

RT-UIE: A Reward-Guided and Task-Oriented Framework for Unified Information Extraction

Wentong Wan^{1,2,3}, Yingxue Qiao^{1,2,3}, Yiwen Zhang^{1,2,3}, Yu Du^{1,2,3}, Song Liu^{1,2,3*}

¹Key Laboratory of Computing Power Network and Information Security, Ministry of Education,
Shandong Computer Science Center (National Supercomputer Center in Jinan),
Qilu University of Technology (Shandong Academy of Sciences), Jinan, China

²Shandong Engineering Research Center of Big Data Applied Technology, Faculty of Computer Science and Technology,
Qilu University of Technology (Shandong Academy of Sciences), Jinan, China

³Shandong Provincial Key Laboratory of Industrial Network and Information System Security,
Shandong Fundamental Research Center for Computer Science, Jinan, China

Email: {10431240080, 10431250164}@stu.qlu.edu.cn, {zywqlu, duyuqlu}@126.com, liusong@qlu.edu.cn

Abstract—Information extraction (IE) faces challenges from task-specific label spaces and heterogeneous data structures. Recent advances in large language models (LLMs) and retrieval-augmented generation (RAG) have enabled unified information extraction across diverse tasks, but existing methods still suffer from several limitations. Conventional RAG does not adequately account for IE-specific preferences, resulting in the underutilization of the background knowledge obtained. Furthermore, most approaches rely on fixed instruction templates and lack adaptive mechanisms for task-aware prompting. In addition, long instruction sequences often cause attention drift, weakening the model’s focus on core extraction targets. To address these problems, we propose RT-UIE, a reward-guided and task-oriented unified information extraction framework. In our framework, we design a task-oriented RAG retrieval and knowledge reconstruction mechanism that structurally reorganizes retrieved background knowledge, improving the consistency between background knowledge and extraction preferences. Moreover, we design a chain-of-thought reward model that dynamically adjusts the depth and quality of the reasoning according to task characteristics, enabling more effective reasoning control. Additionally, a module for instruction and attention calibration is designed to mitigate attention drift and improve the focus of the model on key task instructions. We evaluated our model on the IE_INSTRUCTION, WNUT-17, DocRED, and MAVEN datasets and compared it with other baseline models.

Index Terms—Unified Information Extraction; Large Language Models; Task-Oriented RAG Retrieval and Reconstruction; Chain-of-thought Reward Model; Attention Calibration.

I. INTRODUCTION

Information extraction (IE) [1] is a fundamental research task in natural language processing (NLP) [2], which focuses on the transformation of unstructured text into structured representations. By enabling automatic extraction of structured information, IE provides essential support for a wide range of downstream intelligent applications. In recent years, the

increasing demand for structured knowledge in applications such as intelligent question answering [3], knowledge graph construction [4], and recommendation systems [5] has further highlighted the importance of IE. Consequently, the accurate extraction of structured information from large-scale textual data has become a critical factor in improving the efficiency and decision-making capabilities of these applications.

According to extraction targets, IE can be broadly categorized into several sub-tasks, including Named Entity Recognition (NER) [6], Relation Extraction (RE) [7], and Event Extraction (EE) [8]. NER identifies entities with predefined types from text. Based on these entities, RE extracts semantic relations between entity pairs, while EE identifies events and their corresponding arguments, including event types, triggers, and related entities.

Early IE methods mainly focused on task-specific deep learning models. Lample et al. [9] proposed a BiLSTM-CRF model for end-to-end NER without manual feature engineering. Miwa et al. [10] introduced a neural model that combines sequence and tree representations for joint entity and relation extraction. However, these methods are typically designed for specific tasks or domains, requiring separate models and datasets, which limits their scalability and generalization.

Recent studies have explored large language models (LLMs) [11] for unified information extraction. Lu et al. [12] and Lou et al. [13] decompose IE into sub-tasks via prompt engineering, but they treat each sub-task independently without modeling their interactions. Wang et al. [14] introduce InstructUIE with multi-task instruction tuning to capture task relationships, yet its effectiveness is limited by weak reasoning guidance. Zhan et al. [15] propose the RAMIE framework, which improves multi-task information extraction by enhancing retrieval. However, directly using retrieved examples can introduce irrelevant information and overlook structured output constraints.

Existing unified IE methods with LLMs still suffer from several critical limitations. First, although retrieval-augmented generation (RAG) [16] has been introduced to enhance do-

This work was supported by Key Research and Development Program (Major Innovation Projects) in Shandong Province under Grant 2021CXGC010102 and Key Research and Development Program in Shandong Province (The Project for Enhancing the Innovation Capacity of Technology-based Small and Medium-sized Enterprises) under Grant 2025TSGCCZZB0124.

*Corresponding author: liusong@qlu.edu.cn

main knowledge acquisition, existing RAG-based generation strategies do not adequately account for IE task preferences, leading to semantic misalignment between generated content and extraction objectives. Secondly, reliance on fixed instructions and prompts often leads to feature confusion and error accumulation across tasks, degrading downstream extraction performance. Moreover, LLMs are prone to attention drift when processing long and complex prompts, further exacerbating the challenges of unified information extraction.

To address these challenges, we propose a unified information extraction framework RT-UIE based on reward guidance and task orientation. Unlike existing methods, RT-UIE employs a task-oriented RAG retrieval and reconstruction mechanism, reorganizing retrieved knowledge into enhanced task-specific contextual knowledge to improve task alignment. RT-UIE further introduces a chain-of-thought (CoT) reward model to select optimal reasoning paths, reducing feature confusion and error propagation. In addition, an attention calibration module is designed to model interactions between semantic fragments and generate an attention bias, enabling the model to focus more effectively on core instructions.

In summary, the main contributions of this paper are as follows:

- We propose a task-oriented RAG retrieval and reconstruction mechanism to overcome the limitations of traditional RAG methods in information extraction tasks. Traditional RAG methods only concatenate the retrieved text, failing to capture task-specific preferences. Our proposed mechanism reconstructs knowledge into task-aware background knowledge, thereby reducing the semantic mismatch between background knowledge and information extraction preferences.
- We propose a chain-of-thought reward model to enhance the interpretability and controllability of reasoning in unified information extraction. By leveraging multi-dimensional reward signals to optimize the structure and depth of the reasoning chain, the model effectively guides reasoning behaviors and enables adaptive task-specific optimization.
- To alleviate attention drift and the neglect of core task instructions in long-instruction scenarios, we build a novel attention calibration module. This module incorporates position-aware attention bias and dynamically adjusts it based on task characteristics, thus enhancing the focus of the model on critical instructional components.

II. RELATED WORK

Early information extraction models relied mainly on rule-based and pattern-driven methods. Carlson et al. [17], Kilicoglu et al. [18] proposed handcrafted lexical and syntactic rules for joint entity-relation and event extraction, respectively. Although these approaches achieved some progress on specific benchmarks, their heavy reliance on manual rule engineering resulted in limited scalability and poor generalization.

To alleviate these problems, traditional machine learning methods introduced supervised and distantly supervised

paradigms. Mintz et al. [19] proposed a feature-based supervised relation extraction model, while Quirk and Poon [20] leveraged distant supervision to construct training data from knowledge bases automatically. Despite improved recall and automation, these methods still relied heavily on handcrafted features, which constrained robustness and cross-domain transferability.

With the advancement of deep learning, deep neural networks (DNNs) enabled automatic representation learning for IE tasks. Lample et al. [9] introduced the BiLSTM-CRF architecture for sequence labeling, achieving certain performance on named entity recognition benchmarks. Zhou et al. [21] further applied convolutional neural networks to relation extraction, improving robustness by modeling contextual semantics. However, these models remain limited in their ability to capture long-range dependencies and complex reasoning patterns.

Recently, LLMs, such as LLaMA [22], ChatGPT [23] and GPT-4 [24], have demonstrated a strong generalization ability, allowing a paradigm shift for unified information extraction. Qi et al. [25] proposed ADELIE, which improves IE performance through instruction-aligned fine-tuning, but its unified instruction design is difficult to adapt to various subtasks. Zhan et al. [15] introduced RAMIE by incorporating retrieval-augmented generation, but directly concatenating retrieved examples often introduces irrelevant information and ignores structured output constraints.

To further optimize the unified IE task, Lu et al. [12] proposed the UIE framework with a unified text-to-structure formulation, while Lou et al. [13] introduced USM through lexical linking operations. Wang et al. proposed InstructUIE, which uses instruction tuning to improve zero-shot and cross-task generalization capabilities, alleviating the model's adaptability to different tasks. However, existing unified IE frameworks still suffer from three key limitations: insufficient task-aware utilization of external knowledge, fixed prompt templates lacking task adaptivity, and attention drift under long instruction contexts.

III. METHOD

A. Task Description

Given an input text sequence $X = \{x_1, x_2, \dots, x_n\}$, Unified Information Extraction (UIE) aims to extract structured semantic information from unstructured text, including named entity recognition (NER), relation extraction (RE), and event extraction (EE). For NER task, the model outputs an entity set $\mathcal{E} = \{(s_i, e_i, t_i)\}$, where (s_i, e_i) denotes the entity span and t_i its type. For RE task, given X and $\mathcal{E}$, the model predicts a relation set $\mathcal{R} = \{(e_i, r_{ij}, e_j)\}$, where r_{ij} denotes the relation between entity pair (e_i, e_j) . For EE task, the model outputs an event set $\mathcal{V} = \{(t_k, \mathcal{A}_k)\}$, where t_k is the event type and $\mathcal{A}_k = \{(r_m, a_m)\}$ denotes the argument set with role r_m and corresponding span a_m .

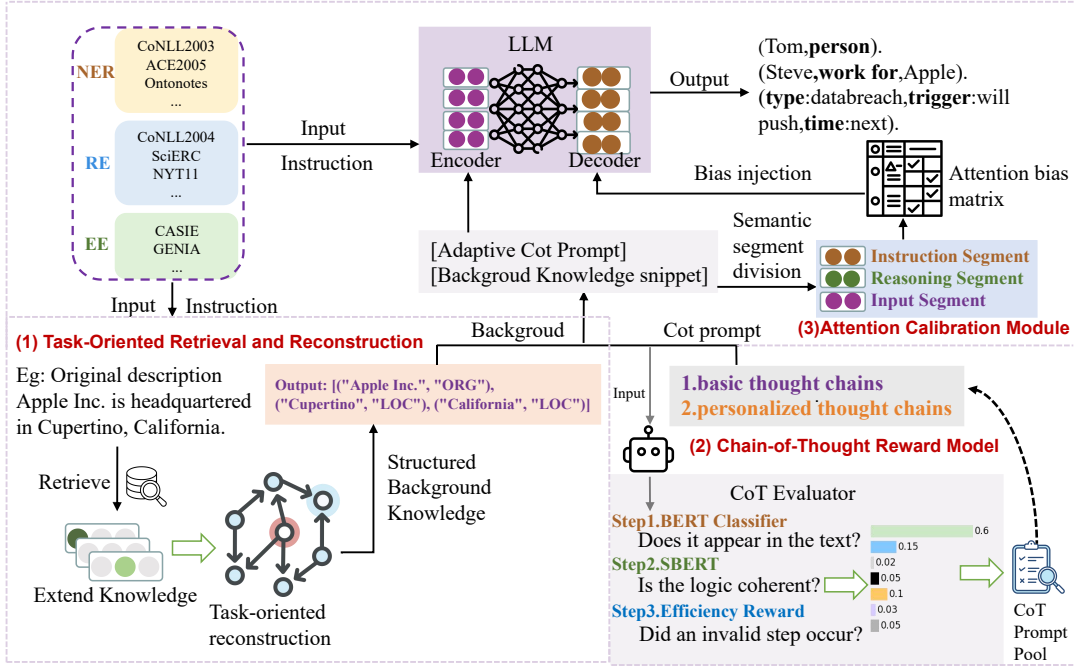


Fig. 1. Overview of the RT-UIE framework. RT-UIE takes the task description, input text, optimized prompts, and reconstructed background knowledge as inputs, generates an attention bias matrix for these inputs, and feeds both the inputs and the matrix into the LLM to produce the information extraction results.

B. Overview

Figure 1 illustrates our proposed LLM-based unified information extraction framework, which comprises three core components: a task-oriented retrieval and reconstruction mechanism, a chain-of-thought reward model, and an attention calibration module. First, the model receives raw text and task instructions as input and constructs the original prompt based on the specific task type, thus forming the initial input context. Next, the initial input context enters the RAG retrieval and reconstruction mechanism, which retrieves candidate knowledge fragments relevant to the current task from external knowledge sources and reconstructs the results based on the task objective, thereby obtaining task-related background knowledge. Then, the chain-of-thought reward model guides the model to optimize the inference path via reward signals, so as to generate chain-of-thought prompt that better meets task requirements. Then, the new input context including the optimized background knowledge and chain-of-thought prompt is fed into the LLM. In addition, Since excessively long inputs may cause attention shifts, we designed an attention calibration module that creates an attention bias matrix for new inputs. During decoding, this bias matrix is injected into the LLM to dynamically adjust the attention distribution of the input context, enabling the LLM to focus more on key information fragments.

C. Task-Oriented RAG Retrieval and Reconstruction

To overcome the problem of insufficient text utilization in information extraction scenarios of traditional RAG frameworks, we propose a task-oriented retrieval and reconstruction

mechanism consisting of two consecutive stages: knowledge retrieval and structured reconstruction. Given the input text X and the task instruction q , a task-aware query is first constructed to retrieve relevant external knowledge. Subsequently, the retrieved knowledge is transformed into structured task-specific background representations using our designed rule-based triples and structured templates. The reconstructed knowledge serves as auxiliary evidence and examples, providing explicit task guidance for downstream extraction models.

Task-Oriented Knowledge Retrieval. Following the Retrieval-Augmented Generation paradigm, we concatenate the task instruction and input text to form a unified retrieval query:

$$\tilde{X} = [q] \parallel [X] \quad (1)$$

where $\parallel$ denotes the concatenation operator. Both query $\tilde{X}$ and each knowledge entry $d_i \in \mathcal{D}$ are mapped into a shared semantic embedding space, resulting in the vector representations $\mathbf{h}_{x,q}$ and $\mathbf{h}_{d_i}$, respectively, where $\mathcal{D}$ denotes the task-oriented IE knowledge base constructed from publicly available information extraction benchmark datasets. We then compute the similarity scores between the query representation and all candidate knowledge entries and select the most relevant entry d for subsequent reconstruction.

$$d = \arg \max_{d_i \in \mathcal{D}} \text{Sim}(\mathbf{h}_{x,q}, \mathbf{h}_{d_i}) \quad (2)$$

Task-oriented knowledge reconstruction. Instead of directly operating on raw input text, we reconstruct the retrieved knowledge d , which provides condensed and semantically aligned evidence while reducing noise from unstructured inputs.

a) *Rule Triple Matching*: For each information extraction task $\tau \in \{\text{NER}, \text{RE}, \text{EE}\}$, we maintain a task-specific rule triple set:

$$\mathcal{R}_\tau = \{r_1^\tau, r_2^\tau, \dots, r_{M_\tau}^\tau\} \quad (3)$$

Each rule triple is defined as:

$$r_k^\tau = \langle P_k^\tau, C_k^\tau, A_k^\tau \rangle \quad (4)$$

where P_k^τ is a factual triple used to match candidate knowledge, C_k^τ enforces type and role consistency through constraints, and A_k^τ maps the matched elements into task-specific structured outputs.

The rule triples are defined according to the output schema of each task and constructed by structural patterns from annotated data, thereby bridging retrieved knowledge with the target structured representations. Given the retrieved knowledge d , the matched rule triples are selected as:

$$\mathcal{T}_\tau = \{r_k^\tau \in \mathcal{R}_\tau \mid d \models r_k^\tau\}, \quad (5)$$

where $\models$ indicates that the retrieved knowledge satisfies both pattern and semantic constraints defined in r_k^τ .

b) *Structured Knowledge Generation*: Each task τ is associated with a predefined structured reconstruction template $\mathcal{S}_\tau$. Based on the task format, the matched triples $\mathcal{T}_\tau$ are mapped via A_k^τ into structured representations, thereby forming structured background knowledge that serves as task-oriented guidance.

$$\mathcal{S}_{\text{NER}} = \{\text{task} = \text{"NER"}, \text{entities} = \{e_i\}_{i=1}^N\}, \quad (6)$$

$$\mathcal{S}_{\text{RE}} = \{\text{task} = \text{"RE"}, \text{relations} = \{n_i\}_{i=1}^M\}, \quad (7)$$

$$\mathcal{S}_{\text{EE}} = \{\text{task} = \text{"EE"}, \text{events} = \{v_i\}_{i=1}^L\}. \quad (8)$$

where $e_i = (\text{text}_i, \text{type}_i)$ denotes an entity and its type, $n_i = (\text{head}_i, \text{tail}_i, \text{rel}_i)$ represents a relation instance, and $v_i = (\text{trigger}_i, \text{args}_i, \text{type}_i)$ denotes an event with its trigger, arguments, and event type. The resulting structured knowledge $\mathcal{K}$ is used as task-oriented background input to guide downstream generation and extraction.

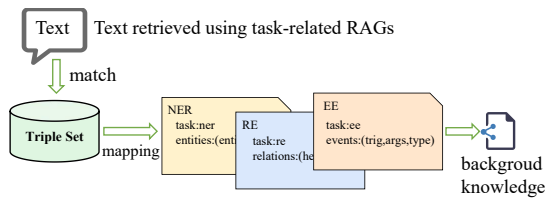


Fig. 2. A diagram illustrating task-oriented RAG reconstruction, where Text represents the relevant knowledge retrieved.

D. Chain-of-Thought Reward Model

Unlike existing unified information extraction methods that rely on implicit black-box reasoning, we propose a thought chain reward model (CoTRM) for explicit and task-aware evaluation of reasoning processes. CoTRM assesses candidate reasoning chains based on accuracy, consistency, and efficiency, and is trained on local input texts. It is then used

during LLM training to score candidate CoTs and select high-quality reasoning paths from a predefined CoT pool. The CoT pool is a collection of reusable reasoning chains that can be combined to provide diverse and structured candidate chains for reasoning.

Accuracy Scorer. Although existing CoT hints can describe intermediate inference steps, their quality varies considerably. Weak semantic alignment between the inference steps and the input can introduce noise and reduce the extraction accuracy. To address this issue, we design an accuracy scorer to explicitly evaluate whether a given CoT is semantically consistent with the input and supports correct extraction. The input to the scorer includes the input text, the CoT, and the corresponding gold labels, which serve as supervision signals. The combined input S is encoded by a BERT-based verifier, and the global representation of the [CLS] token is projected to an accuracy reward score:

$$R_{\text{acc}} = \sigma(W \cdot \text{BERT}(S)_{[\text{CLS}]} + b) \quad (9)$$

where, W and b are learnable parameters, and $\sigma(\cdot)$ denotes the sigmoid function mapping the output to $[0, 1]$. A higher R_{acc} indicates that the inference chain is more likely to support correct information extraction for the given input.

Consistency Scorer. Although CoT reasoning provides explicit intermediate steps, adjacent inference steps may suffer from semantic discontinuity and logical jumps, which reduce the reliability of the reasoning. To mitigate this incoherent transition, we designed a consistency scorer based on Sentence-BERT [26] to measure the semantic smoothness between CoTs. Specifically, given CoTs with N steps, each step is encoded in a semantic representation vector using Sentence-BERT. The semantic consistency between two adjacent steps is measured by cosine similarity:

$$s_i = \frac{\text{SBERT}(\mathbf{e}_i) \cdot \text{SBERT}(\mathbf{e}_{i+1})}{\|\text{SBERT}(\mathbf{e}_i)\| \|\text{SBERT}(\mathbf{e}_{i+1})\|}, \quad i = 1, 2, \dots, N-1. \quad (10)$$

The overall consistency reward is computed as the average similarity across all adjacent step pairs:

$$R_{\text{cons}} = \frac{1}{N-1} \sum_{i=1}^{N-1} s_i, \quad (11)$$

where a higher R_{cons} indicates stronger semantic coherence throughout the CoTs.

Efficiency Scorer. Different information extraction tasks exhibit varying reasoning complexity. For example, NER typically requires fewer reasoning steps than RE or EE tasks with complex entity interactions. To address this, we design an efficiency scorer to align the reasoning depth with task complexity. Given the input text X , a task complexity score $F \in [0, 1]$ is computed based on input length, entity density, and task structure. The ideal number of reasoning steps is then defined as:

$$N_{\text{ideal}} = 2 + \text{round}(6 \cdot F) \quad (12)$$

This formulation allocates a small number of CoTs to low-complexity tasks. In contrast, high-complexity tasks are assigned a larger number of CoTs. Given the actual number of reasoning steps N_{actual} , the efficiency reward is defined as:

$$R_{\text{eff}} = 1 - \frac{|N_{\text{actual}} - N_{\text{ideal}}|}{8} \quad (13)$$

where higher values indicate closer alignment between the actual and ideal depths of reasoning.

Reward Aggregation. Finally, the three reward signals are combined using a weighted sum to obtain the overall reward score:

$$R_{\text{total}} = W_1 \cdot R_{\text{acc}} + W_2 \cdot R_{\text{cons}} + W_3 \cdot R_{\text{eff}} \quad (14)$$

where R_{acc} , R_{cons} , and R_{eff} correspond to accuracy, consistency, and efficiency rewards, respectively. The aggregated reward signal R_{total} is fed back to the LLM to guide the optimization of the CoT selection strategy. We set $W_1 = 0.5$ to prioritize accuracy, $W_2 = 0.3$ to moderately encourage reasoning consistency, and $W_3 = 0.2$ to maintain efficiency. In addition to the general chain-of-thought, we design task-specific personalized CoTs for NER, RE, and EE. These CoTs are distinguished according to task type and domain, and our reward model evaluates them accordingly to select the most suitable CoT for each task.

E. Attention Calibration Module

When multiple modules operate together, the input instruction sequence can become very long and hierarchically structured, including task instructions, retrieved knowledge, CoT steps, and the input text. Such complexity may lead to attention drift and degrade extraction performance.

To address this issue, we propose an attention calibration module. It first performs semantic segmentation to partition the input into functional segments, and then constructs an attention bias matrix to model interactions among these segments. An adaptive bias adjustment strategy is further introduced to modulate the bias strength according to input complexity. Finally, the bias matrix is injected into the LLM decoder to guide attention toward task-relevant tokens, thereby improving focus during generation.

Semantic segmentation. We leverage the annotation information from the prompt generation process to segment the input sequence into three core semantic parts:

- Instruction Segment (Inst): Task definition, output format constraints, and key instruction information.
- Reasoning Segment (Reas): The chain-of-thought generated by the model during inference.
- Input Segment (Inp): The original text from which structured information is extracted.

Attention bias matrix. Based on the semantic segmentation described above, we construct an attention bias matrix for an input sequence of length L :

$$B \in \mathbb{R}^{L \times L} \quad (15)$$

where each element $B_{i,j}$ is used to adjust the attention weight from the i -th token to the j -th token. The bias generation follows the rules:

$$B_{i,j} = \begin{cases} W_{\text{inst} \rightarrow \text{all}} \cdot s, & i \in \text{Inst}, \\ W_{\text{inst} \rightarrow \text{inp}} \cdot s, & i \in \text{Inst} \wedge j \in \text{Inp}, \\ W_{\text{reas} \rightarrow \text{inp}} \cdot s, & i \in \text{Reas} \wedge j \in \text{Inp}, \\ 0, & \text{otherwise.} \end{cases} \quad (16)$$

The bias matrix $B_{i,j}$ is constructed to modulate the attention among different components of the instruction sequence. Specifically, $W_{\text{inst} \rightarrow \text{all}}$ amplifies the influence of instruction tokens on the entire sequence, ensuring that task instructions have a global guiding effect. $W_{\text{inst} \rightarrow \text{inp}}$ strengthens the alignment between the instruction tokens and the input text tokens, facilitating the model's focus on relevant input information. $W_{\text{reas} \rightarrow \text{inp}}$ constrains the reasoning tokens to properly reference the input without dominating the generation process. The scalar s controls the overall magnitude of these bias terms. For all other token pairs not covered by these cases, the bias is set to zero.

Adaptive bias adjustment strategy. To avoid applying a fixed bias strength across all samples, we further introduce an adaptive bias adjustment strategy, in which the coefficient s is dynamically adjusted along three dimensions:

- Task complexity awareness. Different bias strengths are assigned to different information extraction tasks: $s = 0.3$ for NER task, $s = 0.6$ for EE task and $s = 0.8$ for RE task.
- Sequence length gating. When the input sequence length falls below 256 tokens, the attention calibration module is automatically disabled.
- Reasoning-depth modulation. The bias strength is further modulated by the number of reasoning steps: $s = 0.3$ for shallow inference and $s = 0.8$ for multi-hop inference.

Injection of Bias Matrix into LLM. The attention bias matrix is then injected into the Transformer self-attention mechanism of the LLM as follows:

$$\text{Attention}(Q, K, V) = \text{softmax}\left(\frac{QK^\top}{\sqrt{d_k}} + B\right)V \quad (17)$$

where the bias term is added before the softmax operation, this formulation ensures that instruction-related tokens retain high attention weights even under long-range dependency conditions, thereby stabilizing task-oriented generation.

IV. EXPERIMENTS

In this section, we evaluate the proposed model RT-UIE in a supervised learning setting. We adopt FLAN-T5 (13B) [27] as the backbone model and extend it to implement the proposed reward-guided and task-oriented unified information extraction framework. All models are trained with a learning rate of 1×10^{-5} for 5 training epochs. Following standard practice, we use the F_1 score as the primary evaluation metric for all tasks and datasets.

A. Dataset

We conducted supervised experiments on 32 English datasets from the IE_INSTRUCTIONS benchmark datasets [14], which integrates a wide range of widely used datasets spanning three major information extraction tasks: NER task, RE task, and EE task. The benchmark datasets include representative NER datasets such as CoNLL2003 and ACE2005, RE datasets such as NYT and SemEval RE, and EE datasets such as CASIE and PHEE. Besides, to further validate the performance of RT-UIE in complex task environments, we additionally selected three representative complex information extraction datasets for evaluation: WNUT-17 [28], DocRED [29], and MAVEN [30]. These datasets have more complex semantic structures and contextual dependencies, which help to comprehensively evaluate the model’s capabilities in complex scenarios.

B. Baselines

We compare the proposed RT-UIE with the following baseline models:

- **BERT-base** [31] is a widely used pre-trained language model that serves as a strong supervised baseline for downstream tasks.
- **UIE** [12] is a unified text-to-structure generation framework that can universally model different IE tasks and adaptively generate targeted structures.
- **USM** [13] is a unified IE tasks framework, which converts IE tasks to a semantic matching problem.
- **InstructUIE** [14] introduces a unified information extraction framework based on multi-task instruction tuning, enabling a single model to handle heterogeneous extraction tasks.
- **YAYI-UIE** [32] presents a general instruction-tuning-based information extraction framework trained on large-scale conversational data.

C. Evaluation indicators

We use span-based offset **Micro- F_1** as the primary evaluation metric. For the NER task, we employ a span-level evaluation setting in which entity boundaries and types must be correctly predicted. For the RE task, a relation triple is considered correct if the model correctly predicts the boundaries of the subject and object entities, and the entity relation. For the EE task, an event trigger is considered correct if the event type and trigger word are correctly predicted.

D. Experimental Results and Analysis

Tables I, II, III, and IV summarize the performance of different models on named entity recognition, relation extraction, event extraction, and complex extraction tasks, respectively. Overall, RT-UIE consistently outperforms strong baselines such as UIE and InstructUIE across most benchmarks.

Named Entity Recognition: As shown in Table I, RT-UIE achieves an average F_1 score of 86.29% across 20 benchmark

TABLE I
OVERALL F_1 SCORE OF RT-UIE ON THE NAMED ENTITY RECOGNITION TASK.

Dataset	BERT-base	UIE	InstructUIE	Ours
ACE2005	87.30	85.78	86.66	88.42
AnatEM	85.82	77.68	90.89	91.45
bc2gm	80.90	74.77	85.16	83.64
bc4chemd	86.72	82.79	90.30	89.79
bc5cdr	85.28	78.82	89.59	90.89
broadtwitter	58.61	67.02	83.14	83.08
CoNLL03	92.40	92.99	92.94	94.62
FabNER	64.20	73.71	76.20	79.00
FindVehicle	87.13	91.56	89.47	92.40
GENIA-Ent	73.30	67.46	74.71	78.51
HarveyNER	82.26	58.13	88.79	88.56
MIT Movie	88.78	79.56	89.01	89.70
MIT Rest	81.02	81.67	82.55	84.87
multiNERD	91.25	91.75	92.32	92.60
ncbi-disease	80.20	80.13	90.23	89.89
Ontonotes	91.11	86.25	90.19	91.73
polyglot	75.65	68.01	70.15	70.03
tweetNER7	56.59	63.81	64.97	69.48
wikiann	70.60	82.11	85.13	85.39
wikineural	82.78	92.14	91.36	91.73
Avg	80.09	78.81	85.19	86.29

TABLE II
OVERALL F_1 SCORE OF RT-UIE ON RELATION EXTRACTION TASK.

Dataset	UIE	USM	YAYI-UIE	InstructUIE	Ours
ADE corpus	-	-	84.14	82.31	83.14
CoNLL2004	75.00	78.84	79.73	78.48	78.68
GIDS	-	-	72.36	81.98	82.18
kbp37	-	-	59.35	36.14	34.98
NYT	-	-	89.97	90.47	91.15
NYT11 HRL	-	-	57.53	56.06	57.75
SciERC	36.53	37.36	40.94	45.15	45.33
Semeval RE	-	-	61.02	73.23	73.68
Avg	-	-	68.13	67.98	68.36

datasets, outperforming InstructUIE by 1.10%. The best performance is achieved on the CoNLL2003 data set, with RT-UIE achieving a F_1 score of 94.62%. The largest performance gain is observed on the tweetNER7 dataset, with an absolute improvement of 4.51% over InstructUIE. Overall, RT-UIE outperforms InstructUIE on 14 out of 20 datasets and performs well on widely used benchmark datasets such as ACE2005 and CoNLL2003. These results demonstrate that, compared to InstructUIE, RT-UIE incorporates explicit CoT constraints and a task-aware knowledge reconstruction mechanism, which helps to align generated reasoning with task objectives and mitigate bias caused by fixed prompts.

Relation Extraction: Table II indicates that RT-UIE achieves an average F_1 score of 68.36% across eight relation extraction benchmark datasets, outperforming InstructUIE by 0.38%. On the New York Times dataset, RT-UIE achieves a F_1 score of 91.15%, the highest across all datasets. RT-UIE achieves consistent improvements over UIE and USM on benchmarks such as CoNLL2004 and SciERC, demonstrating stronger relational reasoning capability. Although RT-UIE performs slightly worse than YAYI-UIE on some datasets, this may be attributed to differences in pretraining scale and instruction alignment strategies between the two methods.

TABLE III
OVERALL F_1 SCORE OF RT-UIE ON EVENT EXTRACTION TASK.

Dataset	BERT-base	USM	UIE	YAYI-UIE	InstructUIE	Ours
ACE2005	72.5	72.41	72.41	65.00	77.13	82.21
CASIE	68.98	71.73	71.73	63.00	67.80	68.56
PHEE	-	-	-	63.00	70.14	68.36
Avg	-	-	-	63.67	71.69	73.04

TABLE IV
OVERALL F_1 SCORE OF RT-UIE ON COMPLEX TASK.

Dataset	UIE	InstructUIE	Ours
WNUT-17	59.26	59.64	60.24
DocRED	62.14	64.77	65.89
MAVEN	52.68	55.32	57.42
Avg	58.03	59.91	61.18

Nevertheless, the results demonstrate that the proposed reward-guided CoT alignment and attention calibration strategy effectively enhances relation extraction performance while maintaining robust generalization across diverse datasets.

Event Extraction: As shown in Table III, RT-UIE achieves the highest average F_1 score of 73.04% across three event extraction benchmark datasets, outperforming InstructUIE by 1.35%. On the ACE2005 dataset, RT-UIE attains a F_1 score of 82.21%, achieving the best result among all evaluated models. RT-UIE outperforms InstructUIE on both ACE2005 and CASIE datasets and also surpasses UIE across all evaluated benchmarks, demonstrating stronger capability in modeling complex event structures and domain-specific event patterns. These results indicate that the proposed reward-guided reasoning and task-oriented knowledge reconstruction mechanisms effectively constrain the inference process and improve the stability and efficiency of the event extraction performance.

Complex Task Extraction: As summarized in Table IV, RT-UIE consistently outperforms UIE and InstructUIE on complex instruction-based datasets, including WNUT-17, DocRED, and MAVEN. The performance gains indicate that RT-UIE is more effective in handling long-context reasoning and structurally complex extraction scenarios. This improvement mainly benefits from the task-oriented RAG retrieval and reconstruction mechanism and the attention calibration module, which jointly reduce semantic noise and alleviate attention drift under long instruction settings.

E. Ablation experiment

Table V shows the ablation results of RT-UIE on representative entity and relation extraction benchmarks, where **T** denotes the task-oriented RAG retrieval and reconstruction module, **R** denotes the chain-of-thought reward model, and **AC** denotes the attention calibration module.

Task-Oriented RAG Retrieval and Reconstruction. Adding the task-oriented RAG retrieval and reconstruction module (+T) enhances model performance. Compared to the baseline, the +T setting yields notable gains on both entity and relation extraction tasks, especially on NYT, where the F_1 score increases to 92.31%. This improvement demonstrates

that reconstructing retrieved knowledge into task-aligned and structured background information is more effective than directly concatenating raw retrieved text, as it reduces noise and improves task relevance.

Chain-of-Thought Reward Model. With the addition of the chain-of-thought reward module (+R), the model achieves the most significant improvement among the three modules, with F_1 scores on CoNLL2003 and NYT increasing from 92.94% to 93.90% and from 90.47% to 90.67%, respectively. This demonstrates that reward-guided reasoning effectively constrains the reasoning process and promotes the selection of higher-quality CoT, thereby better supporting downstream information extraction.

Attention Calibration Module. After adding the attention calibration module (+AC), the model outperforms the baseline on most datasets, especially on the ACE2005 dataset, where the F_1 score improves from 86.66% to 87.41%. This indicates that attention calibration helps the model maintain focus on core task instructions and input labels when handling complex contextual structures.

TABLE V
 F_1 SCORE OF ENTITY RELATION EXTRACTION DEFUSION EXPERIMENT ON PARTIAL DATASETS.

Settings	Entity		Relation	
	CoNLL03	ACE2005	NYT	SemEval RE
Baseline	92.94	86.66	90.47	73.23
+T	93.68	87.55	92.31	71.53
+R	93.90	87.16	90.67	73.47
+AC	93.45	87.41	87.62	72.68
+R, +T	93.75	87.62	91.20	70.25
+R, +AC	94.02	87.32	90.15	73.15
+T, +AC	93.98	86.98	90.53	73.43
RT-UIE	94.62	87.96	91.15	73.68

When combining multiple modules, we observe apparent complementary effects. The combination of COTRM and the task-oriented RAG retrieval and reconstruction module (+R,+T) yields substantial improvements over single-module variants. Although attention calibration module alone shows limited improvements, combining it with COTRM (+R,+AC) or the task-oriented RAG retrieval and reconstruction module (+T,+AC) yields consistent gains across all datasets. Importantly, the complete RT-UIE model achieves the best overall performance, reaching 94.62% on CoNLL2003, 87.96% on ACE2005, 91.15% on NYT, and 73.68% on SemEval RE. These results demonstrate that the three modules address distinct yet complementary challenges in unified information extraction and that their integration yields additive benefits.

V. CONCLUSION

In this paper, we propose RT-UIE, a reward-guided and task-oriented unified information extraction framework. A task-oriented RAG retrieval and reconstruction module is designed to ensure that the acquired background knowledge closely aligns with the information extraction objective. Furthermore, we design a chain-of-thought reward model to adaptively regulate the quality and depth of CoTs based on the characteristics of the task. Moreover, we construct an attention calibration module to enhance the model's focus on core task instructions and key input features. Our model was tested on the IE_INSTRUCTION benchmark dataset and achieved the best performance on 20 of its subsets, such as ACE2005, AnatEM, Bc5crd, CoNLL2003, WikiNeural, and WikiAnn. Furthermore, our model also demonstrated strong performance on the three benchmark datasets for complex information extraction tasks: WNUT-17, DocRED, and MAVEN. In future work, we plan to leverage richer external knowledge sources and more appropriate post-training to improve our model, and then further apply our framework to fields such as medical and financial information extraction.

REFERENCES

- [1] Z. Zhang, W. You, T. Wu, X. Wang, J. Li, and M. Zhang, "A survey of generative information extraction," in *Proc. 31st Int. Conf. on Computational Linguistics (COLING)*, 2025, pp. 4840–4870.
- [2] T. Young, D. Hazarika, S. Poria, and E. Cambria, "Recent trends in deep learning based natural language processing," *IEEE Computational Intelligence Magazine*, vol. 13, no. 3, pp. 55–75, 2018.
- [3] V. Karpukhin, B. Oguz, S. Min, P. Lewis, L. Wu, S. Edunov, D. Chen, and W. t. Yih, "Dense passage retrieval for open-domain question answering," in *Proceedings of the 2020 Conference on Empirical Methods in Natural Language Processing (EMNLP)*, 2020, pp. 6769–6781.
- [4] S. Ji, S. Pan, E. Cambria, P. Marttinen, and S. Y. Philip, "A survey on knowledge graphs: Representation, acquisition, and applications," *IEEE Trans. Neural Netw. Learn. Syst.*, vol. 33, no. 2, pp. 312–333, 2021.
- [5] S. Zhang, L. Yao, A. Sun, and Y. Tay, "Deep learning based recommender system: A survey and new perspectives," *ACM Comput. Surv.*, vol. 52, no. 1, pp. 1–38, 2019.
- [6] B. Jehangir, S. Radhakrishnan, and R. Agarwal, "A survey on named entity recognition — datasets, tools, and methodologies," *Natural Language Processing Journal*, vol. 3, p. 100017, 2023.
- [7] X. Zhao, Y. Deng, M. Yang, L. Wang, R. Zhang, H. Cheng, W. Lam, Y. Shen, and R. Xu, "A comprehensive survey on relation extraction: Recent advances and new frontiers," *ACM Computing Surveys*, vol. 56, no. 11, pp. 1–39, 2024.
- [8] E. Simon, H. Olsen, H. You, S. Touileb, L. Ovreliid, and E. Velldal, "Generative approaches to event extraction: Survey and outlook," in *Proc. Workshop on the Future of Event Detection (FuturED)*, 2024, pp. 73–86.
- [9] G. Lample, M. Ballesteros, S. Subramanian, K. Kawakami, and C. Dyer, "Neural architectures for named entity recognition," in *Proceedings of the 2016 Conference of the North American Chapter of the Association for Computational Linguistics (NAACL)*, 2016, pp. 260–270.
- [10] M. Miwa and M. Bansal, "End-to-end relation extraction using lstms on sequences and tree structures," *arXiv*, 2016.
- [11] T. B. Brown, B. Mann, N. Ryder, and et al., "Language models are few-shot learners," in *Advances in Neural Information Processing Systems (NeurIPS)*, vol. 33, 2020, pp. 1877–1901.
- [12] Y. Lu, Q. Li, Y. Zhang, and H. Chen, "Unified information extraction via task representation and structure decoupling," in *Proceedings of the 60th Annual Meeting of the Association for Computational Linguistics (ACL)*, 2022, pp. 1254–1268.
- [13] J. Lou, X. Zhang, X. Wang, and Y. Li, "Unified semantic matching for multi-task information extraction," in *Proceedings of the 2023 Conference on Empirical Methods in Natural Language Processing (EMNLP)*, 2023, pp. 10321–10335.
- [14] X. Wang, B. Zhao, Y. Zhang, and X. Li, "Instructuie: Multi-task instruction tuning for unified information extraction," in *Proceedings of the 2023 Conference on Empirical Methods in Natural Language Processing (EMNLP)*, 2023, pp. 1517–1532.
- [15] Z. Zhan, S. Zhou, M. Li, and R. Zhang, "Ramie: Retrieval-augmented multi-task information extraction with large language models on dietary supplements," *Journal of the American Medical Informatics Association*, vol. 32, no. 3, pp. 545–554, 2025.
- [16] P. Lewis, B. Oguz, R. Rinott, S. Riedel, and V. Karpukhin, "Retrieval-augmented generation for knowledge-intensive nlp tasks," in *Proc. NeurIPS*, vol. 33, 2020.
- [17] A. Carlson, J. Betteridge, E. R. H. Jr., and T. M. Mitchell, "Coupling semi-supervised learning of categories and relations," in *Proc. NAACL HLT Workshop on Semi-Supervised Learning for NLP*, 2009, pp. 1–9.
- [18] H. Kilicoglu and S. Bergler, "Syntactic dependency based heuristics for biological event extraction," in *Proc. BioNLP 2009 Workshop*, 2009, pp. 119–127.
- [19] M. Mintz, S. Bills, R. Snow, and D. Jurafsky, "Distant supervision for relation extraction without labeled data," in *Proceedings of ACL-IJCNLP 2009*, 2009, pp. 1003–1011.
- [20] C. Quirk and H. Poon, "Distant supervision for relation extraction beyond the sentence boundary," in *Proceedings of the 2017 Conference of the European Chapter of the Association for Computational Linguistics (EACL)*, 2017, pp. 1171–1182.
- [21] P. Zhou, W. Shi, J. Tian, Z. Qi, B. Li, H. Hao, and B. Xu, "Attention-based bidirectional long short-term memory networks for relation classification," in *Proc. of ACL*, 2016, pp. 207–212.
- [22] H. Touvron, T. Lavril, G. Izacard, X. Martinet, M. A. Lachaux, T. Lacroix, B. Rozière, N. Goyal, E. Hambro, F. Azhar, A. Rodriguez, A. Joulin, E. Grave, and G. Lample, "Llama: Open and efficient foundation language models," *arXiv preprint arXiv:2302.13971*, 2023.
- [23] OpenAI, "Chatgpt: Optimizing language models for dialogue," OpenAI, Tech. Rep., 2022.
- [24] —, "Gpt-4 technical report," OpenAI, Tech. Rep., 2023.
- [25] Y. Qi, H. Peng, X. Wang, B. Xu, L. Hou, and J. Li, "Adelie: Aligning large language models on information extraction," *arXiv preprint arXiv:2405.05008*, 2024.
- [26] N. Reimers and I. Gurevych, "Sentence-bert: Sentence embeddings using siamese bert-networks," in *Proc. EMNLP-IJCNLP*, 2019, pp. 3982–3992.
- [27] T. V. e. a. S. Longpre, L. Hou, "The flan collection: Designing data and methods for effective instruction tuning," in *Proc. ICML*, 2023, pp. 22631–22648.
- [28] L. Derczynski, K. Bontcheva, S. Pyysalo, and H. Uszkoreit, "Results of the wnnt2017 shared task on novel and emerging entity recognition," in *Proceedings of the 3rd Workshop on Noisy User-generated Text (WNUT)*, 2017, pp. 140–147.
- [29] L. Yao, X. Zhu, C. Zhang, Y. Zhang, and M. Sun, "Docred: A large-scale document-level relation extraction dataset," in *Proceedings of the 57th Annual Meeting of the Association for Computational Linguistics (ACL)*, 2019, pp. 764–777.
- [30] Z. Yang, Y. Zhang, and N. Xue, "Maven: A massive general domain event detection dataset," in *Proceedings of the 58th Annual Meeting of the Association for Computational Linguistics (ACL)*, 2020, pp. 3569–3583.
- [31] J. Devlin, M.-W. Chang, K. Lee, and K. Toutanova, "Bert: Pre-training of deep bidirectional transformers for language understanding," in *Proceedings of the 2019 Conference of the North American Chapter of the Association for Computational Linguistics (NAACL)*, 2019, pp. 4171–4186.
- [32] X. Xiao, Y. Wang, N. Xu, Y. Wang, H. Yang, M. Wang, Y. Luo, W. Mao, and D. Zeng, "Yayi-ue: A chat-enhanced instruction tuning framework for universal information extraction," *arXiv preprint arXiv:2312.15548*, 2023.