

Double-Calibration: Towards Reliable LLMs via Calibrating Knowledge and Reasoning Confidence

Yuyin Lu¹, Ziran Liang², Yanghui Rao^{1*}, Wenqi Fan², Fu Lee Wang³ and Qing Li²

¹School of Computer Science and Engineering, Sun Yat-sen University, Guangzhou, China

²Department of Computing, The Hong Kong Polytechnic University, Hong Kong SAR

³School of Science and Technology, Hong Kong Metropolitan University, Hong Kong SAR
luyy37@mail2.sysu.edu.cn, raoyangh@mail.sysu.edu.cn

Abstract

Reliable reasoning in Large Language Models (LLMs) is challenged by their propensity for hallucination. While augmenting LLMs with Knowledge Graphs (KGs) improves factual accuracy, existing KG-augmented methods fail to quantify epistemic uncertainty in both the retrieved evidence and LLMs’ reasoning. To bridge this gap, we introduce DoublyCal, a framework built on a novel double-calibration principle. DoublyCal employs a lightweight proxy model to first generate KG evidence alongside a calibrated evidence confidence. This calibrated supporting evidence then guides a black-box LLM, yielding final predictions that are not only more accurate but also well-calibrated, with confidence scores traceable to the uncertainty of the supporting evidence. Experiments on knowledge-intensive benchmarks show that DoublyCal significantly improves both the accuracy and confidence calibration of black-box LLMs while maintaining low token cost.

1 Introduction

The reliability of Large Language Models (LLMs) is critically undermined by their tendency to hallucinate [Huang *et al.*, 2024], a problem rooted in both intrinsic *epistemic uncertainty* (knowledge gaps) and extrinsic *aleatoric uncertainty* (data ambiguity) [Hüllermeier and Waegeman, 2021]. To mitigate this, Knowledge Graph-augmented Retrieval-Augmented Generation (KG-RAG) has emerged as a leading paradigm [Zhang *et al.*, 2025a]. By augmenting LLM with structured evidence retrieved from external Knowledge Graphs (KGs), KG-RAG is helpful to reduce the model’s internal knowledge gaps and improve the factual accuracy of its responses [Xiang *et al.*, 2025].

However, this KG-augmentation mechanism introduces a new yet critical dependency: *the certainty of the retrieved evidence itself*. Prevailing KG-RAG methods [Luo *et al.*, 2024;

Li *et al.*, 2025a; Liu *et al.*, 2026] often rely on an idealistic assumption that the retrieved evidence is always both sufficient and certain to support correct reasoning for a given query. This assumption is routinely violated in practice due to ambiguous queries, the intrinsic incompleteness of KGs, and imperfections in the retrieval process. Consequently, when provided with partial evidence, LLMs may still produce confidently stated but incorrect predictions [Kalai *et al.*, 2025]. For example, as illustrated in Figure 1, given the partial evidence “Belle is a sibling of Snoopy”, an LLM might incorrectly infer “Belle is Snoopy’s brother”. Thus, the current KG-RAG lacks the ability to assess and control uncertainty at the very source of the reasoning chain.

Concurrently, research on Uncertainty Quantification (UQ) for LLMs aims to calibrate prediction confidence but focuses predominantly on the final output [Xia *et al.*, 2025]. For instance, verbalized UQ methods elicit confidence estimates from black-box LLMs, yet these estimates remain opaque and non-traceable [Tian *et al.*, 2023; Xiong *et al.*, 2024]. It is impossible to discern whether the expressed uncertainty stems from flawed evidence, deficiencies in the model’s own reasoning, or the intrinsic difficulty of the task. Therefore, existing UQ methods cannot synergize effectively with KG-RAG to provide a stepwise-calibrated view of the complete evidence-to-prediction chain.

In summary, a principled solution for systematically managing the propagation of uncertainty in KG-augmented LLMs is lacking. To bridge this gap, we propose a novel **double-calibration** paradigm. Its core lies in moving beyond basic evidence retrieval to the construction of a *calibrated reasoning chain*, where confidence is explicitly estimated and made traceable from the retrieved KG evidence to the final LLM prediction. As shown in Figure 1, this enables the LLM to weigh alternative answers (e.g., correctly favoring “Spike” over “Belle”) based on the calibrated confidence of the supporting evidence, rather than making an overconfident guess.

We instantiate this principle in **DoublyCal**, a framework that implements double-calibrated KG-RAG. DoublyCal grounds the LLM’s reasoning on verifiable KG evidence and performs dual calibration: it first calibrates the confidence of the retrieved evidence, then uses this calibrated evidence to guide and further calibrate the final LLM prediction. Specifically, we formalize KG evidence as constrained relational paths extracted from a KG. We then train a lightweight

*Corresponding Author.

An extended version of this paper can be found at <https://arxiv.org/abs/2601.11956>.

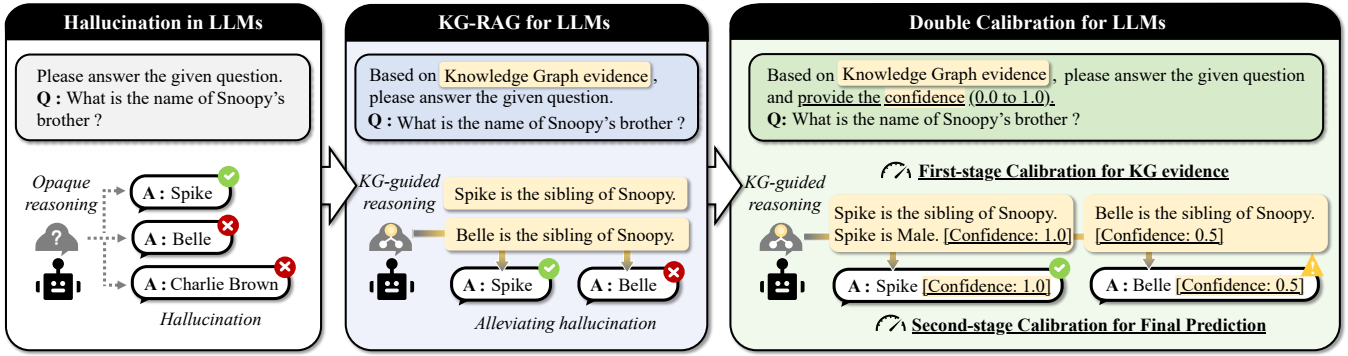


Figure 1: A motivating example of double-calibration against hallucination in KG-augmented LLMs.

proxy model under Bayesian supervision to generate relevant KG evidence alongside a calibrated confidence score for each query. During inference, the primary LLM is prompted with both the KG evidence and its confidence estimate, leading to more accurate and better-calibrated predictions. Crucially, because the evidence confidence explicitly estimates the expected reasoning uncertainty of the LLM when utilizing the provided evidence, the final confidence becomes traceable to the verifiable KG evidence and its calibrated confidence, rather than remaining an opaque global estimate. Our main contributions are summarized as follows:

- We establish the principle of double-calibration for reliable KG-augmented LLMs, which mandates explicit confidence calibration for both the KG evidence and the final LLM predictions.
- We propose DoublyCal¹, a framework that implements this principle via a Bayesian-calibrated proxy model, providing the primary LLM reasoner with KG evidence accompanied by evidence confidence.
- We empirically demonstrate that DoublyCal significantly and consistently improves the accuracy and calibration of diverse black-box LLMs on knowledge-intensive benchmarks in a cost-effective manner.

2 Related Work

2.1 Knowledge-Augmented Generation for Reliable LLMs

Retrieval-Augmented Generation (RAG) reduces the inherent knowledge gaps of LLMs by providing external information, thereby improving the factual accuracy of their responses [Zhang *et al.*, 2025a]. The choice of knowledge source defines a spectrum of RAG variants, ranging from (i) unstructured text in Vanilla RAG [Guo *et al.*, 2025; Sun *et al.*, 2025], to (ii) textual graphs that model latent connections in GraphRAG [He *et al.*, 2024; Li *et al.*, 2025b], and finally to (iii) formal Knowledge Graphs (KGs) with explicit relations in KG-RAG [Luo *et al.*, 2024; Li *et al.*, 2025a]. By providing precise and structured knowledge, KG-RAG offers a rigorous foundation for complex reasoning and has

demonstrated superior performance on knowledge-intensive tasks [Xiang *et al.*, 2025; Pan *et al.*, 2024].

However, the retrieved knowledge itself may be noisy or insufficient, posing a persistent challenge to the reliability of LLM reasoning. To address this, some research dynamically selects external knowledge when it conflicts with the LLM’s internal parametric knowledge [Liu *et al.*, 2026; Zhang *et al.*, 2025b]. Unlike prior work that focus on resolving knowledge conflicts, we introduce a double-calibration principle which explicitly quantifies confidence for both the retrieved evidence and the final LLM prediction, thereby systematically identifying epistemic boundaries.

2.2 Uncertainty Quantification for LLMs

Uncertainty Quantification (UQ) for LLMs aims to calibrate the confidence of model predictions to identify their epistemic boundaries [Huang *et al.*, 2024]. While some studies incorporate uncertainty awareness during training [Stangel *et al.*, 2025], most practical UQ methods often operate post-hoc and are categorized by model access [Xia *et al.*, 2025].

For open-source LLMs, confidence is typically derived from internal states, such as the feature-space distribution of hidden embeddings [Chen *et al.*, 2024; Vazhentsev *et al.*, 2025] and the predictive entropy of the output distribution [Malinin and Gales, 2021]. For black-box LLMs, methods rely on API-based probing. A prevalent strategy generates multiple responses and evaluates their semantic consistency [Manakul *et al.*, 2023; Kuhn *et al.*, 2023; Lin *et al.*, 2024]. A more efficient alternative is verbalized UQ, which directly prompts the LLM to verbalize its own confidence, eliciting introspective uncertainty estimates [Tian *et al.*, 2023; Xiong *et al.*, 2024; Tanneru *et al.*, 2024]. Its plug-and-play nature and low cost make verbalized UQ readily integrable with KG-RAG, forming a strong single-calibration baseline that calibrates only the final output.

A fundamental limitation across prior UQ paradigms is their exclusive focus on the final output, deriving confidence solely from the LLM’s internal states or self-assessment. This makes them vulnerable to the model’s overconfidence biases due to the lack of an objective external anchor [Xiong *et al.*, 2024]. Our double-calibration principle addresses this by first calibrating externally verifiable KG evidence, thereby establishing a traceable foundation for the entire reasoning chain.

¹The source code of DoublyCal is available at <https://github.com/luyy9apples/DoublyCal>.

3 Preliminaries

3.1 Uncertainty and Confidence

We extend the conceptualization of uncertainty and confidence for LLM outputs [Lin *et al.*, 2024] to general predictive systems. Given an input x , a predictive system f produces a probability distribution over possible outputs $P_f(o | x)$. The uncertainty of f regarding x is quantified by the dispersion (e.g., entropy) of this distribution. Conversely, the overall confidence of f can be defined inversely to this uncertainty. For a specific output o_i , its confidence is directly associated with its assigned probability $P_f(o = o_i | x)$.

3.2 Knowledge Graph

A Knowledge Graph (KG) is a graph-structured database representing factual knowledge as a set of triples [Bollacker *et al.*, 2008]. Formally, a KG is denoted as $\mathcal{G} := (\mathcal{V}, \mathcal{R}, \mathcal{E})$, where $\mathcal{V}$ is a set of entities, $\mathcal{R}$ is a set of relations, and $\mathcal{E} := (h, r, t) \subseteq \mathcal{V} \times \mathcal{R} \times \mathcal{V}$ is a set of factual triples. Each triple (h, r, t) represents an atomic fact, stating that relation r holds between head entity h and tail entity t . To mitigate the knowledge gaps of LLMs, KG-RAG provides them with structured evidence extracted from KGs [Xiang *et al.*, 2025].

3.3 Knowledge Graph Question Answering

Knowledge Graph Question Answering (KGQA) is a canonical knowledge-intensive reasoning task [Yih *et al.*, 2016]. Given a natural language question Q involving query entities $\mathcal{V}_Q$, a reasoning system is expected to retrieve relevant evidence from a KG $\mathcal{G}$ and reason over it to produce the correct answer set $\mathcal{A}$. A standard knowledge-augmented pipeline (e.g., KG-RAG) involves two stages: (i) a retriever g that fetches a set of relevant evidence $\mathcal{Z}_Q = g(Q; \mathcal{G})$, and (ii) a reasoner f (e.g., an LLM) that predicts answers $\hat{\mathcal{A}} = f(Q; \mathcal{Z}_Q, \mathcal{G})$. This decomposition naturally highlights two distinct sources of uncertainty that our framework aims to calibrate: the evidence uncertainty in $\mathcal{Z}_Q$, and the reasoning uncertainty in generating $\hat{\mathcal{A}}$ given $\mathcal{Z}_Q$.

4 Methodology

This section introduces **DoublyCal**, a framework designed to establish a calibrated reasoning chain by jointly calibrating both verifiable KG evidence and the final LLM predictions. As illustrated in Figure 2, the proposed DoublyCal framework operates through three core components. Firstly, we formalize KG evidence as constrained relational paths and employ a Bayesian model to estimate a statistically grounded confidence for each evidence (Sec. 4.1). Then, a lightweight proxy model is trained under the supervision of these Bayesian confidence scores to generate KG evidence alongside its calibrated confidence (Sec. 4.2). Finally, the calibrated evidence-confidence pair serves as an objective signal integrated into any black-box LLM to mitigate its inherent overconfidence, thereby enhancing both the calibration and traceability of its final predictions (Sec. 4.3).

4.1 Bayesian Calibration of KG Evidence

KG Evidence Formulation. Effective KG evidence must balance *informativeness* for accurate reasoning with *interpretability* for reliable confidence estimation. While relational paths [Luo *et al.*, 2024] offer step-by-step interpretability, they may lack sufficient context. In contrast, subgraphs [Li *et al.*, 2025a] provide broader context but often introduce redundancy. To resolve this trade-off, we introduce *constrained relational paths* as our primary evidence form. This formulation augments a core relational path with an optional constraint derived from the neighborhood of the candidate answer, thereby enhancing informativeness while preserving interpretability.

Formally, given a KG $\mathcal{G}$ and a question Q , a constrained relational path $\mathcal{P}_c$ is defined as the conjunction of a relational path $\mathcal{P}_r$ and an optional constraint $\mathcal{C}$:

$$\mathcal{P}_c := \mathcal{P}_r [\wedge \mathcal{C}], \quad (1)$$

where $\mathcal{P}_r := \exists v_1, \dots, v_{l-1}. r_1(q, v_1) \wedge r_2(v_1, v_2) \wedge \dots \wedge r_l(v_{l-1}, \hat{a})$ denotes a directed relational path of length l from the query entity $q \in \mathcal{V}_Q$ to a candidate answer $\hat{a}$. Here, each $r_i \in \mathcal{R}$ is a relation, and v_i is an existential variable. The optional constraint $\mathcal{C} := r_c(\hat{a}, c)$ represents a one-hop triple from the candidate answer $\hat{a}$ to a constraint entity $c \in \mathcal{V}$, which serves to filter or refine the candidate set.

Example. Consider the question “What is the name of Snoopy’s brother?” with the query entity $q = \text{Snoopy}$ and the true answer $a = \text{Spike}$. Figure 2 illustrates a relational path evidence z_3 and its constrained counterpart z_2 :

$$z_3 := \text{SiblingOf}(q, \hat{a}) \models_{\mathcal{G}} \{\text{Spike}, \text{Belle}\}, \quad (2)$$

$$z_2 := z_3 \wedge \text{Gender}(\hat{a}, \text{Male}) \models_{\mathcal{G}} \{\text{Spike}\}, \quad (3)$$

where $\models_{\mathcal{G}}$ denotes grounding the evidence in $\mathcal{G}$ to obtain candidate answer entities. The auxiliary constraint on the candidate’s gender effectively identifies $\hat{a} = \text{Spike}$, yielding a more precise and informative evidence for reasoning.

Confidence Estimation with Beta-Bernoulli Model. To estimate a statistically grounded confidence for a given KG evidence z_Q , we model it as a predictive system in accordance with the definition in Sec. 3.1. Formally, the system takes as input a candidate answer $\hat{a}$ drawn from the candidate set $\llbracket z_Q \rrbracket$ obtained by grounding z_Q in $\mathcal{G}$ (i.e., $z_Q \models_{\mathcal{G}} \llbracket z_Q \rrbracket$), and produces a binary output indicating whether $\hat{a}$ is correct ($\hat{a} \in \mathcal{A}$). This defines a Bernoulli distribution for the correctness of a uniformly sampled candidate, characterized by the parameter $p \in [0, 1]$, which is precisely the probability that the candidate is correct.

To obtain a robust estimate of p that accounts for KG incompleteness, we impose a conjugate Beta prior $p \sim \text{Beta}(\alpha, \beta)$, with hyperparameters $\alpha, \beta > 0$. Given Q and its answer set $\mathcal{A}$, we use the posterior mean as the calibrated confidence score p^* , which has the following closed form:

$$p^* = p(\mathcal{A} | z_Q) = \frac{\alpha + |\llbracket z_Q \rrbracket \cap \mathcal{A}|}{\alpha + \beta + |\llbracket z_Q \rrbracket|}, \quad (4)$$

where $|\llbracket z_Q \rrbracket \cap \mathcal{A}|$ counts the number of correct candidates, and $|\llbracket z_Q \rrbracket|$ is the total number of grounded candidates. p^* blends the empirical accuracy of the evidence with prior belief, mitigating the impact of sparse KG grounding.

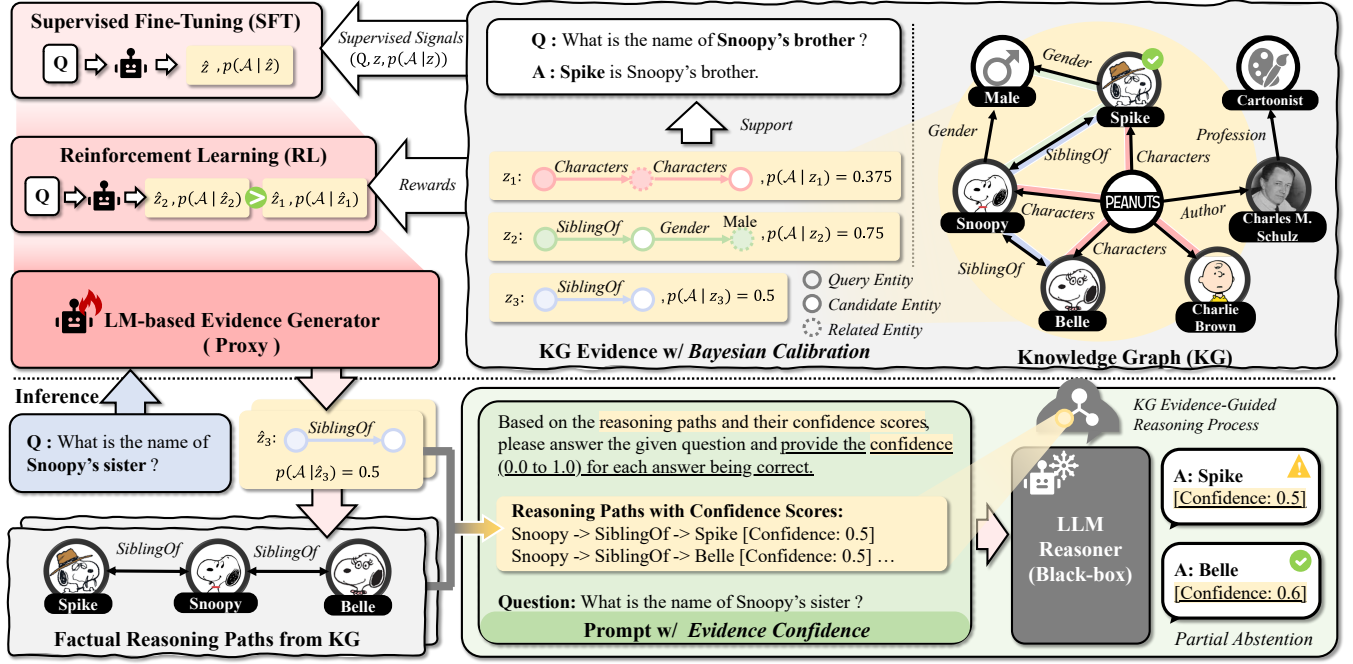


Figure 2: Overview of the DoublyCal framework. **(Top-right)** Bayesian calibration estimates statistically grounded confidence for KG evidence. **(Top-left)** A lightweight proxy model is trained to generate calibrated evidence-confidence pairs. **(Bottom)** During inference, these calibrated pairs guide a black-box LLM reasoner to produce final answers with well-calibrated confidence.

Example (cont.). With a weakly informative prior set to $\alpha = \beta = 0.5$ [Jeffreys, 1998], the statistical confidence for KG evidence z_3 and z_2 is estimated as $p(A|z_3) = 0.5$ and $p(A|z_2) = 0.75$.

4.2 Proxy for Evidence Generation & Calibration

The Beta-Bernoulli model-based confidence provides a statistically grounded but *retrospective* measure of evidence quality. To enable *prospective* evidence retrieval and confidence estimation during inference, we introduce a lightweight **reasoning proxy**. This proxy model is designed to generate high-quality KG evidence alongside well-calibrated confidence estimates for any input question, thereby approximating a reliable reasoning path before the LLM’s final prediction. The proxy model is implemented by an **LM-based Evidence Generator** and trained using the Bayesian confidence scores as its supervisory signal.

Supervised Fine-Tuning (SFT) for Evidence and Confidence Generation. We formalize the dual task of evidence generation and confidence estimation as a sequence-to-sequence problem and conduct the initial training of the proxy model via SFT. This stage equips the proxy with the fundamental ability to identify relevant KG evidence and estimate confidence by mimicking the Bayesian signal.

The SFT training dataset is constructed from triples $(Q, z_Q, p(A|z_Q))$. Each triple is formatted into a structured sequence using a predefined template: the question serves as the instruction, and the target output is the evidence path enclosed in XML-style tags, with the Bayesian confidence score included as an attribute (e.g.,

$\langle \text{PATH confidence} = \dots \rangle \dots \langle / \text{PATH} \rangle$). The proxy model f_θ is trained via standard autoregressive language modeling to generate this target sequence.

Reinforcement Learning (RL) for Evidence Decision. We further refine the proxy by framing evidence generation and calibration as a sequential decision process optimized via RL. This stage transitions the proxy from imitation to strategic decision-making that jointly maximizes both inferential quality and confidence calibration.

For each question Q with a gold evidence set z_Q , we compute a reward for the generated evidence $\hat{z}_Q$ and its predicted confidence $\hat{c}$. Firstly, we define a match score $m(\hat{z}_Q, z_Q) \in [0, 1]$ for each $z_Q \in Z_Q$, which combines Jaccard similarity [Jaccard, 1901] with an order-sensitive Levenshtein ratio [Levenshtein and others, 1966]. The overall reward R is a weighted combination of an inferential quality reward R_{inf} and a calibration alignment reward R_{cal} :

$$R = \lambda \cdot R_{\text{inf}} + (1 - \lambda) \cdot R_{\text{cal}}, \quad (5)$$

$$R_{\text{inf}} = \text{F1}(z_Q) \cdot m(\hat{z}_Q, z_Q), \quad (6)$$

$$R_{\text{cal}} = \max(0, 1 - \xi \cdot |\hat{c} - c|), \quad (7)$$

$$\text{with } c = p(A|z_Q) \cdot m(\hat{z}_Q, z_Q). \quad (8)$$

Here, $\text{F1}(z_Q)$ is a precomputed F1 score assessing the reasoning capability of the gold evidence. Intuitively, a lower match score $m(\hat{z}_Q, z_Q)$ reduces both the inferential quality of $\hat{z}_Q$ and its target confidence c . The weight $\lambda \in (0, 1)$ balances the two objectives, and $\xi > 0$ is a tolerance coefficient. The final reward per generation is the maximum R over Z_Q , followed by transformations to ensure a smooth training signal. The policy π_θ of the proxy model is optimized to max-

Reasoning Method	KG Evidence	+ UQ Method	WebQSP					CWQ				
			Hits	Recall	F1	ECE ↓	ACE ↓	Hits	Recall	F1	ECE ↓	ACE ↓
LLM Reasoner (GPT-3.5-turbo)	<i>No Augmentation</i>	+ Vanilla	74.7	53.1	44.6	27.7	26.6	47.7	40.3	29.5	38.8	38.0
		+ CoT	75.4	53.9	44.4	26.6	25.8	48.2	41.2	29.3	38.4	37.6
		+ Self-Probing	74.1	52.8	50.2	36.5	36.3	43.6	36.6	34.3	48.5	47.9
RoG [Luo <i>et al.</i> , 2024]	<i>Relational Path</i>	+ Vanilla	89.3	77.6	67.1	19.6	20.8	65.3	60.5	43.0	27.4	26.8
		+ CoT	89.9	78.5	68.2	18.9	17.9	65.3	60.4	43.7	27.6	27.0
		+ Self-Probing	87.5	76.6	73.5	13.9	12.8	61.9	56.9	48.7	38.0	37.4
SubgraphRAG [Li <i>et al.</i> , 2025a]	<i>Subgraph</i>	+ Vanilla	88.8	81.3	77.3 [‡]	11.1	9.7	61.5	57.4	52.2 [‡]	39.9	39.7
		+ CoT	89.0	81.0	77.1	10.6	9.5	59.4	55.7	51.4	38.9	38.7
		+ Self-Probing	89.6	80.7	74.9	12.3	12.8	59.9	56.0	50.2	39.1	38.0
SFT-DoublyCal (Ours)	<i>Constrained Relational Path</i>	+ Vanilla	90.0	81.0	72.6	3.1 [†]	3.8 [‡]	68.8	64.3	48.1	17.9	17.5
		+ CoT	90.1 [‡]	81.3	72.1	3.5 [‡]	3.6 [†]	69.0	64.7	47.7	17.8 [‡]	17.4
		+ Self-Probing	88.5	79.4	76.6	7.9	7.1	63.2	58.5	50.9	22.5	22.1
RL-DoublyCal (Ours)	<i>Constrained Relational Path</i>	+ Vanilla	91.5 [†]	84.8 [‡]	76.7	4.5	5.1	70.5 [‡]	66.6 [‡]	50.1	17.9	17.3 [‡]
		+ CoT	91.5 [†]	85.0 [†]	76.8	3.9	4.3	71.3 [†]	67.5 [†]	49.8	17.6 [†]	17.2 [†]
		+ Self-Probing	89.9	83.0	79.3 [†]	6.8	6.8	64.6	60.8	53.0 [†]	23.5	23.1

Table 1: Main results (%) of our DoublyCal and the SingleCal baselines on WebQSP and CWQ datasets. Best, second-best, and worst results are highlighted in red [†], green [‡], and gray, respectively.

imize the expected reward under the Group Relative Policy Optimization (GRPO) [Shao *et al.*, 2024] objective.

4.3 LLM Reasoning with Calibrated Evidence

The trained proxy equips any black-box primary LLM with double-calibration capability in a plug-and-play manner.

For a given question Q , the proxy model generates a set of candidate evidence-confidence pairs, i.e., $\hat{Z}_Q = \{(\hat{z}_Q^{(i)}, \hat{c}^{(i)})\}_{i=1\dots K}$. Each $\hat{z}_Q^{(i)}$ is grounded in $\mathcal{G}$, yielding factual reasoning paths that share the same confidence score $\hat{c}^{(i)}$. These paths and their confidences are verbalized into a natural-language context and integrated into prompts following prior verbalized UQ methods [Tian *et al.*, 2023; Xiong *et al.*, 2024]. Processing this enriched context, the LLM produces a final answer along with a well-calibrated prediction confidence. This design establishes a traceable chain of confidence and achieves double calibration: the proxy first calibrates the external evidence confidence, which then informs and refines the LLM’s final prediction calibration through its own verbalized uncertainty estimation.

5 Experiments

5.1 Experimental Settings

Datasets and Evaluation Metrics. Following [Luo *et al.*, 2024; Li *et al.*, 2025a], we evaluate our framework on two widely-adopted Knowledge Graph Question Answering (KGQA) benchmarks: **WebQSP** [Yih *et al.*, 2016] and **CWQ** [Talmor and Berant, 2018]. To measure prediction accuracy, we report **Hits**, **Recall**, and macro-averaged **F1** score. To assess the reliability of confidence estimates, we report the **Expected Calibration Error (ECE)** [Guo *et al.*, 2017] and the **Adaptive Calibration Error (ACE)** [Nixon *et al.*, 2019]. ECE and ACE measure the expected absolute difference between empirical accuracy and predicted confidence using equal-width and adaptive equal-size bins, respectively.

Baselines. To rigorously evaluate our Double-Calibration mechanism, we construct **Single-Calibration (SingleCal)** baselines by extending reasoning paradigms with verbalized Uncertainty Quantification (UQ) methods, which elicit a self-reported confidence score alongside the predicted answer. We select three state-of-the-art reasoning frameworks: (i) The base **LLM Reasoner** without KG access; (ii) **RoG** [Luo *et al.*, 2024], a KG-RAG method that grounds reasoning in retrieved relational paths; (iii) **SubgraphRAG** [Li *et al.*, 2025a], which retrieves and reasons over KG subgraphs. Each framework is combined with three representative UQ prompting techniques: **Vanilla** [Tian *et al.*, 2023], **CoT** [Kojima *et al.*, 2022], and **Self-Probing** [Xiong *et al.*, 2024].

Implementation Details. All evaluated methods employ GPT-3.5-turbo [Floridi and Chiriatti, 2020] as the primary reasoner unless otherwise specified, ensuring that performance differences are directly attributable to the calibration mechanism rather than the base LLM capability. The evidence proxy in our DoublyCal is implemented with Llama2-7B-Chat [Touvron *et al.*, 2023], which is trained via the SFT then RL pipeline described in Sec. 4.2.

5.2 Main Results

Table 1 summarizes the comparative performance of our DoublyCal against all SingleCal baselines.

Superiority of Double-Calibration. DoublyCal achieves the best overall performance, consistently leading in all prediction metrics (Hits, Recall, F1) and the lowest ECE and ACE. Notably, it establishes a new standard for reliability, reducing the ECE and ACE to levels significantly lower than all SingleCal baselines. This result demonstrates that while KG-RAG methods can enhance LLM factuality, calibrating *both* the external KG evidence and the final prediction is necessary for achieving reliable reasoning. Furthermore, the observed gains in prediction metrics indicate that evidence confidence

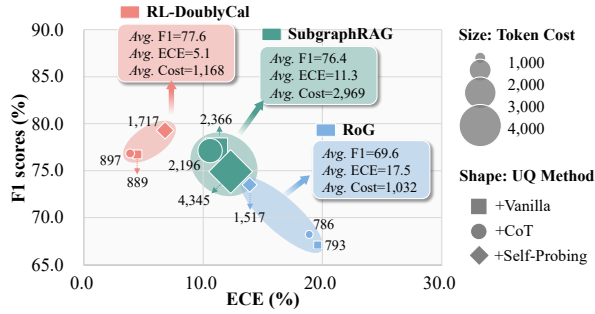


Figure 3: Input token efficiency analysis on WebQSP.

helps improve the quality of the retrieved evidence and further unlocks the LLM’s reasoning potential.

Calibrated Evidence as an Anchor for Verbalized UQ. The efficacy of verbalized UQ methods varies across reasoning backbones, with none proving universally dominant. This inconsistency arises because these methods may be subject to LLMs’ inherent overconfidence. DoublyCal addresses this by supplying KG evidence with calibrated confidence, providing a reliable external anchor that refines the LLM’s uncertainty expression. Consequently, DoublyCal stabilizes all three UQ techniques and achieves the lowest ECE and ACE in nearly every configuration, showing that externally calibrated evidence improves confidence elicitation.

Controlled Enhancement via RL. The RL stage yields a significant performance gain, improving average F1 by ~ 3.0 percentage points over the SFT-only. While prior work notes RL’s risk of harming calibration [Kalai *et al.*, 2025], our Bayesian confidence-aligned reward successfully mitigates this trade-off, resulting in only a minor and controlled variation in ECE and ACE. This confirms that our reward design effectively balances accuracy and calibration.

5.3 Efficiency Analysis

We further analyze the input token efficiency of each method, with results shown in Figure 3.

Superior Cost-Effectiveness of DoublyCal. DoublyCal substantially outperforms RoG (F1 +8.0, ECE -12.4) with only a marginal increase in input token cost, while consuming only about 39% of the input tokens required by SubgraphRAG. This efficiency stems from the high information density of the evidence provided by DoublyCal. Specifically, compared to the simple relational paths in RoG, our constrained relational paths incorporate an optional constraint that yields more precise and informative evidence without compromising conciseness. Moreover, our proxy model is trained through evidence confidence calibration to select more discriminative KG evidence, further enhancing retrieval precision. In contrast, while SubgraphRAG’s subgraphs offer broader context, their lower information density leads to disproportionately high token costs.

Efficiency Across Different UQ Methods. Self-Probing incurs approximately twice the input token cost of Vanilla or CoT due to its two-step prompting design. However, because

Variant	Hits	Recall	F1	ECE ↓	ACE ↓
SFT Only					
DoublyCal	90.0	81.0	72.6	3.1	3.8
SingleCal	90.1 (+0.1)	80.4 (-0.6)	72.5 (-0.1)	21.2 (+18)	20.3 (+16)
Evidence	83.7 (-6.3)	80.2 (-0.8)	62.5 (-10)	21.2 (+18)	21.0 (+17)
With RL					
DoublyCal	91.5	84.8	76.7	4.5	5.1
SingleCal	91.6 (+0.1)	84.2 (-0.6)	75.8 (-0.9)	20.6 (+16)	20.0 (+15)
Evidence	86.4 (-5.1)	84.8	67.8 (-8.9)	21.3 (+16)	21.1 (+16)

Table 2: Ablation study results (%) on the WebQSP dataset with the Vanilla UQ method. Gray marks performance degradation relative to the full model.

DoublyCal and RoG retrieve concise evidence, they can effectively leverage Self-Probing’s reflective “second-thought” process without excessive overhead. Notably, even when equipped with Self-Probing, their total input token cost remains below that of SubgraphRAG+Vanilla.

5.4 Ablation Analysis

We ablate DoublyCal against two variants (Table 2) to examine the contribution of each component: (i) **SingleCal**, which removes the calibrated evidence confidence and applies calibration only to the final LLM output; (ii) **Evidence**, which removes the LLM reasoner and directly outputs the terminal entity of the factual path ($\llbracket z_Q \rrbracket$) as the answer, using the evidence confidence as the final confidence score.

Evidence Confidence is Crucial for Final Prediction Calibration. Ablating from DoublyCal to SingleCal reveals a stark outcome: while predictive accuracy remains stable (e.g., F1 changes within ± 1 point), the calibration error (ECE and ACE) increases drastically from ~ 4 to >20 . This indicates that externally calibrated evidence confidence is essential, as it provides the LLM with a reliable anchor for its self-assessment. Without this first-stage calibration, the black-box LLM cannot reliably judge its own certainty, even when it can identify correct answers using high-quality KG evidence.

The LLM Reasoner Enables Integrative Reasoning. The significant performance gap of the Evidence variant underscores a pivotal design insight: the evidence proxy and the LLM reasoner play distinct yet complementary roles. The proxy specializes in evaluating individual KG evidence, while the LLM reasoner excels at synthesizing an ensemble of such evidence to perform complex reasoning. Consequently, the final prediction confidence is not a simple pass-through of any single evidence confidence, but rather the result of the LLM’s holistic reasoning over the entire set of calibrated evidence. This efficient proxy-reasoner synergy is essential to the framework’s performance.

5.5 Cross-model Compatibility Analysis

To assess generalizability of DoublyCal, we evaluate it across diverse black-box LLMs, including GPT-3.5-turbo, GPT-4o-mini [Achiam *et al.*, 2023], DeepSeek-V3 [Liu *et al.*, 2024], and Gemini-2.5-flash [Comanici *et al.*, 2025].

Performance-Reliability Trade-off in LLM Reasoners. Figure 4 reveals a clear trade-off between accuracy and calibration among standalone LLMs. While GPT-family models

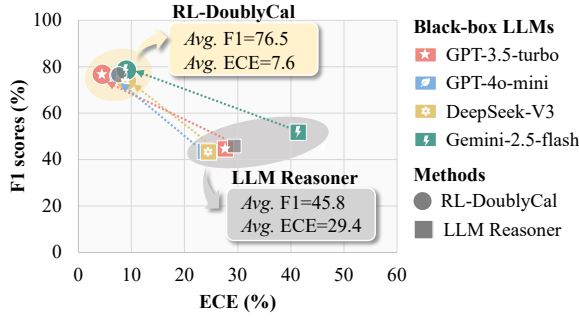


Figure 4: Prediction accuracy (F1) and calibration (ECE) of diverse black-box LLMs with and without DoublyCal.

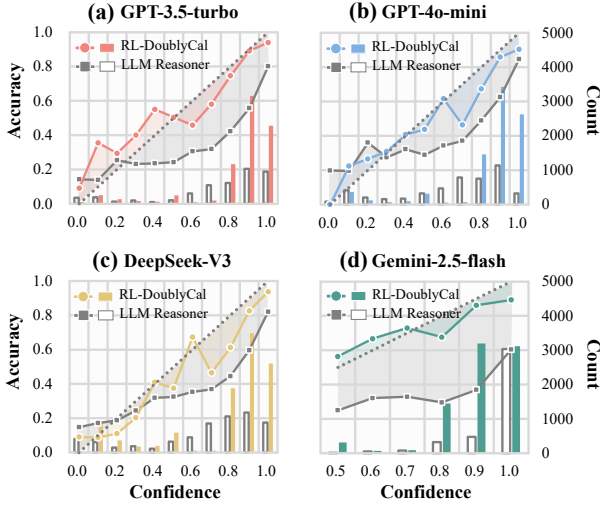


Figure 5: Calibration diagrams (bars: confidence distribution per confidence bin; line: empirical accuracy; dashed: ideal calibration).

and DeepSeek achieve comparable accuracy with moderate calibration errors (F1: 43.3–44.6; ECE: 23.9–27.7), Gemini attains a notably higher F1 (51.8) at the cost of a significantly worse ECE (41.4). This pattern highlights a common pitfall where optimizing purely for accuracy often degrades reliability in standalone LLMs. Confidence distributions (Figure 5) confirms that all models exhibit systematic overconfidence. This issue is most acute in Gemini, where roughly 80% of predictions are made with maximal confidence (1.0), yet the accuracy within this high-confidence group is only about 0.6.

DoublyCal Systematically Decouples the Trade-off.

DoublyCal delivers consistent and substantial improvements across all models, effectively decoupling this trade-off. As shown in Figure 4, it raises the average F1 from 45.8 to 76.5 while reducing the average ECE from 29.4 to 7.6. Crucially, DoublyCal mitigates the overconfidence patterns observed in black-box LLMs (Figure 5), shifting confidence distributions toward well-calibrated and high-accuracy regions. By grounding confidence in externally calibrated evidence, DoublyCal provides a generalizable solution that enhances both accuracy and reliability of diverse black-box LLMs.

Sample	
Question: Where did George W. Bush live as a child?	
Query Entity (q): George W. Bush	Answers: New Haven.
RL-DoublyCal + Self-Probing	
Retrieval:	
- q → people.person.place_of_birth → New Haven [Confidence: 0.8] - q → people.person.place_of_birth → New Haven → location.location.containedby → Connecticut [Confidence: 0.5] - q → people.person.place_of_birth → New Haven → location.location.containedby → United States of America [Confidence: 0.5]	
Predictions: {Connecticut: 0.3}	
SubgraphRAG + Vanilla	
Retrieval:	
(q, people.person.place_of_birth, New Haven) (q, people.person.nationality, United States of America) (m.03prwzr, people.place_lived.location, Midland) (m.02xlp0j, people.place_lived.location, Washington, D.C.) • ...	
Predictions: {Midland: 1.0}	

Table 3: Case study: DoublyCal vs. SingleCal baseline.

Sample	
Question: Where was Martin Luther King, Jr. raised?	
Query Entity (q): Martin Luther King, Jr.	Answers: Atlanta.
RL-DoublyCal (Full) + Vanilla	
Retrieval:	
- q people.person.place_of_birth → Atlanta [Confidence: 0.8] - q → people.deceased.person.place_of_death → Memphis [Confidence: 0.8]	
Predictions: {Atlanta: 0.8, Memphis: 0.1}	
RL-DoublyCal (SingleCal) + Vanilla	
Retrieval: Same factual paths as above, without confidence scores	
Predictions: {Atlanta: 0.8, Memphis: 0.2}	

Table 4: Case study: DoublyCal vs. its SingleCal variant.

5.6 Case Studies

This section presents two case studies that qualitatively illustrate how DoublyCal enhances reliability of LLM predictions.

As shown in Table 3, for a question lacking a direct KG relation (“lived as a child”) to provide accurate support, both methods retrieve the related birthplace fact. However, while SubgraphRAG retrieves scattered evidence leading to an overconfident error (confidence 1.0), DoublyCal presents concise paths with calibrated confidence scores. Its birthplace path receives high confidence (0.8), while less relevant expansions are scored lower (0.5). Consequently, the LLM correctly assigns low confidence (0.3) to the incorrect answer “Connecticut”. Furthermore, the ablation study in Table 4 demonstrates that explicitly providing evidence confidence makes the LLM’s predicted confidence more concentrated on the correct answer. Together, these cases show DoublyCal’s ability to mitigate overconfidence and improve calibration.

6 Conclusion

This paper establishes the principle of double-calibration for constructing a calibrated reasoning chain from KG evidence retrieval to final LLM prediction. We implement this principle in DoublyCal, a reliable KG-RAG framework that integrates plug-and-play verbalized uncertainty quantification, thereby enhancing the traceability and reliability of diverse black-box LLMs. Our work offers a concrete step toward building more reliable and transparent LLM systems, contributing to the advancement of trustworthy AI.

Acknowledgments

We thank the reviewers for their valuable suggestions. This work has been supported by the National Natural Science Foundation of China (62372483), a grant from the Research Grants Council of the Hong Kong Special Administrative Region, China (UGC/FDS16/E23/24), and the Hong Kong Research Grants Council under the Theme-based Research Scheme (T41-517/25-N). This work has also been supported by the Key Special Project of National Natural Science Foundation of China (72442017) and the General Research Funds from the Hong Kong Research Grants Council (PolyU 15207322, 15200023, 15206024, and 15224524).

References

- [Achiam *et al.*, 2023] 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. *CoRR*, abs/2303.08774, 2023.
- [Bollacker *et al.*, 2008] Kurt D. Bollacker, Colin Evans, Praveen K. Paritosh, Tim Sturge, and Jamie Taylor. Freebase: a collaboratively created graph database for structuring human knowledge. In *SIGMOD*, pages 1247–1250, 2008.
- [Chen *et al.*, 2024] Chao Chen, Kai Liu, Ze Chen, Yi Gu, Yue Wu, Mingyuan Tao, Zhihang Fu, and Jieping Ye. IN-SIDE: llms’ internal states retain the power of hallucination detection. In *ICLR*, 2024.
- [Comanici *et al.*, 2025] Gheorghe Comanici, Eric Bieber, Mike Schaeckermann, Ice Pasupat, Noveen Sachdeva, Inderjit Dhillon, Marcel Blistein, Ori Ram, Dan Zhang, Evan Rosen, et al. Gemini 2.5: Pushing the frontier with advanced reasoning, multimodality, long context, and next generation agentic capabilities. *CoRR*, abs/2507.06261, 2025.
- [Floridi and Chiriatti, 2020] Luciano Floridi and Massimo Chiriatti. GPT-3: Its nature, scope, limits, and consequences. *Minds Mach.*, 30(4):681–694, 2020.
- [Guo *et al.*, 2017] Chuan Guo, Geoff Pleiss, Yu Sun, and Kilian Q. Weinberger. On calibration of modern neural networks. In *ICML*, pages 1321–1330, 2017.
- [Guo *et al.*, 2025] Zirui Guo, Lianghao Xia, Yanhua Yu, Tu Ao, and Chao Huang. Lightrag: Simple and fast retrieval-augmented generation. In *Findings of EMNLP*, pages 10746–10761, 2025.
- [He *et al.*, 2024] Xiaoxin He, Yijun Tian, Yifei Sun, Nitesh V. Chawla, Thomas Laurent, Yann LeCun, Xavier Bresson, and Bryan Hooi. G-retriever: Retrieval-augmented generation for textual graph understanding and question answering. In *NeurIPS*, 2024.
- [Huang *et al.*, 2024] Yue Huang, Lichao Sun, Haoran Wang, Siyuan Wu, Qihui Zhang, Yuan Li, Chujie Gao, Yixin Huang, Wenhan Lyu, Yixuan Zhang, Xiner Li, Hanchi Sun, Zhengliang Liu, Yixin Liu, Yijue Wang, Zhikun Zhang, Bertie Vidgen, Bhavya Kailkhura, Caiming Xiong, Chaowei Xiao, Chunyuan Li, Eric P. Xing, Furong Huang, Hao Liu, Heng Ji, Hongyi Wang, Huan Zhang, Huaxiu Yao, Manolis Kellis, Marinka Zitnik, Meng Jiang, Mohit Bansal, James Zou, Jian Pei, Jian Liu, Jianfeng Gao, Jiawei Han, Jieyu Zhao, Jiliang Tang, Jindong Wang, Joaquin Vanschoren, John C. Mitchell, Kai Shu, Kaidi Xu, Kai-Wei Chang, Lifang He, Lifu Huang, Michael Backes, Neil Zhenqiang Gong, Philip S. Yu, Pin-Yu Chen, Quanquan Gu, Ran Xu, Rex Ying, Shuiwang Ji, Suman Jana, Tianlong Chen, Tianming Liu, Tianyi Zhou, William Wang, Xiang Li, Xiangliang Zhang, Xiao Wang, Xing Xie, Xun Chen, Xuyu Wang, Yan Liu, Yanfang Ye, Yinzi Cao, Yong Chen, and Yue Zhao. Position: Trustilm: Trustworthiness in large language models. In *ICML*, 2024.
- [Hüllermeier and Waegeman, 2021] Eyke Hüllermeier and Willem Waegeman. Aleatoric and epistemic uncertainty in machine learning: an introduction to concepts and methods. *Mach. Learn.*, 110(3):457–506, 2021.
- [Jaccard, 1901] Paul Jaccard. Étude comparative de la distribution florale dans une portion des alpes et des jura. *Bull. Soc. Vaudoise. Sci. Nat.*, 37:547–579, 1901.
- [Jeffreys, 1998] Harold Jeffreys. *The theory of probability*. OUP Oxford, 1998.
- [Kalai *et al.*, 2025] Adam Tauman Kalai, Ofir Nachum, Santosh S. Vempala, and Edwin Zhang. Why language models hallucinate. *CoRR*, abs/2509.04664, 2025.
- [Kojima *et al.*, 2022] Takeshi Kojima, Shixiang Shane Gu, Machel Reid, Yutaka Matsuo, and Yusuke Iwasawa. Large language models are zero-shot reasoners. In *NeurIPS*, 2022.
- [Kuhn *et al.*, 2023] Lorenz Kuhn, Yarin Gal, and Sebastian Farquhar. Semantic uncertainty: Linguistic invariances for uncertainty estimation in natural language generation. In *ICLR*, 2023.
- [Levenshtein and others, 1966] Vladimir I Levenshtein et al. Binary codes capable of correcting deletions, insertions, and reversals. *Soviet Physics Doklady*, 10(8):707–710, 1966.
- [Li *et al.*, 2025a] Mufei Li, Siqi Miao, and Pan Li. Simple is effective: The roles of graphs and large language models in knowledge-graph-based retrieval-augmented generation. In *ICLR*, 2025.
- [Li *et al.*, 2025b] Zhuoqun Li, Xuanang Chen, Haiyang Yu, Hongyu Lin, Yaojie Lu, Qiaoyu Tang, Fei Huang, Xianpei Han, Le Sun, and Yongbin Li. Structrag: Boosting knowledge intensive reasoning of llms via inference-time hybrid information structurization. In *ICLR*, 2025.
- [Lin *et al.*, 2024] Zhen Lin, Shubhendu Trivedi, and Jimeng Sun. Generating with confidence: Uncertainty quantification for black-box large language models. *Trans. Mach. Learn. Res.*, 2024, 2024.
- [Liu *et al.*, 2024] Aixin Liu, Bei Feng, Bing Xue, Bingxuan Wang, Bochao Wu, Chengda Lu, Chenggang Zhao, Chengqi Deng, Chenyu Zhang, Chong Ruan, et al. Deepseek-v3 technical report. *CoRR*, abs/2412.19437, 2024.

- [Liu *et al.*, 2026] Shuyi Liu, Yuming Shang, and Xi Zhang. Truthfulrag: Resolving factual-level conflicts in retrieval-augmented generation with knowledge graphs. In *AAAI*, pages 32168–32176, 2026.
- [Luo *et al.*, 2024] Linhao Luo, Yuan-Fang Li, Gholamreza Haffari, and Shirui Pan. Reasoning on graphs: Faithful and interpretable large language model reasoning. In *ICLR*, 2024.
- [Malinin and Gales, 2021] Andrey Malinin and Mark J. F. Gales. Uncertainty estimation in autoregressive structured prediction. In *ICLR*, 2021.
- [Manakul *et al.*, 2023] Potsawee Manakul, Adian Liusie, and Mark J. F. Gales. Selfcheckgpt: Zero-resource black-box hallucination detection for generative large language models. In *EMNLP*, pages 9004–9017, 2023.
- [Nixon *et al.*, 2019] Jeremy Nixon, Michael W Dusenberry, Linchuan Zhang, Ghassen Jerfel, and Dustin Tran. Measuring calibration in deep learning. In *CVPR Workshops*, 2019.
- [Pan *et al.*, 2024] Shirui Pan, Linhao Luo, Yufei Wang, Chen Chen, Jiapu Wang, and Xindong Wu. Unifying large language models and knowledge graphs: A roadmap. *IEEE Trans. Knowl. Data Eng.*, 36(7):3580–3599, 2024.
- [Shao *et al.*, 2024] Zhihong Shao, Peiyi Wang, Qihao Zhu, Runxin Xu, Junxiao Song, Mingchuan Zhang, Y. K. Li, Y. Wu, and Daya Guo. Deepseekmath: Pushing the limits of mathematical reasoning in open language models. *CoRR*, abs/2402.03300, 2024.
- [Stangel *et al.*, 2025] Paul Stangel, David Bani-Harouni, Chantal Pellegrini, Ege Özsoy, Kamilia Zaripova, Matthias Keicher, and Nassir Navab. Rewarding doubt: A reinforcement learning approach to calibrated confidence expression of large language models. *CoRR*, abs/2503.02623, 2025.
- [Sun *et al.*, 2025] Jiashuo Sun, Xianrui Zhong, Sizhe Zhou, and Jiawei Han. Dynamicrag: Leveraging outputs of large language model as feedback for dynamic reranking in retrieval-augmented generation. In *NeurIPS*, 2025.
- [Talmor and Berant, 2018] Alon Talmor and Jonathan Berant. The web as a knowledge-base for answering complex questions. In *NAACL-HLT*, pages 641–651, 2018.
- [Tanneru *et al.*, 2024] Sree Harsha Tanneru, Chirag Agarwal, and Himabindu Lakkaraju. Quantifying uncertainty in natural language explanations of large language models. In *AISTATS*, pages 1072–1080, 2024.
- [Tian *et al.*, 2023] Katherine Tian, Eric Mitchell, Allan Zhou, Archit Sharma, Rafael Rafailov, Huaxiu Yao, Chelsea Finn, and Christopher D. Manning. Just ask for calibration: Strategies for eliciting calibrated confidence scores from language models fine-tuned with human feedback. In *EMNLP*, pages 5433–5442, 2023.
- [Touvron *et al.*, 2023] Hugo Touvron, Louis Martin, Kevin Stone, Peter Albert, Amjad Almahairi, Yasmine Babaei, Nikolay Bashlykov, Soumya Batra, Prajjwal Bhargava, Shruti Bhosale, Dan Bikel, Lukas Blecher, Cristian Canton-Ferrer, Moya Chen, Guillem Cucurull, David Es-
iobu, Jude Fernandes, Jeremy Fu, Wenyin Fu, Brian Fuller, Cynthia Gao, Vedanuj Goswami, Naman Goyal, Anthony Hartshorn, Saghar Hosseini, Rui Hou, Hakan Inan, Marcin Kardas, Viktor Kerkez, Madian Khabsa, Isabel Kloumann, Artem Korenev, Punit Singh Koura, Marie-Anne Lachaux, Thibaut Lavril, Jenya Lee, Diana Liskovich, Yinghai Lu, Yuning Mao, Xavier Martinet, Todor Mihaylov, Pushkar Mishra, Igor Molybog, Yixin Nie, Andrew Poulton, Jeremy Reizenstein, Rashi Rungta, Kalyan Saladi, Alan Schelten, Ruan Silva, Eric Michael Smith, Ranjan Subramanian, Xiaoqing Ellen Tan, Binh Tang, Ross Taylor, Adina Williams, Jian Xiang Kuan, Puxin Xu, Zheng Yan, Iliyan Zarov, Yuchen Zhang, Angela Fan, Melanie Kambadur, Sharan Narang, Aurélien Rodriguez, Robert Stojnic, Sergey Edunov, and Thomas Scialom. Llama 2: Open foundation and fine-tuned chat models. *CoRR*, abs/2307.09288, 2023.
- [Vazhentsev *et al.*, 2025] Artem Vazhentsev, Lyudmila Rvanova, Ivan Lazichny, Alexander Panchenko, Maxim Panov, Timothy Baldwin, and Artem Shelmanov. Token-level density-based uncertainty quantification methods for eliciting truthfulness of large language models. In *NAACL*, pages 2246–2262, 2025.
- [Xia *et al.*, 2025] Zhiqiu Xia, Jinxuan Xu, Yuqian Zhang, and Hang Liu. A survey of uncertainty estimation methods on large language models. In *Findings of ACL*, pages 21381–21396, 2025.
- [Xiang *et al.*, 2025] Zhishang Xiang, Chuanjie Wu, Qinggang Zhang, Shengyuan Chen, Zijin Hong, Xiao Huang, and Jinsong Su. When to use graphs in RAG: A comprehensive analysis for graph retrieval-augmented generation. *CoRR*, abs/2506.05690, 2025.
- [Xiong *et al.*, 2024] Miao Xiong, Zhiyuan Hu, Xinyang Lu, Yifei Li, Jie Fu, Junxian He, and Bryan Hooi. Can llms express their uncertainty? an empirical evaluation of confidence elicitation in llms. In *ICLR*, 2024.
- [Yih *et al.*, 2016] Wen-tau Yih, Matthew Richardson, Christopher Meek, Ming-Wei Chang, and Jina Suh. The value of semantic parse labeling for knowledge base question answering. In *ACL*, 2016.
- [Zhang *et al.*, 2025a] Qinggang Zhang, Shengyuan Chen, Yuanchen Bei, Zheng Yuan, Huachi Zhou, Zijin Hong, Junnan Dong, Hao Chen, Yi Chang, and Xiao Huang. A survey of graph retrieval-augmented generation for customized large language models. *CoRR*, abs/2501.13958, 2025.
- [Zhang *et al.*, 2025b] Qinggang Zhang, Zhishang Xiang, Yilin Xiao, Le Wang, Junhui Li, Xinrun Wang, and Jinsong Su. Faithfulrag: Fact-level conflict modeling for context-faithful retrieval-augmented generation. In *ACL*, pages 21863–21882, 2025.