

APEX-Accounting

Julien Benchek^{1*} Austin Bennett^{1*} Jasmin Kern^{1*} Ryan Stevens² René Sultan² Charis Ching¹ Hayley Popiel¹ Vaibhav Mittal¹ Felix Mercier¹ Brendan Foody¹ Bertie Vidgen¹

¹Mercor ²Ramp

*Equal contribution.

Contact apex@mercort.com for more information about APEX-Accounting.

Abstract

We introduce APEX-Accounting, a benchmark built by Mercor in partnership with Ramp, to assess whether frontier models can do the real work of accountants. Tasks include reconciling accounts, accruing expenses, posting transactions, and producing reports. The private eval set comprises 160 tasks, split across 10 worlds. Each world contains an accounting system, as well as spreadsheets, PDFs, and other files. Every task was authored and solved by experts in accounting and bookkeeping, who also wrote grading rubrics. Across nine frontier models, Claude-Fable-5 (Max) leads with 56.4% Mean Criteria@3, ahead of Muse-Spark-1.1 (xHigh) at 52.6%. No model scores more than 2.6% Pass⁸ (GPT-5.6-Sol (Max+Pro)) and the highest Pass@8 is 21.5% (Muse-Spark-1.1 (xHigh)). We experiment with increasing the token budget from \$1 to \$50 and observe an instance of Simpson’s paradox: scores increase as the token budget increases but within a given budget-constrained harness, scores are lower on tasks where the model spends more tokens. As APEX-Accounting is a closed benchmark, leaderboard evals can be run for any frontier model on request.

1 Introduction

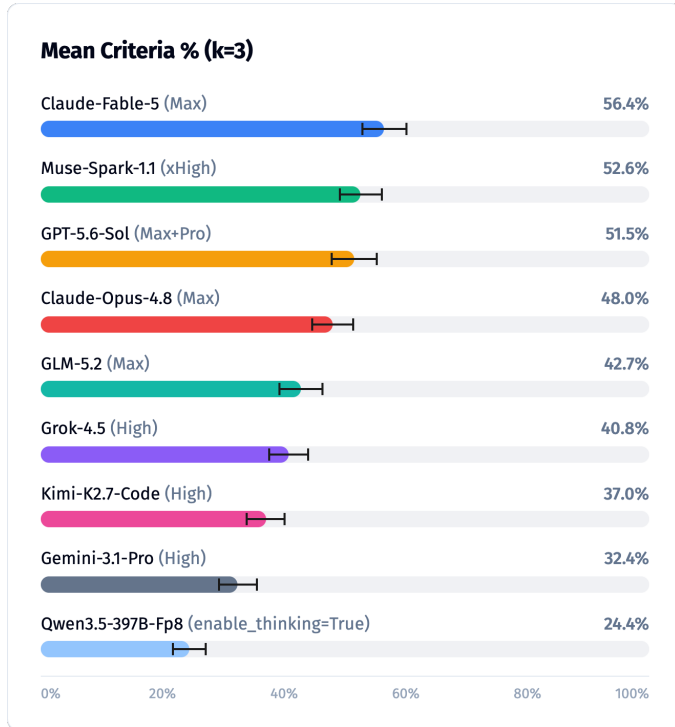
Globally, accounting generates approximately \$700 billion (Hughes, 2024) in revenue and the US alone employed close to 1.6 million accountants in 2024 (National Center for O*NET Development, 2025). Producing correct end-of-year accounts and reporting tax correctly is a requirement for every company, regardless of their size or industry. This sort of knowledge work is skilled and detail-oriented, and requires deep understanding of business context while also being repetitive, procedural, and document-heavy. AI models have long passed the profession’s certification exams: in 2023, GPT-4 with

10-shot prompting averaged 85.1% across the CPA, CMA, CIA, and Enrolled Agent exams (Eulerich et al., 2024). A Stanford survey of 277 accountants found that AI improved the quality of their work, despite them expressing concerns about data security and job stability (Choi and Xie, 2026), and a 2026 Harvard Business School analysis scored accounting at 0.51 on a 0–1 scale of how readily generative AI can substitute human labor for a given occupation (Chen et al., 2024) and (Azpúrua, 2026).

We introduce APEX-Accounting, a benchmark built by Mercor in partnership with Ramp. It comprises 160 tasks across 10 synthetically generated worlds and spans four task types: Reconciliation, Data Entry, Variance Analysis, and Schedules & Accruals. Each world is a self-contained company frozen at month-end close, and every task is authored and solved by

Acknowledgments. We thank all of the experts who contributed to the creation of APEX-Accounting. We also thank Yash Bharti for his support in creating the eval set, Arnav Garg for his research assistance, and members of the Mercor research team for reviewing the paper.

Figure 1: Performance of models on the APEX-Accounting leaderboard using Mean Criteria@3, with 95% task-bootstrap confidence intervals.



accounting experts. Model outputs are graded by an LM judge against task-specific rubrics consisting of binary, outcome-based criteria. We also open-source a public dev set world ($n = 10$ tasks), including prompts, rubrics, and task files.¹

Few benchmarks assess how AI models perform at real accounting work. AuditBench (Wang et al., 2025) pairs real S&P 500 financial statements with synthetic transactions to test inconsistency detection. FinMaster (Jiang et al., 2025) spans the financial pipeline with text-only, simulator-generated tasks, and AccountingBench (Penrose, 2025) evaluates agents on closing the books for one real SaaS company against a human CPA. OpenAI’s GDPval (Patwardhan et al., 2025) has five tasks (of $n = 220$) for accountants’ and auditors’ work in its open-sourced gold subset. None of these test the day-to-day bookkeeping and month-end close work

¹<https://huggingface.co/datasets/mercort/apex-accounting>

that fill most of a bookkeeper or staff accountant’s job at scale.

We evaluated 9 frontier models for the APEX-Accounting leaderboard, using the Loop harness, which implements the canonical while-loop-with-tools agent architecture (Goyal, 2025). On Mean Criteria@3, Claude-Fable-5 is the best-performing model with 56.4%, followed by Muse-Spark-1.1 at 52.6%. The highest score for Pass@8 is 21.5% (Muse-Spark-1.1), and Pass@8 at 2.6% (GPT-5.6-Sol).

We introduce a fine-grained taxonomy of failure categories, adapted to accounting work, and use it to annotate low-scoring trajectories from the three best-performing models. Their profiles are strikingly similar: reasoning failures dominate (59–79% of annotated failures per model). Information gathering and instruction following, which are often improved through better harnesses, account for only about a quarter combined. None of the annotated failures involve tool use. Two ablations point to the same: raising the per-task budget from \$1 to \$50 helps only models whose token appetite makes tight caps binding, and a purpose-built accounting harness shifts Mean Criteria@3 by just +1.2 pp on average.

Section 2 describes how worlds, files, and tasks were designed by accounting experts. Section 3 outlines our evaluation methodology, judge model, and metrics. Section 4 reports the leaderboard, the budget ablations, and the harness experiments. Section 5 reports top-performing models’ failure modes to assist future model and agent development.

2 Benchmark Design

APEX-Accounting comprises 160 tasks across 10 held-out synthetic company worlds, with 16 tasks per world. The public dev set contains one world and 10 tasks.

2.1 Expert Selection

APEX-Accounting was created by 42 accounting experts with a median of 11 years of career experience (mean 12.6). Backgrounds skew toward experience in corporate finance and accounting (35 of 42 experts).

Table 1: APEX–Accounting tasks by category ($n = 160$), with category definitions. File counts are the files required to solve the prompt, whether drawn from the shared world corpus or provided with the task.

Category	Definition	Tasks	Mean Criteria	Median Criteria	Mean Files	Median Files
Reconciliation	Tie two sources together, identify differences, and explain or correct the breaks.	61	14.61	11	7.49	7
Data Entry	Post or update transactions, journal entries, vendor bills, and invoices.	26	16.35	13	8.04	6
Variance Analysis	Compare actuals against budget, prior periods, or expectations and explain the drivers.	28	14.54	11	7.86	7
Schedules & Accruals	Build or update supporting schedules, calculate accruals, and carry the right balances into the close.	45	10.29	8	7.00	6
Overall	—	160	13.66	10	7.51	7

52.4% of experts have worked at a Big Four accounting firm (Deloitte, KPMG, EY, or PwC).

2.2 World Construction

Each world is a self-contained company, frozen at a specific month-end close. Each is based on a realistic business scenario and has its own entity type, chart of accounts, revenue model, and prior-period balances. Across their tasks, worlds draw on 73.1 unique required input files on average (see Table 6). These were scoped in partnership with Ramp to ensure a diverse distribution of company profiles across revenue, headcount, industry, and geography. Every company follows U.S. GAAP accrual-basis accounting.

The worlds were built in four stages. First, we ran a cross-world scoping pass to assign each world a high-level profile, balanced across the full set. Second, we gave each world a detailed specification: a description of every file, every task to be built (including the deliberate traps each is meant to catch), and a trap register cataloging every seeded contradiction. Third, we iterated on a single reference document until its formatting and style were realistic, then derived a per-world style guide. We validated all files against the spec and style guide so that they read as coming from one company. Most world files were synthetically populated, always under expert specification and review. Every document is novel and screened against public sources, so no world can be found online or memorized in advance. On average, 14 experts worked on each world.

2.3 Task Creation

Each task is associated with a single world. It consists of a prompt, a set of required input files, a golden response, and a grading rubric.² The required output for every task is a message sent in the console. On average, there are 7.5 required input files per task that the model needs to find and read. Most of these belong to the shared world file corpus, but about half of the tasks provide additional task-specific files, which are placed in the filesystem alongside the world files when the task begins. Across the held-out set, 83 of 160 tasks (52%) include at least one task-specific file (median 2, mean 2.7 per such task). Task-specific files are most common in Reconciliation (33/61 tasks), Schedules & Accruals (29/45), and Variance Analysis (17/28), and rare in Data Entry (4/26).

Prompts mimic realistic requests that an accountant would receive on the job. They are concise and include a clear expectation of the final output but do not explain the methods where a competent accountant would already know them, nor give the file names unless they are unintuitive or unexpected.³ Each task is in one of four categories: Reconciliation, Data Entry, Variance Analysis, and Schedules & Accruals. Table 1 defines each category and reports its task count and per-task criteria and file statistics.

²Section A shows one dev set task end-to-end — the prompt, two representative rubric criteria, and a condensed evaluated trajectory — to make the task and rubric structure concrete.

³Quality control processes applied during dataset construction, including expert review, automated QC, and independent re-solve baselining, are given in Section B.

Table 2: Requirements every rubric criterion must satisfy.

Requirement	Definition
Self-contained	Gradable without reference to other criteria.
Easy to interpret	One fact or judgment per criterion, so partial credit is never lost to compound requirements.
Aligned with prompt	Grading only what the prompt asks.
Outcome-based	Grading only the final answer, not intermediate steps. We apply acceptable ranges to handle legitimate rounding differences.

Experts created a grading rubric for each task, comprising binary (Pass/Fail) unweighted criteria. On average, there are 13.7 criteria per task. To ensure grading fairness, each criterion is self-contained, easy to interpret, aligned with the prompt, and outcome-based (Table 2). Experts also created a golden response for each task, which represents a top-tier industry-quality output, scoring 100% against the rubric.

3 Experimental Setup

3.1 Models

We evaluate 9 frontier models for the APEX-Accounting leaderboard: Claude-Fable-5, Muse-Spark-1.1, GPT-5.6-Sol, Claude-Opus-4.8, GLM-5.2, Grok-4.5, Kimi-K2.7-Code, Gemini-3.1-Pro, and Qwen3.5-397B-Fp8. Each model independently executes each task 8 times, resulting in a total of 11,520 trajectories. A trajectory comprises the sequence of steps a model takes, where each step is one model turn. It can include output tokens, reasoning tokens, and one or more tool calls. We allow a maximum of 500 steps and 5 million tokens per task.⁴

3.2 Harnesses

We evaluate models using two harnesses: a standard Loop Harness and the Ramp Harness. Both

⁴Provider, context window, maximum output, thinking-effort configuration, and wall-clock time per run for every model are reported in Section D.

harnesses are implemented and executed within Archipelago, our internal framework for running agent evaluations at scale.⁵ Leaderboard results use the Loop Harness. The Loop Harness repeatedly passes context into a model to solve a task. The Ramp Harness, however, expands on a ReAct-style harness (Yao et al., 2023) with parallel subagents. This Ramp Harness was built in partnership with Ramp to mirror their specialized harness for accounting work. Figure 2 shows a schematic of the Ramp Harness. Specifically, there are six properties designed to mirror the constraints of a real accounting deployment:

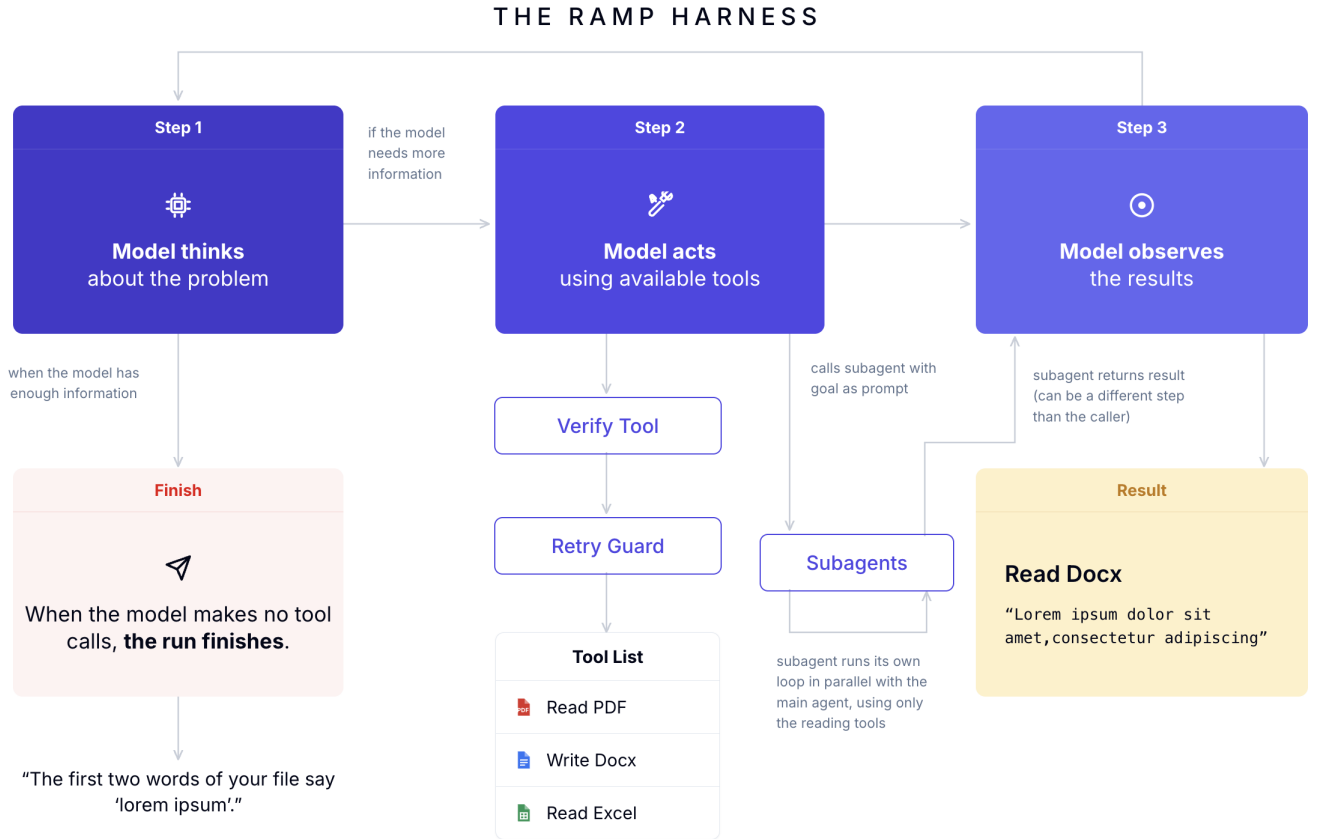
1. *Retry-awareness*: Before executing a tool call, the harness checks whether an identical call has already been made; if so, it blocks the call, avoiding infinite loops and inefficient repeated tool use.
2. *Tool allowlist*: Each agent is exposed to only the set of tools needed to solve the problem.
3. *Tool validation*: Tool arguments are parsed and validated before execution.
4. *Unavailable-tool recovery*: If the model requests a tool it cannot use, the harness returns a model-visible tool error rather than executing anything.
5. *Read/write execution policy*: Read-only tools can run in parallel through subagents, while write tools run serially to avoid conflicting side effects.
6. *Subagent delegation*: Some tool calls can launch a separate agent loop to complete a narrower sub-task, then return its result to the parent loop.

For the canonical leaderboard evals we restrict agents to 5 million tokens and 500 steps. We created versions of the harnesses that tell the model at each step how many steps and tokens remain. If the step or token limit is reached, the model is told to submit a final answer; failure to do so is scored zero. If the token limit is reached we alert the model that the budget is exhausted and give it one additional turn to submit the answer. Any subagents that are currently active are canceled without being allowed to return. This ensures graceful stopping without a significant advantage beyond the original token budget.

Additionally, we conducted an experiment for Claude-

⁵<https://github.com/Mercor-Intelligence/archipelago>

Figure 2: Ramp Harness schematic.



Fable-5, GPT-5.6-Sol, and Gemini-3.1-Pro, where we compared unlimited token spend against the 5 million token cap. For all three models, the 5 million token limit scores were slightly higher, supporting this design choice. The harness reports remaining tokens at every step, and a finite, visible budget pushes the model to consolidate intermediate results and commit to an answer, whereas unlimited runs keep exploring past their best candidate and may dilute the final submission with accumulated context.

3.3 Tool Use

Both harnesses were given 91 operational tools to choose from: filesystem, docx, pdf, excel, mail, code execution, and accounting software. Additionally, the Ramp Harness was given 11 meta-tools to facilitate the run (e.g. subagent operations). Operational tools aid the model’s analysis through read and write functions. The first six categories are stan-

dard tool families for agent harnesses, while the accounting-software tools supplement them, offering a lightweight version of common accounting operations.

3.4 Grading Methodology

DeepSeek-v4-Flash, configured with temperature at 0.1, high reasoning, and a GEPA-optimized (Agrawal et al., 2026) grading prompt, is the grading judge. To grade each criterion the judge receives the task prompt, the criterion text, and the model’s final output, but not the trajectory log. It returns a binary score (Met, Not Met) along with a concise free-text explanation.

To select DeepSeek-v4-Flash as the grading judge, we sampled 40 trajectories from Claude-Fable-5, GPT-5.6-Sol, and Gemini-3.1-Pro. This resulted in a total of 120 trajectories, with 1,687 criteria. Three

Table 3: Confusion matrix for the judge model (DeepSeek-v4-Flash), evaluated against majority-vote human ground truth ($n = 1,687$ criteria, one majority-vote label per criterion).

	Predicted: Met	Predicted: Not Met
Actual: Met	799	21
Actual: Not Met	28	839

experts independently annotated every criterion for each trajectory, and we took the majority vote across the 3 annotators as ground truth for that criterion; raw inter-annotator agreement is 92.8%, with 89.2% of criteria labeled unanimously. Fleiss’ $\kappa = 0.857$, indicating strong agreement beyond chance.⁶

We used the ground truth dataset to assess 10 LM judges, including 7 frontier models (Claude-Opus-4.8 [reasoning = high/low], Gemini-3.1-Pro [reasoning = high/low], Gemini-3-Flash, GPT-5.5 [reasoning = high/low]) and 3 open-weight models (Kimi-K2.6, DeepSeek-v4-Flash [reasoning = high], DeepSeek-v4-Pro [reasoning = high]). DeepSeek-v4-Flash was selected because its performance is within noise of the best frontier judge, Claude-Opus-4.8 (High) ($F_1 = 0.970$ for both), while being open-weight. Table 3 shows a confusion matrix for the 1,687 majority-vote ground truth labels. The judge achieves 97.1% overall accuracy, with 96.6% precision and 97.4% recall on the Met class.⁷

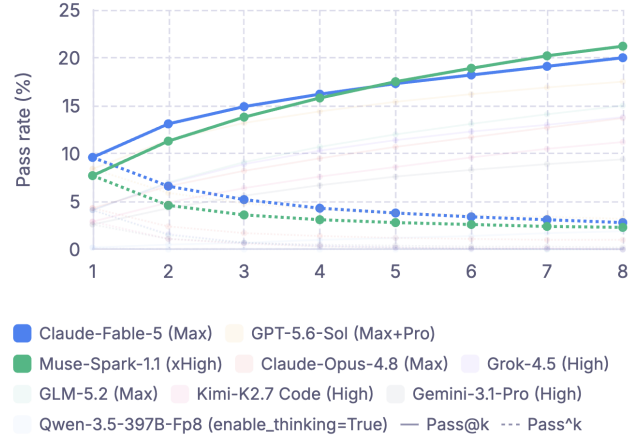
3.5 Metrics

We report four metrics: Mean Criteria@3, Pass@1, Pass@k, and Pass^k (for $k = 1$ through 8). Mean Criteria@3 is our primary leaderboard metric: the percentage of rubric criteria met, averaged across 3 of a model’s 8 runs per task, chosen lexicographically based on task ID strings, which essentially make up a random sample. We use Mean Criteria@3 because per-criterion partial credit separates models that complete most of a task from models that fail outright. Averaging over 3 runs rather than all 8 reduces run-to-run variance at a lower

⁶Table 11 (Section E) reports the full inter-rater reliability values.

⁷Section E reports judge agreement broken out by solver model and shows no per-model inflation; the full comparison across all 10 candidate judges is reported in Table 9 (Section E).

Figure 3: Pass@k versus Pass^k by model, $k = 1$ through 8, capability ceiling versus consistency floor. Claude-Fable-5 (Max) and Muse-Spark-1.1 (xHigh) scores are bolded for visibility.



computational cost.⁸

Beyond the leaderboard, we report Pass@8 as an indicative ceiling on current frontier agent capabilities. It measures whether a model passes a task at least once across 8 attempts. We also report Pass⁸ (Yao et al., 2025), which measures whether a model produces a fully correct output in each of its 8 attempts. We compute 95% confidence intervals via task-level bootstrapping, resampling individual tasks with replacement across 10,000 resamples.

4 Results

4.1 Leaderboard

Figure 1 shows the full APEX-Accounting leaderboard.⁹ Claude-Fable-5 performs best at 56.4%, ahead of Muse-Spark-1.1 (52.6%); GPT-5.6-Sol (51.5%) and Claude-Opus-4.8 (48.0%) follow close behind. The remaining models—GLM-5.2, Grok-4.5, Kimi-K2.7-Code, Gemini-3.1-Pro, and Qwen3.5-397B—trail, with Qwen3.5-397B lowest at 24.4%.

⁸Runs that fail because of the model are scored as zero, while runs that fail for reasons outside the model’s control are excluded from every metric. Section F gives the full convention and its effect on task coverage.

⁹Table 13 (Section H) reports each model’s Mean Criteria@3, Pass@1, Pass@8, and Pass⁸ in full, with 95% bootstrap confidence intervals.

Table 4: Mean Criteria@3 by task category, for all 9 models.

Model	Reconciliation	Data Entry	Variance Analysis	Schedules & Accruals
Tasks (N)	61	26	28	45
Claude-Fable-5 (Max)	59.4%	58.6%	56.8%	51.0%
Muse-Spark-1.1 (xHigh)	52.8%	57.9%	57.2%	46.4%
GPT-5.6-Sol (Max+Pro)	55.3%	52.4%	50.2%	46.5%
Claude-Opus-4.8 (Max)	53.2%	48.7%	51.9%	38.1%
GLM-5.2 (Max)	44.1%	50.0%	51.7%	31.0%
Grok-4.5 (High)	47.4%	38.4%	45.5%	30.2%
Kimi-K2.7-Code (High)	40.5%	40.1%	44.3%	25.9%
Gemini-3.1-Pro (High)	34.6%	28.9%	36.2%	28.9%
Qwen3.5-397B-FP8 (True)	28.5%	20.9%	24.9%	20.6%

Because the leaderboard contains 9 models, we conduct 36 pairwise significance tests; we control the false discovery rate at 5% with the Benjamini–Hochberg procedure over per-task Mean Criteria@3 differences.¹⁰ All gaps spanning two or more leaderboard ranks remain significant after correction; only two adjacent-rank pairs are statistically indistinguishable.

4.2 Performance by Task Category

Table 4 breaks down each model’s Mean Criteria@3 by question category (given in Table 1). Claude-Fable-5 leads three of the four categories (Reconciliation, Data Entry, and Schedules & Accruals) while Muse-Spark-1.1 edges ahead on Variance Analysis (57.2% vs. 56.8%). It sits within a point of Claude-Fable-5 on Data Entry, consistent with the two models’ close aggregate scores. Schedules & Accruals is the hardest category, with every model scoring lowest (Gemini-3.1-Pro ties its Data Entry score) and trailing their highest-scoring categories by 7–21 percentage points, reflecting the category’s multi-step, judgment-heavy nature.

¹⁰Section G reports the full pairwise comparison matrix and which leaderboard gaps remain significant after correction using paired t-tests. Section H compares each model’s leaderboard performance against the public dev set world.

4.3 Pass@8 and Pass^8

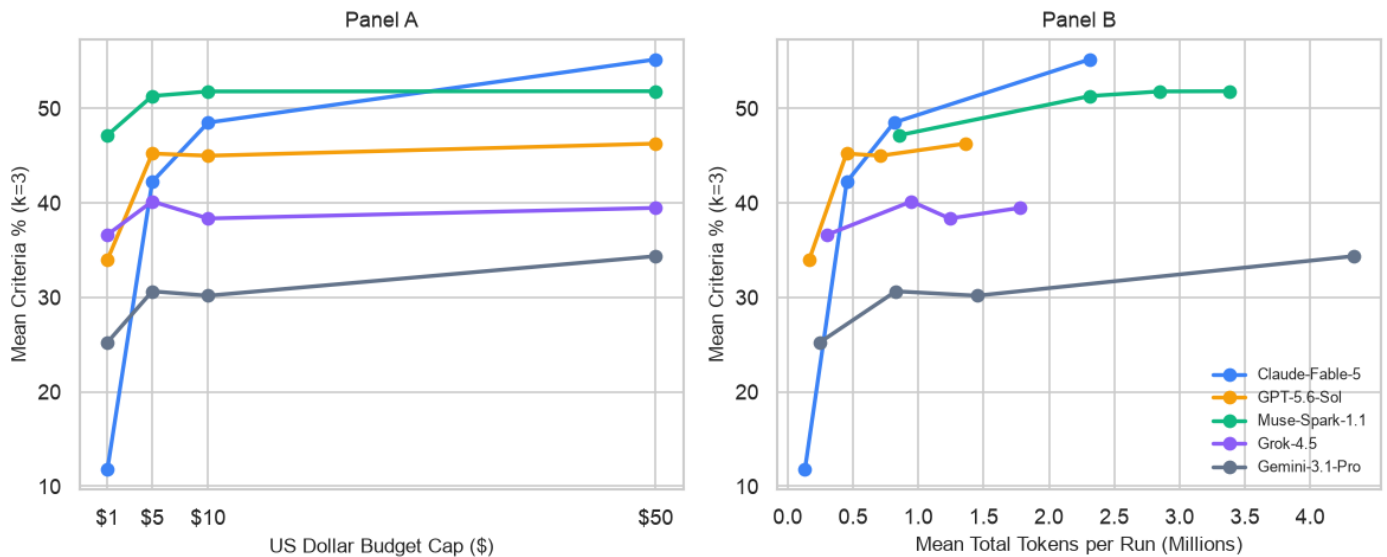
Pass@8 and Pass^8 measure models’ practical capability ceiling and consistency, respectively. The capability ceiling shows the potential maximum performance of models, whereas consistency is often the most important consideration for real-world deployment of AI models. Figure 3 plots Pass@k against Pass^k for $k = 1$ through 8. The two metrics rank models differently: Muse-Spark-1.1 achieves the highest Pass@8 (21.5%), narrowly ahead of Claude-Fable-5 (20.1%), whereas GPT-5.6-Sol scores the highest Pass^8 at 2.6%. Consistency is low across the board, no model exceeds 2.6% on Pass^8, demonstrating how far agents are from reliably completing accounting work end to end.

4.4 Cost Ablation Experiments

We tested five models—Claude-Fable-5, GPT-5.6-Sol, Gemini-3.1-Pro, Grok-4.5, and Muse-Spark-1.1—using the Loop Harness and Mean Criteria@3 at dollar token budgets of \$1, \$5, \$10, and \$50 per task. Standardizing on cost accounts for the different ways that models price and tokenize. To calculate the budget per model, we assume a 90:10 split of input to output tokens using the higher cost tier listed by model providers (prices can vary by token count).¹¹ The harness tells the model at each step

¹¹Section I details this per-model dollar-to-token conversion (Table 15) and shows an example system-prompt excerpt disclosing the resulting token budget to the model (Figure 12).

Figure 4: Panel A. For five models, the Mean Criteria@3 versus the per-task dollar budget cap (set at \$1, \$5, \$10 and \$50). Panel B. The Mean Criteria@3 versus the mean total tokens used at each budget cap. At higher budget caps models do not, on average, fully utilize the tokens available to them.



an updated estimate of its remaining tokens.

Because every task is run at every budget, we measure each model's gain as a paired per-task difference, which removes task difficulty from confounding the comparison.¹² Raising the dollar budget leads to higher scores, but the impact differs sharply across models (Figure 4). The impact is greatest on models whose high token cost means that they have very few tokens at the lower budgets, such as Claude-Fable-5. Claude-Fable-5 scores 11.8% at \$1—scoring exactly zero on 38.8% of its tasks, more than triple any other model—and increases to 55.2% at \$50. This is a gain of +43.4 pp, with significant improvements at every budget step. Muse-Spark-1.1 is the opposite case: it leads at the tightest cap (47.2%) yet gains only +4.7 pp across the whole range. The consequence is that at \$50, Claude-Fable-5 spends \$32.36 per run on average against Muse-Spark-1.1's \$5.25, yet their mean criteria scores are within 4 percentage points.

Outside of the \$1 cap, models typically do not fully utilize the token limits. There is only a moderate

¹²Table 16 reports all paired budget contrasts, with Benjamini-Hochberg-corrected significance.

increase in token usage when going from \$10 to \$50 and, on average, models utilize just 64.7% of the maximum budget available at \$50. This cap constrains just 0% to 16.9% of runs across the five models.¹³

Within a fixed budget, tasks where models use more tokens are not associated with higher scores. After Benjamini-Hochberg correction, every statistically significant within-cap-spend score association is negative. This presents an example of Simpson's Paradox, wherein raising the budget raises a model's scores, yet within any fixed budget harness, the tasks on which a model spends more tokens score lower. This is because task difficulty confounds the result. Harder tasks use more tokens but still score lower. Figures 5 and 6 show the per-task spend-score scatter with fitted within-cap slopes and the slope summary.

¹³Section J reports the per-cell within-cap slopes, Table 17 reports the within-cap levels slope b (change in Mean Criteria@3 per dollar of realized spend) for each model×budget cell.

Figure 5: Per-task Mean Criteria@3 versus realized spend per run (USD), by model, at each dollar budget (90:10 input:output split). Lines are per-model within-cap OLS fits, annotated with the slope b .

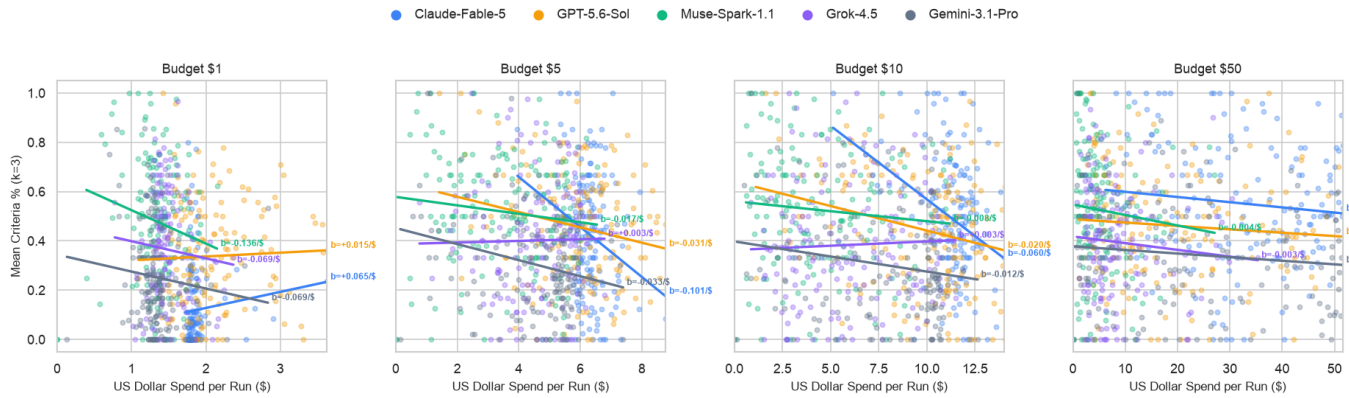
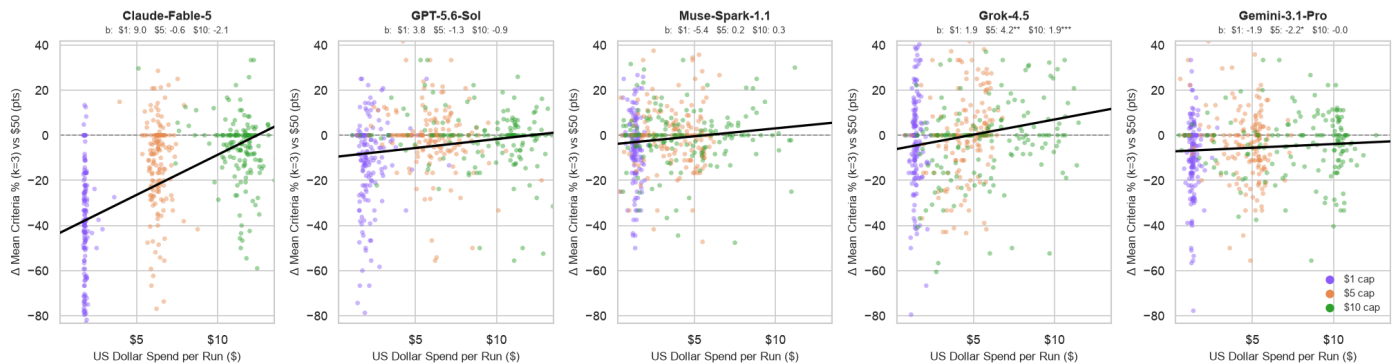


Figure 6: Task-level change in Mean Criteria@3 relative to the \$50 baseline, versus realized spend per run, by model. Points are colored by budget cap (\$1, \$5, \$10); black lines are pooled OLS fits, with per-cap slopes b annotated above each panel. X-axis and Y-axis truncated to middle 98% of values for visual clarity.



4.5 Harness Experiments

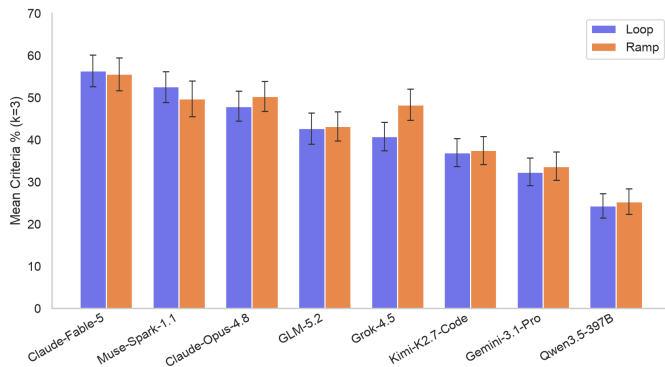
We compare performance of the Loop Harness and the Ramp Harness for each model. Harness choice does not make a substantial impact on Mean Criteria@3. Averaged over the eight models run under both harnesses, the raw result shift (Ramp – Loop) is +1.2 pp, and just +0.3 pp excluding one outlier. Grok-4.5’s gain is the only shift that survives multiple-testing correction; for every other model, the per-harness CIs overlap. Grok-4.5 gains +7.5 pp (40.8 to 48.3), overtaking GLM-5.2 to place fourth among the paired models. The purpose-built harness does not systematically favor the strongest models either: Claude-Fable-5 and Muse-Spark-1.1 both score *lower* under the Ramp Harness (–0.8

and –2.8 pp), while Claude-Opus-4.8 gains +2.4 pp and edges past Muse-Spark-1.1 (based on raw table difference). Switching harness is therefore not a free performance lift: on APEX-Accounting the model had more of an impact than the underlying harness. Harness scores are shown in Figure 7.¹⁴

Only the Ramp Harness allows subagent-delegation tools, and models differ sharply in how much they use them—from Claude-Fable-5 (4.9 subagents per task, 89% of trajectories) to Gemini-3.1-Pro and Qwen3.5-397B (0.2–0.4 per task, 12–13%). Delegation propensity does not line up with harness gains: GLM-5.2 delegates nearly as often as Claude-Fable-5 yet

¹⁴Section K reports the full comparison, including Pass@1 and corrected significance for every paired contrast.

Figure 7: Mean Criteria@3 by model under the Loop Harness versus the Ramp Harness ($k = 3$, $n = 160$). Error bars are task-bootstrap 95% CIs; GPT-5.6-Sol has no Ramp Harness run.



both shift by less than a point, while the one significant Ramp Harness improvement (Grok-4.5) comes from a model that delegates in 69% of its trajectories. Delegation propensity reads more as a model trait than a mechanism that raises scores.¹⁵

5 Failure Analysis

We developed a taxonomy of failure categories for APEX-Accounting, working with accounting experts who reviewed low-scoring trajectories.¹⁶ A substantial part of this analysis was spent verifying that failures were real. For each trajectory, experts manually checked the rubric criteria scores to confirm that the model actually erred. Any issues were corrected during production. Once the taxonomy was finalized, we sampled trajectories from the three best-performing models with low mean scores ($n=74$) and categorized each one using an LM judge.

5.1 Failure Modes of the Best-Performing Models

The three best-performing models fail in very similar ways. Figure 8 compares their L1 failure profiles over annotated failures ($n = 24$ for Claude-Fable-5, 28 for GPT-5.6-Sol, 22 for Muse-Spark-1.1). All 74 failures fall into just four of the taxonomy's seven L1 categories, and reasoning failures dominate every

profile: 79% of annotated failures for Claude-Fable-5, 75% for GPT-5.6-Sol, and 59% for Muse-Spark-1.1. By contrast, the modes that agent harnesses most directly targets are rare: not a single annotated failure involves tool use, and information-gathering and instruction-following failures combined account for 17% of Claude-Fable-5's failures and 21% of GPT-5.6-Sol's. Muse-Spark-1.1 is the outlier at 36% combined, the only model whose failures spread meaningfully into incomplete search and dropped requirements. Top models reliably find the right inputs; they make mistakes when applying accounting judgment and multi-step logic to them.¹⁷

Given that reasoning failures are the modal category, Figure 9 breaks them down by L2 subclass. Two subclasses dominate: non-numeric reasoning failures (22 of the 74 failures) and data handling errors (16 of 74) jointly account for just over half of all annotated failures. We observe one pattern repeatedly: the model retrieves the correct documents, often computes the correct intermediate values, then substitutes the wrong basis, source, or authorization logic downstream. For example, Claude-Fable-5 treats a program director's report as written donor authorization and releases a restricted \$15,000 balance; GPT-5.6-Sol shrinks a fixed room-night denominator by out-of-order rooms; Muse-Spark-1.1 re-normalizes allocation weights after computing the correct allocation. Forgetting the information, deriving a correct intermediate result and then dropping or contradicting it downstream, happens in all three profiles (7 cases) and is GPT-5.6-Sol's third-largest reasoning failure mode.

The skews differ: Claude-Fable-5 has more non-numeric reasoning failures (53% of all its reasoning failures), whereas GPT-5.6-Sol tilts toward data handling and forgetting failures, and Muse-Spark-1.1 splits evenly between data handling and non-numeric reasoning failures. Some tasks cause the same failures in multiple models: on a tenant-onboarding task, Claude-Fable-5 and GPT-5.6-Sol both report the new lease's square footage as the entity-wide total; and on a WIP-to-finished-goods

¹⁵Section K.1 reports per-model spawn rates.

¹⁶The full taxonomy, with definitions for every L1 category and L2 subclass, is in Section M.

¹⁷Full per-model sunburst breakdowns are in Section L, with exact L2 subclass counts in Table 19; Section M defines each label.

Figure 8: L1 failure-mode counts for the three best-performing models, over annotated failures ($n = 24/28/22$ for Claude-Fable-5 (Max), GPT-5.6-Sol (Max+Pro), and Muse-Spark-1.1 (xHigh)). Reasoning failures dominate every profile; no annotated failure involves tool use.

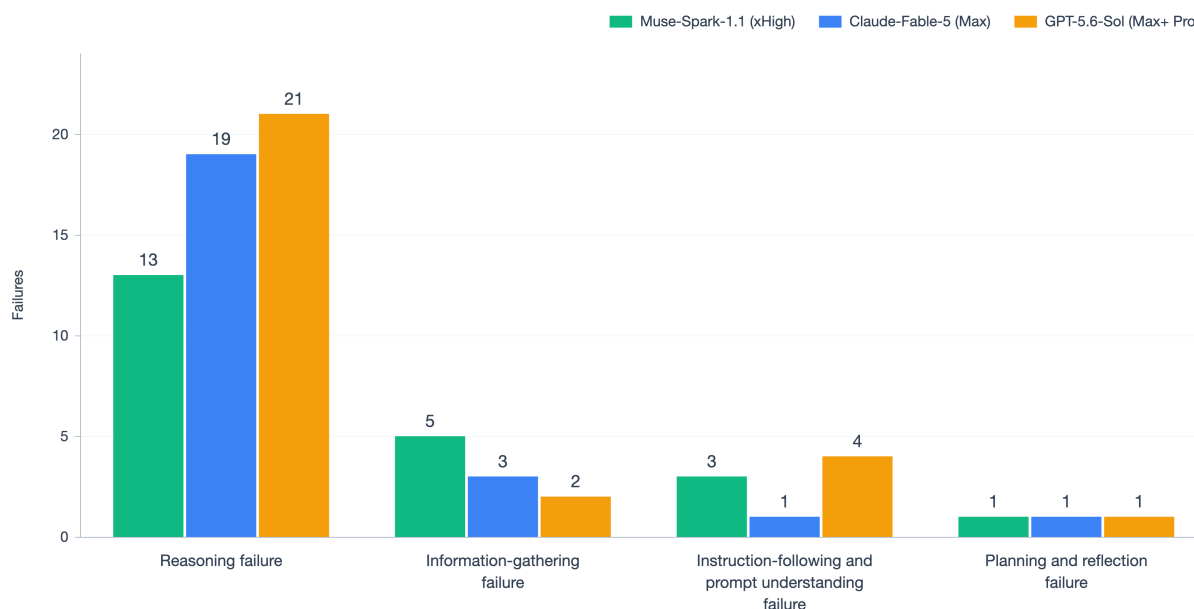
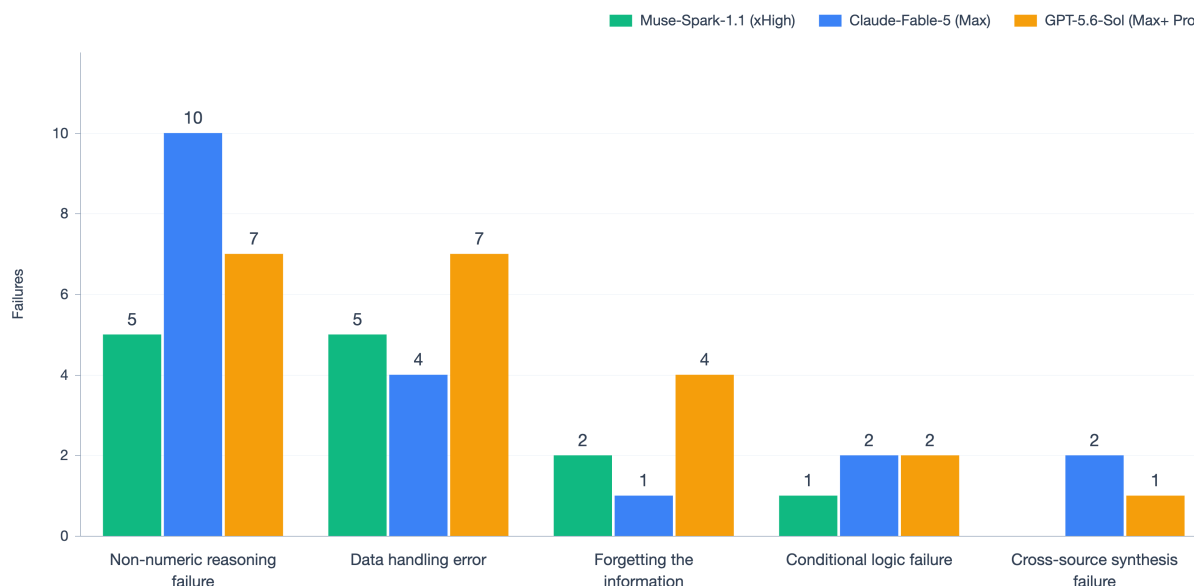


Figure 9: L2 subclass counts within Reasoning failure for the three best-performing models. Non-numeric reasoning failures and data handling errors dominate; forgetting the information recurs across all three.



transfer, GPT-5.6-Sol and Muse-Spark-1.1 both post an incorrect cost variance, taken from the same actual-cost source.

5.2 Per-Task Headroom

We compare models' task win rates head-to-head. For each pair, Figure 10 counts the tasks on which the row model outscores the column model. Because we report gross wins rather than a net difference, the matrix shows headroom in both directions. For example, Claude-Fable-5 outscores Muse-Spark-1.1 on 16 tasks while Muse-Spark-1.1 wins 18 back, a tenth of the benchmark going in each direction. On the other 126 tasks they tie. Wins over the leader extend well down the roster: Grok-4.5, GLM-5.2, and Kimi-K2.7-Code each beat Claude-Fable-5 on 10 tasks despite trailing it by 13+ points on mean score. Muse-Spark-1.1 posts the largest single-pair headroom (32 tasks over Qwen3.5-397B, a fifth of the benchmark) and wins more tasks than it loses against every other model despite trailing Claude-Fable-5 on mean score: Muse-Spark-1.1 wins more tasks, Claude-Fable-5 wins by larger margins. 93 of the 160 tasks (58%) are not fully solved by any model in any run.

6 Limitations

APEX-Accounting has several limitations. Its four task categories (Reconciliation, Data Entry, Variance Analysis, and Schedules & Accruals) cover close-cycle bookkeeping work; we leave out tax, audit, consolidation, multi-entity and multi-currency work, and external reporting. We also designed tasks to be completed end-to-end and excluded any tasks that require a human-in-the-loop. This removes a skill that real staff accountants exercise constantly; knowing when to ask a clarifying question rather than proceed on an assumption. This would increase realism if addressed in the task design. With only 160 held-out tasks across 10 worlds, statistical power is limited. After Benjamini-Hochberg correction, 34 of the 36 pairwise Mean Criteria@3 differences remain significant, but two adjacent-rank pairs are statistically indistinguishable (Section G), so the exact ordering at those boundaries should not be over-read.

Tasks were selected by filtering from a pool of worlds, with tasks already quality controlled for realism and diversity, based on three frontier models (Claude-Opus-4.8, GPT-5.5, and Gemini-3.1-Pro Preview) achieving low scores when graded

(Section B.1). We adopted this design to ensure that only worlds with challenging tasks are selected for the benchmark. However, it introduces a risk that the models used to filter the tasks, as well as others in the same family, have depressed scores; and the effect does not transfer uniformly to models from other families. We do not believe that the leaderboard ordering is an artifact of the filter: we selected tasks based on filtering over all three models rather than indexing on any single family's idiosyncratic failures. Further, Claude-Fable-5, the successor of filtering model Claude-Opus-4.8, ranks first, and GPT-5.6-Sol, the successor of GPT-5.5, ranks third. Thus, being selected against does not by itself determine a family's rank.

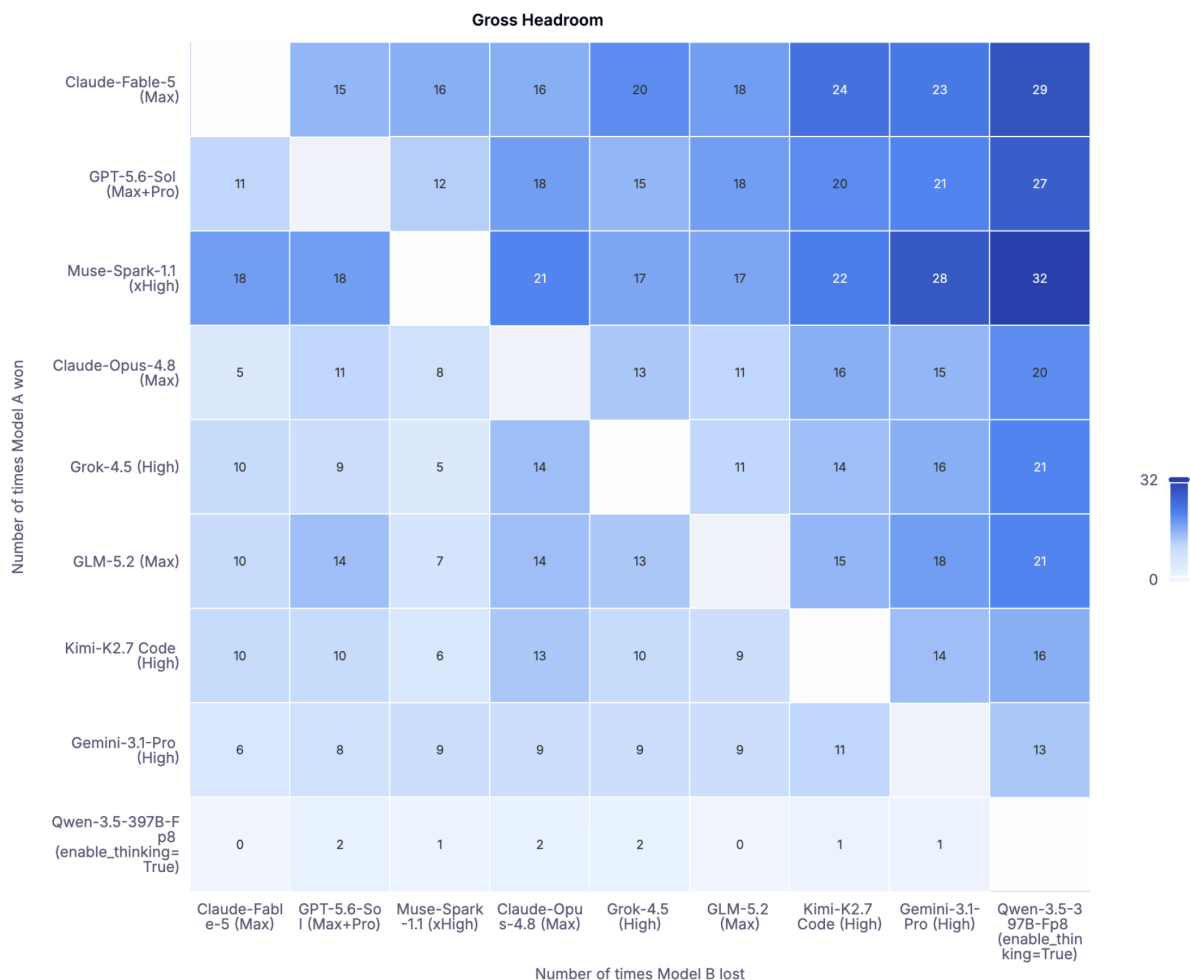
As with any LM-as-a-judge, grading is subject to judge error. We validated DeepSeek-v4-Flash extensively before deployment, benchmarking 10 candidate judges against 1,687 majority-vote expert labels and optimizing its prompt. It reaches 97.1% agreement with the expert ground truth ($F_1 = 0.970$), which is very strong but not perfect fidelity. The judge makes both false-positive and false-negative errors (precision = 0.966, recall = 0.974 on the Met class), so absolute scores differ slightly from expected ground truth labels. Because these residual errors are small and are applied uniformly across all evaluated models, they are unlikely to materially affect relative rankings.

7 Conclusion

APEX-Accounting assesses whether frontier models can perform the real work of accountants: finding and reading the right files, applying accounting judgment, and carrying multi-step analyses through to a final close. The best model reaches 56.4% Mean Criteria@3, no model exceeds 2.6% Pass^8, and 58% of tasks are not fully solved by any model in any run. An agent that is usually correct but not reliably so cannot yet close the books unsupervised.

Our failure analysis shows that the best-performing models can often retrieve the right inputs; they fail by mishandling multi-step reasoning over those inputs: substituting the wrong basis or authorization logic, and dropping correct intermediate results before the

Figure 10: Gross headroom for each model pair: Gross headroom counts tasks where the row model achieved a perfect pass on at least one run but the column model never did; only tasks attempted by the column model are included, partial passes receive no credit, and ties are excluded. Models are ordered by leaderboard rank.



final answer. Tool call errors are not widespread. As such, we anticipate that progress will come less from improving harnesses than improving the models themselves: accounting-specific training that instills the discipline to carry results through, surface contradictions in the documents, and refuse to post an entry the evidence does not support.

References

Lakshya A. Agrawal, Shangyin Tan, Dilara Soylu, Noah Ziems, Rishi Khare, Krista Opsahl-Ong, Arnav Singhvi, Herumb Shandilya, Michael J. Ryan, Meng Jiang, Christopher Potts, Koushik Sen, Alexandros G. Dimakis, Ion Stoica, Dan Klein, Matei Zaharia, and Omar Khattab. 2026. [GEPA:](#)

[Reflective prompt evolution can outperform reinforcement learning.](#) In *International Conference on Learning Representations (ICLR)*.

Ana Elena Azpúrua. 2026. [Enhance or eliminate? How AI will likely change these jobs.](#) Harvard Business School Working Knowledge.

Wilbur Xinyuan Chen, Suraj Srinivasan, and Saleh Zakarinia. 2024. [Displacement or complementarity? The labor market impact of generative AI.](#) Working Paper 25-039, Harvard Business School.

Jung Ho Choi and Chloe L. Xie. 2026. [Human + AI in accounting: Early evidence from the field.](#) *Journal of Accounting Research*, 64(3):1333–1373.

Marc Eulerich, Aida Sanatizadeh, Hamid Vakilzadeh, and David A. Wood. 2024. [Is it all hype? ChatGPT's performance](#)

and disruptive potential in the accounting and auditing industries. *Review of Accounting Studies*, 29(3):2318–2349.

Ankur Goyal. 2025. [The canonical agent architecture: A while loop with tools](#). Braintrust Blog.

Patrick Hughes. 2024. [The takeaway: A Q&A with Patrick Hughes on private equity’s rising interest in accounting firms](#). Houlihan Lokey Insights.

Junzhe Jiang, Chang Yang, Aixin Cui, Sihan Jin, Ruiyu Wang, Bo Li, Xiao Huang, Dongning Sun, and Xinrun Wang. 2025. [Finmaster: A holistic benchmark for mastering full-pipeline financial workflows with llms](#).

National Center for O*NET Development. 2025. [13-2011.00 — accountants and auditors](#). O*NET Online.

Tejal Patwardhan, Rachel Dias, Elizabeth Proehl, Grace Kim, Michele Wang, Olivia Watkins, Simón Posada Fishman, Marwan Aljubei, Phoebe Thacker, Laurance Fauconnet, Natalie S. Kim, Patrick Chao, Samuel Miserendino, Gildas Chabot, David Li, Michael Sharman, Alexandra Barr, Amelia Glaese, and Jerry Tworek. 2025. [GDPval: Evaluating ai model performance on real-world economically valuable tasks](#).

Penrose. 2025. [AccountingBench: Evaluating LLMs on real long-horizon business tasks](#). Blog post.

Rushi Wang, Jiateng Liu, Weijie Zhao, Shenglan Li, and Denghui Zhang. 2025. [AuditBench: A benchmark for large language models in financial statement auditing](#). In *AI for Research and Scalable, Efficient Systems (AI4Research and SEAS Workshops at AAAI 2025)*, volume 2533 of *Communications in Computer and Information Science*, pages 59–81. Springer Nature Singapore. Also available as arXiv:2506.17282 under the title “Automating Financial Statement Audits with Large Language Models”.

Shunyu Yao, Noah Shinn, Pedram Razavi, and Karthik Narasimhan. 2025. τ -bench: A benchmark for tool-agent-user interaction in real-world domains. In *International Conference on Learning Representations (ICLR)*.

Shunyu Yao, Jeffrey Zhao, Dian Yu, Nan Du, Izhak Shafran, Karthik Narasimhan, and Yuan Cao. 2023. [ReAct: Synergizing reasoning and acting in language models](#). In *International Conference on Learning Representations (ICLR)*.

A Worked Task Example

To make task and rubric structure concrete, we walk through one dev set task end to end: the prompt as given to the agent, two representative rubric criteria, and a condensed trajectory from one evaluated model. The example illustrates how binary, outcome-based criteria grade an open-ended deliverable.

A.1 Example Prompt

The prompt below is the complete task input: the agent receives no file list and must locate its own evidence in the world's file system.

Task: World 9_Task 14

Category: Reconciliation

Review the December client cost advance reconciliation. Reconcile the client cost advance ledger to the QBO/GL balance for client cost advances. Using the cost advance ledger methodology and provided support, verify whether hard costs that should be treated as client cost advances were capitalized and whether soft costs were expensed. Identify any discrepancies in cost classification or capitalization treatment that need to be resolved.

Round all amounts to two decimal places. Do not round intermediary values. Provide your answer to the console.

A.2 Example Rubric Criteria

Table 5 shows two of the 24 binary criteria used to grade this task: one qualitative criterion that grades whether the response identifies the task's central seeded error (contract attorney fees expensed rather than capitalized), and a numerical criterion that grades a supporting dollar amount for the same finding. The trajectory excerpted in Section A.3 fails both.

A.3 Example Trajectory

The excerpt below shows six representative turns from one evaluated trajectory (Muse-Spark-1.1, Loop Harness). It omits the remaining turns, tool calls, and

Table 5: Two representative rubric criteria for the worked example task (of 24 total).

Type	Criterion
Qualitative	Identifies the Whitfield M-2024-0018 Contract Attorney Fees Matter Direct (or Account 5000) AP bill lines as a capitalization issue.
Numerical	States that the Whitfield M-2024-0018 account 5000 contract attorney matter direct lines total \$23,500.00 (acceptable value is \$23,500.00).

intervening reasoning. The full trajectory comprises 19 turns containing 41 tool calls; under this harness, an assistant message with no tool calls is treated as the final answer and ends the run. The model met 11 of the task's 24 criteria (45.8%).

Trajectory excerpt — Muse-Spark-1.1, Loop Harness, 11/24 criteria met

Turn 1. "I've mapped the client cost advance task. I'll start by locating your reconciliation files and ledger."
Tool: filesystem_list_files(/)

Turn 3. "I've located the reconciliation workpapers. Now I'll dig into the cost advance methodology and GL balances to uncover the discrepancies."
Tool: excel_read_tab(workpaper_client_cost_advance_ledger_2024.xlsx) ×5 tabs

Turn 4. "I've mapped the cost advance methodology — trust-paid vs. firm-advanced and hard vs. soft. Now I'll reconcile the ledger to GL and verify capitalization treatment."

Tools: excel_read_tab(qbo_closing_trial_balance_2024_12_31.xlsx), excel_read_tab(workpaper_whitfield_settlement_summary.xlsx), filesystem_list_files(/contract_attorney_invoices)

Turn 8. "I've traced the Whitfield settlement and trust-funded costs. Now I'll verify contingency matter capitalization and hard vs. soft cost treatment."

Tool: excel_filter_tab(qbo_general_ledger_detail_2024.xlsx, Matter contains M-2024-0018) ×2

Turn 13 (reasoning written as comments in a code cell). "In our earlier search we only found Whitfield bills (Westfield Legal Services, Jennifer Kao, Marcus Avila). But what about other contingency matters? Meadowbrook and Ferndale have 'matter file maintained off-system' — perhaps there are costs in AP bills that should be capitalized to 1300 for those matters but were expensed?"

Tool: code_execution_run_shell(python3: scan GL detail for AP bills on contingency matters)

Turn 19 (final answer — a message with no tool calls). "QBO/GL balance 1300 = \$0.00 reconciles to Client Cost Advance Ledger = \$0.00. Tie-outs to IOLTA (\$29,400.00) and Clio (\$47,587.09) and settlement (\$45,000.00) pass."

Only classification correction required is re-labeling three Heritage transactions (\$2,950.00) from Soft to Hard in ledger workpaper. No dollar adjustment to GL required at Dec 31, 2024, pending review of off-system Meadowbrook/Ferndale registers."

The model ties the ledger to the GL, verifies the supporting tie-outs, and catches a minor workpaper labeling issue (\$2,950.00), earning 11 of 24 criteria. But it misses the error the task is built around: \$23,500.00 of Whitfield contract attorney fees expensed to account 5000 instead of capitalized to account 1300.

The model retrieved the relevant bills and ledger records but did not open the cost-recovery addendum containing the controlling authorization. We therefore classify the primary failure as incomplete information gathering. A secondary reasoning failure followed: the model treated the zero account-1300 balance as evidence that no capitalization adjustment was required, rather than investigating whether relevant costs had been misclassified.

B Quality Control

This appendix section details the quality-control process behind Section 2. The aim was to verify that every task is solvable from the shipped world files alone and admits a single defensible answer, so that rubric grades measure model ability rather than dataset defects. Quality control was run at both the world and task level (Figure 11). At the world level, four checks were applied:

1. An automated hygiene pass scanned for blatant contradictions and data-hygiene problems.
2. Multiple experts manually reviewed all files for realism and consistency.
3. A validation script confirmed the accounting software data was schema-compatible with the evaluation environment.
4. Three to five experts each completed a handful of tasks against the world files, surfacing contradictions and missing data by having professionals actually do the work.

At the task level, each submission passed through one layer of expert review (with revision cycles as needed) and an automated QC pass. Tasks failing automated QC were audited by the project team, and roughly 5% of approved tasks were manually re-reviewed throughout the project lifecycle.

As a final check, we baselined 20 tasks with independent experts who had no prior exposure to the task, giving them only the prompt and world files and having them re-solve the task from scratch. Comparing each independent answer against the shipped golden response surfaced a single divergent task, on which the independent re-solve disagreed with the shipped golden on 3 of its 19 rubric criteria (5% defect rate by task). Across all 20 baselined tasks this was the only disagreement observed, an error rate of 3 of 276 total criteria (1.1% defect rate by criteria). The divergent task was returned for revision, confirming that the remaining baselined tasks were fully solvable from the shipped files alone and admitted a single defensible answer.

B.1 Difficulty and Selection

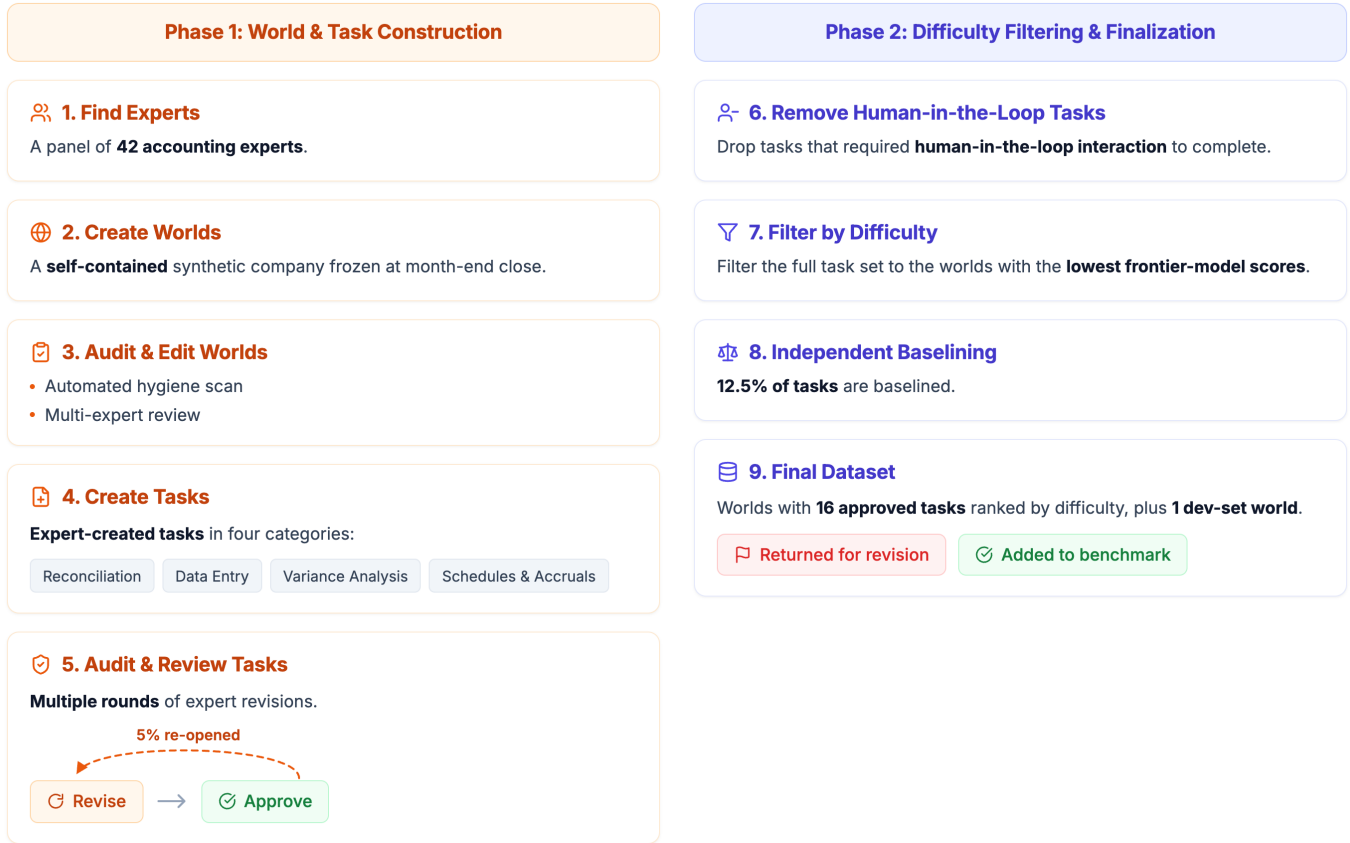
Prior to baselining, the final benchmark was filtered from a larger authored pool of 599 candidate tasks, on which frontier models scored consistently high under the Loop Harness, leaving little headroom to separate them, so we retained only the hardest worlds. We restricted the benchmark to worlds with 16 or more approved tasks and selected the 11 with the lowest average Mean Criteria@3 score among their 16 hardest tasks, using Claude-Opus-4.8, GPT-5.5, and Gemini-3.1-Pro Preview as the filtering models; the 10 most difficult worlds formed the benchmark, and the 11th (a legal-services world) became the public dev set, from which we release 10 tasks.

Before filtering for difficulty, we removed 24 incomplete-information tasks, those where a model is meant to refuse service or ask for clarification, given their small representation in the dataset and the human-in-the-loop support they would require.

C World and Task Statistics

This appendix reports per-world scale statistics (Table 6). We checked whether task complexity (e.g., criteria counts and required files) was spread evenly

Figure 11: Our world construction and audit process consisted of expert and synthetic generation with multiple rounds of expert audits.



across worlds, so that no single world dominates the benchmark.

APEX—Accounting tasks fall into four categories (Table 1); Table 7 reports the same breakdown restricted to the public dev set world (World 9), which has a smaller and less evenly distributed sample per category ($n = 10$ total) than the held-out benchmark. World 9 simulates a boutique law firm completing its December month-end close under U.S. GAAP accrual accounting; its tasks cover law-firm-specific close work such as trust (IOLTA) accounting, client cost advances, payroll and partner-compensation accruals, and realization analysis.

D Model Configurations

This appendix section records the exact model configurations used for the leaderboard runs, for repro-

ducibility and because context limits and thinking-effort settings differ by provider. All models are called via LiteLLM with retry logic (exponential back-off) and a maximum of three attempts per request. Model configurations are described in Table 8.

E Judge Evaluation Details

The question behind this comparison is whether a low-cost model (not on the leaderboard roster) can match frontier-model grading agreement. Table 9 reports full comparison results across all 10 candidate judges with a GEPA-optimized (Agrawal et al., 2026) grading prompt evaluated in Section 3, scored against majority-vote human ground truth. DeepSeek-v4-Flash (the deployed judge, shown in bold) matches the agreement of substantially more expensive proprietary judges—its F_1 of 0.970 matches

Table 6: Task, criteria, and file-type statistics by world.

World	Tasks	Total Criteria	Unique Files	Unique Spreadsheets	Unique PDFs	Unique QuickBooks	Mean Criteria/Task	Mean Files/Task
World 9 (dev set)	10	89	90	34	46	10	8.9	5.6
World 2.2	16	183	64	33	27	2	11.4	6.6
World 3	16	297	79	43	27	5	18.6	6.2
World 5.2	16	190	61	27	23	8	11.9	6.0
World 6.2	16	237	64	40	18	5	14.8	7.0
World 7	16	317	166	85	71	0	19.8	14.1
World 8	16	186	65	30	24	10	11.6	7.6
World 10	16	180	54	28	17	7	11.3	6.0
World 11	16	212	72	32	36	0	13.3	8.4
World 13	16	191	54	30	17	5	11.9	6.7
World 14	16	193	52	32	16	3	12.1	6.5
Mean (w/ dev set)	—	207	74.6	37.6	29.3	5.0	13.4	7.4
Median (w/ dev set)	—	191	64	32	24	5	11.9	6.6
Mean (w/o dev set)	—	219	73.1	38.0	27.6	4.5	13.7	7.5
Median (w/o dev set)	—	192	64	32	23.5	5	12	6.7

Table 7: APEX–Accounting dev set (World 9) tasks by category ($n = 10$).

Category	Tasks	Mean Criteria	Median Criteria	Mean Files	Median Files
Data Entry	1	2	2	5	5
Reconciliation	4	11.5	9	7	7.5
Variance Analysis	2	10.5	10.5	2.5	2.5
Schedules & Accruals	3	6.7	5	6	6
Overall	10	8.9	5.5	5.6	5.5

Claude-Opus-4.8 with high reasoning and exceeds every other individual judge tested.

E.1 Judge Agreement by Solver Model

DeepSeek-v4-Flash is not among the solver models used in this benchmark’s leaderboard (Section 3), so self-preference is not a structural concern. Nonetheless, Table 10 breaks out judge agreement, precision, and recall separately for each solver model’s trajectories as a robustness check.

E.2 Inter-Rater Reliability

To validate the ground truth, Table 11 reports inter-rater reliability among the 3 expert annotators before majority-vote aggregation; Fleiss’ κ of 0.857 indicates strong agreement beyond chance for two-category ratings.

E.3 Self-Preference Stress Test

Because DeepSeek-v4-Flash is not among the solver models used for the benchmark, self-preference is not a structural concern. As a stronger, adversarial check, we additionally ran DeepSeek-v4-Flash as a solver on the full APEX–Accounting benchmark and graded its trajectories with the same DeepSeek-v4-Flash judge, temperature = 0.1, and high reasoning. Graded by itself, DeepSeek-v4-Flash attains a Mean Criteria@3 of 37.1%—essentially matching Kimi-K2.7-Code (37.0%, seventh on the leaderboard) and far below the leader (Claude-Fable-5, 56.4%). A self-favoring judge would push its own solver toward the top of the roster; instead it lands in the lower-middle, consistent with no material self-preference.

F Run Completion and Scoring of Incomplete Runs

Runs that fail because of the model (e.g., no submission, a malformed or empty answer, or the step or token budget exhausted without submitting) are scored as zero on every criterion, so the run contributes 0% to Mean Criteria@3 and is included both the numerator and the denominator. Runs that fail for reasons outside the model’s control, such as a provider API error or a system failure, are excluded from every metric, contributing to neither the nu-

Table 8: Information about the models evaluated on the APEX–Accounting leaderboard: provider, context window, maximum output, thinking configuration, and mean/median wall-clock time per run.

Model	Provider	Context Window	Max Output	Temp.	Thinking	Mean min/run	Median min/run
Claude-Fable-5	Anthropic	1M	128k	Default	adaptive, effort=max	22.0	20.5
Muse-Spark-1.1	Meta	—	—	1.0	reasoning_effort=xHigh	80.9	76.1
GPT-5.6-Sol	OpenAI	1M	—	1.0	reasoning mode=pro, effort=max	54.3	46.6
Claude-Opus-4.8	Anthropic	1M	128k	Default	adaptive, effort=max	44.3	34.9
GLM-5.2	Z.AI	1M	131k	1.0	reasoning_effort=max	43.4	32.9
Grok-4.5	xAI	500k	—	1.0	default (high)	12.0	9.4
Kimi-K2.7-Code	MoonshotAI	256k	—	1.0	reasoning_effort=high	11.6	10.2
Gemini-3.1-Pro	Google	1M	32k	1.0	reasoning_effort=high (+thoughts)	28.4	29.9
Qwen3.5-397B-FP8	Alibaba	—	82k	0.6	enable_thinking=True	60.5	43.4

Table 9: Full LLM-as-a-judge comparison against the majority-vote human ground truth on the Met class. Every judge in this comparison was evaluated with the same GEPA-optimized (Agrawal et al., 2026) grading prompt. The deployed judge, DeepSeek-v4-Flash, is shown in bold. The last three rows are ensemble configurations which combine several judges’ verdicts by majority vote: Open Source aggregates the open-weight judges, Oracle Models the closed frontier judges, and Majority Vote all judges together. (High/Low denote reasoning effort settings.)

Judge	Prec.	Rec.	F_1	Acc.
DeepSeek-v4-Flash (High)	0.966	0.974	0.970	0.971
Claude-Opus-4.8 (High)	0.946	0.994	0.970	0.970
Claude-Opus-4.8 (Low)	0.947	0.993	0.969	0.969
Gemini-3.1-Pro (High)	0.945	0.994	0.969	0.969
Gemini-3.1-Pro (Low)	0.946	0.989	0.967	0.967
DeepSeek-v4-Pro (High)	0.965	0.968	0.967	0.967
Kimi-K2.6	0.973	0.960	0.966	0.967
Gemini-3-Flash	0.938	0.994	0.965	0.965
GPT-5.5 (High)	0.950	0.977	0.963	0.964
GPT-5.5 (Low)	0.948	0.974	0.961	0.962
Open Source (ensemble)	0.975	0.985	0.980	0.981
Oracle Models (ensemble)	0.947	0.991	0.968	0.969
Majority Vote (ensemble)	0.954	0.994	0.974	0.974

erator nor the denominator. Because those runs are dropped rather than scored, a few tasks have fewer than 8 scored runs for a given model, which is why paired comparisons span 156–160 tasks per pair (Section G) rather than all 160.

Table 10: DeepSeek-v4-Flash agreement against majority-vote ground truth, broken out by solver model.

Solver Model	Agreement	Precision	Recall
Claude-Fable-5	0.980	0.984	0.981
GPT-5.6-Sol	0.947	0.918	0.965
Gemini-3.1-Pro	0.984	0.989	0.975
Overall	0.971	0.966	0.974

Table 11: Inter-rater reliability of the 3 expert annotators on the judge-selection sample, prior to majority-vote aggregation.

Metric	Value
Criteria labeled (3 graders)	1,687
Unanimous criteria (all 3 agree)	1,505 (89.2%)
Non-unanimous criteria	182 (10.8%)
Mean pairwise agreement	92.8%
Fleiss’ κ (3 raters, 2 categories)	0.857

G Pairwise Significance Matrix

This appendix section evaluates which leaderboard gaps are statistically reliable once all 36 pairwise comparisons are corrected for multiple testing. For each model pair, we compute the paired per-task difference in Mean Criteria@3 over the tasks with complete runs for both models (156–160 tasks per pair) and correct the resulting p -values with the Benjamini–Hochberg procedure at the 5% false-discovery rate. Table 12 reports the full matrix. 34 of the 36 gaps remain significant after correction. The two exceptions are both adjacent-rank pairs: Muse-Spark-1.1 vs. GPT-

5.6-Sol (+1.2 pp, CI [-2.9, 5.1]) and GLM-5.2 vs. Grok-4.5 (+1.9 pp, CI [-1.6, 5.4]). Three further adjacent-rank gaps survive only narrowly ($0.01 \leq q < 0.05$): Claude-Fable-5 vs. Muse-Spark-1.1 (+3.8 pp), GPT-5.6-Sol vs. Claude-Opus-4.8 (+3.7 pp), and Grok-4.5 vs. Kimi-K2.7-Code (+3.8 pp). Every comparison spanning two or more leaderboard ranks is significant at $q < 0.01$, so the ordering is reliable everywhere except at these adjacent boundaries.

H Full Leaderboard and Dev Set Shift

Table 13 reports the full APEX-Accounting leaderboard: Mean Criteria@3, Pass@1, Pass@8, and Pass^8 for every model, with 95% bootstrap confidence intervals. Table 14 then compares each model's held-out leaderboard performance against the public dev set world, to surface overfitting risk on the released split. Because the dev set world is, by construction, the easiest of the 11 qualifying worlds - it ranked last under the difficulty filter (Section B) - scores are expected to run somewhat higher on the dev set.

I Cost Ablation: Dollar-to-Token Conversion

This appendix section documents how the dollar budgets in the cost ablation are enforced, so the experiment can be reproduced exactly. Section 4.4 enforces each per-task dollar budget (\$1, \$5, \$10, \$50) as a token budget, converting dollars to tokens using each model's own prompt and completion token prices at a 90:10 input-to-output token split, using the higher cost tier listed by model providers (prices can vary by token count). For a model with prompt price p and completion price c (both in dollars per one million tokens), the token allowance for a dollar budget B is

$$\text{tokens}(B) = \frac{B}{(0.9p + 0.1c) / 1,000,000}.$$

Table 15 reports the prompt and completion price for each model in the five-model ablation roster, along with the resulting token allowance at each dollar budget, the exact token counts the Loop Harness enforces directly.

Figure 12 shows an example of how this allowance is surfaced to the model: the converted token budget is

Figure 12: System-prompt excerpt disclosing the converted token budget to the model, shown for Grok-4.5 at the \$5 allowance from Table 15. Models are informed of their token budget the same way they are informed of the per-run step budget (Section 4.4).

Example System Prompt Excerpt (Grok-4.5, \$5 Budget)

```
You are an AI assistant that completes accounting tasks by reasoning and using tools.

## Token Budget
You have a total of 1,041,660 tokens to use throughout the run. If you run out of tokens, submit what you have for partial credit.

## Think Before Acting
Before making tool calls, briefly explain your reasoning in 1-3 sentences: what you learned from the previous step, and what you are doing next and why. Be concise but show your thinking.

[... standard task-file, tool-use, workflow, and output-format instructions omitted ...]
```

stated verbatim in the system prompt's Token Budget section, so the model can pace its own tool calls and reasoning against a concrete, known limit rather than an opaque dollar figure. The remaining token budget is also shared with the model at each step.

J Within-Cap-Spend Score Analysis

This appendix supports the cost ablation in Section 4.4. The paired per-task budget gains are reported in Table 16. Table 17 reports the within-cap levels slope b (change in Mean Criteria@3 per dollar of realized spend) for each model $\times$ budget cell. Because we set the budget cap but not how much of it a model spends on any given task, these slopes describe associations rather than causal effects; the most plausible driver is task difficulty (hard tasks draw more tokens and still score lower), which our design can note but not confirm.

K Harness Comparison Details

This comparison tests whether a purpose-built accounting harness affects model scores once task difficulty is controlled. Table 18 reports Mean Criteria@3 and Pass@1 for every model under the Loop Harness and the Ramp Harness, with task-bootstrap 95% CIs over the shared 160-task set (154-160 paired tasks per model after excluding incomplete runs); GPT-5.6-Sol has no Ramp Harness run. Of the eight

Table 12: Pairwise Mean Criteria@3 differences (row minus column, percentage points, paired per task over the 156–160 tasks with complete runs for both models). Stars are Benjamini–Hochberg-corrected across all 36 comparisons: * $q < 0.05$, ** $q < 0.01$, *** $q < 0.001$; ns = not significant at the 5% false-discovery rate. Models in leaderboard order.

	Muse-Spark-1.1	GPT-5.6-Sol	Claude-Opus-4.8	GLM-5.2	Grok-4.5	Kimi-K2.7-Code	Gemini-3.1-Pro	Qwen3.5-397B
Claude-Fable-5	+3.8*	+5.0**	+8.6***	+13.7***	+15.6***	+19.4***	+24.1***	+32.4***
Muse-Spark-1.1		+1.2(ns)	+5.0**	+9.9***	+11.8***	+15.6***	+20.3***	+28.6***
GPT-5.6-Sol			+3.7*	+8.8***	+10.7***	+14.5***	+19.1***	+27.4***
Claude-Opus-4.8				+5.0**	+6.9***	+10.7***	+15.4***	+23.6***
GLM-5.2					+1.9(ns)	+5.7***	+10.4***	+18.5***
Grok-4.5						+3.8*	+8.4***	+16.6***
Kimi-K2.7-Code							+4.6**	+12.8***
Gemini-3.1-Pro								+8.1***

Table 13: Full APEX–Accounting leaderboard: Mean Criteria@3, Pass@1, Pass@8, and Pass^8 (%) for every model, with 95% task-bootstrap confidence intervals. Best value in each column in bold.

Model	Mean Criteria@3	Pass@1	Pass@8	Pass^8
Claude-Fable-5	56.4 [52.7–60.2]	9.5 [6.0–13.4]	20.1 [13.9–26.6]	1.9 [0.0–4.4]
Muse-Spark-1.1	52.6 [49.0–56.3]	7.7 [4.9–11.0]	21.5 [15.3–28.1]	1.9 [0.0–4.4]
GPT-5.6-Sol	51.5 [47.6–55.4]	8.4 [5.2–12.1]	17.9 [12.2–24.2]	2.6 [0.6–5.2]
Claude-Opus-4.8	48.0 [44.4–51.6]	4.1 [2.1–6.6]	13.7 [8.6–19.5]	0.7 [0.0–2.0]
GLM-5.2	42.7 [39.1–46.5]	4.0 [2.3–5.9]	14.5 [9.4–20.1]	0.0 [0.0–0.0]
Grok-4.5	40.8 [37.4–44.1]	4.1 [2.3–6.0]	13.8 [8.8–19.4]	0.0 [0.0–0.0]
Kimi-K2.7-Code	37.0 [33.7–40.3]	2.9 [1.4–4.7]	11.2 [6.9–16.2]	0.0 [0.0–0.0]
Gemini-3.1-Pro	32.4 [29.2–35.7]	2.6 [1.2–4.3]	9.4 [5.0–13.8]	0.0 [0.0–0.0]
Qwen3.5-397B	24.4 [21.5–27.3]	0.2 [0.0–0.6]	1.9 [0.0–4.5]	0.0 [0.0–0.0]

Table 14: Performance of models on the APEX–Accounting leaderboard ($n = 160$) compared with the public dev set ($n = 10$) using Mean Criteria@3.

Model	Leaderboard Mean Criteria@3	Dev Set Mean Criteria@3	Score Diff
Claude-Fable-5	56.4%	67.9%	+11.5
Muse-Spark-1.1	52.6%	52.4%	−0.2
GPT-5.6-Sol	51.5%	62.3%	+10.8
Claude-Opus-4.8	48.0%	61.8%	+13.8
GLM-5.2	42.7%	44.1%	+1.4
Grok-4.5	40.8%	51.1%	+10.3
Kimi-K2.7-Code	37.0%	38.0%	+1.0
Gemini-3.1-Pro	32.4%	37.0%	+4.6
Qwen3.5-397B	24.4%	27.6%	+3.2

models with both harnesses, only Grok-4.5’s Mean Criteria@3 gain survives Benjamini–Hochberg correction across the 16 paired contrasts (+7.5 pp, CI [5.0, 10.2], $q < 10^{-6}$; 61% of its tasks improve, 25% worsen). Three other shifts exclude zero nominally but do not survive BH correction: Muse-Spark-1.1’s Mean Criteria@3 decline (−3.1 pp, CI [−6.2, −0.2]), Claude-Opus-4.8’s gain (+2.2 pp, CI [0.1, 4.2]), and Muse-Spark-1.1’s Pass@1 gain (+2.5 pp, CI [0.2, 5.0]);

Table 15: Prompt and completion token prices and the resulting per-task token allowance at each dollar budget, assuming a 90:10 input-to-output token split, five-model ablation roster (rounded to the nearest ten).

Model	Prompt Cost (\$/1M)	Completion Cost (\$/1M)	\$1	\$5	\$10	\$50
Claude-Fable-5	\$10.00	\$50.00	71,420	357,140	714,280	3,571,420
GPT-5.6-Sol	\$10.00	\$45.00	74,070	370,370	740,740	3,703,700
Gemini-3.1-Pro	\$4.00	\$18.00	185,180	925,920	1,851,850	9,259,250
Grok-4.5	\$4.00	\$12.00	208,330	1,041,660	2,083,330	10,416,660
Muse-Spark-1.1	\$1.25	\$4.25	645,160	3,225,800	6,451,610	32,258,060

all shifts are paired-task deltas over shared tasks.

K.1 Subagent Delegation

Figure 13 shows how heavily each model delegates work: Claude-Fable-5 spawns subagents in 89% of its trajectories (4.9 per task, CI [4.5, 5.3]); GLM-5.2, Claude-Opus-4.8, and Grok-4.5 follow at 69–89% (2.3–3.2 per task); Kimi-K2.7-Code and Muse-Spark-1.1 delegate in 37–40% (1.1–1.4 per task); and Gemini-3.1-Pro and Qwen3.5-397B delegate rarely (12–13% of trajectories, 0.2–0.4 per task).

Spawn intensity does not line up with harness-level gains: Claude-Fable-5 and GLM-5.2 delegate most heavily yet shift by less than a point under the Ramp Harness (−0.8 and +0.5 pp), while Grok-4.5, the only model with a Benjamini–Hochberg-surviving gain (Table 18), delegates in 69% of its trajectories. Heavy delegation is therefore attributed to model preference rather than performance gains.

Table 16: Paired per-task quality gains from raising the dollar budget (Mean Criteria@3, percentage points), with 95% bootstrap CIs. p -values are Benjamini-Hochberg corrected across the 20 model $\times$ contrast cells; gains significant at the 5% false-discovery-rate level are shown in **bold**.

Model	$\Delta(\$5-\$1)$	$\Delta(\$10-\$5)$	$\Delta(\$50-\$10)$	$\Delta(\$50-\$1)$
Claude-Fable-5	+ 30.4 [+26.8, +34.2]	+ 6.3 [+3.7, +8.9]	+ 6.7 [+4.1, +9.4]	+ 43.4 [+39.1, +47.7]
GPT-5.6-Sol	+ 11.2 [+8.3, +14.3]	-0.2 [-2.5, +2.1]	+1.3 [-0.8, +3.5]	+ 12.3 [+9.2, +15.5]
Gemini-3.1-Pro	+ 5.4 [+3.1, +7.8]	-0.5 [-2.7, +1.8]	+ 4.2 [+2.1, +6.3]	+ 9.1 [+6.5, +11.8]
Muse-Spark-1.1	+ 4.1 [+2.0, +6.3]	+0.5 [-1.5, +2.5]	+0.0 [-2.2, +2.2]	+ 4.7 [+2.3, +7.0]
Grok-4.5	+ 3.5 [+0.7, +6.4]	-1.8 [-4.8, +1.2]	+1.1 [-1.9, +4.2]	+2.8 [-0.2, +5.9]

Table 17: Within-cap levels slope b (Mean Criteria@3 per USD) for each model $\times$ budget cell, 90:10 input:output split, with 95% bootstrap CIs. p -values are Benjamini-Hochberg corrected across the 20 cells; slopes significant at the 5% false-discovery-rate level are shown in **bold**. Slopes are descriptive associations, not causal effects.

Budget	Claude-Fable-5	GPT-5.6-Sol	Gemini-3.1-Pro	Muse-Spark-1.1	Grok-4.5
\$1	0.065 [-0.016, 0.231]	0.015 [-0.027, 0.060]	-0.069 [-0.247, 0.051]	-0.136 [-0.300, 0.013]	-0.069 [-0.202, 0.080]
\$5	-0.101 [-0.167, -0.020]	- 0.031 [-0.052, -0.010]	-0.033 [-0.059, -0.007]	-0.017 [-0.044, 0.009]	0.003 [-0.023, 0.031]
\$10	- 0.060 [-0.084, -0.032]	- 0.020 [-0.033, -0.007]	-0.012 [-0.025, 0.000]	-0.008 [-0.023, 0.006]	0.003 [-0.009, 0.016]
\$50	-0.002 [-0.005, 0.001]	-0.001 [-0.004, 0.001]	-0.001 [-0.003, 0.000]	-0.004 [-0.015, 0.003]	-0.003 [-0.008, 0.002]

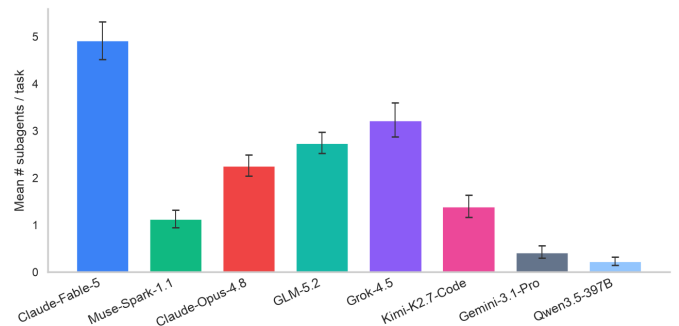
Table 18: Mean Criteria@3 and Pass@1 (%) under the Loop Harness vs. the Ramp Harness ($k = 3$, $n = 160$). Brackets are task-bootstrap 95% CIs; GPT-5.6-Sol has no Ramp Harness run and is omitted. The one Benjamini-Hochberg-surviving gain (Grok-4.5, Mean Criteria@3) is in **bold**.

Model	Mean Criteria@3		Pass@1	
	Loop	Ramp	Loop	Ramp
Claude-Fable-5	56.4 [52.7-60.2]	55.6 [51.7-59.4]	9.5 [6.0-13.4]	10.0 [6.0-14.4]
Muse-Spark-1.1	52.6 [48.9-56.2]	49.8 [45.5-54.0]	7.7 [4.9-11.0]	10.2 [6.5-14.4]
Claude-Opus-4.8	48.0 [44.5-51.5]	50.4 [46.8-53.9]	4.1 [2.1-6.6]	5.8 [3.1-9.0]
GLM-5.2	42.7 [38.9-46.4]	43.2 [39.7-46.7]	4.0 [2.3-5.9]	4.2 [2.1-6.7]
Grok-4.5	40.8 [37.5-44.2]	48.3 [44.7-52.0]	4.1 [2.3-6.0]	5.6 [2.9-8.8]
Kimi-K2.7-Code	37.0 [33.7-40.3]	37.5 [34.2-40.8]	2.9 [1.4-4.7]	2.5 [0.8-4.6]
Gemini-3.1-Pro	32.4 [29.1-35.7]	33.7 [30.4-37.2]	2.6 [1.2-4.3]	3.5 [1.5-6.2]
Qwen3.5-397B	24.4 [21.5-27.3]	25.3 [22.4-28.4]	0.2 [0.0-0.6]	0.2 [0.0-0.6]

L Failure Modes of the Best-Performing Models

Section 5.1 compares the top models at the L1 level and within reasoning failures; Figure 14 gives the complete picture, breaking every annotated failure down to its L2 subclass per model. Inner rings show L1 failure categories; outer rings show the L2 subclasses within each. Beyond the shared concentration in reasoning, the sunbursts show where the profiles diverge: Claude-Fable-5 concentrates hardest in reasoning (79% of its failures, over half of them non-numeric reasoning), GPT-5.6-Sol carries the broadest instruction-following tail (ignoring requirements and

Figure 13: Mean subagents spawned per task by model under the Ramp Harness (task-bootstrap 95% CIs).



implicit-instruction failures), and Muse-Spark-1.1 is the only model whose information-gathering share exceeds a fifth.

L1 L2 Subclass Counts

Table 19 reports the exact per-model counts behind Figure 14. Non-numeric reasoning failures and data handling errors are the two largest subclasses for every model, and no annotated failure involves tool use.

Figure 14: Failure-mode breakdown for the three best-performing models ($n = 24/22/28$ annotated failures). Inner ring: L1 failure categories; outer ring: L2 subclasses.

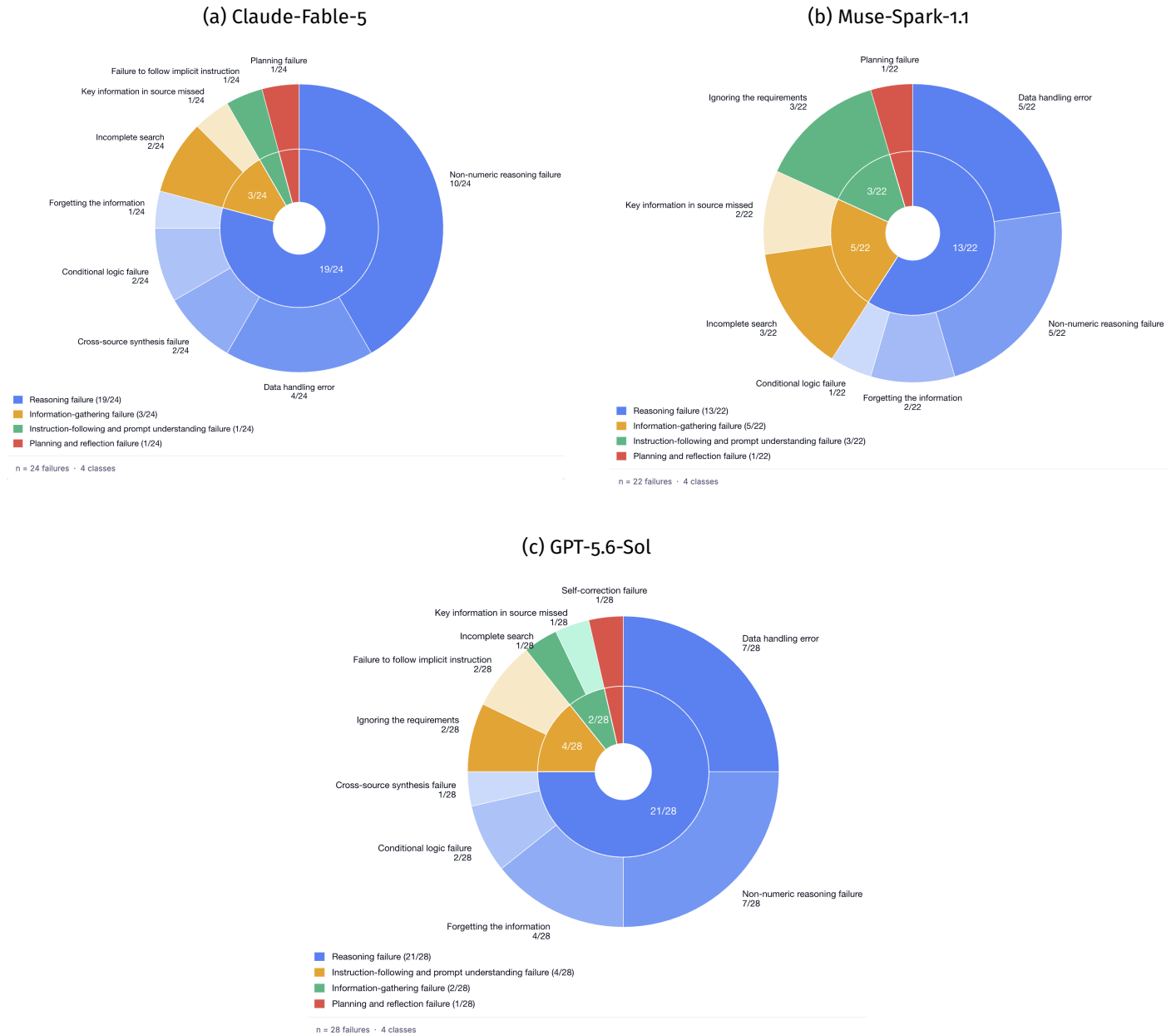


Table 19: Annotated failure counts by L2 subclass for the three best-performing models, grouped by L1 category ($n = 24/28/22$; 74 failures total). The four L1 categories shown are the only ones observed; no annotated failure falls under tool use or any other taxonomy branch.

L2 subclass	Claude-Fable-5	GPT-5.6-Sol	Muse-Spark-1.1	Total
<i>Reasoning failure</i>				
Non-numeric reasoning failure	10	7	5	22
Data handling error	4	7	5	16
Forgetting the information	1	4	2	7
Conditional logic failure	2	2	1	5
Cross-source synthesis failure	2	1	0	3
<i>Information-gathering failure</i>				
Incomplete search	2	1	3	6
Key information in source missed	1	1	2	4
<i>Instruction-following and prompt understanding failure</i>				
Ignoring the requirements	0	2	3	5
Failure to follow implicit instruction	1	2	0	3
<i>Planning and reflection failure</i>				
Planning failure	1	0	1	2
Self-correction failure	0	1	0	1
Total	24	28	22	74

M Failure Taxonomy

Section 5 introduces the taxonomy at a high level. This appendix section lists the taxonomy at the L1/L2 level. 7 L1 categories and 39 L2 subclasses are relevant to accounting work; finer L3 sub-labels are omitted for brevity.

1. Information-gathering failure

The agent fails to find the correct information, or all of the information, required to successfully complete the request.

- **No information gathering.** The agent does not attempt to search or retrieve any files at any point in the task, relying on parametric knowledge to complete the request.
- **Incomplete search.** The agent searches for or retrieves only some of the required information: it never finds the correct source, finds it but does not use it, or finds a similar but poorer-quality source.
- **Unused search result.** The agent finds the correct source in its search results but fails to open or use it.
- **Key information in source missed.** The agent retrieves the correct file(s) but does not find the

relevant or correct information within them.

- **Source reconstruction error.** The agent reconstructs a value from an original source (usually raw inputs) and introduces an inconsistency.
- **Grounding failure – incorrect source attribution.** The agent attributes information to a source that does not actually support or contain it.

2. Instruction-following and prompt understanding failure

The agent fails to understand or apply some or all of the instructions and constraints in the prompt.

- **Ignoring the requirements.** The agent ignores one or more of the prompt’s requirements, constraints, or dimensions, such as a time period, business segment, or entity.
- **Failure to follow implicit instruction.** The agent fails to follow, ignores, or misinterprets an implicit instruction in the prompt.
- **Priority error.** The agent applies an incorrect priority order among multiple satisfiable instructions that compete for attention, resources, or scope.
- **Ignores a negation.** The agent ignores an exclusionary or negating term such as “not,” “except,” or “excluding” and acts as though it were absent.

- **Non-text interpretation failure.** The agent misreads or fails to extract information from non-text content in the prompt, such as tables, graphs, images, or images of text (e.g., scanned documents or screenshots).

3. Reasoning failure

The agent has the correct information available but applies faulty logic, inference, or calculation to it.

- **Numeric reasoning failure.** The agent makes an incorrect calculation or applies faulty quantitative logic to numeric data; the underlying values may be correct and complete, but the arithmetic or quantitative method applied to them is not.
- **Data handling error.** The agent correctly understands the source data's structure and content but selects an incorrect method to filter, join, aggregate, or otherwise process it.
- **Non-numeric reasoning failure.** The agent applies faulty logic or inference to non-numeric information, reaching an incorrect conclusion even though the available information is sufficient.
- **Cross-source synthesis failure.** The agent fails to correctly and completely link or reconcile information across sources into a single accurate value or perspective.
- **Exception blindness.** The agent applies a general rule, default approach, or commonly used value without recognizing or checking for a documented exception, carve-out, or methodology.
- **Conditional logic failure.** The agent identifies an if/then rule but misapplies or ignores it: acting when it should not, failing to act when it should, or applying the wrong outcome to a correctly identified condition.
- **Failure to request clarification.** When the prompt contradicts available evidence, rests on an incorrect premise, requests something infeasible, or includes a mistake by the user, the agent fails to seek clarification before attempting the task.
- **Failure to proactively correct.** The agent fails to resolve confusion, ambiguity, or mistakes in the prompt that could be addressed without input from the user.
- **Forgetting the information.** The agent fails to carry its own correct prior reasoning, calculations, or intermediate conclusions into subsequent steps, contradicting, abandoning, or replacing them without new evidence or reconciling the discrepancy.

- **Disambiguation failure.** The agent fails to disambiguate between similar technical terms with different meanings in context, selecting the wrong meaning or treating the terms as interchangeable.
- **Missed dependency.** The agent fails to identify dependencies between steps, or identifies a dependency but does not sequence or handle it correctly.

4. Tool use failure

The agent fails to use any tools, identifies the correct tool but does not use it, uses an incorrect tool, or uses the correct tool incorrectly.

- **No tool used.** The agent does not identify a tool to use to complete an action, relying solely on parametric knowledge.
- **Tool not used.** The agent identifies the correct tool but does not actually call, load, or execute it.
- **Incorrect tool used.** The agent selects a tool that cannot perform the function the task requires.
- **Tool used incorrectly.** The agent selects the correct tool but operates it in a way that produces a faulty result or no result.
- **Tool output misinterpretation.** The agent fails to correctly understand the results, errors, or status returned by a tool.

5. Planning and reflection failure

The agent fails to appropriately plan, monitor, or adjust its process while completing the request.

- **Planning failure.** The agent fails to plan, or to course-correct, a viable course of action for completing the request.
- **Failure to complete task.** The agent enters a doom loop, repeating the same or functionally similar steps, and cannot exit the loop to complete the task.
- **Verification failure.** The agent does not verify its final output before submission, or verifies it only partially; this applies at the point of the final response. It does not apply at intermediate steps.
- **Self-correction failure.** The agent makes an error while completing the request but fails to identify or correct it.
- **Failure to retry.** The agent does not retry or opt

for an appropriate, available fallback when a step fails.

6. Communication and presentation failure

The agent makes a style, grammar, formatting, or other communication error.

- **Final output mistake.** The agent correctly completes all reasoning and analysis but misstates the results in its output.
- **Tone failure.** The agent fails to adapt its tone or formality to the context, audience, or a specific user request.
- **Language failure.** The agent makes a typographical, grammar, or sentence-construction error in its final output.
- **Excessively verbose response.** The agent's output contains unnecessary length or repetition, egregious enough to make the output considerably harder to read.
- **Overly concise response.** The agent's output omits key information it had already gathered or derived, leaving the response too brief to be useful or complete.
- **Overconfidence.** The agent presents its response as a certainty when the available evidence does not support that level of confidence.
- **Output format failure.** The agent's output is not structured or presented in the format(s) requested in the prompt.

7. Other valid failure

The agent fails in a way that is not captured by any other failure label.

N LLM Usage

We used LLMs to assist with drafting and refinement of this paper.