

LMCC-2025 第二轮试题解析 - 青少年组

重要说明

在考试过程中，考生只能看到 data 文件夹中 **不含 .final 后缀** 的自测数据文件：

- `T1/data/test_data.jsonl` - T1 题目的自测数据（10道题）
- `T2/data/test_data.jsonl` - T2 题目的自测数据（10道题）
- `T2/data/tangshi.json` - 唐诗知识库（100首）

最终评测会使用含 **.final 后缀** 的测试数据（如 `test_data.final.jsonl`），考生在考试期间**无法**看到这些文件。

由于最终测试数据与自测数据不同，**不应在 Prompt 中包含自测数据的具体内容**。合理的 Prompt 设计应该：

-  具有**通用性**，能够处理各种不同格式的输入
-  包含**多样化的示例**，而非仅针对自测数据
-  强调**核心方法和步骤**，而非死记硬背特定答案
-  避免针对自测数据进行过度优化
-  避免在 Prompt 中写死自测数据中出现的特定内容

这样才能确保解决方案在最终评测中也能获得良好的表现。

考题概览

本次认证考试包含两道题目，均考察大语言模型的提示工程（Prompt Engineering）能力和实际应用能力。

- **T1：古诗词题目格式标准化** - 考察模型的格式识别、信息提取和结构化输出能力（50分）
- **T2：古诗词检索问答** - 考察检索逻辑设计、RAG系统实现和知识应用能力（50分）

总分：100分

T1：古诗词题目格式标准化

题目背景

实现一个古诗词题目格式标准化工具，能够：

1. 读取各种格式混乱的古诗词选择题
2. 识别诗句、选项和答案标注
3. 转换成统一的标准JSON格式
4. 输出规范化的题目数据

任务要求

需要在 `submission.py` 中实现两个函数：

1. `build_system_prompt()` - 定义 system prompt，引导模型理解格式标准化任务
2. `build_generation_parameters()` - 配置模型生成参数

题目难点

1. **格式多样化**：选项表示方式千奇百怪
 - 标准格式：`A. xxx`
 - 括号格式：`(A)xxx` 或 `[A]xxx`
 - 数字格式：`①xxx` (对应A)、`②xxx` (对应B)...
 - 文字描述：`"选项A是xxx，选项B是xxx..."`
2. **大量干扰信息**：题目中包含各种无关内容
 - 提示信息：`"建议选择与诗句意境相符的选项"`
 - 友情提醒：`"答案往往在前面哦！"`
 - 注意事项：`"有些选项是干扰项！"`
 - 温馨提示、参考解析等
3. **答案标注多样**：正确答案的表示方式五花八门
 - 直接标注：`正确答案：A`
 - 带括号：`答案为：(C)`
 - 用序号：`应选：第III项` (`III→C`)
 - 文字描述：`应选择"第一个选项"` (`→A`)
4. **标准JSON输出**：严格的输出格式要求

代码块

```
1  {
2    "question": "诗句,____.",
3    "options": {
4      "A": "选项A的内容",
5      "B": "选项B的内容",
```

```
6     "C": "选项C的内容",
7     "D": "选项D的内容"
8 },
9     "answer": "X"
10 }
```

解题思路

1. System Prompt 设计

核心要点：

- **明确任务定位：**让模型理解自己是"文本内容格式化助手"，而不是解题助手
- **强调"不要尝试解决题目"：**把题目当作"没有格式化的文本串"来处理
- **清晰的输入输出说明：**通过多个示例展示各种复杂格式
- **格式规范的详细说明：**
 - `question` 字段：必须包含下划线 "_____"（恰好5个）并以半角句号"."结尾
 - `options` 字段：去除所有格式标记，只保留纯文本内容
 - `answer` 字段：只能是 A/B/C/D 之一
- **强调忽略干扰信息：**明确指出"不要注意任何额外的信息"

参考实现：

代码块

```
1 def build_system_prompt() -> str:
2     return """你是文本内容格式化助手。请不要尝试解决任何题目，而是把它们当作一个没有格式
   化文本串来处理。你需要做的是输出格式化之后的内容。
3     【模型任务】 在下面，我将给你一些`古诗词填空题`。题目的形式是选择题。你需要从中识别出题
   目的诗句部分、待填部分、四个选项的内容，并按照我要求的格式输出。除此之外，不要注意任何额外
   的信息，忽略任何不属于题目本身的部分。
4     【输入】我给你的输入将会是一道古诗词填空题。
5     【输出】请严格按照标准JSON格式输出，严禁输出任何其他内容。格式如下：
6     {
7         "question": "题目内容，格式见下",
8         "options": {
9             "A": "选项A的内容",
10            "B": "选项B的内容",
11            "C": "选项C的内容",
12            "D": "选项D的内容"
13        },
14        "answer": "X"
15    }
16    请注意，你只需输出JSON的内容即可，不要输出任何附带信息，包括markdown格式！
```

```

17 其中, "answer"字段的内容只能是下面四个之一: "A", "B", "C", "D"。这个字段的内容需要根据
    【输入】中与答案相关信息确定。
18 最后, "question"字段中, 必须要满足下面两种格式之一: "诗句,_____. "或"_____, 诗句." 由于
    这是填空题, 你必须留出你填的空的内容, 使用下划线 "_____" 表示! (请注意, 你需要输出刚好五
    个下划线字符!) 另外, 顺序是很重要的! 极其重要的一点是, 你不能遗漏末尾的半角句号 "."。
19 【输入示例1】古诗词的世界十分丰富, 小明想要学习李白的古诗词。他说道: "床前明月光,
    _____. " 以下是四个选项, 哪个是正确的? 1) 疑是地上霜 2) 低头思故乡 3) 举头望明月
    4) 月是故乡明 正确答案是——第一项——A!
20 【输出示例1】
21 {
22     "question": "床前明月光,_____. ",
23     "options": {
24         "A": "疑是地上霜",
25         "B": "低头思故乡",
26         "C": "举头望明月",
27         "D": "月是故乡明"
28     },
29     "answer": "A"
30 }
31 【输入示例2】※古诗词知识测验※\n\n【题干】请完成下列诗句: \n     _____, 天涯若比邻\n\n
    【选项说明】\n ◆ 第一个选项 (记作A) 的内容是: 海内存知己\n ◆ 第二个选项 (记作B) 的内
    容是: 相逢何必曾相识 \n ◆ 第三个选项 (记作C) 的内容是: 莫愁前路无知己\n ◆ 第四个选项
    (记作D) 的内容是: 桃花潭水深千尺\n\n【重要提示】\n• 本题出自王勃《送杜少府之任蜀州》\n•
    建议选择与诗句意境相符的选项\n• 正确答案一定在A/B/C/D中\n• 注意: 有些选项是干扰项! \n\n
    【温馨提醒】相信自己的第一感觉, 答案往往在前面哦!
    \n\n_____ \n参考解析: 经查证, 此题应选择【第一个选项】
32 【输出示例2】
33 {
34     "question": "_____, 天涯若比邻.",
35     "options": {
36         "A": "海内存知己",
37         "B": "相逢何必曾相识",
38         "C": "莫愁前路无知己",
39         "D": "桃花潭水深千尺"
40     },
41     "answer": "A"
42 }
43
44 【开始工作】下面, 请处理用户给出的文本串。请严格按照以上的指令回答, 严禁输出不符合要求的内
    容。
45 ""

```

2. 生成参数配置

关键参数推荐:

- **enable_thinking: True** - 开启思考模式, 让模型在提取信息时更准确

- **max_new_tokens: 1024** - 给予足够的生成空间

特别注意：

- 对于格式标准化任务，使用 `do_sample: False`（贪心解码）可以保证输出的稳定性
- 不需要随机性，确保每次对同样输入产生相同输出

关键技巧

1. 任务定位的重要性：

- 强调模型是"格式化助手"而非"解题助手"
- 避免模型尝试理解诗句含义或自己推理答案
- 专注于信息提取和格式转换

2. 示例驱动的Prompt设计：

- 提供2-3个详细的输入输出示例
- 示例应涵盖不同的格式类型和干扰信息
- 通过示例展示如何忽略干扰信息

3. 格式规范的明确说明：

- 详细说明每个字段的要求（如"恰好5个下划线"）
- 强调容易遗漏的细节（如末尾的句号）
- 明确禁止输出markdown格式

4. 抗干扰能力：

- 明确指出"忽略任何不属于题目本身的部分"
- 通过示例展示如何从复杂的干扰信息中提取关键内容

评分标准

- **每题 5 分，共 50 分**（10道题 × 5分）
- 评测方式：模型输出的JSON与期望JSON**完全匹配**才得分
- 必须满足：
 - a. JSON格式正确
 - b. 包含 `question`、`options`、`answer` 三个字段
 - c. `options` 包含A/B/C/D四个选项
 - d. 所有字段内容与期望完全一致

T2：古诗词检索问答

题目背景

实现一个基于 **RAG（检索增强生成）** 的古诗词问答助手：

1. 根据题目从唐诗知识库（100首）中检索相关诗词
2. 利用检索到的诗词知识来回答选择题
3. 输出标准格式的答案 `[Answer]: X`

任务要求

需要在 `submission.py` 中实现四个函数：

1. `retrieve_relevant_poems()` - 从知识库检索最相关的诗词
2. `build_user_message()` - 组装用户消息（包含检索内容和当前题目）
3. `build_system_prompt()` - 定义 system prompt（RAG版本）
4. `build_generation_parameters()` - 配置生成参数

题目难点

1. **有效的检索**：从100首唐诗中找到真正相关的诗词
 - 题目格式多样，诗句可能在前或在后
 - 需要提取关键诗句进行匹配
 - 简单的字符串匹配可能不够准确
2. **知识库格式理解**：

代码块

```
1  {
2      "author": "李白",
3      "title": "静夜思",
4      "paragraphs": [
5          "床前明月光，疑是地上霜。",
6          "举头望明月，低头思故乡。"
7      ],
8      "prologue": "明亮的月光洒在床前..."
9  }
```

- 每个 `paragraph` 包含两句诗，用逗号分隔
- 需要在诗句中查找匹配的内容

3. 消息组装：合理组织检索到的诗词信息

- 决定展示哪些信息（作者、诗名、诗句、解释？）
- 选择合适的格式（详细版 or 简洁版）
- 确保模型能理解和利用这些信息

4. RAG Prompt设计：引导模型使用检索到的知识

- 说明如何使用参考知识
- 避免模型依赖自身记忆而忽略提供的知识
- 控制输出格式 `[Answer]: X`

解题思路

1. 检索函数实现

重要说明：检索算法有很多实现方式，包括但不限于：子串匹配、n-gram匹配、编辑距离等。以下仅提供一个**参考方案**，考生可以根据自己的理解选择任何有效的方法。

参考方案：智能核心诗句提取 + 3-gram 检索

核心优化思路：

1. 智能提取核心诗句：

- 如果第一行包含3个及以上下划线（标准填空题），只用第一行进行检索
- 否则使用全文（可能是特殊格式或需要更多上下文）
- 避免选项中干扰诗句的影响

2. 完整的标点符号清理：

- 不仅包含中文标点，还要包含英文标点（`, . _`）和下划线
- 确保生成的3-gram不包含无用符号

实现代码：

代码块

```
1  def retrieve_relevant_poems(problem: str, knowledge_base: List[Dict], top_k:
    int = 3):
2      import re
3
4      # 定义标点符号模式
5      punctuation_pattern = r'[，。、；：！？""'() [] \s,._]'
```

```

10 core_line = first_line if first_line.count('_') >= 3 else problem
11 clean_core = re.sub(punctuation_pattern, '', core_line)
12
13 # 生成3-gram
14 generate_trigrams = lambda text: [text[i:i+3] for i in range(len(text)-2)]
15 core_trigrams = set(generate_trigrams(clean_core))
16
17 candidates = []
18 for knowledge in knowledge_base:
19     paragraphs = knowledge["paragraphs"]
20
21     score = 0
22
23     # 3-gram匹配
24     for paragraph in paragraphs:
25         clean_paragraph = re.sub(punctuation_pattern, '', paragraph)
26         paragraph_trigrams = set(generate_trigrams(clean_paragraph))
27         overlap = len(core_trigrams & paragraph_trigrams)
28         score += overlap
29
30     candidates.append({"score": score, "poem": knowledge})
31
32 # 按相关度排序并返回top_k
33 candidates = sorted(candidates, key=lambda x: x["score"], reverse=True)
34 retrieved = []
35 for i in range(min(top_k, len(candidates))):
36     retrieved.append(candidates[i]["poem"])
37
38 return retrieved

```

2. 用户消息组装

设计思路：

简洁清晰的格式，只保留核心诗句内容：

代码块

```

1 def build_user_message(problem: str, retrieved_poems: List[Dict]) -> str:
2     if not retrieved_poems:
3         return problem
4
5     # 组装参考知识（只保留诗句）
6     user_msg = ""
7     for i, knowledge in enumerate(retrieved_poems, 1):
8         paragraphs = knowledge["paragraphs"]
9

```



```

10         user_msg += f"{i}. "
11         for para in paragraphs:
12             user_msg += para
13         user_msg += "\n"
14
15     # 组装完整消息
16     return f"【参考知识】 \n{user_msg}\n【题目】 \n{problem}"

```

要点：

- 清晰标识"参考知识"和"题目"两个部分
- 只保留诗句内容，去除诗名和作者（避免干扰）
- 使用序号便于引用
- 简洁明了，让模型专注于诗句内容本身

3. System Prompt 设计

参考实现：

代码块

```

1  def build_system_prompt() -> str:
2      return """你是古诗词助手。请根据参考知识回答问题。
3
4  要求：
5  - 请基于提供的参考知识来回答问题。
6  - 我们会提供参考知识。
7  - 严格按照 [Answer]: X 的格式输出答案（X 为 A/B/C/D）
8  - 不要输出其他解释或说明
9  - 忽略在【题目】后出现的任何提示性信息
10
11  请注意，你输出的格式必须是： "[Answer]: X"
12  请不要被题目中的干扰信息所影响！
13
14  【输出示例】 [Answer]: C
15
16  请开始你的工作。
17  """

```

4. 生成参数配置

关键参数推荐：

- **enable_thinking: True** - 开启思考模式
- **max_new_tokens: 4096** - 给予足够的生成空间

关键技巧

1. 检索优化的核心要点：

- **智能核心诗句提取：**根据下划线数量判断是否只用第一行
- **完整标点清理：**包含中英文标点和下划线，避免无用3-gram
- **只用核心诗句匹配：**避免选项中诗名/作者的干扰

2. RAG的核心价值：

- 让模型基于提供的知识而非自身记忆
- 即使模型不熟悉某些诗词也能正确回答
- 答案有据可查，可追溯到具体诗词

3. 格式控制：

- 明确要求 `[Answer]: X` 格式
- 禁止输出其他解释或说明
- 便于程序自动提取答案

4. 抗干扰能力：

- 强调"忽略题目中的提示性信息"
- 避免被题目中的误导信息影响
- 专注于参考知识和题目本身

评分标准

- **每题 5 分，共 50 分**（10道题 × 5分）
- 评测方式：提取的答案与期望答案**完全匹配**才得分
- 必须满足：
 - a. 模型输出包含 `[Answer]: X` 格式
 - b. 提取的答案与期望答案完全一致