

Fane 在 SemEval-2025 任务 10: 零样本实体框架与大型语言模型

Enfa Fane* Mihai Surdeanu* Eduardo Blanco* Steven R. Corman

*University of Arizona Arizona State University

{ enfageorge, msurdeanu, eduardoblanco } @arizona.edu
{ steve.corman } @asu.edu

Abstract

理解新闻叙事如何构建实体框架对于研究媒体对社会事件认知的影响至关重要。本文中，我们评估了大型语言模型 (LLMs) 在零样本情况下的分类角色框架能力。通过系统实验，我们评估了输入上下文、提示策略和任务分解的影响。我们的研究结果表明，首先识别广泛角色然后进行细粒度角色分类的层次方法优于单步分类。我们还表明，不同任务级别的最佳输入上下文和提示是不同的，这突出了需要针对子任务的特定策略。我们达到了 89.4 % 的主要角色准确率和 34.5 % 的精确匹配率，展示了我们方法的有效性。我们的研究结果强调了为提高 LLM 在实体框架中的表现而量身定制提示设计和优化输入上下文的重要性。

1 介绍

新闻的生产和消费规模空前庞大，然而新闻业中对中立性和客观性的期望 (Schudson, 2001) 常常与现实产生对比，即媒体选择性地决定什么才具有新闻价值 (Eliaz and Spiegler, 2024)。通过包含、遗漏或强调特定细节，记者塑造公众对事件的解释 (Iyengar, 1990)，推广特定建议 (Entman, 1993)，进而影响公众舆论和政策决策 (Ziems and Yang, 2021)。

当这种选择性的框架用于塑造对某一实体的认知时，称为实体构框 (van den Berg et al., 2020)。SemEval-2025 任务 10 引入了一个用于分析实体构框的基准 (Piskorski et al., 2025; Stefanovitch et al., 2025)，研究从词汇选择到叙述结构的语言选择如何影响认知。

图 1 中的示例展示了如何通过策略性的语言选择将比尔·盖茨描绘为一个对立者。尽管提到了他在气候变化项目上的大量投资，但叙述着重于他的高碳足迹，突出了虚伪。诸如“精心策划的采访”以及关于他与爱泼斯坦关系的引用，进一步加强了欺骗和逃避责任的形象。这种选择性的强调塑造了读者对盖茨作为骗子和腐败人物的认知，展示了框架如何微妙地影响解读。

Gates claimed that because he continues to spend billions of dollars on climate change activism, his carbon footprint isn't an issue [...] Elsewhere during the carefully constructed interview, Gates said he was surprised that he was targeted by 'conspiracy theorists' for pushing vaccines during the pandemic [...] While the BBC interview was set up to look like Gates was being challenged or grilled, he wasn't asked about his close friendship with the elite pedophile Jeffrey Epstein

Figure 1: 来自数据集的一个示例文档摘录。作者策略性地使用偏激的语言、对比框架和选择性强调来塑造读者对比尔·盖茨腐败和欺骗的看法。通过强调他在公众宣传中的矛盾之处，对其媒体露面的真实性提出质疑，并引用有争议的关联，该文本强化了负面框架。

近年来，大型语言模型 (LLMs) 在文本分类任务中显示出了优越的性能，常常在复杂任务上超越传统的机器学习方法 (Kostina et al., 2025)。在本文中，我们探讨这种优势是否也适用于实体框架分析。仅依赖提示，我们试图了解 LLMs 能否准确识别新闻文章中分配给实体的框架角色。鉴于 LLM 预测对提示变化的敏感性 (Zhuo et al., 2024)，我们评估了提示工程策略如何影响分类性能，例如基于角色的提示，结合标签定义等任务相关信息，以及理由生成。

为此，我们开发了一种完全依赖提示的模块化零样本方法。我们的方法将任务分解为两个阶段，首先预测广义的叙述角色（主角、反派、无辜者），然后细化为精细角色。我们系统地变化输入上下文，比较完整文章、以实体为中心的摘录，以及保持框架的摘要，并设计提示模板，结合专家角色和任务相关信息。

本文的主要贡献是：

- (1) 对大型语言模型在零样本实体框架中的系统评估：我们评估大型语言模型在无需微调的情况下如何从隐含的叙述线索中推断框架角色。
- (2) 多步骤提示策略：我们表明，将任务分解为两个阶段，首先预测广泛的角色，然后是细粒度的角色，比一步法效果更好。

- (3) 战略提示设计：我们证明了有效的提示设计，结合基于角色的指导和任务定义，可以改进分类。精心设计的提示可以让一个较小的模型表现得与一个较大的模型相当。

我们报告了在 SemEval 2025 任务 10 子任务 1 中的英语主角准确率为 89.4 %，精确匹配率为 34.5 %，在参与的团队中排名第六。我们将我们的代码库和相关提示公开发布。¹

2 背景：任务 & 数据集

SemEval 2025 任务 10：来自网络新闻的多语言叙事元素的表征与抽取，引入了三个子任务，用于分析五种语言的新闻文章中的叙事元素：保加利亚语、英语、印地语、欧洲葡萄牙语和俄语。本文聚焦于子任务 1：英语中的实体框架。

实体框架的目标是确定一个命名实体在新闻文章中的框架方式。框架在两个层面上表示：主要角色和细粒度角色。主要角色广泛地将实体分类为主人公、反派或无辜者，而细粒度角色则提供更详细的标签。² 一个实体可能会被框架为多个细粒度标签，这使得这是一个多类、多标签的分类问题。

有关数据集的更多详细信息，包括领域覆盖范围、语料库统计数据和注释方法，请参见 (Piskorski et al., 2025; Stefanovitch et al., 2025)。我们现在描述我们利用大型语言模型 (LLM) 进行实体框架分类的方法。

3 相关工作

语言通过选择性地强调现实的某些方面来塑造解读 (Entman, 1993)。此前的研究识别出了两种影响框架效应的关键语言机制：命名惯例，其中尊重或正式的称谓与积极情绪有很强的关联 (van den Berg et al., 2019, 2020)；以及叙事框架，反映了媒体对事件的描绘中党派差异 (Ziems and Yang, 2021)。然而，这些研究主要将框架视为与其他任务相关的隐含现象，而不是明确地分类或预测特定的框架角色或标签。

除了文本方面，在多模态领域，Sharma et al. (2022) 引入了 HVVMemes，一个涵盖 COVID-19 和美国政治的表情包数据集，并用粗粒度的框架角色进行标注：英雄、恶棍、受害者或无。相比之下，我们的工作专注于新闻文本，并对实体进行更细粒度的层次化角色分类。

大型语言模型 (LLM) 最近在文本分类任务中表现出色，通常在复杂数据集上超过传统的

```
{ expert_phrasing }
You will be provided with { input_format_phrase } .
Your task is to { task_definition } .
{ task_instructions }
{ output_format_with_example }
{ input_context }
```

Figure 2: 提示模板的结构。花括号中的内容会根据实验设置动态替换（具体措辞见附录 D）。该模板旨在最大限度地减少变异性，确保观察到的差异仅源于针对性提示的修改。

机器学习方法 (Kostina et al., 2025)。然而，仍然存在显著挑战，特别是关于对提示措辞和输入结构的敏感性 (Zhuo et al., 2024)。结构化提示策略，如角色条件 (Kong et al., 2024) 和指令重构 (Mishra et al., 2022)，已被提出。然而，这些策略的有效性仍然高度依赖于任务，并且通常需要根据具体问题设置进行仔细调整 (Atreja et al., 2024)。

在这些见解的指导下，我们开发并系统地测试了多种分层实体框架的提示策略，旨在改进粗粒度和细粒度的实体框架分类。

为了系统地评估不同提示策略对框架分类的影响，我们采用了一种结构化的、基于模板的方法 (图 2)。大型语言模型的决策过程通常是不透明的。通过使用固定模板，我们确保输出的任何变化都来自我们控制的提示修改，而非随机性。该模板由适应于不同实验条件的模块化组件组成。

我们的方法围绕三个关键组成部分进行构建：(1) 输入语境，其确定为分类提供的文本语境，具体在 input_context 中指定。(2) 提示设计，研究提示策略，如角色提示，在 task_definition、task_instructions 和 output_format_with_example 中结构化。(3) 推理策略，定义分类是作为单一任务进行，还是分解为两个步骤，从而反映框架的层次性质。以下章节将详细讨论每个组成部分。

3.1 输入上下文变化

框架涉及有选择地强调某些细节，同时省略其他细节，以突出所感知现实的特定方面 (Entman, 1993)。因此，优化上下文对于最大化叙述信号至关重要。为了评估上下文粒度对框架分类的影响，我们定义了五种输入设置，包括提取和基于总结的方法，³ 每种设置在分类可用的上下文信息量上有所不同。

- 全文 (FT)：整篇文章。
- 实体-句子 (Ent-Sent)：仅提及实体的句子。

¹<https://github.com/beingenfa/semeval2025task10>

²详细标签列在附录 A 中

³有关总结的详细信息，请参见附录 C。

- 实体-句子邻居 (Ent-Neigh): 涉及实体的句子加上前面一句和后面一句。
- 中立总结 (Neutral-Sum): 通过一个大型语言模型生成的总结, 提示其中立地关注实体的参与、行动和框架。
- 框架保留摘要 (FP-Sum): 由大型语言模型生成的摘要, 旨在保留文章的原始框架, 强调正面或负面的行为。

3.2 提示设计

在已经提出的各种提示策略中 (Gu et al., 2023), 我们专注于以下内容:

- (1) 角色/人格提示: 在 LLM 提示中分配角色或人格可能会影响其性能。虽然有研究发现基于角色的提示可以提高任务性能 (Kong et al., 2024), 另一项研究则表明其效果高度依赖于任务, 并且在不同环境中有所变化 (Zheng et al., 2024)。为了评估其对实体框架的影响, 我们将中性提示与明确指定专家角色的人格提示进行比较。
- (2) 与任务相关的信息: 我们测试了两种配置: 一种仅提供任务定义, 另一种同时包括任务和标签定义。
- (3) 输出理由: 我们比较带有和不带有模型生成解释的预测, 以评估要求理由是否提高或降低分类准确性。

3.3 推理策略

我们比较两种推理策略, 这两种策略决定分类是作为单一任务进行还是分解成两个阶段, 反映了框架的层次结构。单步预测是在单个推理步骤中联合预测主要角色和细粒度角色, 而多步预测则首先推断主要角色, 然后将其纳入提示以预测细粒度角色。

在确立了我们的提示策略后, 我们现在展示用于系统评估其有效性的实验设置。

4 实验设置

我们首先讨论研究中的一个关键假设, 然后是评估指标和 LLM 设置。

虽然 SemEval 任务将实体框架定义为一个多标签分类问题, 但我们观察到实体提及通常只获得一个细粒度角色, 在英语训练集划分中平均每个提及获得 1.08 个细粒度角色, 这一趋势在不同语言和主要角色中一致 (附录 B)。此外, 在我们的实验中, 当允许模型预测仅一个主要角色时, 我们的模型仍然始终预测两个主要角色。鉴于这些模式, 我们采用单标签近似, 为每个提及分配一个细粒度角色。

我们使用两个指标来评估性能: 主角色准确率 (MRA), 即预测的主角色与金标准标签匹

Metric	Baseline	+ EP	+ LD	+ RA
Main Role Accuracy				
FT	0.93	0.92	0.89	0.91
Ent-Sent	0.90	0.93	0.92	0.91
Ent-Neigh	0.92	0.93	0.92	0.93
Neutral-Sum	0.69	0.73	0.69	0.71
FP-Sum	0.95	0.95	0.93	0.93
Exact Match Ratio				
FT	0.29	0.35	0.30	0.29
Ent-Sent	0.32	0.33	0.35	0.31
Ent-Neigh	0.30	0.34	0.32	0.31
Neutral-Sum	0.22	0.21	0.21	0.24
FP-Sum	0.33	0.32	0.34	0.33

Table 1: 在同时提示主要角色和细化角色时, 不同输入上下文和提示设计策略的性能比较。EP = 专家角色, LD = 标签定义, RA = 理由。对于主要角色, 保留框架的总结是最强的输入上下文, 实现了最高的准确度 (0.95)。对于细化角色, 全文以及仅包含实体的句子是最强的, 当与提示策略一起使用时。

配的实例比例, 以及精确匹配率 (EMR), 即主角色和细化角色都完全匹配的实例比例, 没有部分匹配的计分。我们通过 OpenAI API 使用 GPT-4o (gpt-4o-2024-08-06)。

接下来, 我们研究在两种推理设置中使用不同提示策略和输入上下文所获得的结果。

5 结果

我们展示了在开发集上的提示实验结果, 接下来是官方的 SemEval 结果和赛后分析。关键发现将在第 6 节中讨论。

5.1 发展

我们系统地评估了在单步骤和多步骤环境中跨不同输入上下文的单个提示修改的影响, 使用附录 D 中详述的基线提示。为了确保效应的精确归因, 我们单独分析提示修改, 而不是评估组合策略。尽管在第 5.2 节中包含了一个组合设置, 但提示策略组合的系统评估超出了本文的范围。

单步法 表 1 显示了在单步设置中不同输入上下文和提示工程策略的性能。当使用保留框架的摘要时, 主角色的准确性最高 (0.95), 额外的提示策略没有进一步的提升。相比之下, 细粒度的角色分类从提示策略中获益, 其中全篇文本 + 专家角色和仅实体句子 + 标签定义实现了最高的精确匹配率 (0.35)。

这些研究结果表明, 单一输入策略可能无法最佳地支持这两个任务, 这促使我们进一步探讨多步骤方法, 其中主要角色和细粒度角色分别预测。此外, 由于一贯较低的性能, 中性色总结设置被排除在后续实验之外。

Metric	Baseline	+ EP	+ LD	+ RA
FT	0.89	0.91	0.91	0.87
Ent-Sent	0.84	0.86	0.87	0.82
Ent-Neigh	0.91	0.92	0.91	0.88
FP-Sum	0.92	0.93	0.96	0.93

Table 2: 不同输入上下文和提示工程策略在多步骤主角色预测中的性能比较，报告单位为准确率。EP = 专家角色，LD = 标签定义，RA = 理由。保持框架的摘要仍然是最强的输入上下文，达到了最高的准确率 (0.96)，现在从标签定义 (LD) 策略中获益，不同于单步设置。

Metric	Baseline	+ EP	+ LD	+ RA
FT	0.33	0.36	0.31	0.35
Ent-Sent	0.36	0.44	0.35	0.37
Ent-Neigh	0.36	0.38	0.34	0.37
FP-Sum	0.29	0.29	0.35	0.31

Table 3: 在细粒度角色分类的多步骤设置中，使用精确匹配率 (Exact Match Ratio) 指标对不同输入上下文和提示工程策略的性能进行比较。EP = 专家角色，LD = 标签定义，RA = 理由。与单步预测相比，多步骤方法在大多数设置中提高了 EMR，其中实体-句子 + 专家角色 (EP) 的 EMR 最高 (0.44)。

多步法 在多步设置中，模型首先预测主角色，然后用其来指导细粒度分类。

主角色预测：表 2 报告了独立预测时的主角色准确率。框架保持的摘要仍然是最强的设置，达到了 0.96 的准确率。

细粒度角色：从性能最好的设置中获取的主角色预测用于指导细粒度角色分类的提示。此外，我们仅将可能的输出标签集限制为对预测的主角色有效的细粒度角色。这减少了歧义并提高了分类准确性。

如表 3 所示，多步方法在大多数情况下优于单步预测。表现最佳的策略是实体句子 + 专家角色 (EP)，其精确匹配率 (EMR) 达到 0.44，相较于最佳单步结果 (0.35) 有显著提升。实体句子仍然是细粒度角色的最强输入上下文 (0.44 EMR)，而完整文本则略逊一筹。

这些研究结果进一步支持了将主角色预测与细粒度角色预测分开的有效性，并通过限制有效的角色标签来提高分类准确性。

5.2 官方 SemEval 结果

我们开发实验中的发现突出了不同输入背景和提示策略的有效性。然而，在我们正式的 SemEval 提交中，我们使用了“全文 + 专家角色 + 多步设置”的组合，并使用了 O1 (o1-2024-12-17) 作为我们的模型，因为该配置是基于赛前实验选择的。

如表 4 所示，我们的系统整体排名第六，达

Rank	Team	EMR (Δ)	MRA (Δ)
1	DUTIR	0.41	0.95
2	PAteam	0.38 (-0.03)	0.89 (-0.06)
3	DEMON	0.37 (-0.04)	0.92 (-0.03)
4	gowithnlp	0.37 (-0.04)	0.94 (-0.01)
5	TartanTritons	0.36 (-0.05)	0.72 (-0.23)
6	Ours	0.34 (-0.07)	0.89 (-0.06)
27	Baseline	0.04 (-0.37)	0.29 (-0.66)

Table 4: 官方 SemEval 测试集结果。团队按准确匹配率 (EMR) 排名，同时也报告了主角色准确率 (MRA)。我们的系统排名第六，达到了 0.34 的 EMR，仅比排名第一的系统 (DUTIR) 低 0.07。 Δ 值表示与排名第一的系统的差异。

Approach	EMR	MRA	Model	Price/1M Tokens	
				Input	Output
SemEval	0.345	0.894	O1	\$ 15.00	\$ 60.00
Improved	0.349	0.894	GPT-4o	\$ 2.50	\$ 10.00

Table 5: 我们的官方 SemEval 提交与改进方法的比较。尽管性能几乎相同，但改进的方法值得注意，因为它使用了 GPT-4o，这是一种显著更小且更具成本效益的模型来实现这些结果。这凸显了优化提示设计的重要性，它使得较小的模型能够匹配或超越更大且更昂贵的替代品的性能。

到了 0.34 的准确匹配率 (EMR)，比表现最好的系统 (DUTIR) 落后 0.07。虽然具有竞争力，但我们的后 SemEval 消融研究揭示了一种更有效的策略，我们将在下一节讨论。

5.3 后 SemEval

在提交至 SemEval 之后，我们进行了结构化的消融研究来优化我们的方法，如在第 5.1 节中讨论的那样。如表 5 所示，新方法实现了 0.349 的 EMR，相较于我们正式提交时的 0.345，MRA 则保持不变为 0.894。这个结果具有重要意义，因为它是使用 GPT-4o 这个较小且成本显著降低的模型实现的。

6 洞察

我们展示了实验中的关键发现，研究了不同方法的影响。

如表 1 所示，以事实性和公正方式呈现实体的中性摘要，其表现始终不如保留框架的摘要。这与现有的框架理论一致，进一步证明框架不仅涉及事实选择，还涉及通过选择性强调来塑造解释 (Entman, 1993)。

解释说明不能提高分类性能 我们的实验表明，要求模型为其预测提供理由并不会提高分类准确性。性能主要由提示中的信息决定。

更多信息并不总是更好 更长的上下文并不总是有用。主要角色分类受益于建立总体叙述的广泛上下文，而细致的角色需要专注于实体特定的细节。过多输入，例如全文上下文，可能会稀释关键框架信号，而浓缩且框架保留的摘要可以提高准确性。这突出了根据具体分类任务调整上下文粒度的重要性。

我们发现，一个结构化的多步骤分类方法能显著提高性能。通过首先预测主要角色，然后基于该预测进行细粒度分类，模型在每个阶段都能从更清晰的上下文中获益。这个方法减少了歧义，并确保每个分类步骤针对其特定的框架级别进行了优化。

最后，我们的实验强调了精心设计的提示的有效性。精心制作的提示可以让较小的模型与较大的模型相媲美。对于 GPT-4o，经过良好优化的提示可以达到或超过在更大、更昂贵的 o1 模型上未调优提示的效果。这些结果表明，在扩展到更大架构之前，改进较小模型的提示设计可以在效率和准确性上带来显著提升。

在本文中，我们探讨了大语言模型 (LLMs) 在零样本环境中进行实体框架分类的有效性。尽管 LLMs 在广泛角色分类上表现良好，但细粒度分类仍然具有挑战性。我们采用的多步骤方法在每个阶段结合了不同的输入上下文和提示策略，显著提高了整体性能。

7

致谢 这项研究得到了美国海军研究办公室 (N00014-22-1-2596) 的资助。

References

- Shubham Atreja, Joshua Ashkinaze, Lingyao Li, Julia Mendelsohn, and Libby Hemphill. 2024. Prompt design matters for computational social science tasks but in unpredictable ways. *arXiv preprint arXiv:2406.11980*.
- Kfir Eliaz and Ran Spiegler. 2024. [News media as suppliers of narratives \(and information\)](#). Preprint, *arXiv:2403.09155*.
- Robert M Entman. 1993. Framing: Toward clarification of a fractured paradigm. *Journal of communication*, 43(4):51–58.
- Jindong Gu, Zhen Han, Shuo Chen, Ahmad Beirami, Bailan He, Gengyuan Zhang, Ruotong Liao, Yao Qin, Volker Tresp, and Philip Torr. 2023. A systematic survey of prompt engineering on vision-language foundation models. *arXiv preprint arXiv:2307.12980*.
- Shanto Iyengar. 1990. Framing responsibility for political issues: The case of poverty. *Political behavior*, 12:19–40.
- Aobo Kong, Shiwan Zhao, Hao Chen, Qicheng Li, Yong Qin, Ruiqi Sun, Xin Zhou, Enzhi Wang, and Xiaohang Dong. 2024. [Better zero-shot reasoning with role-play prompting](#). In Proceedings of the 2024 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies (Volume 1: Long Papers), pages 4099–4113, Mexico City, Mexico. Association for Computational Linguistics.
- Arina Kostina, Marios D. Dikaiakos, Dimosthenis Stefanidis, and George Pallis. 2025. [Large language models for text classification: Case study and comprehensive review](#). Preprint, *arXiv:2501.08457*.
- Swaroop Mishra, Daniel Khashabi, Chitta Baral, Yejin Choi, and Hannaneh Hajishirzi. 2022. [Reframing instructional prompts to GPTk’s language](#). In Findings of the Association for Computational Linguistics: ACL 2022, pages 589–612, Dublin, Ireland. Association for Computational Linguistics.
- Jakub Piskorski, Tarek Mahmoud, Nikolaos Nikolaidis, Ricardo Campos, Alípio Jorge, Dimitar Dimitrov, Purificação Silvano, Roman Yangarber, Shivam Sharma, Tanmoy Chakraborty, Nuno Guimarães, Elisa Sartori, Nicolas Stefanovitch, Zhuohan Xie, Preslav Nakov, and Giovanni Da San Martino. 2025. SemEval-2025 task 10: Multilingual characterization and extraction of narratives from online news. In Proceedings of the 19th International Workshop on Semantic Evaluation, SemEval 2025, Vienna, Austria.
- Michael Schudson. 2001. [The objectivity norm in american journalism](#). *Journalism*, 2(2):149–170.
- Shivam Sharma, Tharun Suresh, Atharva Kulkarni, Himanshi Mathur, Preslav Nakov, Md. Shad Akhtar, and Tanmoy Chakraborty. 2022. [Findings of the CONSTRAINT 2022 shared task on detecting the hero, the villain, and the victim in memes](#). In Proceedings of the Workshop on Combating Online Hostile Posts in Regional Languages during Emergency Situations, pages 1–11, Dublin, Ireland. Association for Computational Linguistics.
- Nicolas Stefanovitch, Tarek Mahmoud, Nikolaos Nikolaidis, Jorge Alípio, Ricardo Campos, Dimitar Dimitrov, Purificação Silvano, Shivam Sharma, Roman Yangarber, Nuno Guimarães, Elisa Sartori, Ana Filipa Pacheco, Cecília Ortiz, Cláudia Couto, Glória Reis de Oliveira, Ari Gonçalves, Ivan Koychev, Ivo Moravski, Nicolo Faggiani, and 19 others. 2025. Multilingual Characterization and Extraction of Narratives from Online News: Annotation Guidelines. Technical Report JRC141322, European Commission Joint Research Centre, Ispra (Italy).
- Esther van den Berg, Katharina Korfhage, Josef Ruppenhofer, Michael Wiegand, and Katja Markert. 2019. [Not my president: How names and titles frame political figures](#). In Proceedings of the Third Workshop on Natural Language Processing and Computational Social Science, pages 1–6, Minneapolis, Minnesota. Association for Computational Linguistics.
- Esther van den Berg, Katharina Korfhage, Josef Ruppenhofer, Michael Wiegand, and Katja Markert. 2020. [Doctor who? framing through names and titles in German](#). In Proceedings of the Twelfth Language Resources and Evaluation Conference, pages 4924–4932, Marseille, France. European Language Resources Association.
- Mingqian Zheng, Jiaxin Pei, Lajanugen Logeswaran, Moontae Lee, and David Jurgens. 2024. [When “a helpful assistant” is not really helpful: Personas in system prompts do not improve performances of large language models](#). In Findings of the Association for Computational Linguistics: EMNLP 2024, pages 15126–15154, Miami, Florida, USA. Association for Computational Linguistics.
- Jingming Zhuo, Songyang Zhang, Xinyu Fang, Haodong Duan, Dahua Lin, and Kai Chen. 2024. [ProSA: Assessing and understanding the prompt sensitivity of LLMs](#). In Findings of the Association for Computational Linguistics: EMNLP 2024, pages 1950–1976, Miami, Florida, USA. Association for Computational Linguistics.
- Caleb Ziems and Diyi Yang. 2021. [To protect and to serve? analyzing entity-centric framing of police violence](#). In Findings of the Association for Computational Linguistics: EMNLP 2021, pages 957–976, Punta Cana, Dominican Republic. Association for Computational Linguistics.

A 实体框架标签

实体框架分类法由三个主要角色和 22 个细化角色组成。有关详细的定义和注释指南，请参见 (Piskorski et al., 2025; Stefanovitch et al., 2025)。

- 主角：守护者，殉道者，和事佬，反叛者，失败者，美德者。
- 反派：煽动者，阴谋者，暴君，外国敌人，叛徒，间谍，破坏者，腐败，无能，恐怖分子，欺骗者，偏执者。
- 无辜者：被遗忘、被利用、受害者、替罪羊。

B 跨语言的平均子角色

	EN	BG	PT	HI	RU
All	1.08	1.13	1.05	1.16	1.06
Protagonist	1.06	1.02	1.01	1.24	1.00
Antagonist	1.10	1.19	1.09	1.12	1.11
Innocent	1.02	1.01	1.00	1.06	1.01

Table 6: 在训练数据中的不同语言（英语 (EN)、保加利亚语 (BG)、欧洲葡萄牙语 (PT)、印地语 (HI) 和俄语 (RU)）和主要角色中，每个实例的子角色平均数量。数值接近于 1，证实了我们在细粒度分类中使用单标签近似的方法。

C 摘要生成提示

C.1 中性总结

Summarize the following article with a specific focus on { entity } . Write the summary as a standalone description, ensuring that the entity and its role are clearly introduced without referring to 'the article' or assuming prior context. Clearly state their involvement, actions, and framing within the event. Maintain a factual and neutral tone.

Figure 3: 用于生成中性摘要的提示。说明指导 LLM 确保实体的角色被明确介绍，同时保持事实和公正的语气。

C.2 框架保留摘要

Write a standalone summary that clearly reflects how the author of the article frames entity and their actions. Do not present the entity neutrally—mirror the language and implicit bias of the article itself. If the author portrays the entity favorably, highlight their positive actions, successes, and beneficial impact. If the author is critical, emphasize the entity's negative actions, failures, or harmful consequences.

If the framing is mixed or subtle, encode the contrast in tone and nuance. Strongly reflect the author's framing through:

- Loaded or emotionally charged language (if present in the article)
- Emphasis on the entity's perceived intentions and motivations
- Who is affected by their actions and how the consequences are framed
- Any direct or implied judgments made by the author.

Maintain the style and tone of the article, ensuring the framing is explicit in how the entity's role and impact are described. Do not add external information or neutralize the bias. The summary should feel as though it was written by the original author.

Table 7: 用于生成保持框架的摘要的提示。说明指导模型反映文章的原始框架，强调偏见和语气，同时避免中立或外部信息。

D 提示模板详情

我们的系统提示遵循之前介绍的结构化格式。输入上下文由用户提示提供，而其余部分则仍然是系统提示的一部分。

Prompt Type	Template
System Prompt	{ expert_phrasing } You will be provided with { input_format_phrase } . Your task is to { task_definition } . { task_instructions } { output_format_with_example }
User Prompt	{ input_context }

Table 8: 系统和用户提示的结构。系统提示提供指令，而用户提示为模型提供上下文输入。

我们在下面概述了这些元素在不同环境中采取的值。

D.1 输入格式短语 & 输入上下文

输入上下文的格式为 DOCUMENT: { content } 。 请注意，在 Entity-Sentences 和 Entity-Sentences Neighbors 设置中，单个句子或句子组是由 [...] 分隔的。

在系统提示中，输入格式短语指定了提供给模型的输入上下文类型。格式遵循以下结构：a { document_type_str } in the following format- DOCUMENT: { document_type_str } 。 每种输入上下文类型对应的 document_type_str 在下表中提供：

Input Context Type	document_type_str
Full Text	news article
Entity-Sentences or Entity-Sentences Neighbors	excerpt of a news article
Neutral Summary or Framing-Preserved Summary	news article summary

Table 9: 输入上下文类型与相应输入短语之间的映射。

D.2 提示设计

D.2.1 专家角色

在角色/人格提示的上下文中，如果没有指定的角色，人格 { expert_phrasing } 保持为空。否则，它将被以下文本替换：

```
You are an expert in analyzing how a specific named entity is portrayed in a given text. Read the text carefully and focus on everything said about { entity } .
```

D.2.2 输出原理

在下面的模板中，我们通过包含三个变量来引入提示进行论证：

- { ask_reasoning } : "with a reasoning for your prediction"
- { reasoning_in_json } : ', "reasoning" : "your reasoning here"
- { reasoning_example } : , "reasoning" : "The article frames Kremlin propagandists as instigators and deceivers, highlighting their role in spreading falsehoods and promoting extreme measures."

当设置不需要用输出标签进行验证时，这些变量保持为空。

D.2.3 任务相关信息

Setting	task_definition	task_instructions	output_format_with_example
Single Step	classify the narrative framing of the { entity } in the document based on the taxonomy that follows in the format [(broad role,[list of valid fine grained roles/])]. The taxonomy is { taxonomy } .	Instructions: 1. Assign exactly one broad role from: Protagonist, Antagonist, or Innocent. 2. Determine one or a maximum of two corresponding fine grained role from the taxonomy. 3. Order the sub-roles by likelihood, with the most likely fine-grained role listed first. { label_definitions }	4. Finally, return your conclusion { ask_reasoning } as a single JSON object with no extra text, in this format: { "main_role": "<most_likely_main_role>", "fine_grained_roles": ["<Most likely sub-role>", "<Second most likely sub-role if relevant>"] { reasoning_in_json } } Example Output: { "main_role": "Antagonist", "fine_grained_roles": ["Conspirator"] { reasoning_example }
Multi-Step Main Role	: classify the narrative framing of the entity in the document as either Protagonist, Antagonist or Innocent.	Instructions: 1. Assign exactly one main role from: Protagonist, Antagonist, or Innocent. { label_definitions }	2.Return your conclusion { ask_reasoning } as a single JSON object with no extra text, in this JSON format: { "main_role": "<most_likely_main_role>" { reasoning_in_json } } Example Output: { "main_role": "Antagonist" { reasoning_example }
Multi-Step Fine-grained Role	classify the narrative framing of the { entity } in the document. The taxonomy is { taxonomy } . { label_definitions }	You have previously identified the broader narrative frame to be main_role_candidate. Instructions: 1. Determine one or a maximum of two corresponding fine grained role from the taxonomy. 2. Order the sub-roles by likelihood, with the most likely fine-grained role listed first.	3. Finally, return your conclusion { ask_reasoning } as a single JSON object with no extra text, in this format: { "fine_grained_roles": ["<Most likely sub-role>", "<Second most likely sub-role if relevant>"] { reasoning_in_json } } Example Output: { "fine_grained_roles": ["Conspirator"] { reasoning_example }

Table 10: 不同设置、任务定义和叙事框架分类对应说明之间的映射。

分类 提示中共享的分类法与附录 A 中相同。提供的标签定义来自官方论文，并在表 11 中共享。

Role Level	{ label_definitions }
Main Role	Protagonist: The central figure or party in a news article, often portrayed as a hero or positive force driving the narrative. Antagonist: The opposing figure or force in a news article, often depicted as the source of conflict or challenge to the protagonist. Innocent: An individual or group portrayed as untainted or blameless in the context of the news, typically victimized or wronged.
Fine Grained Roles Protagonist	Guardians: Heroes or guardians who protect values or communities, ensuring safety and upholding justice. They often take on roles such as law enforcement officers, soldiers, or community leaders (e.g., climate change advocacy community leaders). Martyr: Martyrs or saviors who sacrifice their well being, or even their lives, for a greater good or cause. These individuals are often celebrated for their selflessness and dedication. This is mostly in politics, not in Climate Change. Peacemaker: Individuals who advocate for harmony, working tirelessly to resolve conflicts and bring about peace. They often engage in diplomacy, negotiations, and mediation. This is mostly in politics, not in Climate Change. Rebel: Rebels, revolutionaries, or freedom fighters who challenge the status quo and fight for significant change or liberation from oppression. They are often seen as champions of justice and freedom. Underdog: Entities who are considered unlikely to succeed due to their disadvantaged position but strive against greater forces and obstacles. Their stories often inspire others. Virtuous: Individuals portrayed as virtuous, righteous, or noble, who are seen as fair, just, and upholding high moral standards. They are often role models and figures of integrity.
Fine Grained Roles Antagonist	Conspirator: Those involved in plots and secret plans, often working behind the scenes to undermine or deceive others. They engage in covert activities to achieve their goals. Instigator: Individuals or groups initiating conflict, often seen as the primary cause of tension and discord. They may provoke violence or unrest. Deceiver: Deceivers, manipulators, or propagandists who twist the truth, spread misinformation, and manipulate public perception for their own benefit. They undermine trust and truth. Incompetent: Entities causing harm through ignorance, lack of skill, or incompetence. This includes people committing foolish acts or making poor decisions due to lack of understanding or expertise. Their actions, often unintentional, result in significant negative consequences. Corrupt: Individuals or entities that engage in unethical or illegal activities for personal gain, prioritizing profit or power over ethics. This includes corrupt politicians, business leaders, and officials. Tyrant: Tyrants and corrupt officials who abuse their power, ruling unjustly and oppressing those under their control. They are often characterized by their authoritarian rule and exploitation. Foreign Adversary: Entities from other nations or regions creating geopolitical tension and acting against the interests of another country. They are often depicted as threats to national security. This is mostly in politics, not in Climate Change. Terrorist: Terrorists, mercenaries, insurgents, fanatics, or extremists engaging in violence and terror to further ideological ends, often targeting civilians. They are viewed as significant threats to peace and security. This is mostly in politics, not in Climate Change. Bigot: Individuals accused of hostility or discrimination against specific groups. This includes entities committing acts falling under racism, sexism, homophobia, Antisemitism, Islamophobia, or any kind of hate speech. This is mostly in politics, not in Climate Change. Saboteur: Saboteurs who deliberately damage or obstruct systems, processes, or organizations to cause disruption or failure. They aim to weaken or destroy targets from within. Traitor: Individuals who betray a cause or country, often seen as disloyal and treacherous. Their actions are viewed as a significant breach of trust. This is mostly in politics, not in Climate Change. Spy: Spies or double agents accused of espionage, gathering and transmitting sensitive information to a rival or enemy. They operate in secrecy and deception. This is mostly in politics, not in Climate Change.
Fine Grained Roles Innocent	Victim: People cast as victims due to circumstances beyond their control, specifically in two categories: (1) victims of physical harm, including natural disasters, acts of war, terrorism, mugging, physical assault, etc., and (2) victims of economic harm, such as sanctions, blockades, and boycotts. Their experiences evoke sympathy and calls for justice, focusing on either physical or economic suffering. Scapegoat: Entities blamed unjustly for problems or failures, often to divert attention from the real causes or culprits. They are made to bear the brunt of criticism and punishment without just cause. Exploited: Individuals or groups used for others' gain, often without their consent and with significant detriment to their wellbeing. They are often victims of labor exploitation, trafficking, or economic manipulation. Forgotten: Marginalized or overlooked groups who are often ignored by society and do not receive the attention or support they need. This includes refugees, who face systemic neglect and exclusion.

Table 11: 标记定义在提示中共享。这些来自任务定义 (Piskorski et al., 2025) 。