets (0.56 percent, $p = 0.303$; 1.60 percent, $p = 0.134$; 3.00 percent, $p = 0.104$), so the effect is fragile, reaching significance only in a middle band of budgets and never approaching the pre-registered 30 percent at any budget. One clarification on how to read these figures: the successful-bad-action counts are an audit-simulation counterfactual, the number of labeled malicious steps a strategy fails to select at a given budget, and are distinct from the empirical injection-success rate of 70 of 690 tasks (about 10 percent); every percentage reduction we report is a relative change in the former, not in the latter.

Discussion

The headline is a control-evaluation caution rather than a new defense. A timing-adaptive attacker that delays its injected action to a low-suspicion moment measurably defeats suspicion-based auditing (a 17.0-point drop in attack-step suspicion, with suspicion-topk recovery falling from 6.8 percent under a naive attacker to essentially zero under delay), and the uncertainty-triggered auditing we proposed to fill that gap only weakly recovers: at most 20.9 percent reduction versus random from disagreement-alone, and 4.76 percent pooled from the combined signal, far short of the 30 percent we pre-registered. Control evaluations that assume random or static attack timing may therefore overstate the safety a monitoring-plus-auditing protocol delivers (Greenblatt et al. 2023; Wen et al. 2024), because the operating point they measure is one an adaptive agent can move away from, and the uncertainty signals that survive the move are weak.

The mechanism is legible. Action-entropy fails because the injected payload is deterministic: an agent complying with an injected instruction emits a stable tool call, so its semantic entropy over resampled actions stays low precisely at the attack step, which is the opposite of what an entropy-triggered audit needs. This is the concern the reference note on semantic entropy anticipated (Farquhar et al. 2024; Kuhn et al. 2023): instability tracks a model not knowing its answer, whereas a hijacked agent confidently knows the malicious action. Monitor disagreement, the variance of suspicion across monitor resamples, is the only signal that partially survives, because a well-timed attack can hold the monitor’s mean suspicion low while still leaving it internally inconsistent about the step. Combining the two by taking the more extreme of the z-scored signals does not help and hurts at higher budgets, because the uninformative entropy component displaces disagreement-selected steps from the audit set; a protocol that used disagreement alone would have been strictly better than

the combined rule we pre-registered. The large oracle headroom (79.9 percent reduction at a 10 percent budget under delay) shows the attacks are localizable in principle, so the gap is a weakness of these particular cheap signals, not evidence that budgeted auditing is hopeless.

Limitations and reviewer responses

The most important limitation, which we concede rather than hide, is that our adaptive attacker is only partially adaptive. The delayed condition instructs the agent to minimize monitor suspicion by timing, but it does not optimize against the action-entropy or disagreement signals we actually propose. An attacker aware of those signals could plausibly suppress disagreement as well, and our result therefore bounds robustness to a suspicion-lowering attacker, not to a fully signal-aware one; a signal-aware adaptive-attack condition is the natural and necessary follow-up, and until it is run our disagreement result should be read as an upper bound on what these signals buy. Ctrl-Z-style resampling-based control (Bhatt et al. 2025) is a candidate defense in that stronger threat model that we do not evaluate here.

Three narrower limitations qualify the numbers. We study a single untrusted agent (Qwen2.5-7B-Instruct) monitored by a single independent model family (Llama-3.1-8B-Instruct), so the magnitudes may not transfer to other capability gaps or monitor pairings. The workspace environment produced too few successful bad actions to support a per-environment estimate, so our per-environment claims rest mainly on banking and travel. And the combined-versus-suspicion comparison, while clearly failing to reach significance ($p = 0.559$), has a wide confidence interval; more seeds would tighten it, though the point estimate gives no reason to expect the 30 percent bar would be cleared with more data. We report the honest negative because it is the load-bearing result: uncertainty-triggered auditing, as instantiated with entropy and monitor disagreement, is not a reliable answer to timing-adaptive injected attacks, and disagreement’s weak survival is the one thread a follow-up should pull.

Three further points raised by reviewers deserve an explicit response. First, combined and disagreement auditing recover more under the delayed attacker than under the naive one (for instance 11.8 versus near-zero percent reduction at a 5 percent budget), which is counterintuitive for a defense described as stressed by delay. We read this as a headroom-and-mechanism effect rather than evidence the defense prefers stronger attackers: under the naive attacker the malicious step coincides with high, confident monitor suspicion, so the random and suspicion baselines already catch much of it and the monitor’s low resample variance leaves disagreement little to discriminate, whereas under delay the attack moves to low-suspicion steps where the monitor is internally uncertain and disagreement becomes comparatively more informative. Because the naive and delayed pools contain different numbers of successful bad actions, cross-condition reductions are not directly comparable, so this pattern should be read as suggestive, a base-rate confound we flag rather than resolve. Second, the workspace environment yields only about three successful bad actions per audit, so its per-environment reductions are noise; they enter the pooled headline only through the paired bootstrap over task instances, in which workspace contributes few malicious-step observations and therefore little weight, but we report its instability plainly rather than rest anything on its point estimates. Third, we test a single combination rule, the maximum of the z-scored signals, so the finding that combining hurts at higher budgets may be specific to that rule; a weighted or learned combination could in principle recover disagreement’s advantage without the entropy dilution, and we leave that to the signal-aware follow-up.

Reproducibility

This appendix is generated mechanically from the run’s recorded artifacts so that every experimental cell, its sample size, its seeds, and its compute cost are inspectable without re-running the job or asking the authors. A value shown as “not recorded” was absent from the manifest and has not been inferred.

The run comprises 1 recorded experiments. Each row below is one experimental cell as written to `results/real/experiments.json`.

Experiment	Model	Dataset	Mode	n (per cell)	Seed(s)	Key hyperparameters	GPU minutes
gen_trajectory	Qwen2.5-7B-Instruct	AgentDojo (workspace, banking, travel)	inference	690	not recorded	params=7B	280.85

Seed policy. No RNG seeds are recorded in the experiment manifest.

Cross-validation. No cross-validation fold fields are recorded in the manifest.

Statistical tests. Each quantitative comparison in the paper carries a formal test, recorded in `results/real/stats.json`.

Claim	Test	Statistic	p	95% CI	n	Seeds
combined-vs-random reduction at 0.5% audit budget	paired bootstrap over task instances	0.005610060886342	86342832	[-0.005025125628140741, 0.02061597830934082]	690	3
combined-vs-suspicion-topk at 0.5% audit budget	paired bootstrap over task instances	1.0600000000000002	0000023	[-1.8017500000000011, 4.5423749999999979]	690	3
combined-vs-random reduction at 1% audit budget	paired bootstrap over task instances	0.015979904868734655	68734655	[-0.004426285325161736, 0.03359374125045763]	690	3
combined-vs-suspicion-topk at 1% audit budget	paired bootstrap over task instances	1.8799999999999995	9999985	[-3.4302500000000067, 7.0]	690	3
combined-vs-random reduction at 2% audit budget	paired bootstrap over task instances	0.032285930245117864	45117864	[0.0042393769310184291, 0.058178843392367646]	690	3
combined-vs-suspicion-topk at 2% audit budget	paired bootstrap over task instances	3.7599999999999991	999991	[-4.0, 9.0]	690	3
combined-vs-random reduction at 5% audit budget	paired bootstrap over task instances	0.047619047619047526	19047526	[0.010426003638672764, 0.09948776631542065]	690	3

Claim	Test	Statistic	p	95% CI	n	Seeds
combined-vs-suspicion-topk at 5% audit budget	paired bootstrap over task instances	3.0	0.559	[-6.0, 13.704999999999961]	690	3
combined-vs-random reduction at 10% audit budget	paired bootstrap over task instances	0.029982363305096616	0.506	[-0.008140504481967896, 0.11267484239558692]	690	3
combined-vs-suspicion-topk at 10% audit budget	paired bootstrap over task instances	-11.0	0.345	[-21.104999999999993, 6.799999999999984]	690	3
delayed entropy-topk-vs-random reduction at 5% audit budget	paired bootstrap over delayed-condition task instances	-0.01469890943575165	0.488	[-0.044471644226846206, 0.038319594976198725]	345	3
manipulation check: delayed-vs-naive attack-step monitor suspicion	bootstrap over attack-step observations	-16.974982674982705	0.001	[-26.24349965349968, -7.553853083853118]	189	3

Compute. Total recorded GPU time across all experiments is 4.6808 GPU-hours on RTX 5090.

Code and data availability. A public code repository URL is not recorded in `project.yaml` ([links.github](#)). The per-example data of record that backs every reported number is provided under `results/real/` in the project repository: `curve.csv`, `steps.csv`, `tasks.csv`, `experiments.json`.

References

- Bhatt, Aryan, Cody Rushing, Adam Kaufman, et al. 2025. “Ctrl-z: Controlling AI Agents via Resampling.” *arXiv Preprint arXiv:2504.10374*.
- Bowman, Samuel R, Jeeyoon Hyun, Ethan Perez, et al. 2022. “Measuring Progress on Scalable Oversight for Large Language Models.” *arXiv Preprint arXiv:2211.03540*.
- Debenedetti, Edoardo, Jie Zhang, Mislav Balunovic, Luca Beurer-Kellner, Marc Fischer, and Florian Tramèr. 2024. “AgentDojo: A Dynamic Environment to Evaluate Attacks and Defenses for LLM Agents.” *arXiv Preprint arXiv:2406.13352*.
- Dubey, Abhimanyu, Abhinav Jauhri, Abhinav Pandey, Abhishek Kadian, et al. 2024. “The Llama 3 Herd of Models.” *arXiv Preprint arXiv:2407.21783*.

- Farquhar, Sebastian, Jannik Kossen, Lorenz Kuhn, and Yarin Gal. 2024. “Detecting Hallucinations in Large Language Models Using Semantic Entropy.” *Nature* 630 (8017): 625–30.
- Greenblatt, Ryan, Buck Shlegeris, Kshitij Sachan, and Fabien Roger. 2023. “AI Control: Improving Safety Despite Intentional Subversion.” *arXiv Preprint arXiv:2312.06942*.
- Greshake, Kai, Sahar Abdelnabi, Shailesh Mishra, Christoph Endres, Thorsten Holz, and Mario Fritz. 2023. “Not What You’ve Signed up for: Compromising Real-World LLM-Integrated Applications with Indirect Prompt Injection.” *ACM Workshop on Artificial Intelligence and Security (AISec)*.
- Hubinger, Evan, Carson Denison, Jesse Mu, et al. 2024. “Sleeper Agents: Training Deceptive LLMs That Persist Through Safety Training.” *arXiv Preprint arXiv:2401.05566*.
- Kuhn, Lorenz, Yarin Gal, and Sebastian Farquhar. 2023. “Semantic Uncertainty: Linguistic Invariances for Uncertainty Estimation in Natural Language Generation.” *International Conference on Learning Representations (ICLR)*.
- Liu, Yi, Gelei Deng, Yuekang Li, et al. 2023. “Prompt Injection Attack Against LLM-Integrated Applications.” *arXiv Preprint arXiv:2306.05499*.
- Team, Qwen. 2024. “Qwen2.5 Technical Report.” *arXiv Preprint arXiv:2412.15115*.
- Weij, Teun van der, Felix Hofstätter, Ollie Jaffe, Samuel F Brown, and Francis Rhys Ward. 2024. “AI Sandbagging: Language Models Can Strategically Underperform on Evaluations.” *arXiv Preprint arXiv:2406.07358*.
- Wen, Jiaxin, Vivek Hebbar, Caleb Larson, et al. 2024. “Adaptive Deployment of Untrusted LLMs Reduces Distributed Threats.” *arXiv Preprint arXiv:2411.17693*.
- Yao, Shunyu, Jeffrey Zhao, Dian Yu, et al. 2023. “ReAct: Synergizing Reasoning and Acting in Language Models.” *International Conference on Learning Representations (ICLR)*.
- Zheng, Lianmin, Wei-Lin Chiang, Ying Sheng, et al. 2023. “Judging LLM-as-a-Judge with MT-Bench and Chatbot Arena.” *Advances in Neural Information Processing Systems (NeurIPS)*.