

A Reward-Guided Semantic Decoupling Framework for Zero-Shot Multi-Intent Detection

1st Chuanwang Xiong

*School of Artificial Intelligence and Computer Science
Jiangnan University
Wuxi, Jiangsu, China*

2nd Hongwei Ge*

*School of Artificial Intelligence and Computer Science
Jiangnan University
Wuxi, Jiangsu, China*

Abstract—Unsupervised disentanglement of high-dimensional semantic spaces remains a critical bottleneck in zero-shot multi-intent detection. While Large Language Models (LLMs) possess vast knowledge, they suffer from attentional dilution and semantic hallucinations when processing dense, intertwined intents. Existing approaches often rely on heuristic prompt engineering, treating models as opaque black boxes and lacking explicit inference verification. To transcend these limitations, this paper proposes Reward-Guided Semantic Decoupling (RGSD), a parameter-efficient inference framework designed to unlock the complex reasoning potential of open-source models without fine-tuning. Departing from standard end-to-end generation, RGSD restructures the task into a transparent "Locate-Decouple-Reason-Verify" cognitive cycle: (1) A lightweight reward model constructs a semantic saliency map to minimize information entropy and lock onto core features; (2) Structured subspace projection enforces semantic orthogonality, decomposing mixed contexts into independent units to physically block feature interference; (3) Parallel reasoning is conducted within isolated subspaces, followed by a closed-loop global consistency check to audit logical self-consistency. Experiments on MixATIS and MixSNIPS demonstrate that RGSD, utilizing only a 14B parameter model, achieves zero-shot performance comparable to GPT-4. This approach validates the "Explicit Decoupling + Guided Reasoning" paradigm, offering a generalizable theoretical solution for enhancing the interpretability and robustness of large models in complex structured prediction.

Index Terms—Intent Detection, Slot Filling, Multi-view learning, Spoken Language Understanding

I. INTRODUCTION

Intent detection serves as the cornerstone of Natural Language Understanding (NLU), tasked with mapping unstructured user utterances into precise semantic goals. However, real-world interactions rarely align with idealized single-intent scenarios; users frequently articulate multiple, deeply coupled intents within a single, brief utterance. This semantic entanglement creates a fundamental bottleneck: when systems fail to explicitly decouple these intertwined demands, the resulting high-entropy ambiguity causes traditional single-intent paradigms to collapse, inevitably leading to downstream execution failures. Consequently, achieving fine-grained, robust disentanglement in complex semantic scenarios has shifted from a mere optimization problem to a critical architectural imperative for advanced NLU.

Existing research [1] [2] [3] has primarily focused on optimizing neural architectures under fully supervised settings, employing techniques ranging from advanced encoders (BERT [4], RoBERTa [5]) to graph neural networks (AGIF [6], GLGIN [7]). While methods like joint modeling ([3] and [8]) and local perspective characterization, such as CLID [9] and Uni-MIS [10] have improved performance, they fundamentally remain data-driven discriminative fitting paradigms. These models are brittle when facing data scarcity or domain shifts, struggling to generalize without massive, stable annotated corpora. The structural rigidity of these fully supervised approaches limits their adaptability, highlighting the urgent need for frameworks that can function under resource-constrained conditions.

The advent of Large Language Models (LLMs) offers a potential breakthrough via zero-shot reasoning. However, naively deploying general-purpose LLMs for multi-intent detection reveals inherent structural weaknesses. Traditional end-to-end generation lacks an explicit computational trajectory, treating the mapping from input to label as an opaque "black box." In scenarios involving long texts or dense intents, models suffer from attentional dispersion, leading to the loss of fine-grained features. Furthermore, the unconstrained associative power of LLMs often induces semantic hallucinations, generating plausible but incorrect labels. Relying solely on heuristic prompt engineering fails to impose the necessary logical constraints required for rigorous structured prediction.

To overcome these boundaries, this study explores a mechanism-driven approach to elevate the performance of small-to-medium open-source models to match closed-source giants like GPT-4 [11]. We propose the Reward-Guided Semantic Decoupling (RGSD) framework. Unlike recent decomposition frameworks such as DSCP that rely purely on heuristic logical chains, RGSD restructures the inference process into a transparent, hierarchical cognitive cycle: 1). Semantic Anchor Localization: Instead of passive reading, we employ a reward model to maximize information gain, filtering background noise to construct a high-signal-to-noise ratio semantic saliency map. 2). Structured Subspace Projection: We introduce a dynamic orthogonality constraint, projecting the complex mixed context into independent subspaces. This physically isolates semantic units, blocking feature interference between intents. 3). Local Sub-Intent Reasoning: Inference is conducted

*Corresponding author.

within these purified subspaces, ensuring the model focuses on targeted recall without contextual distraction. 4).Global Consistency Verification: We implement a closed-loop control mechanism, where local predictions are mapped back to the global context for a bidirectional logical audit, effectively suppressing over-generation.

To ensure universality, we instantiate RGSD with two complementary strategies: Poswise (Prior-based Static Weighting), which injects linguistic rules to filter noise, and Promptive (Semantics-based Dynamic Adaptation), which elicits deep semantic knowledge for adaptable anchor positioning. Experimental results on MixATIS and MixSNIPS confirm that the RGSD framework, utilizing only a 14B parameter model, achieves zero-shot performance comparable to GPT-4. This work validates that shifting from "black-box generation" to "Explicit Decoupling + Guided Reasoning" offers a robust, interpretable theoretical solution for complex semantic understanding tasks.

II. RGSD

We formalize the proposed Reward-Guided Semantic Decoupling (RGSD) as a multi-stage probabilistic inference process. As illustrated in Figure 1, the RGSD framework comprises four hierarchical phases: (1) Semantic Anchor Localization, (2) Structured Subspace Projection, (3) Local Sub-Intent Reasoning, and (4) Global Consistency Verification.

A. Semantic Anchor Localization

In zero-shot multi-intent detection, input utterances often contain significant background noise, leading to attentional dispersion in large models. The objective of this phase is to identify a set of semantic anchors, $K = \{k_1, k_2, \dots, k_t\}$, that maximally represent the core intent within the input utterance $S = \{w_1, w_2, \dots, w_n\}$. We model this process as a reward-based saliency evaluation problem. We aim to learn a reward function $R(\cdot)$ that quantifies the information gain r_i of each token w_i relative to the final intent determination. Based on these scores, we apply a Top- k sampling strategy to filter for the highest semantic density:

$$r_i = R(w_i|S) \quad (1)$$

$$K = \text{Top-}k(\{w_i, r_i\}_{i=1}^n; \rho) \quad (2)$$

where n is the utterance length, and ρ represents the semantic retention rate, a hyperparameter controlling the granularity of information disentanglement. To accommodate varying computational constraints and prior knowledge availability, we design two complementary instantiation strategies for the reward function: Poswise (Prior-based Static Reward) and Promptive (Semantics-based Dynamic Reward).

1) *Poswise*: The Poswise strategy constructs a lightweight, domain-agnostic noise filter. It operates on the linguistic hypothesis that, absent domain supervision, syntactic structures (e.g., Part-of-Speech) provide reliable priors for semantic contribution—content words (verbs, nouns) typically carry

core intents, while function words serve merely as structural support.

We parameterize the reward function as a lightweight evaluation network, E_θ . Specifically, the architecture consists of a pre-trained BERT encoder followed by a fully connected linear layer that projects the token-level hidden states into scalar reward scores. To avoid introducing task-specific bias, this model is not fine-tuned on multi-intent datasets but is pre-trained via self-supervision on the English Wikipedia corpus. The training objective minimizes the Mean Squared Error (MSE) between the predicted score and a "prior importance distribution" constructed from linguistic rules:

$$\mathcal{L}_{MSE} = \frac{1}{n} \sum_{i=1}^n (E_\theta(w_i) - s_i)^2 \quad (3)$$

Here, s_i denotes the prior score mapped from the Part-of-Speech (POS) tag of word w_i :

$$s_i = \phi(\text{POS}(w_i)) \quad (4)$$

where ϕ is a mapping function assigning differential weights to POS tags. We assign higher weights to content words and suppress function words. By learning this task-agnostic "fine-grained POS-importance" mapping, Poswise effectively assigns higher rewards to critical content words even in unseen downstream datasets, thereby enhancing zero-shot performance.

2) *Promptive*: In contrast, the Promptive approach is entirely training-free, leveraging the intrinsic semantic modeling capabilities of LLMs without requiring explicit scoring functions or annotated data. To bias the extraction towards information-dense words, we inject POS-related priors into the instructional constraints. Here, the "reward" is not an explicit score but is implicit within the LLM's conditional probability distribution.

We reframe anchor localization as a conditional subset generation task. Given an utterance S and a constraint I_{pos} containing POS priors, the LLM is guided to return a subset K that maximizes the probability under the semantic condition:

$$K = \arg \max_{K' \subseteq S} P(K'|S, I_{pos}) \quad (5)$$

where I_{pos} encodes the selection criteria based on linguistic priors. In practice, the LLM directly generates the optimal subset K following this directive.

B. Structured Subspace Projection

Once the anchor set K is established, direct reasoning over the full text may still suffer from feature entanglement between multiple intents. To physically block this interference, we introduce a Structured Subspace Projection mechanism. Unlike rigid splitting based on punctuation, we view this as a semantic reconstruction process. The goal is to dynamically project the original utterance S into a set of semantically independent subspaces $C = \{c_1, c_2, \dots, c_m\}$ based on the anchor distribution. Each subspace c_j represents not merely a

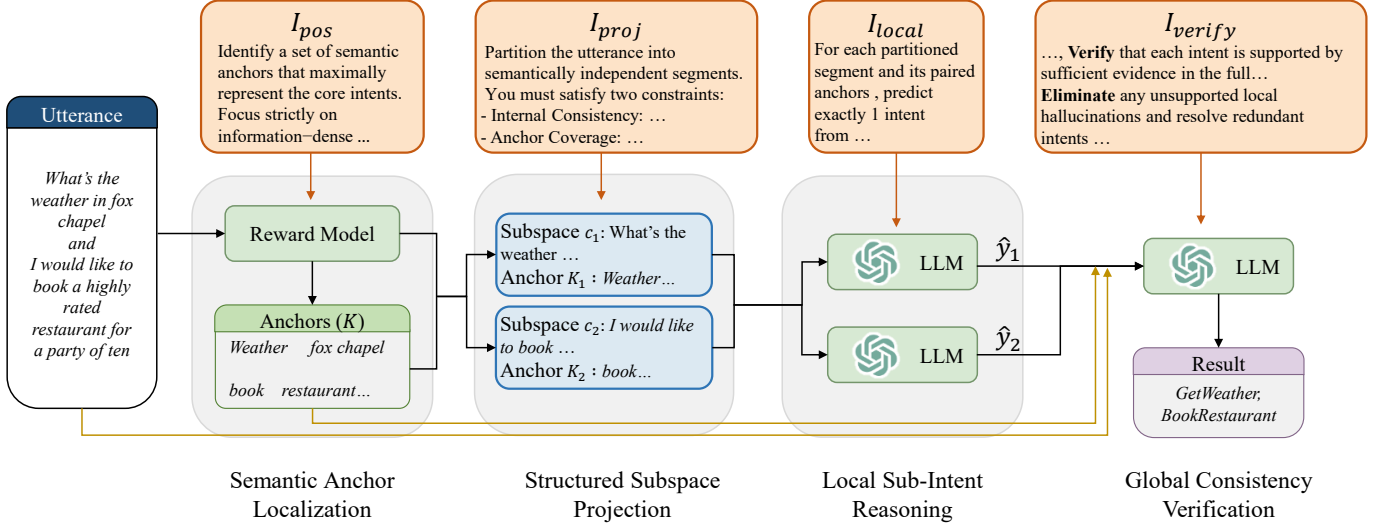


Fig. 1. The overview of the Reward-Guided Semantic Decoupling (RGSD) framework.

text fragment, but a semantic context window centered around a specific cluster of anchors.

We leverage the LLM’s context-awareness to model this projection as an anchor-driven structural prediction problem. Given S , anchors K , and instruction I_{proj} , the LLM seeks an optimal partition scheme C that satisfies two constraints: 1) Internal Consistency: c_j must contain exactly one primary intent and its modifiers; 2) Anchor Coverage: every identified anchor must be assigned to the subspace that best explains its semantic role. Formally, we maximize the following conditional probability:

$$C = \arg \max_{C'} P(C'|S, K, I_{proj}) \quad (6)$$

C. Local Sub-Intent Reasoning

The core task of this phase is to execute parallel local inference within each purified subspace, transforming sub-intent recognition into a series of low-noise, single-target decision processes. We define this not as discriminative classification, but as evidence-based contextual reasoning. Here, the subspace c_j provides the semantic context, while its paired anchor subset $K_j \subseteq K$ serves as the critical evidentiary support. This dual-input mechanism effectively resolves reference ambiguities common in fragmented expressions.

We formalize local sub-intent reasoning as a posterior probability maximization problem over a predefined intent space L . Given fragment c_j , its anchor subset K_j , and instruction I_{local} , the predicted intent $\hat{y}_j$ is:

$$\hat{y}_j = \arg \max_{y \in L} P(y|c_j, K_j, I_{local}) \quad (7)$$

Notably, this inference is parallelizable. This design not only reduces the computational complexity associated with single long-context inputs but, crucially, forces the model to focus

on local features, fundamentally preventing “hallucinatory correlations” common in multi-intent scenarios.

D. Global Consistency Verification

Following local sub-intent reasoning, we obtain a set of candidate sub-intents $Y_{cand} = \{\hat{y}_1, \dots, \hat{y}_m\}$. However, due to the independence assumption in subspace projection, a simple aggregation of local optima does not guarantee a global optimum and may contain logical conflicts. Therefore, we introduce a bidirectional verification mechanism to establish a logical closed loop from local details back to the global context.

Specifically, RGSD scrutinizes the full utterance S , global anchors K , and fragment predictions Y_{cand} . Under the verification constraint I_{verify} , the model first validates whether each candidate $\hat{y}_j$ is supported by sufficient evidence in the original context S (eliminating local hallucinations) and then resolves redundancies. The final intent set Y_{final} is derived by maximizing the joint probability, representing the most credible intent combination under a global view:

$$Y_{final} = \arg \max_{Y' \subseteq Y_{cand}} P(Y'|S, K, Y_{cand}, I_{verify}) \quad (8)$$

III. EXPERIMENTS

A. Experimental Setup

We optimize the supervised reward model using the Adam optimizer with hyperparameters set to $\alpha = 0.9$, $\beta = 0.999$. To ensure a rigorous and unbiased benchmarking protocol, consistent with DSCP [12], we evaluate our method on the standard multi-intent benchmarks: MixATIS and MixSNIPS. Performance is measured using three core metrics: Intent Accuracy, Intent Macro-averaged F1-score, and Intent Micro-averaged F1-score. All experiments are evaluated over 5 independent runs to account for generation randomness, reporting the mean $\pm$ standard deviation. The temperature coefficient

for LLM generation is sampled from the interval $[0, 1]$. To accelerate the inference pipeline, we employ vLLM [13] as the inference engine, with a maximum generation length constraint of 4096 tokens. All experiments were conducted on a single NVIDIA 4090D GPU.

B. Main Results

Table I presents a systematic comparison of the Qwen3-14B-powered RGSD against various prompt engineering paradigms based on GPT-3.5 and GPT-4 across the MixATIS and MixSNIPS datasets. The primary findings are as follows: (1) Substantial Superiority over GPT-3.5 Baselines: Despite being deployed on a significantly smaller model architecture (Qwen3-14B: ~ 14 B parameters vs. GPT-3.5: ~ 175 B parameters), RGSD consistently surpasses all GPT-3.5-based methods across key metrics. On the more challenging MixATIS dataset, RGSD achieves an Intent Accuracy of 53.15%, markedly outperforming the best GPT-3.5 method, Inter-DSCP (41.11%). On MixSNIPS, RGSD attains an Intent Accuracy of 93.12%, again significantly exceeding the GPT-3.5 baseline (Inter-DSCP: 74.9%). Furthermore, RGSD demonstrates comprehensive leadership in both Macro-F1 and Micro-F1 scores compared to all GPT-3.5 baselines.

(2) Competitive Performance against GPT-4: Remarkably, RGSD achieves parity with, and in several instances outperforms, GPT-4-based methods. On MixATIS, RGSD surpasses all GPT-4 baselines in Intent Accuracy (53.15%) and Micro-F1 (83.00%). On MixSNIPS, RGSD outperforms all GPT-4 prompting strategies—including Least-to-Most, VMID, Plan-and-Solve, DSCP, and Inter-DSCP—in both Intent Accuracy and Macro-F1, falling only slightly behind Inter-DSCP in Micro-F1.

Experiments based on Mistral-7B [14] (Table II) further corroborate the robustness of the RGSD framework. It consistently outperforms baselines across both datasets and all evaluation metrics. On MixATIS, RGSD achieves an Intent Accuracy of 19.2%, a 4.48% improvement over Inter-DSCP (14.72%); Macro-F1 and Micro-F1 reach 58.08% and 61.91%, reflecting gains of 11.12% and 15.31%, respectively. On MixSNIPS, RGSD attains an Intent Accuracy of 66.35%, a dramatic 25.1% increase over Inter-DSCP (41.25%). Its Macro-F1 and Micro-F1 scores reach 87.98% and 87.67%, whereas all baseline methods fail to exceed 71%.

These results provide compelling evidence for the efficacy of RGSD. By empowering medium-scale open-source models (e.g., Qwen3-14B) to not only significantly outperform GPT-3.5 but also rival GPT-4, RGSD demonstrates that explicit structural decoupling is a viable pathway to high-performance semantic understanding in resource-constrained environments.

IV. ANALYSIS AND DISCUSSION

To provide a comprehensive understanding of the RGSD framework, this section dissects the internal mechanisms driving its performance. We structure the discussion around six pivotal questions: (1) What is the theoretical necessity of Semantic Anchor Localization? (2) How does Structured

Subspace Projection fundamentally enhance decoding? (3) What role does Global Consistency Verification play in the control loop? (4) How is the optimal information sparsity determined? (5) How do different reward modeling paradigms compare?

A. Rationalizing the Necessity of Semantic Anchor Localization

Semantic Anchor Localization (Step 1) serves as the cornerstone of the RGSD framework. Its primary mission is to resolve the signal-to-noise ratio (SNR) bottleneck inherent in long-context processing. By proactively filtering non-informative functional tokens and redundant modifiers, this module compresses the raw, high-entropy natural language stream into a set of high-density feature indices, constructing a purified input space for subsequent deep reasoning.

The quantitative results from the ablation study (Table III) empirically validate the critical nature of this mechanism. When the Semantic Anchor Localization module is removed (w/o Step 1), the model suffers a catastrophic performance drop: on MixATIS, Intent Accuracy plummets by 11.24% (from 53.15% to 41.91%), while Macro-F1 and Micro-F1 shrink by 17.01% and 12.52%, respectively. A similar degradation is observed on MixSNIPS (10.94% drop in Accuracy). This severe regression exposes the fragility of the “direct reasoning” paradigm: without the guidance of anchors, the model is forced to conduct a blind search within a context saturated with high-frequency functional noise. This leads to attentional dispersion, where intent representations are diluted by irrelevant tokens, preventing key decision cues from emerging. Conversely, by introducing reward-driven anchor localization, RGSD establishes a focused and interpretable reasoning path. These explicitly extracted anchors serve not only as data compression but as high-confidence semantic landmarks, guiding the model to allocate computational resources to the core intent regions.

B. Mechanism of Performance Enhancement via Structured Subspace Projection

Structured Subspace Projection (Step 2) is the core disentanglement mechanism designed to counter semantic entanglement. By physically isolating the complex composite intent space into multiple independent single-intent subspaces, this module fundamentally shifts the inference paradigm.

The ablation results (w/o Step 2 in Table III) reveal the indispensability of this module. Removing Step 2 triggers the most disastrous collapse in performance: on MixATIS, Intent Accuracy precipitous falls from 53.15% to 22.71%. This absolute loss of nearly 30% constitutes the largest drop across all ablation settings. The synchronous decline in Macro-F1 and Micro-F1 confirms that this degradation is systemic.

This phenomenon underscores the limitations of “direct multi-intent decoding” in zero-shot settings. Without projection-based decoupling, the model is compelled to perform end-to-end multi-label classification within a raw, high-dimensional mixed semantic space, leading to attentional com-

TABLE I

PERFORMANCE COMPARISON OF METHODS BASED ON GPT-3.5, GPT-4, AND QWEN3 [15] ON THE MIXATIS AND MIXSNIPS DATASETS. ALL BASELINE RESULTS ARE REPORTED FROM DSCP [12]. **BOLD** NUMBERS INDICATE THE BEST OVERALL PERFORMANCE. * MARKS RESULTS THAT OUTPERFORM ALL GPT-3.5-BASED METHODS. ALL RESULTS ARE REPORTED AS MEAN $\pm$ STANDARD DEVIATION OVER 5 INDEPENDENT RUNS. IMPROVEMENTS OVER ALL BASELINES ARE STATISTICALLY SIGNIFICANT ($p < 0.05$ UNDER T-TEST).

Model	MixATIS			MixSNIPS		
	Intent Acc(%)	Macro F1(%)	Micro F1(%)	Intent Acc(%)	Macro F1(%)	Micro F1(%)
GPT-3.5						
Zero-CoT [16]	26.11	68.77	66.11	63.21	87.79	87.38
Least-to-Most [17]	25.69	68.99	65.45	66.48	89.90	89.47
VMID [18]	21.94	67.40	66.01	63.12	87.87	87.80
Plan-and-Solve [19]	27.08	65.32	62.59	69.12	90.35	90.04
DSCP [12]	29.17	73.75	79.05	72.90	91.14	90.79
Inter-DSCP [12]	41.11	75.04	75.41	74.90	91.50	91.30
GPT-4						
Zero-CoT [16]	40.97	78.62	77.40	87.77	95.46	95.47
Least-to-Most [17]	48.33	82.26	80.67	87.27	94.60	94.71
VMID [18]	32.64	76.74	74.02	88.45	92.81	92.90
Plan-and-Solve [19]	36.67	79.44	76.68	84.45	92.81	92.90
DSCP [12]	50.69	84.57	81.97	89.68	95.69	95.80
Inter-DSCP [12]	52.22	83.99	81.86	92.00	96.59	96.66
RGSD (14B)	53.15* ± 0.02	82.69* ± 0.03	83.00* ± 0.02	93.12* ± 0.01	96.62* ± 0.01	95.71* ± 0.01

TABLE II

MAIN PERFORMANCE FOR MISTRAL-7B ON MIXATIS AND MIXSNIPS DATASETS. THE OPTIMAL RESULTS ARE SHOWN IN **BOLD**.

Model	MixATIS			MixSNIPS		
	Intent Acc(%)	Macro F1(%)	Micro F1(%)	Intent Acc(%)	Macro F1(%)	Micro F1(%)
Zero-CoT [16]	8.19	43.59	42.40	32.98	69.61	69.79
Least-to-Most [17]	6.67	32.44	32.97	29.00	61.91	62.65
VMID [18]	5.97	39.46	37.88	32.87	66.35	67.48
Plan-and-Solve [19]	7.93	34.25	34.78	28.99	61.78	62.27
DSCP [12]	11.21	42.20	41.81	36.65	69.44	71.28
Inter-DSCP [12]	14.72	46.96	46.60	41.25	69.62	70.71
RGSD	19.20	58.08	61.91	66.35	87.98	87.67

petition and feature cross-contamination. By introducing structured subspace projection, RGSD implements a “Divide and Conquer” strategy: 1) It reduces the complex N -to- N mapping problem into a set of parallel 1-to-1 simple reasoning tasks, eliminating label interference; 2) Each projected subspace and its anchors constitute a low-entropy, high-SNR local sub-intent reasoning unit, significantly lowering the cognitive load on the LLM.

C. The Role of Global Consistency Verification

The Global Consistency Verification phase (Step 4) acts as both a “Quality Assurance Auditor” and a “Hallucination Filter” within the RGSD framework. Ablation results clearly demonstrate the criticality of this phase. Removing it (w/o Step 4 in Table III) leads to significant degradation on both datasets: Intent Accuracy on MixATIS drops from 53.15% to 40.34%, and Micro-F1 slides from 83.00% to 70.28%.

This finding highlights the potential risk that “local optimality does not imply global optimality.” While the preceding

steps effectively reduce reasoning difficulty, they introduce a risk of “local myopia”—where the model, operating within an isolated subspace, might over-interpret specific features to generate high-confidence pseudo-correlations or redundant predictions. Without a global audit, these “hallucinated intents,” which lack sufficient evidentiary support in the full original context, would contaminate the final output. Step 4 ensures logical self-consistency by mapping local decisions back to the global context.

D. Sensitivity Analysis of Semantic Sparsity

The semantic retention rate, ρ , is the core hyperparameter controlling the granularity of information disentanglement in RGSD; it directly determines the semantic sparsity of the input space. To investigate the non-linear impact of this parameter on zero-shot performance, we conducted a sensitivity analysis. Figure 2 reveals a distinct inverted U-shaped relationship between ρ and model performance, reflecting a delicate trade-off between “information completeness” and “noise suppression”:

TABLE III

ABLATION EXPERIMENTS FOR QWEN3-14B ON THE MIXATIS AND MIXSNIPS DATASETS. STEP 1, STEP 2, AND STEP 4 DENOTE SEMANTIC ANCHOR LOCALIZATION, STRUCTURED SUBSPACE PROJECTION, AND GLOBAL CONSISTENCY VERIFICATION, RESPECTIVELY.

Model	Intent Acc(%)	Macro F1(%)	Micro F1(%)
<i>MixATIS</i>			
w/o Step 1	41.91	65.68	70.48
w/o Step 2	22.71	62.45	65.61
w/o Step 4	40.34	65.56	70.28
RGSD	53.15	82.69	83.00
<i>MixSNIPS</i>			
w/o Step 1	82.18	93.53	93.62
w/o Step 2	84.40	92.29	92.45
w/o Step 4	85.36	90.66	91.08
RGSD	93.12	96.62	95.71

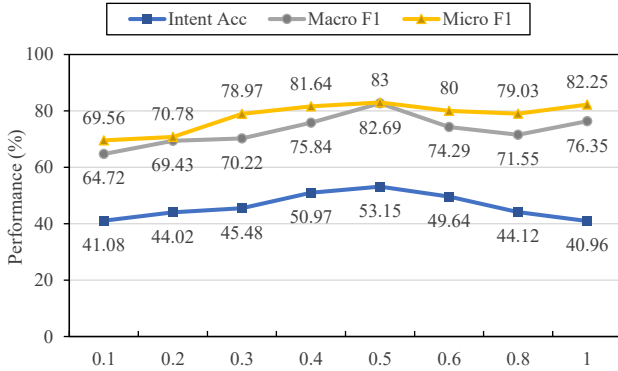


Fig. 2. Performance on MixATIS with Qwen3-14B under different values of ρ .

- 1) **Information Loss at Low Ratios ($\rho < 0.3$):** When retaining only the top 10% of tokens, performance degrades significantly (Intent Acc: 41.08%). This collapse is attributed to over-compression, where secondary contextual cues supporting intent judgment are deleted, leading to fragmented representations. This causes high miss rates for intents relying on long-range dependencies.
- 2) **Optimal Signal-to-Noise Balance ($\rho \in [0.3, 0.5]$):** Performance peaks at $\rho = 0.5$ (Intent Acc 53.15%). In this range, the reward model prioritizes core intent words while retaining sufficient auxiliary context, achieving an ideal balance between information coverage and attentional focus.
- 3) **Noise Re-introduction at High Ratios ($\rho > 0.6$):** As ρ exceeds 0.5, the re-inclusion of low-relevance words dilutes the core signal, causing Intent Acc to regress. Notably, at $\rho = 1.0$ (full text), while Micro-F1 recovers slightly, Intent Acc does not improve. This divergence indicates that while full context aids common intents, the amplified noise severely hampers the resolution of complex, long-tail multi-intent combinations, validating the necessity of semantic decoupling.

E. Comparative Analysis of Reward Modeling Paradigms

To evaluate the efficacy of different reward mechanisms, we systematically compared the Poswise (Prior-based) and Promptive (Semantics-based) strategies on MixATIS. Figure 3 illustrates the results.

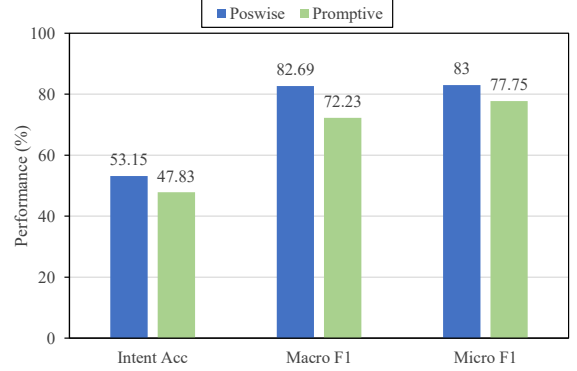


Fig. 3. Performance difference under Reward Methods on MixATIS.

Overall, Poswise outperforms Promptive across all metrics, achieving an Intent Acc of 53.15% compared to 47.83%—a gain of 5.32%. The superiority of Poswise stems from its explicit utilization of linguistic structural priors. By hard-coding syntactic rules (e.g., up-weighting nouns/verbs), Poswise constructs a deterministic, high-frequency semantic filter. This is inherently advantageous for intent detection, where intents are often triggered by specific verb-object structures.

In contrast, while Promptive leverages the semantic power of LLMs, it exposes limitations in fine-grained token-level scoring. LLMs tend to focus on global coherence rather than individual token contribution (calibration error) and are sensitive to prompt phrasing (instability). The results strongly suggest that for zero-shot tasks sensitive to local cues, injecting explicit inductive biases (Poswise) is more robust than relying on the implicit “black box” reasoning of LLMs (Promptive).

V. CONCLUSION

Addressing the challenges of attentional dilution and semantic hallucination in zero-shot multi-intent detection, this paper proposes the Reward-Guided Semantic Decoupling (RGSD) mechanism. RGSD reconstructs the inference process into a transparent cognitive cycle: Semantic Anchor Localization, Structured Subspace Projection, Local Sub-Intent Reasoning, and Global Consistency Verification. Experimental results on MixATIS and MixSNIPS demonstrate that RGSD, powered by a 14B open-source model, achieves performance comparable to, or exceeding, GPT-4 based methods. Ablation and sensitivity analyses further confirm the critical role of explicit decoupling and bidirectional verification in enhancing robustness and suppressing divergence. RGSD provides a highly interpretable and resource-efficient theoretical solution for high-quality semantic parsing under constrained conditions. Future work will explore the generalization of this mechanism to a broader range of complex structured understanding tasks.

ACKNOWLEDGMENT

This work is supported by the National Natural Science Foundation of China under Grants 52374155, Anhui Provincial Natural Science Foundation under Grant No. 2308085MF218, and PAPD of Jiangsu Higher Education Institutions.

REFERENCES

- [1] T. Pham, C. Tran, and D. Q. Nguyen, “MISCA: A joint model for multiple intent detection and slot filling with intent-slot co-attention,” in *Findings of the Association for Computational Linguistics: EMNLP 2023*. Singapore: Association for Computational Linguistics, Dec. 2023, pp. 12 641–12 650.
- [2] Z. Zhu, W. Xu, X. Cheng, T. Song, and Y. Zou, “A dynamic graph interactive framework with label-semantic injection for spoken language understanding,” in *ICASSP 2023-2023 IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP)*. IEEE, 2023, pp. 1–5.
- [3] N. A. Tu, H. T. T. Uyen, T. M. Phuong, and N. X. Bach, “Joint multiple intent detection and slot filling with supervised contrastive learning and self-distillation,” in *ECAI 2023*. IOS Press, 2023, pp. 2370–2377.
- [4] 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: human language technologies, volume 1 (long and short papers)*, 2019, pp. 4171–4186.
- [5] Y. Liu, M. Ott, N. Goyal, J. Du, M. Joshi, D. Chen, O. Levy, M. Lewis, L. Zettlemoyer, and V. Stoyanov, “Roberta: A robustly optimized bert pretraining approach,” *ArXiv*, vol. abs/1907.11692, 2019.
- [6] L. Qin, X. Xu, W. Che, and T. Liu, “AGIF: An adaptive graph-interactive framework for joint multiple intent detection and slot filling,” in *Findings of the Association for Computational Linguistics: EMNLP 2020*, Nov. 2020, pp. 1807–1816.
- [7] L. Qin, F. Wei, T. Xie, X. Xu, W. Che, and T. Liu, “GL-GIN: Fast and accurate non-autoregressive model for joint multiple intent detection and slot filling,” in *Proceedings of the 59th Annual Meeting of the Association for Computational Linguistics and the 11th International Joint Conference on Natural Language Processing (Volume 1: Long Papers)*, Aug. 2021, pp. 178–188.
- [8] B. Liu and I. Lane, “Attention-based recurrent neural network models for joint intent detection and slot filling,” in *Interspeech 2016*, 2016, pp. 685–689.
- [9] H. Huang, P. Huang, Z. Zhu, J. Li, and P. Lin, “Clid: A chunk-level intent detection framework for multiple intent spoken language understanding,” *IEEE Signal Processing Letters*, vol. 29, pp. 2123–2127, 2022.
- [10] S. Yin, P. Huang, and Y. Xu, “Uni-mis: United multiple intent spoken language understanding via multi-view intent-slot interaction,” in *Proceedings of the AAAI Conference on Artificial Intelligence*, vol. 38, 2024, pp. 19 395–19 403.
- [11] J. Achiam, S. Adler, S. Agarwal, L. Ahmad, I. Akkaya, F. L. Aleman, D. Almeida, J. Altschmidt, S. Altman, S. Anadkat *et al.*, “Gpt-4 technical report,” *arXiv preprint arXiv:2303.08774*, 2023.
- [12] L. Qin, Q. Chen, J. Zhou, J. Wang, H. Fei, W. Che, and M. Li, “Divide-solve-combine: An interpretable and accurate prompting framework for zero-shot multi-intent detection,” in *Proceedings of the AAAI Conference on Artificial Intelligence*, vol. 39, 2025, pp. 25 038–25 046.
- [13] W. Kwon, Z. Li, S. Zhuang, Y. Sheng, L. Zheng, C. H. Yu, J. Gonzalez, H. Zhang, and I. Stoica, “Efficient memory management for large language model serving with pagedattention,” in *Proceedings of the 29th symposium on operating systems principles*, 2023, pp. 611–626.
- [14] A. Q. Jiang, A. Sablayrolles, A. Mensch, C. Bamford, D. S. Chaplot, D. de Las Casas, F. Bressand, G. Lengyel, G. Lample, L. Saulnier, L. R. Lavaud, M.-A. Lachaux, P. Stock, T. L. Scao, T. Lavril, T. Wang, T. Lacroix, and W. E. Sayed, “Mistral 7b,” *ArXiv*, vol. abs/2310.06825, 2023.
- [15] A. Yang, A. Li, B. Yang, B. Zhang, B. Hui, B. Zheng, B. Yu, C. Gao, C. Huang, C. Lv, C. Zheng, D. Liu, F. Zhou, F. Huang, F. Hu, H. Ge, H. Wei, H. Lin, J. Tang, J. Yang, J. Tu, J. Zhang, J. Yang, J. Yang, J. Zhou, J. Zhou, J. Lin, K. Dang, K. Bao, K. Yang, L. Yu, L.-C. Deng, M. Li, M. Xue, M. Li, P. Zhang, P. Wang, Q. Zhu, R. Men, R. Gao, S.-Q. Liu, S. Luo, T. Li, T. Tang, W. Yin, X. Ren, X. Wang, X. Zhang, X. Ren, Y. Fan, Y. Su, Y.-C. Zhang, Y. Zhang, Y. Wan, Y. Liu, Z. Wang, Z. Cui, Z. Zhang, Z. Zhou, and Z. Qiu, “Qwen3 technical report,” *ArXiv*, vol. abs/2505.09388, 2025.
- [16] T. Kojima, S. S. Gu, M. Reid, Y. Matsuo, and Y. Iwasawa, “Large language models are zero-shot reasoners,” *Advances in neural information processing systems*, vol. 35, pp. 22 199–22 213, 2022.
- [17] D. Zhou, N. Schärli, L. Hou, J. Wei, N. Scales, X. Wang, D. Schuurmans, C. Cui, O. Bousquet, Q. Le *et al.*, “Least-to-most prompting enables complex reasoning in large language models,” *arXiv preprint arXiv:2205.10625*, 2022.
- [18] W. Pan, Q. Chen, X. Xu, W. Che, and L. Qin, “A preliminary evaluation of chatgpt for zero-shot dialogue understanding,” *arXiv preprint arXiv:2304.04256*, 2023.
- [19] L. Wang, W. Xu, Y. Lan, Z. Hu, Y. Lan, R. K.-W. Lee, and E.-P. Lim, “Plan-and-solve prompting: Improving zero-shot chain-of-thought reasoning by large language models,” in *Proceedings of the 61st Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, Jul. 2023, pp. 2609–2634.