

Adaptive Stopping for Multi-Turn LLM Reasoning

Xiaofan Zhou

Department of Computer Science
University of Illinois Chicago
Chicago, IL 60607, USA
xzhou77@uic.edu

Huy Nguyen

Data science
Augustana College
Rock Island, IL 61201, USA
huynguyen23@augustana.edu

Bo Yu

Department of Civil and Environmental Engineering
University of Utah
Salt Lake City, UT 84112, USA
bo.yu@utah.edu

Chenxi Liu

Department of Civil and Environmental Engineering
University of Utah
Salt Lake City, UT 84112, USA
chenxi.liu@utah.edu

Lu Cheng

Department of Computer Science
University of Illinois Chicago
Chicago, IL 60607, USA
lucheng@uic.edu

Abstract

Large Language Models (LLMs) increasingly rely on multi-turn reasoning and interaction, such as adaptive retrieval-augmented generation (RAG) and ReAct-style agents, to answer difficult questions. These methods improve accuracy by iteratively retrieving information, reasoning, or acting, but introduce a key challenge: **When should the model stop?** Existing approaches rely on heuristic stopping rules or fixed turn budgets and provide no formal guarantees that the final prediction still contains the correct answer. This limitation is particularly problematic in high-stakes domains such as finance and healthcare, where unnecessary turns increase cost and latency, while stopping too early risks incorrect decisions. Conformal prediction (CP) provides formal coverage guarantees, but existing LLM-CP methods only apply to a single model output and cannot handle multi-turn pipelines with adaptive stopping. To address this gap, we propose **Multi-Turn Language Models with Conformal Prediction (MiCP)**, the first CP framework for multi-turn reasoning. MiCP allocates different error budgets across turns, enabling the model to stop early while maintaining an overall coverage guarantee. We demonstrate MiCP on adaptive RAG and ReAct, where it achieves the target coverage on both single-hop and multi-hop question answering benchmarks while reducing the number of turns, inference cost, and prediction set size. We further introduce a new metric that jointly evaluates coverage validity and answering efficiency.

1 Introduction

Recent Large Language Models (LLMs) have become one of the most transformative technologies in scientific research (Rane et al., 2023), industry (Wang et al., 2025), and daily life assistance (Achiam et al., 2023). Beyond producing a direct answer in a single forward pass, modern LLM systems increasingly rely on multi-turn interaction and reasoning. When the initial context is insufficient, the model may iteratively retrieve additional evidence, generate intermediate reasoning steps, or interact with external tools before deciding on a final answer. Representative examples include adaptive Retrieval-Augmented Generation (RAG) (Moskvoretskii et al., 2025), which repeatedly retrieves documents until enough evidence has been gathered, and ReAct-style agents (Yao et al., 2022), which interleave reasoning and actions in a multi-turn loop.

These multi-turn strategies create a new challenge: deciding when to stop. If the model stops too early, it may fail to gather enough information and return an incorrect answer. If it continues for too many turns, the additional retrievals or reasoning steps increase latency, computation cost, and sometimes even introduce more errors. Existing approaches therefore rely on heuristic stopping rules, such as confidence thresholds, fixed retrieval budgets, or manually designed criteria. However, these heuristics provide no formal guarantee that the resulting prediction still contains the correct answer.

The need for principled uncertainty quantification becomes especially important in high-stakes applications such as finance and healthcare. Conformal prediction (CP) (Angelopoulos & Bates, 2021; Zhou et al., 2025) has recently emerged as a powerful framework to provide finite-sample coverage guarantees for LLM outputs. By calibrating a nonconformity score on held-out data, CP constructs a prediction set that contains the correct answer with a user-specified confidence. Existing LLM-CP methods have successfully been applied to short-answer tasks (Su et al., 2024a; Li et al., 2024) and long-form generation (Mohri & Hashimoto, 2024). Nevertheless, all prior approaches assume a single-shot prediction: the model produces one final answer, and CP is applied only after the entire process is complete.

In contrast, multi-turn LLM pipelines may terminate after different numbers of turns depending on the difficulty of the question. A conformal predictor for such systems must therefore answer two questions simultaneously: (1) whether the model should stop at the current turn, and (2) whether the resulting prediction set still satisfies the desired coverage guarantee. Naively applying standard CP independently at each turn fails to solve this problem. First, it does not provide a guarantee on the overall coverage across all possible stopping times. Second, the sampling cost grows combinatorially across turns. For example, if three candidate answers are sampled at each of five turns, then a naive approach would require 3^5 sampled trajectories. Finally, practical systems often need to operate under a limited computation budget, requiring the model to abstain or return a “cannot answer” decision when sufficient confidence cannot be achieved.

To address these challenges, we propose **Multi-Turn Language Models with Conformal Prediction (MiCP)**, a multi-level CP framework for multi-turn LLMs reasoning. MiCP assigns distinct error budgets to different turns, enabling the model to stop early when sufficient evidence has been accumulated while still maintaining a rigorous overall coverage guarantee. Intuitively, early turns are encouraged to answer easy questions quickly, whereas later turns reserve additional error budget for more difficult questions that require further retrieval or reasoning. The framework naturally supports both adaptive RAG and ReAct-style agents, and can additionally incorporate a rejection option for questions that remain too uncertain within the allowed budget. We further introduce a new metric that explicitly reflects the trade-off between reliability and efficiency, making it particularly suitable for evaluating multi-turn reasoning systems.

In summary, our work makes the following contributions:

- We propose MiCP, the first CP framework for multi-turn LLMs. MiCP provides formal coverage guarantees while simultaneously enabling early stopping in iterative pipelines.
- We develop a multi-level error allocation strategy that distributes the overall conformal error budget across turns. This allows LLMs to decide whether to stop or continue at

each turn while preserving the desired end-to-end coverage guarantee, even under a fixed computation budget and optional rejection mechanism.

- We introduce a new evaluation metric for multi-turn LLMs that jointly measures coverage validity and efficiency. The metric explicitly rewards methods that answer correctly with fewer turns and penalizes unnecessary retrieval or reasoning steps.
- Extensive experiments on both adaptive RAG and ReAct benchmarks demonstrate that MiCP consistently achieves the target coverage level while substantially reducing the number of turns, inference cost, and prediction set size on both single-hop and multi-hop question answering tasks.

2 Related Work

2.1 Multi-turn LLMs

In multi-turn settings, LLMs can iteratively request additional information before committing to a final answer. A prominent category is adaptive RAG, where models dynamically determine whether to retrieve and how much text to retrieve from external sources based on their own output signals, rather than answering directly in a single pass. IRCoT (Trivedi et al., 2023) interleaves retrieval with chain-of-thought reasoning, fetching additional passages when the reasoning steps so far are insufficient to produce the answer. Adaptive-RAG (Jeong et al., 2024) trains a T5-large classifier to predict whether retrieval is needed and how many turns to perform. FLARE (Jiang et al., 2023) triggers retrieval whenever token-level generation probability falls below a threshold. DRAGIN (Su et al., 2024b) leverages attention weights to identify salient keywords for formulating retrieval queries. Rowen (Ding et al., 2024) uses cross-lingual answer consistency to decide whether retrieval is necessary, while SeaKR (Yao et al., 2024) relies on internal hidden states for the same purpose. QucoRAG (Min et al., 2025) detects hallucination by measuring how frequently generated tokens overlap with training data. Moskvoretskii et al. (2025) provide a comprehensive study on how uncertainty estimation affects adaptive RAG performance. Beyond retrieval-augmented settings, ReAct (Yao et al., 2022) enables LLMs to interleave reasoning with action steps, generating queries to external tools when additional information is needed, and Reflexion (Shinn et al., 2023) allows LLMs to iteratively self-reflect on their own outputs to refine answers across multiple turns. In this work, we primarily evaluate our framework under the adaptive RAG setting and additionally experiment with the ReAct paradigm to demonstrate its generalizability.

2.2 CP for LLM

CP has become an important tool for guaranteeing the factuality of predictions from LLMs in high-stakes scenarios (Zhou et al., 2025; Campos et al., 2024), and has been applied to various LLM tasks (Sheng et al., 2025; Kumar et al., 2023; Ye et al., 2024). In early applications, (Quach et al., 2023) repeatedly sample outputs until the prediction set is confident enough to contain the ground truth, which is computationally expensive (Su et al., 2024a). Several subsequent methods (Su et al., 2024a; Li et al., 2024) instead draw a fixed number of samples and use frequency to measure uncertainty, improving efficiency. Orthogonally, (Schuster et al., 2022) allow the transformer to exit early at the cost of a bounded performance drop. However, the nearly infinite output space of LLMs makes achieving the target coverage intractable. A common strategy is to replace uncertain answers with more general ones. (Zhang et al., 2024) propose a conformal structure prediction framework for tasks whose labels lie in a directed acyclic graph. (Mohri & Hashimoto, 2024; Rubin-Toles et al., 2025), which can be viewed as a special variant of this approach, remove uncertain statements from long-form generations to maintain factuality or coherence. Another variant by (Li et al., 2024) assigns a “Can’t Answer” label to questions that the LLM cannot reliably answer within the fixed sample set. Our method extends the “Can’t Answer” mechanism to multi-turn LLM interactions, allowing the model to early stop when sufficient confidence is reached while preserving the coverage guarantee in the no-early-stop setting.

3 Preliminary

Foundations of Split CP. CP is a distribution-free framework for constructing prediction sets with finite-sample coverage guarantees (Shafer & Vovk, 2008; Angelopoulos & Bates, 2021). It operates as a post-hoc wrapper around any pre-trained model, without requiring any additional pretraining. More specifically, Split CP separates model training from threshold calibration. Given a dataset $\{(x_i, y_i)\}_{i=1}^n$, it is partitioned into a training set $\mathcal{D}_{\text{train}} = \{(x_i, y_i)\}_{i=1}^{n_{\text{train}}}$, a calibration set $\mathcal{D}_{\text{cal}} = \{(x_i, y_i)\}_{i=1}^{n_{\text{cal}}}$, and a test set $\mathcal{D}_{\text{test}} = \{(x_i, y_i)\}_{i=1}^{n_{\text{test}}}$, where x_i is the input question, y_i is its corresponding answer, n_* denotes the size of each respective subset, and $n = n_{\text{train}} + n_{\text{cal}} + n_{\text{test}}$. Let $S(f, (x, y))$ denote a nonconformity score that measures how unusual a label y is for input x under a model f trained on $\mathcal{D}_{\text{train}}$. In a classification setting, for instance, a common choice is

$$S(f, (x, y)) = 1 - f_y(x), \quad (1)$$

where $f_y(x)$ is the predicted probability that f assigns label y to input x . Using the calibration set $\{(x_i, y_i)\}_{i=1}^{n_{\text{cal}}}$, one evaluates the nonconformity scores $\{S(f, (x_i, y_i))\}_{i=1}^{n_{\text{cal}}}$. A threshold is then obtained by computing a quantile over these calibration scores at a user-specified error rate level $\alpha \in (0, 1)$:

$$\hat{q}_\alpha = \text{Quantile}\left(\{s_i\}_{i=1}^{n_{\text{cal}}}; \frac{[(n_{\text{cal}} + 1)(1 - \alpha)]}{n_{\text{cal}}}\right). \quad (2)$$

For a new test input x_{n+1} and each candidate label $y \in \mathcal{Y}$, the nonconformity score is computed as $s_{n+1}(y) = S(f, (x_{n+1}, y))$. The prediction set is then defined as

$$\mathcal{C}_\alpha(x_{n+1}) = \{y : s_{n+1}(y) < \hat{q}_\alpha\}, \quad (3)$$

which satisfies the marginal coverage guarantee $\mathbb{P}(y_{n+1} \in \mathcal{C}_\alpha(x_{n+1})) \geq 1 - \alpha$. CP is not only expected to achieve the coverage guarantee but also keep its prediction sets as concise as possible so that the outputs are useful.

CP for LLMs. Applying CP to LLMs is challenging because outputs are open-ended text rather than fixed label sets. Recent work (Quach et al., 2023; Su et al., 2024a) samples multiple responses from the LLM, clusters semantically equivalent answers, and uses the cluster frequency as the basis for nonconformity scores. Given M sampled answers clustered into groups $\{C_1, \dots, C_K\}$, the nonconformity score for a cluster C_j is defined as: $S(C_j) = 1 - f(C_j)$, where $f(C_j) = |C_j|/M$ is the fraction of samples assigned to cluster C_j . A prediction set is then constructed by including all clusters whose frequency exceeds a conformally calibrated threshold, providing finite-sample coverage guarantees without requiring access to internal model logits. In CP for LLMs, a concentration problem arises because frequency-based nonconformity scores are discrete (Su et al., 2024a). To address this, (Su et al., 2024a) propose combining frequency with negative entropy (NE) and semantic similarity as tie-breaking mechanisms. NE is calculated as:

$$\text{NE}_t(x) = -\frac{1}{\log M} \sum_{j=1}^{K_t} f(C_j) \log f(C_j), \quad (4)$$

where $\text{NE}_t(x) \in [0, 1]$, with values near 0 indicating that samples concentrate on a single answer (high confidence) and values near 1 indicating a uniform spread across clusters (high uncertainty). The NE penalized frequency used in this work is formulated as:

$$f(C_i) = f(C_i) - \eta \cdot \text{NE}, \quad (5)$$

where η is a hyperparameter.

4 Method

We propose MiCP, a conformal framework for multi-turn LLMs that consists of three components: turn-level confidence estimation, adaptive early stopping, and final prediction

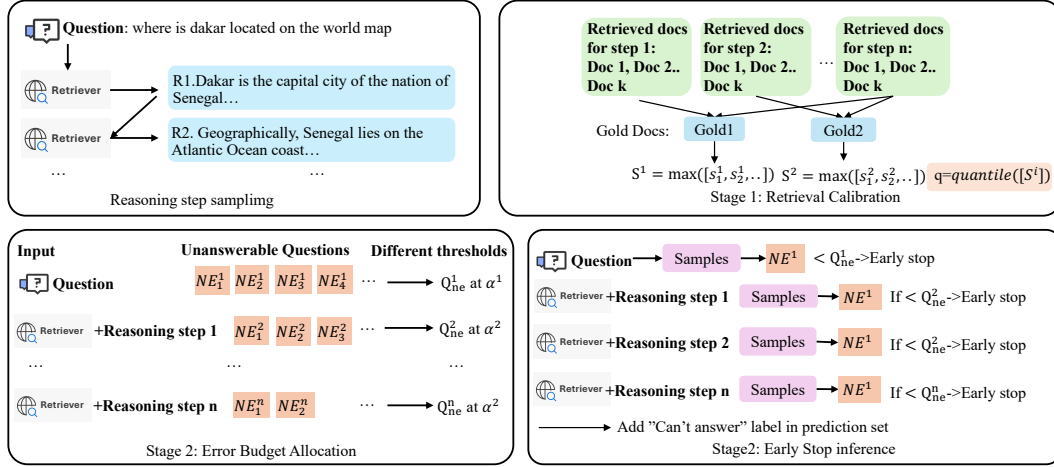


Figure 1: MiCP pipeline.

set construction. At each turn, MiCP first evaluates the reliability of the current prediction set using a calibrated confidence criterion. It then decides whether to stop or continue by comparing the turn-specific score against conformal thresholds with different error budgets across iterations. If additional information or reasoning is needed, the model performs another retrieval, action, or reasoning step, as in adaptive RAG or ReAct. Once the model stops, MiCP constructs a final prediction set over the sampled answers, with an optional "Can't Answer" label when no reliable prediction can be made within the allotted budget.

4.1 Problem Setup and Notation

We consider a single reasoning trajectory per example. For question x_i , the model iteratively generates a sequence of reasoning steps $\mathbf{r}^i = (r_0^i, r_1^i, \dots, r_{T-1}^i)$, where r_0^i is the original question and each subsequent state is conditioned on prior reasoning and any retrieved passages or external actions. For example, in adaptive RAG, r_t^i corresponds to a retrieval query, and in ReAct, r_t^i may include reasoning and tool use.

MiCP calibration proceeds in two stages: turn-level retrieval filtering, followed by prediction set construction with adaptive early stopping.

Stage 1: Retrieval Filtering. Let $\{(x_i, \mathcal{P}_i^*)\}_{i=1}^{n_{\text{cal}}}$ be the calibration set, where

$$\mathcal{P}_i^* = \{p \in \mathcal{G}_i \mid \exists t, p \in \text{Top-}K(r_t^i)\} \quad (6)$$

is the subset of gold passages $\mathcal{G}_i$ retrievable from at least one reasoning step. We calibrate a threshold $\hat{q}^{\text{ret}}$ and retain only passages with relevance score $s_t(p) \geq \hat{q}^{\text{ret}}$, where $s_t(p)$ denotes the relevance score between passage p and the query at turn t , ensuring approximately $1 - \alpha_{\text{ret}}$ of retrievable gold passages are kept while filtering low-relevance evidence.

Stage 2: Prediction Set and Early Stopping. Let $\{(x_i, y_i, \{\mathcal{A}_t^i\}_{t=0}^T)\}_{i=1}^{n_{\text{cal}}}$ be the calibration set, where y_i is the gold answer and $\mathcal{A}_t^i$ is the set of M sampled answers at turn t using filtered context from Stage 1. We jointly learn: (1) turn-specific stopping thresholds $\{\hat{q}_t\}_{t=0}^T$ that determine whether the model should stop after turn t , and (2) a final prediction set

$$\mathcal{C}(x_i) \subseteq \bigcup_{t=0}^T \mathcal{A}_t^i \cup \{\text{Can't Answer}\}. \quad (7)$$

At each turn t , MiCP computes a calibrated confidence score from the sampled answers and stops whenever the score exceeds $\hat{q}_t$. If the model stops at turn $t < T$, the prediction

set is constructed from $\bigcup_{t=0}^T \mathcal{A}_t^i$; after the final turn T , it may additionally include “Can’t Answer”. The prediction set is calibrated to satisfy

$$\mathbb{P}\left(y_i \in \mathcal{C}(x_i) \text{ if } y_i \in \bigcup_{t=0}^T \mathcal{A}_t^i, \text{ Can't Answer} \in \mathcal{C}(x_i) \text{ if } y_i \notin \bigcup_{t=0}^T \mathcal{A}_t^i\right) \geq 1 - \alpha. \quad (8)$$

That is, MiCP guarantees the prediction set either contains the correct answer when it appears among the sampled answers, or abstains via “Can’t Answer” when no sampled answer is correct.

4.2 Retrieval Threshold Calibration

A key challenge in multi-turn pipelines is that each retrieval turn returns a fixed set of top- K passages, many of which may be irrelevant. As retrieval proceeds, these irrelevant passages accumulate and pollute the LLM’s context, degrading both answer quality and stopping decisions. Prior work has applied CP to filter irrelevant passages in single-round retrieval (Chakraborty et al., 2025); we extend this to the multi-turn setting, applying retrieval filtering at every turn while preserving coverage over the entire reasoning process.

Because a gold passage p may be retrieved at multiple turns with different relevance scores, we define its optimistic relevance score as the maximum across all turns:

$$s^*(p, x_i) = \max_{t: p \in \text{Top-}K(r_t^i)} s_t(p). \quad (9)$$

We aggregate these scores across the calibration set: $\mathcal{S}_{\text{ret}} = \{s^*(p, x_i) \mid p \in \mathcal{P}_i^*, i = 1, \dots, n_{\text{cal}}\}$, and define the conformal retrieval threshold as

$$\hat{q}_{\text{ret}} = \text{Quantile}\left(\mathcal{S}_{\text{ret}} \mid \frac{(n-1)\alpha_{\text{ret}}}{n}\right). \quad (10)$$

At test time, for turn t , MiCP retains only passages whose relevance scores exceed the threshold:

$$\mathcal{C}_t^{\text{ret}}(x) = \{p \in \text{Top-}K(r_t) \mid s_t(p) > \hat{q}_{\text{ret}}\}. \quad (11)$$

This guarantees that at least a $1 - \alpha_{\text{ret}}$ fraction of retrievable gold passages are retained, while adaptively removing low-relevance passages that would otherwise hinder subsequent reasoning.

4.3 Sequential NE Calibration with Error Budget Allocation

The multi-turn retrieval loop introduces a sequential decision structure: at each turn t , the system must decide whether the question can be answered with sufficient confidence or requires further retrieval. We use the normalized entropy (NE) of the LLM’s sampled answer distribution as the uncertainty signal.

A naive approach uses one global error rate for all turns, but this ignores differences in question difficulty. Some questions can be answered after one retrieval step, while others require multiple hops. Using the same budget everywhere wastes it on easy, early cases and leaves less for harder, later ones.

To address this, we decompose the total error budget α across turns, allowing each turn to operate with its own calibrated NE threshold. We allocate a per-turn budget α_t for $t = 0, \dots, T$, subject to:

$$\sum_{t=0}^T (1 - c_{\text{ans}}^t) \cdot \alpha_t \leq (1 - c_{\text{ans}}^{\text{final}}) \cdot \alpha, \quad (12)$$

where c_{ans}^t is the fraction of questions not early-stopped that are answerable at turn t , and $c_{\text{ans}}^{\text{final}}$ is the fraction of all calibration questions for which the gold answer appears in any turn’s samples. The term $(1 - c_{\text{ans}}^t) \cdot \alpha_t$ accounts for unanswerable questions early-stopped at turn t , whose prediction sets cannot include the gold answer regardless of the threshold.

Let $\mathcal{U}_t$ be the subset of questions still active at turn t for which the gold answer has not appeared in any sampled answers up to and including t . The per-turn NE threshold is:

$$\hat{q}_t = \text{Quantile}\left(\{\text{NE}_t(x_i) \mid x_i \in \mathcal{U}_t\}; \frac{(n_{\text{cal}} - 1)\alpha_t}{n_{\text{cal}}}\right). \quad (13)$$

4.4 Composite Evaluation Metric

We introduce a composite metric to evaluate adaptive stopping in multi-turn LLMs. Some basic metrics provide misleading optimization incentives. The average number of turns encourages the model to stop as early as possible. Answer rate ignores that the proportion of answerable and unanswerable questions may vary across turns.

To balance these objectives, we propose:

$$\mathcal{L} = \gamma \cdot \bar{T} - \frac{\sum_{t=0}^T n_{\text{corr}}^t / c_{\text{ans}}^t}{\sum_{t=0}^T n_{\text{wrong}}^t / (1 - c_{\text{ans}}^t)} \quad (14)$$

where $\bar{T}$ is the average number of turns used, n_{corr}^t and n_{wrong}^t are the numbers of correctly and incorrectly answered examples that stop at turn t , and γ controls the efficiency–quality trade-off. The second term measures how well the stopping policy distinguishes answerable from unanswerable questions at each turn, rewarding correct answers at turns where answering is feasible while penalizing incorrect answers at turns where most questions remain unresolved. Although introduced in the context of MiCP, this metric is general and can evaluate stopping policies in adaptive RAG, ReAct, and other multi-turn LLM systems.

Grid Search Optimization We perform a grid search over $\alpha_0, \dots, \alpha_{T-1}$, deriving α_T from Eq. (12), and select the allocation minimizing $\mathcal{L}$ on a held-out optimization set $\mathcal{D}_{\text{opt}}$.

4.5 Final Answer Prediction Set Confidence Threshold Calibration

For each calibration sample whose gold answer appears in at least one sampled answer before stopping, let t_i^* denote the stopping turn (or T if no early stopping occurs). At each turn $t \leq t_i^*$, we cluster the sampled answers and compute a confidence score for each cluster based on its frequency, optionally adjusted by the NE. A high-confidence cluster corresponds to many consistent samples and indicates that the model may already have sufficient evidence.

The confidence of the gold-answer cluster is:

$$f(C_{\text{gold}}^i) = \max_{0 \leq t \leq t_i^*} \max_{C_j^t: y_i \in C_j^t} f(C_j^t), \quad (15)$$

where the outer maximum ranges over turns up to the stopping point, and the inner maximum ranges over all answer clusters C_j^t containing the gold answer y_i .

The conformal threshold is:

$$\hat{q}_{\text{freq}} = \text{Quantile}\left(\{f(C_{\text{gold}}^i)\}_{i \in \mathcal{D}_{\text{cal}}^{\text{ans}}}, \frac{(n_{\text{cal}} - 1)\alpha}{n_{\text{cal}}}\right), \quad (16)$$

where $\mathcal{D}_{\text{cal}}^{\text{ans}}$ denotes calibration samples whose gold answer appears in the sampled answers before stopping.

4.6 Prediction Set Construction

At test time, for a test sample x , MiCP stops at the first turn whose highest-confidence cluster exceeds the stopping threshold: $t^* = \min\{t : f(C_{\text{max}}^t) \geq \hat{q}_t\}$. If no threshold is satisfied, the model continues through all T turns.

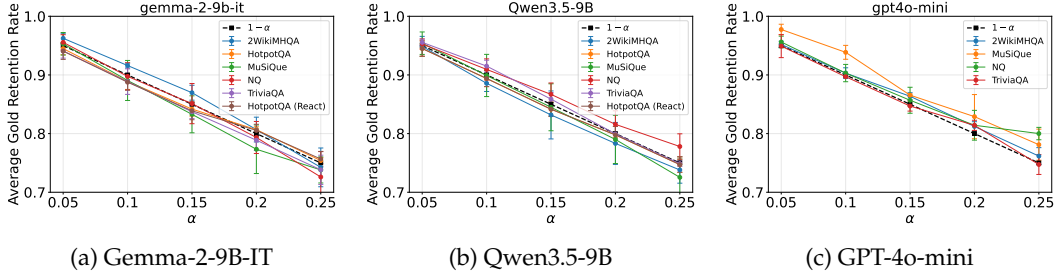


Figure 2: Empirical gold retention rate vs. the target $1 - \alpha$ across five datasets on adaptive RAG and ReAct for retrieval calibration. All models maintain coverage at or above the $1 - \alpha$ guarantee for all tested error rates.

For every turn $t \leq t^*$, sampled answers are clustered using the procedure in Section 3. The final prediction set is:

$$\mathcal{C}(x) = \begin{cases} \bigcup_{t=0}^{t^*} \{C_j^t \mid f(C_j^t) \geq \hat{q}_{\text{freq}}\}, & \text{if } t^* < T, \\ \bigcup_{t=0}^T \{C_j^t \mid f(C_j^t) \geq \hat{q}_{\text{freq}}\} \cup \{\text{Can't Answer}\}, & \text{otherwise.} \end{cases} \quad (17)$$

If the model reaches sufficient confidence before the final turn, MiCP returns only the high-confidence answer clusters accumulated up to that point. Otherwise, after exhausting all turns, it additionally includes “Can’t Answer” to preserve the coverage guarantee.

5 Experiment

We investigate the following research questions:

- **RQ1:** Can our proposed method achieve the desired coverage guarantee for both retrieval and prediction sets in multi-turn reasoning settings?
- **RQ2:** How does the efficiency of our method, in terms of both prediction set size and number of inference steps, compare with existing baselines?
- **RQ3:** To what extent does grid search improve the optimization of our proposed objective?

5.1 Experimental Setup

We evaluate MiCP on five question answering benchmarks spanning both single-hop and multi-hop reasoning: Natural Questions (NQ) (Kwiatkowski et al., 2019), TriviaQA (Joshi et al., 2017), HotpotQA (Yang et al., 2018), MuSiQue (Trivedi et al., 2022), and 2WikiMulti-HopQA (2WikiMHQA) (Ho et al., 2020). To demonstrate the generality of MiCP beyond adaptive RAG, we additionally evaluate on the multi-turn reasoning framework ReAct (Yao et al., 2022), where the model alternates between reasoning and information-seeking actions. For each dataset, we randomly sample 900 examples and split them into an optimization set ($|\mathcal{D}_{\text{opt}}| = 300$), a calibration set ($|\mathcal{D}_{\text{cal}}| = 300$), and a test set ($|\mathcal{D}_{\text{test}}| = 300$). We use three popular LLMs of comparable scale: Gemma-2-9B-IT (Team et al., 2024), Qwen3.5-9B (Qwen Team, 2026), and GPT-4o-mini (Achiam et al., 2023), with answer clustering performed by the all-MiniLM-L6-v2 sentence transformer using agglomerative clustering at a cosine similarity threshold of 0.9. Other settings can be found in Appendix A.1.

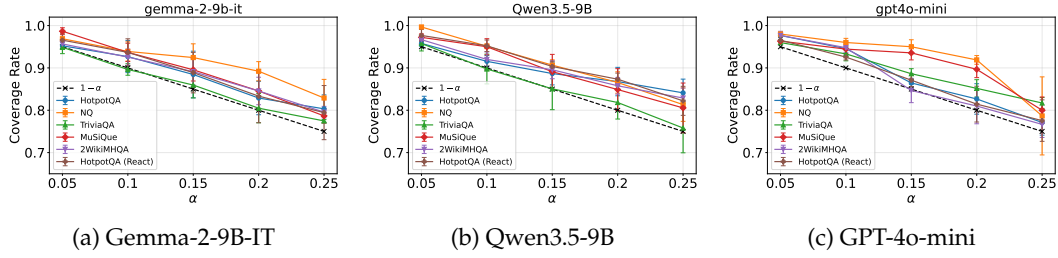


Figure 3: Empirical coverage rate vs. the target $1 - \alpha$ across five datasets on adaptive RAG and ReAct for the final prediction set. All models maintain coverage at or above the $1 - \alpha$ guarantee for all tested error rates.

Table 1: Average number of retrieval turns (Avg #Turns), prediction set size, and composite objective across five datasets, averaged over all α values. **Bold** indicates the best result per row. Lower is better for all metrics. w/o ES means no early stop.

Dataset	Qwen3.5-9B				Gemma-2-9B-IT				GPT-4o-mini			
	Ours	w/o ES	Avg Iter	Answer Rate	Ours	w/o ES	Avg Iter	Answer Rate	Ours	w/o ES	Avg Iter	Answer Rate
Avg #Turns												
HotpotQA	2.562	3.000	2.331	2.433	2.782	3.000	2.481	2.854	2.497	3.000	2.173	2.607
HotpotQA (ReAct)	2.522	3.000	2.313	2.557	2.752	3.000	2.472	2.757	2.587	3.000	2.167	2.877
NQ	2.621	3.000	2.459	2.770	2.840	3.000	2.597	2.856	3.000	3.000	2.849	2.924
TriviaQA	2.528	3.000	2.356	2.587	2.898	3.000	2.687	2.806	3.000	3.000	3.000	3.000
MuSiQue	2.705	3.000	2.519	2.861	2.795	3.000	2.488	2.681	2.870	3.000	2.278	2.936
2WikiMHQA	1.778	3.000	1.472	1.477	2.553	3.000	2.275	2.790	2.293	3.000	2.293	2.767
Avg Set Size												
HotpotQA	10.75	11.45	10.56	10.57	7.99	8.10	7.68	7.88	3.503	2.738	2.469	4.382
HotpotQA (ReAct)	11.94	12.85	11.50	11.79	6.82	6.85	6.59	6.69	4.426	3.781	2.576	4.776
NQ	19.63	21.16	19.02	19.82	9.38	9.66	8.77	9.27	5.753	5.753	5.683	5.740
TriviaQA	5.64	5.77	5.47	5.50	3.26	3.22	3.13	3.14	5.017	5.019	5.019	5.019
MuSiQue	43.25	42.02	40.09	42.44	24.87	25.46	22.50	23.34	19.568	25.121	10.207	25.066
2WikiMHQA	7.32	8.62	6.54	6.57	9.09	8.89	8.73	8.70	4.033	4.164	4.033	4.078
Ours Obj.												
HotpotQA	-5.95	-1.96	-5.17	-4.84	-1.95	-1.34	-2.47	-1.78	-2.097	-1.843	-2.083	-2.045
HotpotQA (ReAct)	-6.95	-2.68	-7.28	-4.42	-2.58	-1.74	-4.07	-2.16	-2.932	-1.518	-1.993	-1.902
NQ	-8.67	-3.23	-5.20	-4.30	-3.20	-1.93	-3.27	-2.50	-0.368	-0.368	-0.364	-0.361
TriviaQA	-3.61	-2.13	-3.31	-2.84	-2.06	-0.97	-2.26	-1.37	0.030	0.030	0.030	0.030
MuSiQue	-5.17	-2.08	-4.37	-2.99	-2.43	-2.41	-2.29	-2.37	-2.785	-1.953	-1.787	-2.701
2WikiMHQA	-11.79	-3.01	-8.16	-8.19	-2.51	-1.58	-3.21	-1.73	-1.920	-1.334	-1.920	-1.449

5.2 RQ1

Figure 2 and Figure 3 present the empirical coverage results for retrieval filtering and prediction set construction, respectively. For retrieval, the average gold retention rate closely tracks the $1 - \alpha$ target across all five datasets and all models, confirming that the conformal retrieval threshold $\hat{q}_{\text{ret}}$ successfully preserves retrievable gold passages while filtering irrelevant ones. For the prediction set, the coverage rate similarly stays at or above $1 - \alpha$ for nearly all α values and datasets. It shows a little overcoverage problem, which might be caused by the concentration problem coming from the frequency (Su et al., 2024a). These results confirm that MiCP maintains the desired $1 - \alpha$ coverage guarantee across both the retrieval and prediction stages, even with early stopping enabled, which comes from our multi-hop retrieval calibration and error budget allocation.

5.3 RQ2

Table 1 compares three variations of MiCP—Ours (optimizing the composite objective in Eq. (14)), Avg Iter (optimizing for average turns), and Answer Rate (optimizing for the ratio of questions without the “Can’t Answer” label)—against the No Early Stop baseline where all questions run for T turns. We highlight two findings.

First, MiCP successfully improves the inference efficiency (avg #turns), with better and more recent models (e.g., qwen3.5-9B, gpt-4o-mini), our method is able to get better inference efficiency. Especially on 2wikiMHQA, it is easier to distinguish answerable and unanswerable questions as it is a fully template-generated multi-hop QA dataset. MiCP also consistently improves the prediction efficiency (prediction set size). It achieves this by filtering answers from unnecessary late turns, which may not include the correct answer and may even bring noisy, wrong answers.

Second, we find that using average turns as the optimization objective consistently yields the fewest turns and smallest prediction set sizes. This is because optimizing for fewer turns encourages the system to allocate more error budget to early turns, causing the LLM to stop earlier and thereby avoiding the accumulation of noisy, low-confidence answers from unnecessary later turns.

5.4 RQ3

As shown in Table 1, grid search over Qwen3.5-9B and gpt-4-mini successfully optimizes our proposed objective, yielding more rational early-stopping decisions. Moreover, all variants of MiCP consistently improve upon the baseline objective, demonstrating the necessity of early stopping in multi-turn question answering. On Gemma-2-9B-IT, the improvements are smaller but still consistent, confirming that the grid search generalizes across models with different uncertainty characteristics. Noticing that for the Gemma model, using avg-turn even achieves the best performance on the objective without optimizing it. An explanation can be that it is hard for Gemma-2-9B to distinguish the difference between answerable and unanswerable questions, which makes the optimization not stable.

6 Conclusion

We presented MiCP, the first conformal framework for multi-turn LLMs, which jointly calibrates retrieval filtering, adaptive early stopping, and final prediction set construction while maintaining an overall $1 - \alpha$ coverage guarantee. MiCP applies broadly to multi-turn pipelines such as adaptive RAG and ReAct, where the model may retrieve, reason, or act over multiple turns before answering. By allocating the error budget across turns and optimizing turn-specific thresholds through grid search, MiCP preserves the same coverage guarantee as a strategy without early stopping while substantially improving efficiency. Experiments on five single-hop and multi-hop QA benchmarks with three popular LLMs show that MiCP achieves the target coverage at both the retrieval and answer stages, while significantly reducing the number of turns, retrieval cost, and maintaining competitive prediction set sizes.

References

- Josh Achiam, Steven Adler, Sandhini Agarwal, Lama Ahmad, Ilge Akkaya, Florencia Leoni Aleman, Diogo Almeida, Janko Altschmidt, Sam Altman, Shyamal Anadkat, et al. Gpt-4 technical report. *arXiv preprint arXiv:2303.08774*, 2023.
- Anastasios N Angelopoulos and Stephen Bates. A gentle introduction to conformal prediction and distribution-free uncertainty quantification. *arXiv preprint arXiv:2107.07511*, 2021.
- Margarida Campos, António Farinhas, Chrysoula Zerva, Mário AT Figueiredo, and André FT Martins. Conformal prediction for natural language processing: A survey. *Transactions of the Association for Computational Linguistics*, 12:1497–1516, 2024.
- Debashish Chakraborty, Eugene Yang, Daniel Khashabi, Dawn Lawrie, and Kevin Duh. Principled context engineering for rag: Statistical guarantees via conformal prediction. *arXiv preprint arXiv:2511.17908*, 2025.
- Hanxing Ding, Liang Pang, Zihao Wei, Huawei Shen, and Xueqi Cheng. Retrieve only when it needs: Adaptive retrieval augmentation for hallucination mitigation in large language models. *CoRR*, abs/2402.10612, 2024. doi: 10.48550/ARXIV.2402.10612. URL <https://doi.org/10.48550/arXiv.2402.10612>.
- BV Elasticsearch. Elasticsearch. *software*, version, 6(1), 2018.
- Xanh Ho, Anh-Khoa Duong Nguyen, Saku Sugawara, and Akiko Aizawa. Constructing a multi-hop QA dataset for comprehensive evaluation of reasoning steps. In *Proceedings*

- of the 28th International Conference on Computational Linguistics, pp. 6609–6625, Barcelona, Spain (Online), 2020. International Committee on Computational Linguistics.
- Soyeong Jeong, Jinheon Baek, Sukmin Cho, Sung Ju Hwang, and Jong Park. Adaptive-rag: Learning to adapt retrieval-augmented large language models through question complexity. In Kevin Duh, Helena Gómez-Adorno, and Steven Bethard (eds.), *Proceedings of the 2024 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies (Volume 1: Long Papers)*, NAACL 2024, Mexico City, Mexico, June 16-21, 2024, pp. 7036–7050. Association for Computational Linguistics, 2024. doi: 10.18653/V1/2024.NAACL-LONG.389. URL <https://doi.org/10.18653/v1/2024.naacl-long.389>.
- Zhengbao Jiang, Frank Xu, Luyu Gao, Zhiqing Sun, Qian Liu, Jane Dwivedi-Yu, Yiming Yang, Jamie Callan, and Graham Neubig. Active retrieval augmented generation. In Houda Bouamor, Juan Pino, and Kalika Bali (eds.), *Proceedings of the 2023 Conference on Empirical Methods in Natural Language Processing*, pp. 7969–7992, Singapore, December 2023. Association for Computational Linguistics. doi: 10.18653/v1/2023.emnlp-main.495. URL <https://aclanthology.org/2023.emnlp-main.495>.
- Mandar Joshi, Eunsol Choi, Daniel S. Weld, and Luke Zettlemoyer. TriviaQA: A large scale distantly supervised challenge dataset for reading comprehension. In *Proceedings of the 55th Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, pp. 1601–1611, Vancouver, Canada, 2017.
- Bhawesh Kumar, Charlie Lu, Gauri Gupta, Anil Palepu, David Bellamy, Ramesh Raskar, and Andrew Beam. Conformal prediction with large language models for multi-choice question answering. *arXiv preprint arXiv:2305.18404*, 2023.
- Tom Kwiatkowski, Jennimaria Palomaki, Olivia Redfield, Michael Collins, Ankur Parikh, Chris Alberti, Danielle Epstein, Illia Polosukhin, Jacob Devlin, Kenton Lee, Kristina Toutanova, Llion Jones, Matthew Kelcey, Ming-Wei Chang, Andrew M. Dai, Jakob Uszkoreit, Quoc Le, and Slav Petrov. Natural questions: A benchmark for question answering research. *Transactions of the Association for Computational Linguistics*, 7:453–466, 2019.
- Shuo Li, Sangdon Park, Insup Lee, and Osbert Bastani. TRAQ: Trustworthy retrieval augmented question answering via conformal prediction. In Kevin Duh, Helena Gomez, and Steven Bethard (eds.), *Proceedings of the 2024 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies (Volume 1: Long Papers)*, pp. 3799–3821, Mexico City, Mexico, June 2024. Association for Computational Linguistics. doi: 10.18653/v1/2024.naacl-long.210. URL <https://aclanthology.org/2024.naacl-long.210/>.
- Dehai Min, Kailin Zhang, Tongtong Wu, and Lu Cheng. Quco-rag: Quantifying uncertainty from the pre-training corpus for dynamic retrieval-augmented generation. *arXiv preprint arXiv:2512.19134*, 2025.
- Christopher Mohri and Tatsunori Hashimoto. Language models with conformal factuality guarantees. *arXiv preprint arXiv:2402.10978*, 2024.
- Viktor Moskvoretskii, Maria Marina, Mikhail Salnikov, Nikolay Ivanov, Sergey Pletenev, Daria Galimzianova, Nikita Krayko, Vasily Konovalov, Irina Nikishina, and Alexander Panchenko. Adaptive retrieval without self-knowledge? bringing uncertainty back home. In *Proceedings of the 63rd Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, pp. 6355–6384, 2025.
- Victor Quach, Adam Fisch, Tal Schuster, Adam Yala, Jae Ho Sohn, Tommi S Jaakkola, and Regina Barzilay. Conformal language modeling. *arXiv preprint arXiv:2306.10193*, 2023.
- Qwen Team. Qwen3.5: Towards native multimodal agents, February 2026. URL <https://qwen.ai/blog?id=qwen3.5>.

- Nitin Liladhar Rane, Abhijeet Tawde, Saurabh P Choudhary, and Jayesh Rane. Contribution and performance of chatgpt and other large language models (llm) for scientific and research advancements: a double-edged sword. *International Research Journal of Modernization in Engineering Technology and Science*, 5(10):875–899, 2023.
- Maxon Rubin-Toles, Maya Gambhir, Keshav Ramji, Aaron Roth, and Surbhi Goel. Conformal language model reasoning with coherent factuality. *arXiv preprint arXiv:2505.17126*, 2025.
- Tal Schuster, Adam Fisch, Jai Gupta, Mostafa Dehghani, Dara Bahri, Vinh Tran, Yi Tay, and Donald Metzler. Confident adaptive language modeling. *Advances in Neural Information Processing Systems*, 35:17456–17472, 2022.
- Glenn Shafer and Vladimir Vovk. A tutorial on conformal prediction. *Journal of Machine Learning Research*, 9(3), 2008.
- Huanxin Sheng, Xinyi Liu, Hangfeng He, Jieyu Zhao, and Jian Kang. Analyzing uncertainty of llm-as-a-judge: Interval evaluations with conformal prediction, 2025. URL <https://arxiv.org/abs/2509.18658>.
- Noah Shinn, Federico Cassano, Ashwin Gopinath, Karthik Narasimhan, and Shunyu Yao. Reflexion: Language agents with verbal reinforcement learning. *Advances in neural information processing systems*, 36:8634–8652, 2023.
- Jiayuan Su, Jing Luo, Hongwei Wang, and Lu Cheng. API is enough: Conformal prediction for large language models without logit-access. In Yaser Al-Onaizan, Mohit Bansal, and Yun-Nung Chen (eds.), *Findings of the Association for Computational Linguistics: EMNLP 2024*, pp. 979–995, Miami, Florida, USA, November 2024a. Association for Computational Linguistics. doi: 10.18653/v1/2024.findings-emnlp.54. URL <https://aclanthology.org/2024.findings-emnlp.54/>.
- Weihang Su, Yichen Tang, Qingyao Ai, Zhijing Wu, and Yiqun Liu. Dragin: Dynamic retrieval augmented generation based on the information needs of large language models, 2024b. URL <https://arxiv.org/abs/2403.10081>.
- Gemma Team, Thomas Mesnard, Cassidy Hardin, Robert Dadashi, Surya Bhupatiraju, Shreya Pathak, Laurent Sifre, Morgane Rivière, Mihir Sanjay Kale, Juliette Love, et al. Gemma: Open models based on gemini research and technology. *arXiv preprint arXiv:2403.08295*, 2024.
- Harsh Trivedi, Niranjan Balasubramanian, Tushar Khot, and Ashish Sabharwal. MuSiQue: Multihop questions via single-hop question composition. *Transactions of the Association for Computational Linguistics*, 10:539–554, 2022.
- Harsh Trivedi, Niranjan Balasubramanian, Tushar Khot, and Ashish Sabharwal. Interleaving retrieval with chain-of-thought reasoning for knowledge-intensive multi-step questions. In Anna Rogers, Jordan L. Boyd-Graber, and Naoaki Okazaki (eds.), *Proceedings of the 61st Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, ACL 2023, Toronto, Canada, July 9-14, 2023, pp. 10014–10037. Association for Computational Linguistics, 2023. doi: 10.18653/v1/2023.ACL-LONG.557. URL <https://doi.org/10.18653/v1/2023.acl-long.557>.
- Jiaqi Wang, Enze Shi, Huawen Hu, Chong Ma, Yiheng Liu, Xuhui Wang, Yincheng Yao, Xuan Liu, Bao Ge, and Shu Zhang. Large language models for robotics: Opportunities, challenges, and perspectives. *Journal of Automation and Intelligence*, 4(1):52–64, 2025.
- Zhilin Yang, Peng Qi, Saizheng Zhang, Yoshua Bengio, William W. Cohen, Ruslan Salakhutdinov, and Christopher D. Manning. HotpotQA: A dataset for diverse, explainable multi-hop question answering. In *Conference on Empirical Methods in Natural Language Processing (EMNLP)*, 2018.
- Shunyu Yao, Jeffrey Zhao, Dian Yu, Nan Du, Izhak Shafran, Karthik R Narasimhan, and Yuan Cao. React: Synergizing reasoning and acting in language models. In *The eleventh international conference on learning representations*, 2022.

Zijun Yao, Weijian Qi, Liangming Pan, Shulin Cao, Linmei Hu, Weichuan Liu, Lei Hou, and Juanzi Li. Seakr: Self-aware knowledge retrieval for adaptive retrieval augmented generation. *CoRR*, abs/2406.19215, 2024. doi: 10.48550/ARXIV.2406.19215. URL <https://doi.org/10.48550/arXiv.2406.19215>.

Fanghua Ye, Mingming Yang, Jianhui Pang, Longyue Wang, Derek F. Wong, Emine Yilmaz, Shuming Shi, and Zhaopeng Tu. Benchmarking LLMs via uncertainty quantification. In *The Thirty-eight Conference on NIPS Datasets and Benchmarks Track*, 2024. URL <https://openreview.net/forum?id=L0oSfTroNE>.

Botong Zhang, Shuo Li, and Osbert Bastani. Conformal structured prediction. *arXiv preprint arXiv:2410.06296*, 2024.

Xiaofan Zhou, Baiting Chen, Yu Gui, and Lu Cheng. Conformal prediction: A data perspective. *ACM Computing Surveys*, 2025.

A Appendix

A.1 Experiment Setting

For adaptive RAG and ReAct, retrieval is performed using an Elasticsearch-based retriever (Elasticsearch, 2018) over Wikipedia, returning the top- K passages at each turn with $K = 10$. We set the maximum number of turns to $T = 3$, draw $M = 15$ stochastic answer samples per turn using temperature 1.4, and use $G = 20$ grid steps per dimension when optimizing the allocation of turn-specific error budgets. We evaluate target error rates $\alpha \in \{0.05, 0.10, 0.15, 0.20, 0.25\}$ and report results averaged over three random seeds for the dataset split. In Eq. 5, we set $\eta = 0.1$, and throughout all experiments we fix the retrieval error budget to $\alpha_{\text{ret}} = 0.1$.

A.2 Proof of Retrieval Calibration Coverage Guarantee

Theorem 1 (Retrieval Coverage Guarantee) Let $\{(x_i, \mathcal{P}_i^*)\}_{i=1}^{n_{\text{cal}}}$ be the calibration set, where $\mathcal{P}_i^* = \{g \in \mathcal{G}_i \mid \mathcal{T}(g, x_i) \neq \emptyset\}$ is the set of retrievable gold passages for question x_i . Let $s^*(g, x_i) = \max_{t \in \mathcal{T}(g, x_i)} s_t(g)$ be the optimistic relevance score. Collect all such scores into $\mathcal{S}_{\text{ret}} = \{s^*(g, x_i) \mid g \in \mathcal{P}_i^*, i = 1, \dots, n_{\text{cal}}\}$ with $|\mathcal{S}_{\text{ret}}| = n$. Define the conformal threshold as:

$$\hat{q}_{\text{ret}} = \text{Quantile}\left(\mathcal{S}_{\text{ret}}; \frac{(n-1)\alpha}{n}\right). \quad (18)$$

Then for a new exchangeable test example, any retrievable gold passage $g \in \mathcal{P}_{n+1}^*$ satisfies:

$$\mathbb{P}(s^*(g, x_{n+1}) \geq \hat{q}_{\text{ret}}) \geq 1 - \alpha. \quad (19)$$

Proof. Under the exchangeability assumption, the n calibration scores and the test score $s_{n+1} = s^*(g, x_{n+1})$ form an exchangeable sequence of size $n+1$. By the split conformal prediction guarantee, setting the threshold at the $\frac{(n-1)\alpha}{n}$ -quantile of the n calibration scores yields:

$$\mathbb{P}(s_{n+1} < \hat{q}_{\text{ret}}) \leq \frac{(n-1)\alpha}{n}. \quad (20)$$

Since $\frac{(n-1)\alpha}{n} < \alpha$ for all $n \geq 1$, we have:

$$\mathbb{P}(s^*(g, x_{n+1}) \geq \hat{q}_{\text{ret}}) \geq 1 - \frac{(n-1)\alpha}{n} > 1 - \alpha. \quad (21)$$

Since this holds marginally for each retrievable gold passage, the expected gold retention rate satisfies:

$$\mathbb{E}\left[\frac{|\{g \in \mathcal{P}_{n+1}^* \mid \exists t : s_t(g) \geq \hat{q}_{\text{ret}}\}|}{|\mathcal{P}_{n+1}^*|}\right] \geq 1 - \frac{(n-1)\alpha}{n} \geq 1 - \alpha. \quad (22)$$

□

Remark 2 We calibrate on $\mathcal{P}_i^*$ (retrievable gold passages) rather than $\mathcal{G}_i$ (all annotated gold passages), avoiding inflation of the threshold with unretrievable passages.

A.3 Proof of Prediction Set Coverage Guarantee

Theorem 3 (Prediction Set Coverage Guarantee) Let $\{(x_i, y_i, \{A_t^i\}_{t=0}^T)\}_{i=1}^{n_{\text{cal}}}$ be the calibration set. Define the answerable subset $\mathcal{D}_{\text{cal}}^{\text{ans}} = \{i \mid y_i \in \bigcup_{t=0}^T A_t^i\}$ with $|\mathcal{D}_{\text{cal}}^{\text{ans}}| = m$. For each $i \in \mathcal{D}_{\text{cal}}^{\text{ans}}$, let $f(C_{\text{gold}}^i)$ be the maximum frequency score of the gold-containing cluster (Eq. 15). Define:

$$\hat{q}_{\text{freq}} = \text{Quantile}\left(\{f(C_{\text{gold}}^i)\}_{i \in \mathcal{D}_{\text{cal}}^{\text{ans}}}; \frac{(m-1)\alpha}{m}\right). \quad (23)$$

Then for a new exchangeable test example (x_{n+1}, y_{n+1}) , the prediction set $C(x_{n+1})$ satisfies:

$$\mathbb{P}\left(y_{n+1} \in C(x_{n+1}) \text{ if } y_{n+1} \in \bigcup_t A_t^{n+1}, \text{ "Can't Answer"} \in C(x_{n+1}) \text{ if } y_{n+1} \notin \bigcup_t A_t^{n+1}\right) \geq 1 - \alpha. \quad (24)$$

Proof. We analyze the two cases separately.

Case 1: The gold answer is never sampled ($y_{n+1} \notin \bigcup_{t=0}^T A_t^{n+1}$).

The NE thresholds $\{\hat{q}_t\}$ are calibrated on the unanswerable subset U_t at each turn (Eq. 13), with per-turn error budgets $\{\alpha_t\}$ satisfying:

$$\sum_{t=0}^T (1 - c_{\text{ans}}^t) \cdot \alpha_t \leq (1 - c_{\text{ans}}^{\text{final}}) \cdot \alpha. \quad (25)$$

This bounds the total probability of incorrectly early-stopping an unanswerable question before it reaches the final turn. When a question is not early-stopped and exhausts all T turns, the "Can't Answer" label is always appended by construction (Eq. 15). Therefore:

$$\mathbb{P}\left(\text{"Can't Answer"} \in C(x_{n+1}) \mid y_{n+1} \notin \bigcup_t A_t^{n+1}\right) \geq 1 - \alpha. \quad (26)$$

Case 2: The gold answer appears in some turn's samples ($y_{n+1} \in \bigcup_{t=0}^T A_t^{n+1}$).

By exchangeability, the m calibration scores $\{f(C_{\text{gold}}^i)\}_{i \in \mathcal{D}_{\text{cal}}^{\text{ans}}}$ and the test score $f(C_{\text{gold}}^{n+1})$ form an exchangeable sequence of size $m+1$. By the split conformal guarantee at quantile level $\frac{(m-1)\alpha}{m}$:

$$\mathbb{P}\left(f(C_{\text{gold}}^{n+1}) < \hat{q}_{\text{freq}}\right) \leq \frac{(m-1)\alpha}{m} < \alpha. \quad (27)$$

Therefore:

$$\mathbb{P}\left(y_{n+1} \in C(x_{n+1}) \mid y_{n+1} \in \bigcup_t A_t^{n+1}\right) \geq 1 - \frac{(m-1)\alpha}{m} > 1 - \alpha. \quad (28)$$

Combining both cases. Let E denote the event $y_{n+1} \in \bigcup_t A_t^{n+1}$. By the law of total probability:

$$\begin{aligned} \mathbb{P}(\text{coverage holds}) &= \mathbb{P}(E) \cdot \mathbb{P}(y_{n+1} \in C(x_{n+1}) \mid E) + \mathbb{P}(\bar{E}) \cdot \mathbb{P}(\text{"Can't Answer"} \in C(x_{n+1}) \mid \bar{E}) \\ &\geq \mathbb{P}(E) \cdot (1 - \alpha) + \mathbb{P}(\bar{E}) \cdot (1 - \alpha) \\ &= 1 - \alpha. \end{aligned}$$

□

Remark 4 (Role of the error budget decomposition) The NE thresholds $\{\hat{q}_t\}_{t=0}^T$ control which turn each question stops at. The error budget constraint (Eq. 12) ensures that unanswerable questions are not prematurely early-stopped, preserving the "Can't Answer" mechanism, while allowing efficient per-hop early stopping without violating the overall $1 - \alpha$ coverage guarantee.

Remark 5 (Conditional vs. marginal coverage) *The guarantees above are marginal, holding on average over the randomness of calibration and test data. Conditional coverage for specific subgroups would require group-conditional calibration, which we leave for future work.*