

Do We Always Need Query-Level Workflows? Rethinking Agentic Workflow Generation for Multi-Agent Systems

Anonymous ACL submission

Abstract

Multi-Agent Systems (MAS) built on large language models typically solve complex tasks by coordinating multiple agents through workflows. Existing approaches generate workflows either at task level or query level, but their relative costs and benefits remain unclear. After rethinking and empirical analyses, we show that query-level workflow generation is not always necessary, since a small set of top-K best task-level workflows together already covers equivalent or even more queries. We further find that exhaustive execution-based task-level evaluation is both extremely token-costly and frequently unreliable. Inspired by the idea of self-evolution and generative reward modeling, we propose a low-cost task-level generation framework **SCALE**, which means **S**elf prediction of the optimizer with few shot **C**ALibration for **E**valuation instead of full validation execution. Extensive experiments demonstrate that **SCALE** maintains competitive performance, with an average degradation of just 0.61% compared to existing approach across multiple datasets, while cutting overall token usage by up to 83%.

1 Introduction

Large Language Model (LLM)-based multi-agent systems (MAS) have recently emerged as a powerful paradigm for solving complex reasoning, coding, and decision-making tasks (Zhang et al., 2024b,a; Zhuge et al., 2024; Niu et al., 2025). By decomposing a task into multiple interacting agents and organizing their collaboration through agentic workflows, MAS can substantially extend the capabilities of a single agent.

Based on the granularity of workflow construction, agentic workflow generation methods fall into two categories: task-level and query-level approaches. Task-level approaches, such as search-based Aflow (Zhang et al., 2024b) and learning-based GPTSwarm (Zhuge et al., 2024) and

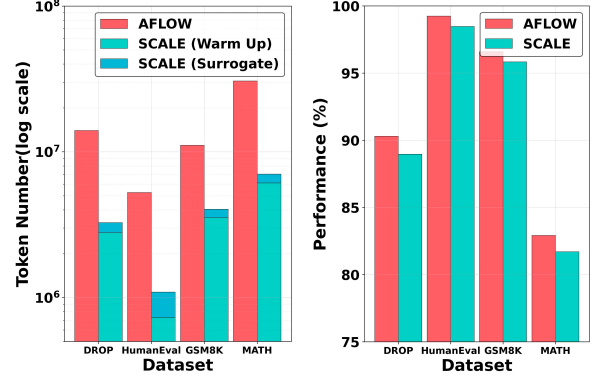


Figure 1: Comparison of Aflow and our rethought task-level workflow generation framework. Left: total token number during workflow generation (log-scale axis). Right: final test performance. Our method **SCALE** achieves comparable performance while significantly reducing token number.

AgentPrune (Zhang et al., 2024a), generate a single workflow intended to perform well across an entire dataset or task distribution. However, this generality comes at a high cost: evaluation dominates computation, as each candidate requires full execution over the validation set. As shown in Figure 1, the whole Aflow’s generation process on four benchmarks consumes approximately 10^6 – 10^8 LLM tokens.

In parallel, query-level approaches generate a separate workflow for each input query (Ye et al., 2025; Wang et al., 2025a; Gao et al., 2025). For every query, the system constructs a customized multi-agent workflow, allowing the agent roles and interaction patterns to adapt to the specific problem. This design aims to better handle heterogeneous queries and can yield strong per-query performance. However, this adaptivity comes with clear costs. A new workflow must be generated for every query, which introduces substantial inference overhead. For many simple or similar inputs, such query-level generation may be unnecessary, providing limited gains relative to its train time computational cost

and test time generation cost.

From the characteristics of these two paradigms, we raise two fundamental questions that have not been systematically examined as shown in Figure 2. First, *Is query-level workflow generation always necessary?* Second, *Is high-cost evaluation in task-level workflow generation necessary?* For the first question, we show that a small set of top-k task-level workflows already achieves strong query coverage comparing to query-level method’s performance. It indicates that query-level workflow generation is not always necessary in practice. For the second, we find that exhaustive execution-based evaluation of task-level workflows is both extremely expensive and frequently unreliable.

Inspired by the idea of self-evolution and generative reward modeling, we propose a low-cost task-level generation framework **SCALE**, which means **S**elf prediction of the optimizer with few shot **CAL**ibration for **E**valuation instead of full validation execution. By leveraging the inherent evaluative ability of LLM-based optimizers, **SCALE** makes self predictions in a generative manner and calibrates them using few shot executions, thereby achieving highly reliable predictions with minimal token cost. Experimental analysis further shows that the calibrated self predictions in **SCALE** closely approximate true execution scores, it achieves a low MAE of 0.16 and maintain consistent ranking with a Pearson correlation of 0.52 (range: $[-1, 1]$), further validating reliability. Overall, our contributions are threefold as shown below:

- **Rethinking Insights:** We present a new empirical rethinking of workflow generation in multi-agent systems. Our analysis yields two main findings: (1) Query-level methods is not always necessary in practice. (2) Exhaustive execution-based evaluation in task-level approaches is both costly and unreliable.
- **Improved Framework:** Motivated by these observations and inspired by self-evolution, we develop a low-cost and effective framework **SCALE** for task-level workflow generation. Instead of exhaustively executing candidate workflows on the full validation set, our approach combines the LLM-based optimizer’s self prediction with few shot calibration to evaluate workflows efficiently.
- **Empirical Validation:** Extensive experiments demonstrate that **SCALE** maintains

competitive performance, with an average degradation of just 0.61% compared to existing approach across multiple datasets, while cutting overall token usage by up to 83%.

2 Preliminaries

2.1 Agentic MAS Workflow

We consider a task $\mathcal{T}$ given as a dataset of queries $q \in \mathcal{D}$, and an agentic MAS workflow $W \in \mathcal{W}$ that orchestrates a LLM-based MAS to produce an answer y . Formally, we formalizes the workflows space as:

$$\mathcal{W} = \{(P_1, \dots, P_n, E, O_{\theta_1}, \dots, O_{\theta_n}) \mid P_i \in \mathcal{P}, E \in \mathcal{E}, O_{\theta_i} \in \mathcal{O}\} \quad (1)$$

where n is the number of agents, $\mathcal{P}$ is the prompt space, $\mathcal{E}$ is the information control flow that governs the execution order, data dependencies, and communication among these agents, which can be represented in various forms: as executable code (e.g. Hu et al., 2024; Zhang et al., 2024b; Xu et al., 2025) or as directed acyclic graphs (e.g. Zhuge et al., 2024; Zhang et al., 2024a; Wang et al., 2025b; Zhang et al., 2025). $\mathcal{O}$ is a set of predefined LLM-based agents parameterized by O_{θ_i} . The prompt P_i and parameters θ_i enable the agent to adapt its behaviors to the task at hand, such as *Review* (Yao et al., 2022), *Ensemble* (Liang et al., 2024), or *Self-Correction* (Shinn et al., 2023).

At a high level, a workflow is a series of agent calls to answer a certain query $y = W(q)$. Given a task-specific evaluator $s(\cdot, \cdot)$ (e.g., exact match, pass@1), the performance of W on a query q is measured by $s(W(q), q) \in [0, 1]$.

2.2 Agentic MAS Workflow Generation

A number of recent methods have been proposed for agentic workflow generation, which can be grouped into two paradigms according to the granularity at which workflows are generated: task-level approaches and query-level approaches.

2.2.1 Task-level workflow generation

Task-level methods aim to generate a single workflow W^* that performs well on a distribution of queries from the same task as shown in (1)A and (1)B of Figure 2. Given a validation set $\mathcal{D}_{\text{val}} = \{q_i\}_{i=1}^N$ sampled from the task $\mathcal{T}$, the objective is:

$$W_{\text{task}}^* = \arg \max_{W \in \mathcal{W}} S^{\text{exec}}(W, \mathcal{D}_{\text{val}}) \quad (2)$$

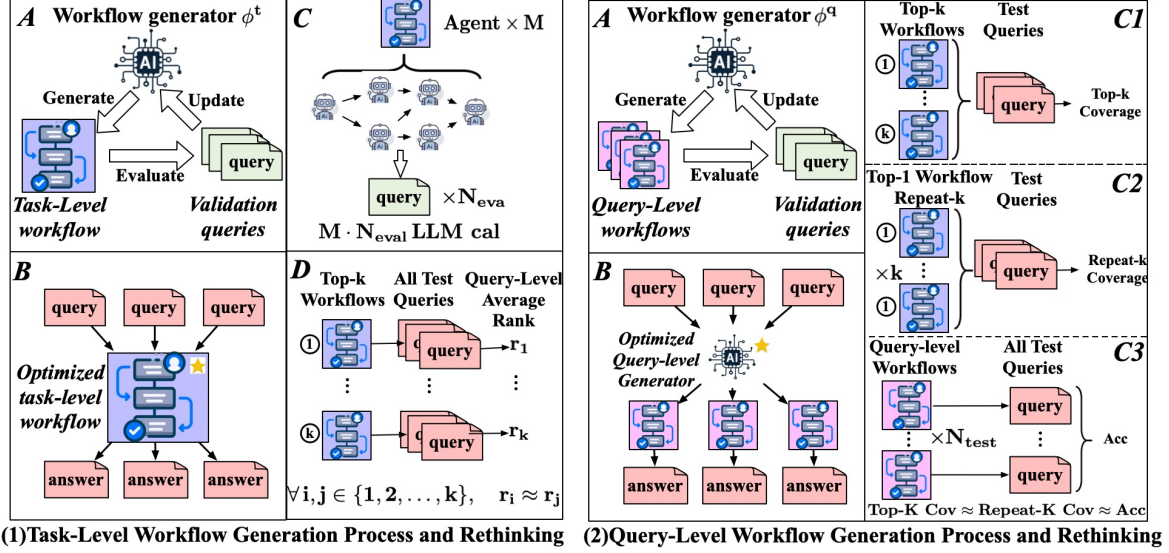


Figure 2: Task-level vs. Query-level workflow generation on their process and rethinking. (1) Task-level generation. (1)A shows searching/training: the generator generates a single workflow using validation queries; (1)B shows inference: the optimized workflow is reused for all test queries. (1)C–D present our rethinking: (1)C shows that repeated full-set evaluations is very token-costly, and (1)D shows top-k workflows have very similar query-level ranks. (2) Query-level generation. (2)A shows training: a workflow is generated per query; (2)B shows inference: producing customized workflows for each input. (2)C1–C3 summarize our rethinking on query-level workflows: top-k task-level workflows, repeat-k runs of the top-1 workflow, and true query-level generation yield comparable coverage/performance.

where $\mathcal{W}$ is the workflow search space, and $S^{exec}(W, \mathcal{D}_{val}) = \frac{1}{|\mathcal{D}_{val}|} \sum_{q \in \mathcal{D}_{val}} s(W(q), q)$ denotes the average execution score of workflow W on dataset $\mathcal{D}_{val}$.

Despite implementation differences, these methods share the same closed-loop paradigm. First, generating candidate workflows through a task-level workflow generator ϕ^t . Second, evaluating the generated workflow on $\mathcal{D}_{val}$. Finally, updating the model ϕ^t using evaluation as a feedback. Aflow (Zhang et al., 2024b) uses a LLM as optimizer ϕ^t and improve the initial workflow through an MCTS-style loop. AgentPrune (Zhang et al., 2024a) use a graph model as workflow generator ϕ^t and learn it through reinforcement learning methods to generate better workflows.

Despite good performance, their evaluation is extremely costly, since each candidate’s evaluation requires the MAS workflow to execute on the full validation set. The token number scales with the validation dataset size and agent numbers resulting in a sharp increase as the loop continues.

2.2.2 Query-level workflow generation

As shown in (2)A and (2)B of Figure 2, query-level methods instead learn a query-level workflow generator ϕ^q that maps each query q to its own workflow $W_q = \phi^q(q)$ and the optimization objec-

tive is to maximize expected performance over the validation set:

$$\phi^{q,*} = \arg \max_{\phi^q} \frac{1}{|\mathcal{D}_{val}|} \sum_{q \in \mathcal{D}_{val}} [s(W_q(q), q)] \quad (3)$$

Existing approaches differ mainly in how ϕ^q is trained. MAS-GPT (Ye et al., 2025) adopts supervised fine-tuning on a curated dataset of query–workflow pairs. ScoreFlow (Wang et al., 2025a) optimizes the workflow generator ϕ^q using a preference-based optimization approach which enhances the original DPO (Rafailov et al., 2023). For many simple or structurally similar inputs, such query-level workflow generation may be unnecessary, providing limited gains relative to its test-time generation cost.

3 Rethinking Agentic Workflow Generation for Multi-Agent Systems

Our rethinking is twofold. First, as shown in (2)C1–C3 of Figure 2, we rethink the necessity of query-level workflow generation. Second, as shown in (1)C and (1)D of Figure 2, we rethink task-level methods, arguing that execution on validation set for evaluation is both token-costly and unreliable.

Dataset	Aflow				S.Flow	
	Top-1 Perf	Top-5 Perf	Top-5 Cov	Repeat-5 Perf	Repeat-5 Cov	All Perf
DROP	90.30	89.81	93.87	90.04	92.45	91.48
HumanEval	99.24	98.17	100.00	97.82	99.24	98.91
GSM8K	96.58	95.89	97.35	96.17	96.87	97.79
MATH	82.92	79.84	87.04	82.81	83.39	84.35

Table 1: Comparison of task-level and query-level workflow effectiveness. **Perf** denotes average test performance (%). **Cov** denotes coverage (%) of test queries. **S.Flow** denotes the query-level method ScoreFlow.

3.1 Is Query-level Workflow Generation Always Necessary?

Query-level methods offers fine-grained adaptivity, but it introduces considerable test-time generation cost that task-level methods don’t have. This raises a fundamental question:

Is query-level workflow generation always necessary to achieve strong performance?

For each dataset, we evaluate the following settings. (1) Task-level Top-1: We report the test performance of Aflow’s (Zhang et al., 2024b) best generated workflow. (2) Task-level Top-5: We report average test performance and coverage of Aflow’s top-5 workflows. *Coverage* is defined as the fraction of test queries that are covered by these workflows. A query is counted as covered if at least one of these workflows answers it correctly. (3) Repeat-5 of Top-1: We execute the Aflow’s top-1 workflow on test set 5 times. We report the average performance and coverage on test queries. This isolates the effect of the test-time execution stochasticity. (4) Query-level workflows: A dedicated workflow is generated for each query using a representative query-level method ScoreFlow (Wang et al., 2025a), and we report its average test performance.

As shown in Table 1, we obtain three main observations. First, a single Top-1 task-level workflow already performs strongly, indicating substantial structural sharing across queries. Second, Top-5 gains a clear increase in query coverage even more than query-level method’s performance, meaning that improvements can come from gathering few candidate task-level workflows rather than generating single strictly better workflows. Third, Repeat-5 achieves coverage comparable to query-level method, showing that much of the gain by query-level method compared to task-level can be covered by stochastic execution rather than diverse workflow structures.

Taken together, these results suggest that most benefits attributed to query-level method can already be achieved by a small pool of reusable task-

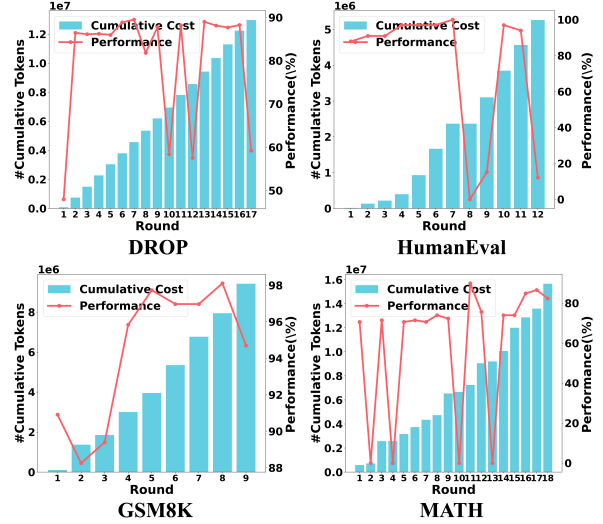


Figure 3: Cumulative Token Number v.s. Performance during Aflow’s task-level workflow generation process.

level workflows or even repeated execution of a single strong task-level workflow. The main advantage can come from coverage and stochasticity, not necessarily in query-level workflows.

3.2 Is High-Cost Evaluation in Task-Level Workflow Generation Necessary?

In this subsection, we revisit the evaluation in task-level workflow generation from two complementary perspectives: (1) how many tokens actually incurs in exhaustive execution-based evaluation, and (2) whether such costly evaluation truly leads to meaningfully different workflows.

3.2.1 How many evaluation tokens are incurred during task-level workflow generation?

We analyze the cumulative evaluation token number incurred during Aflow’s workflow generation. Figure 3 illustrates the relationship between evaluation token number and performance across each generated candidate workflow. For each benchmark, we plot the test performance achieved by candidate workflows, together with the cumulative token number which is defined as the total tokens used to evaluate all candidate workflows generated up to and including the current round.

Figure 3 shows that cumulative evaluation token number grows fast continuously with search rounds, while test performance quickly saturates and yields only marginal or negative gains afterward. This mismatch is evident: the cost is exploding but the corresponding performance improvement is minimal or even negative. These results indicate that the prevailing task-level evaluation paradigm is both

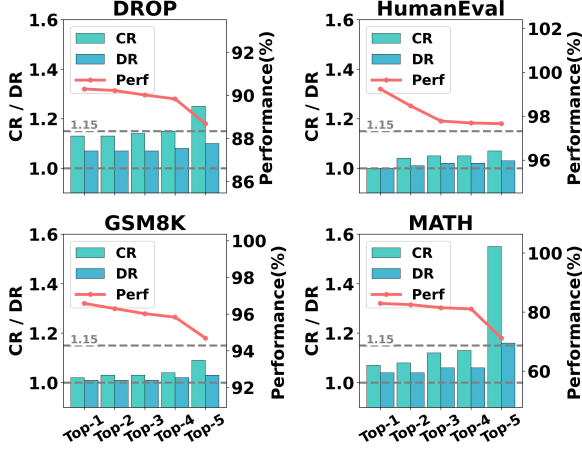


Figure 4: Performance and ranking statistics of the top-5 task-level workflows generated by Aflow across four benchmarks. **Perf** denotes average test performance. **CR** and **DR** denote average competition rank and dense rank, respectively, computed over test queries.

expensive and unreliable in the high-performance regime, motivating the need for cheaper and more reliable task-level workflow evaluation.

3.2.2 Do high-cost task-level evaluations actually distinguish better workflows?

To further examine the efficacy of Aflow’s costly evaluation process, we analyze the Top-5 workflows, defined as the five candidates with the highest validation performance across all candidate workflows produced in one complete run.

Beyond test performance, we also analyze query-level ranking. Concretely, all Top-5 workflows are executed and ranked for each test query. Then Top-5 workflows’ competition rank (**CR**) which allows ties and dense rank (**DR**) are averaged across queries. This reveals how consistently one workflow beats another. If the evaluation are strongly discriminative, these ranks would show clear separation among Top-5 workflows.

Figure 4 shows that the Top-5 (especially Top-4) task-level workflows obtained by Aflow exhibit very similar performance. Their performances vary only slightly across all benchmarks, while both CR and DR remain close to 1 with minimal variation, which indicates that expensive full validation provides limited benefit in identifying substantially better workflows.

3.3 Rethinking Results

In Section 3.1, we observed that query-level workflow generation is not always necessary, because a small set of task-level workflows or even repeated executions of a single workflow already covers

more queries. In Section 3.2, we further showed that exhaustive execution-based task-level evaluation is extremely costly while providing limited discriminative benefit.

Together, these findings show that current agentic workflow methods waste a lot of computation either generating unnecessary query-level workflows or evaluating task-level workflows that have almost the same performances. This motivates a new framework for task-level workflow generation that avoids both query-level generation and high-cost full-execution-based evaluation.

4 Methodology

4.1 Motivation

We propose a low-cost task-level generation framework **SCALE**, which means **S**elf prediction with few shot **C**ALibration for **E**valuation. Inspired by the self-evolution paradigm and generative reward modeling in agentic systems, we treat the workflow optimizer as a self-predictor. In other words, the same LLM that generates workflows is also prompted to estimate the candidate workflow’s expected performance.

To reduce overconfidence, we separate generation and evaluation prompts and obtain scores in a dedicated evaluation context. We calibrate self predictions using execution results from few shot queries, typically 1–3% of the validation set. As a result, **SCALE** enables task-level evaluation without full validation execution, while maintaining reliable workflow scoring and ranking.

4.2 Overall framework

Our method **SCALE** operates in two stages: a short warm-up stage with full execution for evaluation, followed by a surrogate evaluation stage. Specifically, we use Aflow (Zhang et al., 2024b) for a few steps as the warm-up stage, and then instead of repeatedly computing $S^{\text{exec}}(W, \mathcal{D}_{\text{val}})$, we estimate workflow quality using a self prediction score calibrated by few shot execution as the surrogate evaluation stage.

4.2.1 Warm-up Stage

We begin with M warm-up rounds. Concretely, starting from a single question-answering LLM agent call as the initial workflow W_1 with execution score S_1^{exec} , we run an MCTS-style loop for $t = \{1, \dots, M\}$ consisting of four steps:

1. Selection. Given the existing workflows and scores $\{S_i^{\text{exec}}\}_{i=1}^t$, we select a parent workflow

using a soft mixed policy:

$$P_i = \lambda \cdot \frac{1}{t} + (1-\lambda) \cdot \frac{\exp(\alpha(S_i^{\text{exec}} - S_{\max}^{\text{exec}}))}{\sum_{j=1}^t \exp(\alpha(S_j^{\text{exec}} - S_{\max}^{\text{exec}}))}, \quad (4)$$

where $S_{\max}^{\text{exec}} = \max_j S_j^{\text{exec}}$, and λ, α control the exploration–exploitation trade-off.

2. Expansion. A new workflow is generated by editing W_i using the LLM-based optimizer:

$$W_{t+1} = \phi(W_i; P_{t+1}^{\text{optimizer}}) \quad (5)$$

The optimizer prompt $P_{t+1}^{\text{optimizer}}$ is dynamically built from local experience E_i^{local} and global experience E^{global} introduced in backpropagation step.

3. Evaluation. We execute W_{t+1} on the validation set $S_{t+1}^{\text{exec}} = \frac{1}{|\mathcal{D}_{\text{val}}|} \sum_{q \in \mathcal{D}_{\text{val}}} s(W_{t+1}(q), q)$.

4. Backpropagation The evaluation result updates both local and global experience. For W_i the local experience is $E_i^{\text{local}} \leftarrow E_i^{\text{local}} \cup (e_i^{t+1})$ where $e_i^{t+1} = ((W_i, S_i), \Delta_i^{t+1}, (W_{t+1}, S_{t+1}))$ and Δ_i^{t+1} denotes the optimizer’s natural language description of the edit from W_i to W_{t+1} . We also maintain a global experience: $E^{\text{global}} \leftarrow E^{\text{global}} \cup \{(W_{t+1}, S_{t+1}^{\text{exec}})\}$. These experience is reused to update future optimizer prompts and the selection policy.

4.2.2 Surrogate Evaluation Stage

After warm-up stage, we continue the loop but the subsequent workflows are evaluated through self prediction with few shot execution calibration. At iteration $t > M$, we select and expand using Equation 4 and Equation 5 to get the newly generated workflow W_{t+1} . We evaluate it with our method:

Self prediction Firstly, we query the optimizer itself under a dynamically dedicated evaluation prompt $P_{t+1}^{\text{optimizer}}$ to obtain the prediction:

$$S_{t+1}^{\text{pred}} = S^{\text{pred}}(W_{t+1}) = \phi(W_{t+1}; P_{t+1}^{\text{eval}}) \quad (6)$$

where ϕ is the same LLM-based optimizer used for workflow expansion in Equation 5. The evaluation prompt P_{t+1}^{eval} is separated from the optimization prompt $P_{t+1}^{\text{optimizer}}$ to reduce overconfidence and prompt entanglement. The full template of P^{eval} is provided in Appendix A.1.

Few shot execution calibration To reduce bias, we execute W_{t+1} only on a small subset $\mathcal{D}_{\text{few}} \subset$

$\mathcal{D}_{\text{val}}$ with $|\mathcal{D}_{\text{few}}| \ll |\mathcal{D}_{\text{val}}|$:

$$S_{t+1}^{\text{few}} = \frac{1}{|\mathcal{D}_{\text{few}}|} \sum_{q \in \mathcal{D}_{\text{few}}} s(W_{t+1}(q), q) \quad (7)$$

The sampling of $\mathcal{D}_{\text{few}}$ is guided by the full-execution statistics collected during warm-up. With warming up workflows $\{W_m\}_{m=1}^M$, for each validation query $q \in \mathcal{D}_{\text{val}}$, we compute its empirical warm-up score $\bar{s}(q) = \frac{1}{M} \sum_{m=1}^M s(W_m(q), q)$ which reflects how well warm-up workflows already solve q . We then partition the range of $\bar{s}(q)$ into K bins $\{B_k\}_{k=1}^K$ (e.g., $K = 10$), where $B_k = \{q \in \mathcal{D}_{\text{val}} \mid \bar{s}(q) \in I_k\}$ and $\{I_k\}_{k=1}^K$ are disjoint score intervals covering $[0, 1]$. Let $n_k = |B_k|$ be the number of queries in bin k . We define a bin-level sampling distribution by a softmax over bin counts $p_k = \frac{\exp(\gamma n_k)}{\sum_{j=1}^K \exp(\gamma n_j)}$, $k = 1, \dots, K$, where $\gamma > 0$ is the sampling temperature.

Given a target budget $|\mathcal{D}_{\text{few}}| = \rho |\mathcal{D}_{\text{val}}|$ with $\rho \in [0.01, 0.03]$, we sample queries without replacement by first sampling a bin index k from the categorical distribution $\{p_k\}$, and then sampling a query q uniformly from B_k . Repeating this procedure until $|\mathcal{D}_{\text{few}}|$ is reached yields a few shot subset that preserves the warm-up difficulty distribution while covering both easy and hard queries.

Calibrated surrogate score The final score used to evaluate the workflow is:

$$\hat{S}_{t+1} = (1 - \alpha_{t+1}) S_{t+1}^{\text{pred}} + \alpha_{t+1} S_{t+1}^{\text{few}} \quad (8)$$

where $\alpha_{t+1} \in [0, 1]$ controls how strongly we trust the few shot execution relative to self prediction.

In practice, α_{t+1} is set adaptively based on the discrepancy between S_{t+1}^{pred} and S_{t+1}^{few} and the few shot sampling ratio. Let $\epsilon_{t+1} = |S_{t+1}^{\text{pred}} - S_{t+1}^{\text{few}}|$, and let $\tau > 0$ be a calibration tolerance. We also define the few shot ratio $\psi = \frac{|\mathcal{D}_{\text{few}}|}{|\mathcal{D}_{\text{val}}|}$, and an upper bound $\alpha_{\max} \in (0, 1]$ on the calibration strength. We set α_{t+1} as:

$$\alpha_{t+1} = \begin{cases} 0, & \text{if } \epsilon_{t+1} \leq \tau, \\ \min\left(\frac{\epsilon_{t+1}}{\tau} \psi, \alpha_{\max}\right), & \text{otherwise.} \end{cases} \quad (9)$$

Intuitively, when S_{t+1}^{pred} and S_{t+1}^{few} agree within the tolerance τ , we keep $\hat{S}_{t+1}$ equal to the self prediction ($\alpha_{t+1} = 0$). When the discrepancy exceeds τ , we increase α_{t+1} proportionally to how many tolerance units ϵ_{t+1} spans and to the few shot ratio

Method	DROP		HotpotQA		GSM8K		MATH		HumanEval		MBPP	
	Perf	Cost	Perf	Cost	Perf	Cost	Perf	Cost	Perf	Cost	Perf	Cost
ScoreFlow	91.48%	2.06e7	76.85%	2.72e7	97.79%	2.83e7	84.35%	4.03e7	98.91%	3.86e6	89.63%	3.93e6
AgentPrune	89.22%	6.65e6	76.73%	2.81e7	95.42%	1.07e7	81.53%	2.82e7	98.03%	1.65e6	88.37%	1.34e6
Aflow	90.30%	1.40e7	77.57%	3.11e7	96.58%	1.11e7	82.92%	3.06e7	99.24%	5.26e6	89.74%	3.92e6
SCALE	88.96%	3.27e6	77.41%	1.43e7	95.83%	4.04e6	81.70%	7.04e6	98.47%	1.09e6	90.32%	6.83e5
Δ	-1.34%	$\downarrow$ 76%	-0.16%	$\downarrow$ 54%	-0.75%	$\downarrow$ 63%	-1.22%	$\downarrow$ 77%	-0.77%	$\downarrow$ 79%	+0.58%	$\downarrow$ 83%
SCALE_ S^{pred}	78.61%	4.45e6	75.40%	1.50e7	92.51%	7.69e6	76.33%	7.55e6	09.16%	2.30e5	90.62%	5.75e5
SCALE_ S^{few}	86.06%	2.85e6	73.05%	1.55e7	94.22%	6.28e6	78.10%	5.49e6	96.95%	4.13e5	91.20%	6.17e5
SCALE_ S^{conf}	00.00%	1.78e6	00.00%	1.04e7	0.00%	4.43e6	57.61%	8.48e6	97.71%	3.15e5	88.86%	1.24e6

Table 2: Test performance **Perf** and token number **Cost** comparison across six benchmarks. Δ reports the performance change and cost reduction of **SCALE** relative to **Aflow**.

ψ , but cap it by $\alpha_{\max}$. This realizes the heuristic that large disagreements are more likely due to prediction error and should be corrected more aggressively toward the few shot estimate, while still respecting the limited size of $\mathcal{D}_{\text{few}}$.

To this end, we replace full validation execution with the calibrated surrogate score $\hat{S}_i$ in Equation 8. The surrogate scores of low-cost stage $\{\hat{S}_i\}_{i>M}$ are used for selection, expansion, and experience update together with the full-execution scores in warm up stage $\{S_i^{\text{exec}}\}_{i=1}^M$.

Overall, **SCALE** eliminates full validation runs in the main search phase: only few shot execution is performed to calibrate the optimizer’s self prediction. As a result, token number scales with $|\mathcal{D}_{\text{few}}|$ instead of $|\mathcal{D}_{\text{val}}|$, cutting the token cost substantially while maintaining the performance.

5 Experiments

We empirically evaluate our framework on multiple benchmarks, aiming to answer these two questions: **Q1**: Can **SCALE** reduce cost while maintaining performance? **Q2**: Why does calibrated prediction approximate full execution score well?

5.1 Experimental Setup

Agent and Optimizer In all experiments, we adopt Qwen-Plus (Hui et al., 2024) as base models $\{O_{\theta_i}\}_{i=1}^n$ of executor agents and Qwen3-8B (Yang et al., 2025) as the workflow optimizer ϕ .

Benchmarks We evaluate on six benchmarks spanning diverse domains: DROP (Dua et al., 2019) and HotpotQA (Yang et al., 2018) (multi-hop reasoning), GSM8K (Cobbe et al., 2021) and MATH (Hendrycks et al., 2021) (mathematical reasoning), HumanEval (Chen, 2021) and MBPP (Austin et al., 2021) (program synthesis). The dataset splits follow Zhang et al. (2024b).

Baselines We compare **SCALE** with both task-level methods Aflow (Zhang et al., 2024b), Agent-Prune (Zhang et al., 2024a) and query-level method ScoreFlow (Wang et al., 2025a). We also include internal ablations that vary only in the surrogate score: **SCALE_ S^{pred}** uses uncalibrated self prediction S^{pred} ; **SCALE_ S^{few}** uses few shot score S^{few} ; **SCALE_ S^{conf}** uses self-confidence S^{conf} , defined as $S_{t+1}^{\text{conf}} = \phi(W_i, \hat{P}_{t+1}^{\text{optimizer}})$, where $\hat{P}_{t+1}^{\text{optimizer}}$ appends “Output your confidence on the answer” to the optimizer prompt $P_{t+1}^{\text{optimizer}}$. Unlike S^{pred} , S^{conf} isolates the effect of our dedicated evaluation prompt P^{eval} by contrasting it with a minimal modification of the generation prompt.

Metrics The main metrics reported are: test performance and overall LLM token number incurred excluding test-time execution. Each benchmark’s performance metric is the same as in (Zhang et al., 2024b). Importantly, the token number is computed differently across methods to reflect their distinct optimization paradigms: for *task-level* approaches, it includes all tokens consumed in evaluating candidate workflows on the validation set to select the single best task-level policy; for *query-level* method, it includes tokens used in generating data for training the optimizer ϕ , evaluating query-level workflows during validation, and generating the query-level workflow at test time.

5.2 Main Results and Ablation Study

Across all six benchmarks, the results in Table 2 show that our method substantially reduces token number while maintaining test performance. Compared with Aflow, our method **SCALE** yields an average test performance drop of only 0.61%, while reducing total token number by 54% to 83% across all benchmarks. This demonstrates that full validation execution is not necessary for discovering high-quality workflows. Relative to

AgentPrune (Zhang et al., 2024a), which lowers token cost through structural pruning but still relies on repeated execution-based evaluation, our method achieves further token number reduction while delivering comparable performance, indicating that replacing the evaluation paradigm itself yields greater savings than modifying the search structure alone.

The ablation variants further clarify where these gains come from. **SCALE** $_S^{\text{pred}}$ drastically reduces cost but exhibits noticeable performance degradation, suggesting that self prediction alone suffers from model bias. **SCALE** $_S^{\text{few}}$ improves robustness but still requires larger execution budgets. **SCALE** $_S^{\text{conf}}$ performs very bad across tasks, confirming that naive confidence signals are unreliable surrogates for workflow quality. In contrast, using calibrated prediction as surrogate evaluation **SCALE** strikes a stable balance between cost and performance by combining model-based prediction with few shot execution signals.

Overall, these results answer **Q1** affirmatively: **SCALE** maintains the test performance while reducing searching token number by up to 83%.

5.3 Comparing Different Surrogate Evaluation Methods

To answer **Q2**, For every workflow generated in a full Aflow’s run, we log and compare S^{exec} together four surrogate scores. Surrogate scores are intended to take place of $S_{i>M}^{\text{exec}}$ to guide selection and expansion, so a good surrogate should satisfy two properties: First, the value should be close to S^{exec} . If the scales differ, the search will unfairly favor one side of the two-stages. Second, it should induce a ranking consistent with S^{exec} to reliably distinguish workflows.

Let $\{x_t\}_{t=1}^T$ denote the sequence of S^{exec} along the search progress, and $\{y_t\}_{t=1}^T$ the corresponding surrogate scores. We quantify agreement between x and y with:

(1) Pearson correlation: $\text{Pearson}(x, y) = \frac{\sum_t (x_t - \bar{x})(y_t - \bar{y})}{\sqrt{\sum_t (x_t - \bar{x})^2} \sqrt{\sum_t (y_t - \bar{y})^2}}$, which measures linear ranking consistency. The larger pearson correlation means the better linear ranking consistency between two metrics.

(2) First-order difference cosine similarity: $\text{DiffCos}(x, y) = \frac{\sum_{t=2}^T \Delta x_t \Delta y_t}{\sqrt{\sum_{t=2}^T \Delta x_t^2} \sqrt{\sum_{t=2}^T \Delta y_t^2}}$, where $\Delta x_t = x_t - x_{t-1}$ and $\Delta y_t = y_t - y_{t-1}$. It captures whether the direction of round-to-round changes is aligned between two workflow’s measure metrics.

Method	Pearson $\uparrow$	DiffCos $\uparrow$	MAE $\downarrow$
S^{conf}	-0.0576	0.0377	0.0802
S^{pred}	0.0517	0.1275	0.0511
S^{few}	0.6827	0.6192	0.2160
$\hat{S}$	0.5217	0.5545	0.1634

Table 3: Agreement between surrogate evaluation metrics and full execution. $\hat{S}$ strikes a good balance between value accuracy and ranking consistency.

Similar to pearson correlation, the larger DiffCos means the two metrics are better aligned.

(3) Mean absolute error (MAE): $\text{MAE}(x, y) = \frac{1}{T} \sum_{t=1}^T |x_t - y_t|$ which directly measures value-level approximation quality of the surrogate evaluation methods.

Table 3 reports these metrics for different surrogates. S^{conf} performs poorly on all metrics. S^{pred} achieves lowest MAE but almost zero correlation, indicating that it roughly matches the average scale of S^{exec} yet fails to order different workflows. S^{few} shows strong Pearson correlation and DiffCos but suffers from large MAE due to high variance from small sample size. $\hat{S}$ combines the strengths of both: it improves correlation over self prediction while reducing MAE compared with few shot execution.

Overall, the answer to **Q2** is that the calibrated prediction $\hat{S}$ serves as the most effective surrogate for the full-execution score S^{exec} . As a *surrogate method*, it aligns best with S^{exec} in both value agreement and ranking consistency; as a *evaluation score* for task-level workflow generation, it maintains a competitive test performance while substantially reducing token cost.

6 Conclusion

We revisited agentic workflow generation and arrived at two main insights: (1) query-level workflow generation is not always necessary, as a small pool of task-level workflows already achieves strong coverage, and (2) execution-based task-level evaluation is extremely costly while providing limited benefit. Motivated by these findings, we developed a task-level workflow generation framework **SCALE** to replace costly evaluation with calibrated self prediction. Across multiple benchmarks, it maintains competitive test performance with an average degradation of just 0.61% compared to existing approach while reducing overall token number by up to 83%.

7 Limitations

This work still has some limitations. Although our method avoids the main drawbacks of both query-level and task-level workflow generation methods, it remains primarily based on task-level workflow generation and has not yet fully leveraged the generalization capability of task-level approaches together with the fine-grained adaptability of query-level methods. We also do not assess cross-domain generalization in this work.

References

Jacob Austin, Augustus Odena, Maxwell Nye, Maarten Bosma, Henryk Michalewski, David Dohan, Ellen Jiang, Carrie Cai, Michael Terry, Quoc Le, and 1 others. 2021. Program synthesis with large language models. *arXiv preprint arXiv:2108.07732*.

Mark Chen. 2021. Evaluating large language models trained on code. *arXiv preprint arXiv:2107.03374*.

Karl Cobbe, Vineet Kosaraju, Mohammad Bavarian, Mark Chen, Heewoo Jun, Lukasz Kaiser, Matthias Plappert, Jerry Tworek, Jacob Hilton, Reiichiro Nakano, and 1 others. 2021. Training verifiers to solve math word problems. *arXiv preprint arXiv:2110.14168*.

Dheeru Dua, Yizhong Wang, Pradeep Dasigi, Gabriel Stanovsky, Sameer Singh, and Matt Gardner. 2019. Drop: A reading comprehension benchmark requiring discrete reasoning over paragraphs. *arXiv preprint arXiv:1903.00161*.

Hongcheng Gao, Yue Liu, Yufei He, Longxu Dou, Chao Du, Zhijie Deng, Bryan Hooi, Min Lin, and Tianyu Pang. 2025. Flowreasoner: Reinforcing query-level meta-agents. *arXiv preprint arXiv:2504.15257*.

Dan Hendrycks, Collin Burns, Saurav Kadavath, Akul Arora, Steven Basart, Eric Tang, Dawn Song, and Jacob Steinhardt. 2021. Measuring mathematical problem solving with the math dataset. *arXiv preprint arXiv:2103.03874*.

Shengran Hu, Cong Lu, and Jeff Clune. 2024. Automated design of agentic systems. *arXiv preprint arXiv:2408.08435*.

Binyuan Hui, Jian Yang, Zeyu Cui, Jiayi Yang, Dayiheng Liu, Lei Zhang, Tianyu Liu, Jiajun Zhang, Bowen Yu, Keming Lu, and 1 others. 2024. Qwen2. 5-coder technical report. *arXiv preprint arXiv:2409.12186*.

Tian Liang, Zhiwei He, Wenxiang Jiao, Xing Wang, Yan Wang, Rui Wang, Yujiu Yang, Shuming Shi, and Zhaopeng Tu. 2024. Encouraging divergent thinking in large language models through multi-agent debate. In *Proceedings of the 2024 conference on empirical methods in natural language processing*, pages 17889–17904.

Boye Niu, Yiliao Song, Kai Lian, Yifan Shen, Yu Yao, Kun Zhang, and Tongliang Liu. 2025. Flow: Modularized agentic workflow automation. *arXiv preprint arXiv:2501.07834*.

Rafael Rafailov, Archit Sharma, Eric Mitchell, Christopher D Manning, Stefano Ermon, and Chelsea Finn. 2023. Direct preference optimization: Your language model is secretly a reward model. *Advances in neural information processing systems*, 36:53728–53741.

Noah Shinn, Federico Cassano, Ashwin Gopinath, Karthik Narasimhan, and Shunyu Yao. 2023. Reflexion: Language agents with verbal reinforcement learning. *Advances in Neural Information Processing Systems*, 36:8634–8652.

Yinjie Wang, Ling Yang, Guohao Li, Mengdi Wang, and Bryon Aragam. 2025a. Scoreflow: Mastering llm agent workflows via score-based preference optimization. *arXiv preprint arXiv:2502.04306*.

Zhexuan Wang, Yutong Wang, Xuebo Liu, Liang Ding, Miao Zhang, Jie Liu, and Min Zhang. 2025b. Agentdropout: Dynamic agent elimination for token-efficient and high-performance llm-based multi-agent collaboration. *arXiv preprint arXiv:2503.18891*.

Shengxiang Xu, Jiayi Zhang, Shimin Di, Yuyu Luo, Liang Yao, Hanmo Liu, Jia Zhu, Fan Liu, and Min-Ling Zhang. 2025. Robustflow: Towards robust agentic workflow generation. *arXiv preprint arXiv:2509.21834*.

An Yang, Anfeng Li, Baosong Yang, Beichen Zhang, Binyuan Hui, Bo Zheng, Bowen Yu, Chang Gao, Chengen Huang, Chenxu Lv, and 1 others. 2025. Qwen3 technical report. *arXiv preprint arXiv:2505.09388*.

Zhilin Yang, Peng Qi, Saizheng Zhang, Yoshua Bengio, William Cohen, Ruslan Salakhutdinov, and Christopher D Manning. 2018. Hotpotqa: A dataset for diverse, explainable multi-hop question answering. In *Proceedings of the 2018 conference on empirical methods in natural language processing*, pages 2369–2380.

Shunyu Yao, Jeffrey Zhao, Dian Yu, Nan Du, Izhak Shafran, Karthik R Narasimhan, and Yuan Cao. 2022. React: Synergizing reasoning and acting in language models. In *The eleventh international conference on learning representations*.

Rui Ye, Shuo Tang, Rui Ge, Yaxin Du, Zhenfei Yin, Siheng Chen, and Jing Shao. 2025. Mas-gpt: Training llms to build llm-based multi-agent systems. *arXiv preprint arXiv:2503.03686*.

Guibin Zhang, Luyang Niu, Junfeng Fang, Kun Wang, Lei Bai, and Xiang Wang. 2025. Multi-agent architecture search via agentic supernet. *arXiv preprint arXiv:2502.04180*.

Guibin Zhang, Yanwei Yue, Zhixun Li, Sukwon Yun, Guancheng Wan, Kun Wang, Dawei Cheng, Jeffrey Xu Yu, and Tianlong Chen. 2024a. Cut the crap: An economical communication pipeline for llm-based multi-agent systems. *arXiv preprint arXiv:2410.02506*.

Jiayi Zhang, Jinyu Xiang, Zhaoyang Yu, Fengwei Teng, Xionghui Chen, Jiaqi Chen, Mingchen Zhuge, Xin Cheng, Sirui Hong, Jinlin Wang, and 1 others. 2024b. Aflow: Automating agentic workflow generation. *arXiv preprint arXiv:2410.10762*.

Mingchen Zhuge, Wenyi Wang, Louis Kirsch, Francesco Faccio, Dmitrii Khizbullin, and Jürgen Schmidhuber. 2024. Gptswarm: Language agents as optimizable graphs. In *Forty-first International Conference on Machine Learning*.

A Appendix

A.1 Self Prediction Prompt Template

The following prompt template is used to guide an LLM-based evaluator in performing *self-prediction*, i.e., estimating the expected accuracy of a candidate workflow over the entire evaluation dataset before actual execution. The prompt is carefully structured to enforce rigorous static analysis while leveraging historical execution feedback for calibration. Key components include: (1) a clear task definition requiring both justification and a calibrated probability score; (2) contextual information about the dataset, few-shot examples, and the workflow’s code structure; (3) explicit descriptions of allowed LLM-based operators to validate correct usage; (4) reference experiences from prior rounds to enable prediction refinement through error reflection; and (5) strict validation rules (e.g., package imports, prompt definitions, operator interfaces) that trigger immediate failure (score = 0.0) if violated. The output format enforces structured reasoning via `<reason>` and `<box>` tags to ensure parseable and consistent responses.

```
1 SELF_PREDICTION_PROMPT = """
2 You are an expert evaluator of workflows
3 .
4 Your task is to predict the probability
5 that a given workflow will correctly
6 execute on the WHOLE DATASET,
7 which represents your estimation of its
8 overall accuracy.
9 Respond with a brief explanation first,
10 followed by a single floating-point
11 number between 0.0 and 1.0.

12 Dataset Description:
13 <dataset>
14 {dataset_description}
15 </dataset>
```

```
16 Few-shot samples of the Dataset (jsonl
17 format):
18 {dataset_few_shots}
19
20 Workflow to evaluate (python code):
21 <workflow>
22 {workflow}
23 </workflow>
24
25 Prompt used in the workflow (python code
26 ):
27 <prompt>
28 {prompt}
29 </prompt>
30
31 The workflow is Python code; the key
32 function is __call__(question),
33 which produces the workflow's
34 response.
35
36 The workflow may call LLM-based
37 operators described below:
38
39 <operator_description>
40 {operator_description}
41 </operator_description>
42
43 Reference experiences:
44 During the warm-up rounds, several
45 workflows have been executed and
46 evaluated.
47
48 Each record includes the round (the
49 iteration number of the workflow),
50 the score (the actual reward
51 obtained after execution), the
52 prediction (the reward you predicted
53 in the previous round), and the
54 python code of the workflow and
55 prompt. (these workflows use the
56 same operators as shown above in the
57 <operator_description></
58 operator_description>)
59
60 These experiences are provided to help
61 you calibrate your future
62 predictions by comparing your past
63 predicted rewards with the actual
64 scores, you can adjust your
65 estimation strategy to make your
66 predicted rewards as close as
67 possible to the real execution
68 results.
69
70 <experience>
71 {experience}
72 </experience>
73
74 **General Instructions for evaluation:**
75 1. Step by step, carefully check for
76 critical errors that could prevent
77 execution:
78
79 - Package check (VERY IMPORTANT): The
80 workflow code imports the required
81 packages (for example: import numpy,
82 asyncio and other commonly used
83 Python packages). If any package is
84 used in the workflow but missing or
85 commented out, output 0.0 directly
86 and do not continue other checks.
87
88 - Prompt check (VERY IMPORTANT): The
89 workflow code uses prompts written
90 in Python format.
```

```

844 45 If the workflow uses no prompts then
845 just continue other checks.
846 46 If the workflow uses prompts,
847 47 For every prompt referenced in the
848 workflow, you must verify that this
849 prompt is properly defined in the
850 prompt.py file (commonly imported as
851 prompt_custom).
852 48 A prompt is considered properly
853 defined only if: It appears in the
854 prompt.py file without being
855 commented out AND the prompt name
856 matches exactly (including
857 capitalization, underscores, and
858 punctuation).
859 49 If any prompt used in the workflow is
860 missing, misspelled, or commented
861 out in prompt.py, you must
862 immediately output 0.0.
863 50 Check ALL prompts used in the
864 workflow following the same rule.
865 51
866 52 - Operator check (VERY IMPORTANT):
867 The operator is provided in text
868 description format. If the workflow
869 uses an operator, it must be among
870 the operators defined in <
871 operator_description>. If the
872 workflow uses an undefined operator
873 (including mismatched names,
874 incorrect parameters, or improper
875 usage) OR the parameters passed when
876 using an operator do not comply
877 with the interface requirements
878 defined in operator_description,
879 output 0.0 directly and do not
880 continue other checks.
881 53
882 54 - Workflow check (VERY IMPORTANT):
883 The __call__ function must return
884 the output string of the workflow
885 and the token usage only, more or
886 less is totally wrong. The input of
887 the workflow are only the input
888 string, more or less is totally
889 wrong.
890 55
891 56 2. Step by step, Analyze whether the
892 workflow can logically solve the
893 queries in the WHOLE DATASET. Please
894 carefully analyze the function of
895 each operator and whether the
896 position of each operator can
897 smoothly promote the resolution of
898 the problem.
899 57
900 58 3. Consider potential hallucinations
901 from the operators, for now, we use
902 {backbone_model} as backbone model
903 in operators.
904 59
905 60 4. Evaluate carefully and
906 comprehensively across all query
907 types in the WHOLE DATASET.
908 61
909 62 5. Be fair and rational: do not easily
910 assign 0.0 unless there is a severe
911 problem, and avoid scoring 1.0 with
912 overconfidence.
913 63

```

```

64 6. There's no need to focus heavily on
65 output details, such as formatting
66 inconsistencies, since extraction is
67 handled simply or by specialized
68 subsequent steps.
69
70 Output format:
71 - Provide a brief explanation of your
72 reasoning in a <reason> tag.
73 - Wrap your final probability in a <box>
74 tag **after** the <reason>.
75
76 For example:
77 <reason>The workflow correctly calls all
78 operators and uses only defined
79 prompts.</reason>
80 <box>0.85</box>
81 """

```

Listing 1: The prompt template for self prediction.

A.2 Workflows Generated By SCALE

Here we show the best workflows generated by our method SCALE across six datasets.

```

1 from typing import Literal
2 import workspace_SCALE.DROP.workflows_43
3 import workspace_SCALE.DROP.workflows_43
4 from scripts.async_llm import
5 create_llm_instance
6
7 from scripts.evaluator import
8 DatasetType
9
10 class Workflow:
11     def __init__(
12         self,
13         name: str,
14         llm_config,
15         dataset: DatasetType,
16     ) -> None:
17         self.name = name
18         self.dataset = dataset
19         self.llm = create_llm_instance(
20             llm_config)
21         self.answer_generate = operator.
22 AnswerGenerate(self.llm)
23         self.custom = operator.Custom(
24             self.llm)
25         self.sc_ensemble = operator.
26 ScEnsemble(self.llm)
27         self.verify_output = operator.
28 Custom(self.llm) # Reusing Custom
29 as VerifyOutput
30
31 async def __call__(self, problem:
32 str):
33     """
34     Implementation of the workflow
35     """
36     # Step-by-step generation using
37 AnswerGenerate
38     initial_solution = await self.
39 answer_generate(input=problem)

```

```

977 30     initial_thought =
978     initial_solution['thought']
979 31     initial_answer =
980     initial_solution['answer']
981
982 32     # Refine using custom method
983     with clearer instruction and context
984 34     refined_solutions = []
985 35     for _ in range(5): # Increase
986     number of refinements for better
987     consensus
988 36     refined_sol = await self.
989     custom(
990 37         input=f"{problem}\n\
991     nInitial Thought: {initial_thought}\
992     nInitial Answer: {initial_answer}",
993 38         instruction=
994     prompt_custom.
995     REFINE_WITH_CONTEXT_AND_ACCURATE_PERCENTAGE_CALCULATION_PROMPT
996
997 39     )
998 40     refined_solutions.append(
999     refined_sol['response'])
1000
1001 41     # Use self-consistency ensemble
1002 42     to select the best solution
1003 43     ensemble_result = await self.
1004     sc_ensemble(solutions=
1005     refined_solutions)
1006 44     raw_final_response =
1007     ensemble_result['response']
1008
1009 45     # Verification step to ensure
1010 46     final output format compliance
1011 47     verified_solution = await self.
1012     verify_output(
1013 48         input=raw_final_response,
1014 49         instruction=prompt_custom.
1015     VERIFY_OUTPUT_FORMAT_PROMPT
1016 50     )
1017 51
1018 52     return verified_solution['
1019     response'], self.llm.
1020     get_usage_summary()["total_tokens"]

```

Listing 2: The best workflow generated by SCALE for DROP

```

1021 1 VERIFY_OUTPUT_FORMAT_PROMPT = """Ensure
1022     the response is properly formatted
1023     with the answer() wrapper. If the
1024     answer is not wrapped in answer(),
1025     add it around the final result. For
1026     example:
1027 2 - If the answer is a number: answer(42)
1028 3 - If the answer is text: answer(Wilson)
1029 4 - If the answer is a vector: answer(\
1030     begin{pmatrix} 1 \\\ 2 \end{
1031     pmatrix})
1032 5 - If the answer is a range: answer(1-10)
1033 6 - If the answer is a percentage: answer
1034     (95%)
1035
1036 7 The response should contain only the
1037     final answer wrapped in answer()
1038     with no additional text or
1039     explanation."""
1040
1041 9 VERIFY_MATH_REASONING_PROMPT = """Check
1042     the mathematical reasoning in the

```

```

11     solution. Verify that:
12 1. All calculations are correct
13 2. The logic follows mathematical
14     principles
15 3. The steps lead to the correct
16     conclusion
17 4. The final answer is consistent with
18     the reasoning
19
20 If any errors are found, correct them
21 and provide the corrected solution.
22 Ensure the final answer is wrapped
23 in answer() with no additional text.
24 """

```

Listing 3: The prompt used in the executor agents of best workflow generated by SCALE for DROP

```

1 from typing import Literal
2 import workspace_SCALE.HotpotQA.
3     workflows.template.operator as
4     operator
5
6 import workspace_SCALE.HotpotQA.
7     workflows.round_15.prompt as
8     prompt_custom
9
10 from scripts.async_llm import
11     create_llm_instance
12
13
14
15 from scripts.evaluator import
16     DatasetType
17
18
19 class Workflow:
20     def __init__(
21         self,
22         name: str,
23         llm_config,
24         dataset: DatasetType,
25     ) -> None:
26         self.name = name
27         self.dataset = dataset
28         self.llm = create_llm_instance(
29             llm_config)
30         self.custom = operator.Custom(
31             self.llm)
32         self.answer_generate = operator.
33             AnswerGenerate(self.llm)
34         self.sc_ensemble = operator.
35             ScEnsemble(self.llm)
36         self.refine = operator.Custom(
37             self.llm) # New operator for post-
38             ensemble refinement
39
40     async def __call__(self, problem:
41         str):
42         """
43         Implementation of the workflow
44         """
45         # Generate multiple reasoning
46         paths
47         solutions = []
48         for _ in range(3):
49             solution = await self.
50             answer_generate(input=problem)
51             solutions.append(solution[
52                 'answer'])
53

```



```

1107 34      # Self-consistency ensemble to
1108      choose most frequent answer
1109 35      ensemble_response = await self.
1110      sc_ensemble(solutions=solutions)
1111 36      selected_answer =
1112      ensemble_response['response']
1113 37
1114 38      # Verify ensemble result for
1115      validity and confidence before
1116      refining
1117 39      verify_response = await self.
1118      custom(
1119 40          input=f"Question: {problem}\
1120          nCandidate Answer: {selected_answer}"
1121      ,
1122          instruction=prompt_custom.
1123      VERIFY_ENSEMBLE_CONFIDENCE_PROMPT
1124      )
1125 43      is_valid = "answer(valid)" in
1126      verify_response['response']
1127 44
1128 45      if not is_valid:
1129 46          fallback_response = await
1130      self.custom(
1131 47          input=problem,
1132 48          instruction=
1133      prompt_custom.
1134      FALLBACK_ANSWER_GENERATION_PROMPT
1135      )
1136 50      selected_answer =
1137      fallback_response['response']
1138 51
1139 52      # Refine the selected answer
1140      with a stricter categorical/entity-
1141      focused prompt
1142 53      refined_response = await self.
1143      refine(
1144 54          input=f"Question: {problem}\
1145          nSelected Answer: {selected_answer}"
1146      ,
1147          instruction=prompt_custom.
1148      REFINE_TO_ENTITY_OR_CATEGORY_PROMPT
1149      )
1150 57      refined_answer =
1151      refined_response['response']
1152 58
1153 59      # Check if refined answer is
1154      valid; if not, trigger fallback
1155 60      if "answer(None)" in
1156      refined_answer or not refined_answer
1157      .strip():
1158 61          fallback_response = await
1159      self.custom(
1160 62          input=problem,
1161 63          instruction=
1162      prompt_custom.
1163      FALLBACK_ANSWER_GENERATION_PROMPT
1164      )
1165 65      refined_answer =
1166      fallback_response['response']
1167 66
1168 67      # Final formatting verification
1169      verified_result = await self.
1170      custom(
1171 69          input=f"Question: {problem}\
1172          nCandidate Answer: {refined_answer}"
1173      ,
1174          instruction=prompt_custom.
1175      FINAL_FORMAT_VERIFICATION_PROMPT
1176      )

```

```

72      final_output = verified_result['
73      response']
74      return final_output, self.llm.
      get_usage_summary()["total_tokens"]

```

Listing 4: The best workflow generated by SCALE for HotpotQA

```

1 VERIFY_ENSEMBLE_CONFIDENCE_PROMPT = """
    You are given a question and a
    candidate answer derived via
    ensemble. Determine whether the
    candidate answer is logically
    consistent with the question and
    shows sufficient confidence. If the
    answer is relevant and confident,
    respond with answer(valid).
    Otherwise, respond with answer(
    invalid). Do not explain or rephrase
    just evaluate confidence and
    relevance."""
2
3 REFINE_TO_ENTITY_OR_CATEGORY_PROMPT = """
    You are given a question and a
    candidate answer. Your task is to
    reformulate the candidate answer
    into a precise categorical label or
    named entity that best fits the
    question. Avoid explanatory or vague
    language. Return only the most
    accurate concise form, such as a
    person's name, a location, a
    historical event, or a categorical
    label. Do not add punctuation or
    quotation marks. Always wrap your
    final output in answer(...).
    Examples: answer(Polish independence
    ), answer(William Shakespeare),
    answer(no), answer(chronological
    collection of critical quotations).
    """
4
5 FINAL_FORMAT_VERIFICATION_PROMPT = """You
    are tasked with extracting and
    formatting the final answer from a
    candidate answer such that it
    precisely matches the expected
    format. Avoid including any
    descriptive or explanatory text.
    Focus on named entities, binary
    responses, or categorical labels as
    appropriate. For named entities (
    people, places, works), return only
    the name. For yes/no questions,
    return exactly "yes" or "no". For
    categorical responses, return the
    exact category. Always wrap your
    final response in answer(...).
    Examples: answer(Limbo), answer(no),
    answer(Southern Isles). If the
    candidate answer contains multiple
    possibilities, choose the most
    likely one based on the question. If
    the candidate is unclear, make a
    best-guess effort to extract the
    intended answer. Do not add quotes
    or extra punctuation. Do not explain
    your choice."""
6

```

```

1243 7 Fallback_ANSWER_GENERATION_PROMPT = """
1244     Given the original question, please
1245     generate a concise and direct answer
1246     focusing strictly on the key entity
1247     or fact requested. Avoid
1248     explanations or additional
1249     commentary. Always wrap your final
1250     output in answer(...). Example:
1251     Question: Who wrote Pride and
1252     Prejudice? Output: answer(Jane
1253     Austen)"""

```

Listing 5: The prompt used in the executor agents of best workflow generated by SCALE for HotpotQA

1254

```

1255 1 from typing import Literal
1256 2 import workspace_calibrated_prediction.
1257   GSM8K.workflows.template.operator as
1258   operator
1259 3 import workspace_calibrated_prediction.
1260   GSM8K.workflows.round_7.prompt as
1261   prompt_custom
1262 4 from scripts.async_llm import
1263   create_llm_instance
1264 5
1265 6
1266 7 from scripts.evaluator import
1267   DatasetType
1268 8
1269 9 class Workflow:
1270 10     def __init__(
1271 11         self,
1272 12         name: str,
1273 13         llm_config,
1274 14         dataset: DatasetType,
1275 15     ) -> None:
1276 16         self.name = name
1277 17         self.dataset = dataset
1278 18         self.llm = create_llm_instance(
1279 19             llm_config)
1280 20         self.custom = operator.Custom(
1281 21             self.llm)
1282 22         self.programmer = operator.
1283   Programmer(self.llm)
1284 23         self.sc_ensemble = operator.
1285   ScEnsemble(self.llm)
1286 24
1287 25     async def __call__(self, problem:
1288   str):
1289 26         """
1290 27         Implementation of the workflow
1291 28         """
1292 29         # Step 1: Generate multiple
1293   solutions using Programmer for
1294   diverse computation paths
1295 30         solutions = []
1296 31         for _ in range(3):
1297 32             solution = await self.
1298   programmer(problem=problem, analysis
1299   ="Solve the following math problem
1300   precisely. Return only the final
1301   numeric result.")
1302 33             solutions.append(solution['
1303   output'])
1304 34
1305 35         # Step 2: Use ScEnsemble to
1306   select the most consistent solution
1307   among the generated ones

```

```

34         ensemble_result = await self.
sc_ensemble(solutions=solutions,
problem=problem)
35
36         # Step 3: Format the selected
result properly using Custom to
ensure it meets expected structure
37         formatted_solution = await self.
custom(
38             input=f"Problem: {problem}\
nComputed Result: {ensemble_result['
response']}",
39             instruction=prompt_custom.
FORMAT_ANSWER_PROMPT
40         )
41
42         # Step 4: Validate that the
result makes sense in context (e.g.,
not negative where inappropriate,
correct order of magnitude)
43         validated_solution = await self.
custom(
44             input=f"Problem: {problem}\
nFormatted Result: {
formatted_solution['response']}",
45             instruction=prompt_custom.
VALIDATE_NUMERIC_RESULT_PROMPT
46         )
47
48         # Step 5: Final formatting
verification to ensure answer is
boxed correctly
49         final_result = await self.custom
(
50             input=f"Problem: {problem}\
nValidated Result: {
validated_solution['response']}",
51             instruction=prompt_custom.
FINAL_BOXING_CHECK_PROMPT
52         )
53
54         return final_result['response'],
self.llm.get_usage_summary()["
total_tokens"]
1308
1309
1310
1311
1312
1313
1314
1315
1316
1317
1318
1319
1320
1321
1322
1323
1324
1325
1326
1327
1328
1329
1330
1331
1332
1333
1334
1335
1336
1337
1338
1339
1340
1341
1342
1343
1344
1345
1346
1347
1348
1349
1350
1351

```

Listing 6: The best workflow generated by SCALE for GSM8K

```

1 FORMAT_ANSWER_PROMPT = """You are given a
math problem and its computed
numeric result. Your task is to
format the result in a standardized
way by placing it inside \boxed{}.
Only return the final formatted
answer without any additional text
or explanation. For example, if the
result is 123, return \boxed{123}.
"""
2
3
4 VALIDATE_NUMERIC_RESULT_PROMPT = """You
are given a math word problem and a
formatted numeric result. Check
whether the result is logically
reasonable in the context of the
problem (e.g., not negative when
expecting a count, correct magnitude
). If it seems incorrect, estimate a
plausible value and return it in
the same \boxed{} format. Otherwise
1352
1353
1354
1355
1356
1357
1358
1359
1360
1361
1362
1363
1364
1365
1366
1367
1368
1369
1370
1371
1372
1373

```

```

1374         , return the original result in \\
1375         boxed{} format. Only return the
1376         final result in \\boxed{}."""
1377
1378
1379 7 FINAL_BOXING_CHECK_PROMPT = """You are
1380     given a math problem and a validated
1381     numeric result. Ensure that the
1382     final result is enclosed in \\boxed
1383     {} and represents a clean numeric
1384     answer without any extra commentary
1385     or formatting issues. Return only
1386     the properly boxed result."""

```

Listing 7: The prompt used in the executor agents of best workflow generated by SCALE for GSM8K

```

1387
1388
1389 1 from typing import Literal
1390 2 import workspace_SCALE.MATH.workflows.
1391   template.operator as operator
1392 3 import workspace_SCALE.MATH.workflows.
1393   round_20.prompt as prompt_custom
1394 4 from scripts.async_llm import
1395   create_llm_instance
1396
1397 5
1398 6
1399 7 from scripts.evaluator import
1400   DatasetType
1401
1402 8
1403 9 class Workflow:
1404 10     def __init__(
1405 11         self,
1406 12         name: str,
1407 13         llm_config,
1408 14         dataset: DatasetType,
1409 15     ) -> None:
1410 16         self.name = name
1411 17         self.dataset = dataset
1412 18         self.llm = create_llm_instance(
1413 19             llm_config)
1414 20
1415 21         self.custom = operator.Custom(
1416 22             self.llm)
1417 23
1418 24         self.verify_format = operator.
1419   Custom(self.llm)
1420 25
1421 26         self.verify_math = operator.
1422   Custom(self.llm)
1423 27
1424 28     async def __call__(self, problem:
1425   str):
1426 29         """
1427 30         Implementation of the workflow
1428 31         """
1429 32         solution = await self.custom(
1430 33             input=problem, instruction="")
1431 34
1432 35         # Verify mathematical reasoning
1433 36         verified_solution = await self.
1434   verify_math(input=solution['response
1435   '], instruction=prompt_custom.
1436   VERIFY_MATH_REASONING_PROMPT)
1437
1438 37
1439 38         # Verify and format the output
1440 39         to ensure it has answer() wrapper
1441 40         formatted_solution = await self.
1442   verify_format(input=
1443   verified_solution['response'],

```

```

1438         instruction=prompt_custom.
1439         VERIFY_OUTPUT_FORMAT_PROMPT)
1440
1441 35
1442 36         return formatted_solution['
1443   response'], self.llm.
1444         get_usage_summary()["total_tokens"]

```

Listing 8: The best workflow generated by SCALE for MATH

```

1444 1 VERIFY_OUTPUT_FORMAT_PROMPT = """Ensure
1445   the response is properly formatted
1446   with the answer() wrapper. If the
1447   answer is not wrapped in answer(),
1448   add it around the final result. For
1449   example:
1450 2 - If the answer is a number: answer(42)
1451 3 - If the answer is text: answer(Wilson)
1452 4 - If the answer is a vector: answer(\\
1453   begin{pmatrix} 1 \\\\ 2 \\end{
1454   pmatrix})
1455 5 - If the answer is a range: answer(1-10)
1456 6 - If the answer is a percentage: answer
1457   (95%)
1458 7
1459 8 The response should contain only the
1460   final answer wrapped in answer()
1461   with no additional text or
1462   explanation."""
1463 9
1464 10 VERIFY_MATH_REASONING_PROMPT = """Check
1465   the mathematical reasoning in the
1466   solution. Verify that:
1467 11 1. All calculations are correct
1468 12 2. The logic follows mathematical
1469   principles
1470 13 3. The steps lead to the correct
1471   conclusion
1472 14 4. The final answer is consistent with
1473   the reasoning
1474 15
1475 16 If any errors are found, correct them
1476   and provide the corrected solution.
1477   Ensure the final answer is wrapped
1478   in answer() with no additional text.
1479   """

```

Listing 9: The prompt used in the executor agents of best workflow generated by SCALE for MATH

```

1480
1481
1482 1 from typing import Literal
1483 2 import workspace_SCALE.HumanEval.
1484   workflows.template.operator as
1485   operator
1486 3 import workspace_SCALE.HumanEval.
1487   workflows.round_11.prompt as
1488   prompt_custom
1489 4 from scripts.async_llm import
1490   create_llm_instance
1491
1492 5
1493 6
1494 7 from scripts.evaluator import
1495   DatasetType
1496
1497 8
1498 9 class Workflow:
1499 10     def __init__(
1500 11         self,
1501 12         name: str,

```

```

1499 13         llm_config,
1500 14         dataset: DatasetType,
1501 15     ) -> None:
1502 16         self.name = name
1503 17         self.dataset = dataset
1504 18         self.llm = create_llm_instance(
1505         llm_config)
1506 19         self.custom = operator.Custom(
1507         self.llm)
1508 20         self.custom_code_generate =
1509         operator.CustomCodeGenerate(self.llm
1510         )
1511 21         self.sc_ensemble = operator.
1512         ScEnsemble(self.llm)
1513 22         self.test = operator.Test(self.
1514         llm)
1515 23
1516 24     async def __call__(self, problem:
1517     str, entry_point: str):
1518 25         """
1519 26         Implementation of the workflow
1520 27         Custom operator to generate
1521         anything you want.
1522 28         But when you want to get
1523         standard code, you should use
1524         custom_code_generate operator.
1525 29         """
1526 30         # Rephrase the problem for
1527         clarity
1528 31         rephrased_problem = await self.
1529         custom(input="", instruction=
1530         prompt_custom.
1531         REPHRASE_PROBLEM_PROMPT + problem)
1532 32         clarified_problem =
1533         rephrased_problem['response']
1534 33
1535 34         # Generate multiple solutions
1536         for ensemble
1537 35         solutions = []
1538 36         tested_solutions = []
1539 37         for _ in range(5):
1540 38             solution = await self.
1541             custom_code_generate(problem=
1542             clarified_problem, entry_point=
1543             entry_point, instruction="")
1544 39
1545 40             # Reflect on the generated
1546             solution to improve it
1547 41             reflection_prompt = f"Review
1548             the following code solution and fix
1549             any logical or syntax errors:\n\n
1550             {solution['response']}"
1551 42             reflected_solution = await
1552             self.custom(input=clarified_problem,
1553             instruction=reflection_prompt)
1554 43
1555 44             solutions.append(
1556             reflected_solution['response'])
1557 45
1558 46             # Pre-test each refined
1559             solution to filter valid ones early
1560 47             test_result = await self.
1561             test(problem=problem, solution=
1562             reflected_solution['response'],
1563             entry_point=entry_point)
1564 48             if test_result['result']:
1565 49                 tested_solutions.append(
1566                 test_result['solution'])
1567 50
1568 51         # Prioritize validated solutions

```

```

52         ; fallback to all if none pass
53         if tested_solutions:
54             ensemble_input =
55             tested_solutions
56         else:
57             ensemble_input = solutions
58
59         # Use ScEnsemble to select the
60         most consistent solution
61         ensemble_result = await self.
62         sc_ensemble(solutions=ensemble_input
63         , problem=problem)
64         final_solution = ensemble_result
65         ['response']
66
67         return final_solution, self.llm.
68         get_usage_summary()["total_tokens"]

```

Listing 10: The best workflow generated by SCALE for HumanEval

```

1 REPHRASE_PROBLEM_PROMPT = """Please
2 rephrase the following programming
3 problem in clearer terms, making
4 sure to highlight the key
5 requirements and expected output
6 format.Problem: """

```

Listing 11: The prompt used in the executor agents of best workflow generated by SCALE for HumanEval

```

1592
1593 1 from typing import Literal
1594 2 import workspace_SCALE.MBPP.workflows.
1595     template.operator as operator
1596 3 import workspace_SCALE.MBPP.workflows.
1597     round_7.prompt as prompt_custom
1598 4 from scripts.async_llm import
1599     create_llm_instance
1600 5
1601 6
1602 7 from scripts.evaluator import
1603     DatasetType
1604 8
1605 9 class Workflow:
1606 10     def __init__(
1607 11         self,
1608 12         name: str,
1609 13         llm_config,
1610 14         dataset: DatasetType,
1611 15     ) -> None:
1612 16         self.name = name
1613 17         self.dataset = dataset
1614 18         self.llm = create_llm_instance(
1615         llm_config)
1616 19
1617 20         self.custom = operator.Custom(
1618         self.llm)
1619 21         self.custom_code_generate =
1620         operator.CustomCodeGenerate(self.llm
1621         )
1622 22         self.test = operator.Test(self.
1623         llm)
1624 23
1625 24     async def __call__(self, problem:
1626     str, entry_point: str):
1627 25         """
1628 26         Implementation of the workflow

```



```

1629 27     Custom operator to generate
1630     anything you want.
1631 28     But when you want to get
1632     standard code, you should use
1633     custom_code_generate operator.
1634 29     """
1635 30     # Generate the initial solution
1636 31     solution = await self.
1637     custom_code_generate(problem=problem
1638     , entry_point=entry_point,
1639     instruction="Generate Python code
1640     that solves the given problem. Make
1641     sure to return the result of the
1642     function, not just print it.")
1643 32
1644 33     # Verify that the solution has
1645     proper format and return statements
1646 34     verified_solution = await self.
1647     custom(input=f"Problem: {problem}\n
1648     Entry point: {entry_point}\n
1649     Solution:\n{solution['response']}",
1650     instruction=prompt_custom.
1651     VERIFY_CODE_FORMAT_PROMPT)
1652 35
1653 36     # Test the verified solution to
1654     ensure it works correctly
1655 37     test_result = await self.test(
1656     problem=problem, solution=
1657     verified_solution['response'],
1658     entry_point=entry_point)
1659 38
1660 39     final_solution = test_result['
1661     solution'] if test_result['result']
1662     else verified_solution['response']
1663 40
1664 41     return final_solution, self.llm.
1665     get_usage_summary()["total_tokens"]

```

Listing 12: The best workflow generated by SCALE for MBPP

```

1666 1 VERIFY_CODE_FORMAT_PROMPT = """Verify
1667     that the provided code solution has
1668     the correct function signature and
1669     includes proper return statements.
1670     The solution must:
1671 2 1. Contain the function with the exact
1672     name specified in the entry_point
1673 3 2. Include a return statement that
1674     returns the result of the function (
1675     not just print it)
1676 4 3. Follow proper Python syntax
1677 5
1678 6 If the code is missing the function
1679     signature or return statement,
1680     please fix it. Return the corrected
1681     code."""

```

Listing 13: The prompt used in the executor agents of best workflow generated by SCALE for MBPP

1682