

赛题五：基于 AI 的智能旅行规划

一、赛题背景

（一）行业趋势与问题痛点

近年来，中国旅游市场规模持续扩大，2023 年国内旅游人次突破 45 亿，个性化、定制化旅行需求快速增长。然而，传统旅行规划工具普遍存在三大痛点：

- 需求理解局限：现有工具依赖固定模板或简单标签筛选，难以深度解析用户自然语言中的复杂意图（如“想带孩子体验非遗文化，但预算有限且需避开人群”）。
- 动态适配不足：无法实时整合交通、天气、景区限流等动态数据，导致方案可靠性差。
- 知识更新滞后：旅游信息迭代快（如网红打卡点更替、政策变化），传统系统维护成本高且覆盖不全。

（二）学术研究价值

大语言模型（LLM）的突破为旅行规划带来新机遇，但面临关键挑战：

- 多渠道信息集成与决策：需融合 POI 数据、用户画像、时空约束等多维度信息
 - 可信可靠旅行方案生成：平衡旅行体验与可行性（如路线合理性、时间成本）
- 本赛题将推动智能体技术在真实世界复杂决策场景的落地，促进大语言模型与神经符号系统、智能体技术的融合创新，为 AI 垂直领域应用提供重要范式。

二、赛题应用场景

（一）典型应用场景

1. 自由行深度定制

1) 用户需求：“十一想带父母从上海出发，5 天行程要兼顾古镇慢游和现代都市，父亲腿脚不便需减少步行”

2) 智能体响应：自动规避阶梯地形景点，设计古镇（如乌镇）与城市文化地标（上海天文馆）的组合路线，推荐无障碍交通接驳方案

2. 地方文化沉浸体验

1) 用户需求：“带孩子研学旅行，希望深度学习当地历史文化，全程总预算不超过 8000 元”

2) 智能体响应：匹配当地博物馆、名人故居等历史文化景点，自动计算门票+住宿套餐优惠，动态调整餐饮标准控制预算

3. 特殊需求响应

1) 用户需求：“下周北京出差需参加 3 场会议，间隙想体验本地特色餐饮，晚上

要处理邮件"

2) 智能体响应：基于会议地点智能聚类酒店餐厅选址，推荐高效通勤路线，穿插预订可快速就餐的本地特色餐厅

(二) 场景技术关联

1. 多源旅行信息整合：需整合景点、餐厅、住宿等地理位置、开放时间、人均开销等多维度信息，生成旅行方案。

2. 神经符号决策规划：结合大语言模型需求理解和符号推理可信可靠的优势，提高智能体方案对于用户需求的满足能力，确保方案可行可靠。

通过本赛题，参赛者将直面真实应用场景中的复杂约束，推动对话式 AI 从"信息检索"向"可信决策"跨越，为旅游产业数字化升级提供关键技术支撑。

三、出题信息

本赛题由大赛组委会组织专家团队出题，主办单位为南京大学、华为，主办方出题人为：郭兰哲、李宇峰、邵杰晶、杨晓文、张博闻、韩思予、于坤杨、周植。

四、赛题任务

本次竞赛将为参赛者提供一个旅行沙盒以及配套的用户需求数据集。

该旅行沙盒包括中国热门旅游目的地和城市中的具体的旅行资源信息。

用户需求数据集包括自然语言形式的合成需求和来自问卷调研的真实需求，旨在还原旅行场景中的复杂用户需求。

(一) 训练流程

在训练阶段，用户将基于提供的沙盒环境和用户需求数据，设计智能体算法并与沙盒环境交互以获取旅行相关信息，最终生成符合要求的旅行规划方案。用户可对算法进行调优和模型测试，以优化规划效果。除提供的数据外，用户可自行合成数据或使用开源互联网数据进行模型训练，但需在复赛、决赛阶段提交相应数据并在最终技术报告中详细说明数据合成方法或标注外部数据来源。

(二) 测试流程

在测试阶段，系统需通过自然语言交互准确理解用户的旅行需求，并自动生成满足各项约束条件的旅行规划方案。规划效果的评估将以约束满足通过率作为核心指标，重点考察系统在以下关键挑战中的鲁棒性表现：时空细粒度规划的准确性、多约束条件下的复杂需求处理能力和旅行偏好的适应性。

最终目标是构建一个在真实场景下具备高可靠性、强泛化能力的智能旅行规划系统。

五、数据集及数据说明

本次竞赛提供的离线数据集基于旅行沙盒生成。其中初赛数据集仅用于参赛者的

算法验证与调试，复赛、决赛提供全新数据集（数据格式与初赛保持一致）。这样是为了能模拟现实需求中“模型上线需面对新增用户的不同需求”的挑战，筛选出那些更具泛化性、更可靠的算法模型。数据集详细信息如下：

旅行沙盒包括中国 10 个知名城市：北京、成都、重庆、广州、杭州、南京、上海、深圳、苏州、武汉，城市间的航线、火车各 700、5700 余项。沙盒为每个城市提供平均 300 余个景点、400 余餐厅和 400 余酒店的基础信息。

数据集包括合成数据和用户真实数据，每条数据包含一条自然语言询问，代表用户的旅行需求。

示例数据可从 <https://box.nju.edu.cn/d/8941a6d3cddc4c7591e4/> 获取。

六、性能指标要求

本次比赛以需求满足率和旅行偏好满足度为主要评价指标。参赛者需提交策略模型，测试服务器将基于该模型与旅行沙盒进行交互式测试。最终，系统将根据所有数据上的约束通过率，将其作为模型的最终评分结果。

评价指标包含下述几项：

1. 环境约束：环境约束评价了输出规划方案中的信息是否与提供的沙盒环境信息一致，度量了规划方案的可行性。

$$EPR - micro = \frac{\sum_{p \in P} \sum_{c \in Env} 1_{passed(c, p)}}{|P| * |Env|}$$

$$EPR - macro = \frac{\sum_{p \in P} \prod_{c \in Env} 1_{passed(c, p)}}{|P|}$$

2. 条件逻辑约束：条件逻辑约束评价了输出规划方案中在满足环境约束的前提下对用户个性化需求的满足程度。

$$C - LPR = \frac{\sum_{p \in P} 1_{passed(Env, p)} \cdot \sum_{c \in C_p} 1_{passed(c, p)}}{\sum_{p \in P} |P|}$$

3. 最终约束通过率：最终约束通过率表达了输出规划方案中满足所有环境约束和逻辑约束的比例。

$$C - LPR = \frac{\sum_{p \in P} 1_{passed(Env, p)} \cdot \sum_{c \in C_p} 1_{passed(c, p)}}{\sum_{p \in P} |P|}$$

4. 偏好评估：

TPC 比赛中，我们提供了三个旅行中常见的偏好指标：

每天访问的景点数量尽可能多, Daily Average Attractions Visited, DAV，数值归一化到[0,4]作为分数

$$DAV-score = (DAV - 0) / 4$$

平均交通时间尽可能少，Averaged Transportation Time, ATT，数值归一到[15,120]（分钟）作为分数

$$ATT-score = \max(\min((120 - ATT) / (120 - 15), 1), 0)$$

每天餐饮推荐数量尽可能多, Daily Dining Recommendations, DDR，数值归一到[0,3] 作为分数。

$$DDR-score = \min((DDR - 0) / (3 - 0), 1)$$

5. 最终得分：

$$\text{Overall Score} = 10\% * \text{EPR-micro} + 10\% * \text{EPR-macro} + 25\% * \text{C-LPR} + 40\% * \text{FPR} + 5\% \text{ DAV-Score} + 5\% \text{ ATT-Score} + 5\% \text{ DDR-Score}$$

七、功能要求

1. 算法设计要求：

训练方式：基于提供的数据完成模型调优

2. 算法优化：

（1）旅行规划助手通过与旅行沙盒环境进行交互，并整合获取的信息来提供切实有效，符合沙盒信息的旅行规划方案。

（2）旅行规划助手应通过理解用户的自然语言需求，确保最终提供的旅行规划方案确切符合用户需求。

（3）旅行规划助手应具备具备良好的泛化能力，能泛化到新用户的新需求。

八、开发环境

参赛者需使用 Python 进行开发，建议参赛者确保其开发环境与提供测试环境配置一致，以避免因环境差异导致的运行问题。

九、成绩评价

本赛题分为初赛、复赛和决赛三个阶段：

1. 初赛：初赛提供数据集仅供参赛者在前期进行算法验证与调试，选手提交的算法与方案结果不计入最终决赛总分。

2. 复赛：复赛提供新的数据集（格式和初赛保持一致），选手提交算法模型及技术报告，赛事方验证算法运行结果，根据性能评估指标与提交材料完整性、质量等进行综合打分。

决赛：线上决赛提供决赛验证数据集（格式与初赛复赛一致）。线下决赛最终综合成绩由客观评分和主观评分构成，比例为 70%和 30%。（国二、国三成绩不涉及 30%主观评分部分）

（1）客观评分：基于经过标准化处理后的机器评测得分；

（2）主观评分：依据经过标准化处理后的答辩得分。答辩评价将综合考察参赛者的答辩表现，以及所提交的技术方案和代码文档。

十、解题思路

1. 语言模型提示调优：通过提示调优和上下文学习，利用大语言模型对自然语言需求进行推理，提高旅行规划助手对用户需求的理解程度

2. 语言模型微调提升：通过大语言模型后训练，提高大语言模型对自然语言需求的理解能力

3. 程序 workflow 优化：通过程序级 workflow 优化，提高旅行规划助手的方案求解效率

重点待解决的技术挑战：

1) 针对用户复杂多样的自然语言表达，设计需求理解算法，调优语言模型，提高对自然语言表达的形式化约束提取能力。

2) 基于符号验证器，设计方案修正算法，通过未满足约束的提示信息，帮助语言模型进行方案修正

3) 针对神经符号系统求解效率低下，设计调优 workflow，提高求解效率；

十一、参考资源

旅行规划、大模型相关论文及开源代码库，例如：

ChinaTravel: A Real-World Benchmark for Language Agents in Chinese Travel Planning

To the Globe (TTG) : Towards Language-Driven Guaranteed Travel Planning

TravelPlanner: A Benchmark for Real-World Planning with Language Agents

ITINERA: Integrating Spatial Optimization with Large Language Models for Open-domain Urban Itinerary Planning

十二、提交要求

本次竞赛要求选手在给定测试框架下提交算法推理代码和模型，测试服务器将利用该模型与旅行沙盒进行交互，对数据集生成一批规划结果。最终，计算这批规划结果的约束满足率均值作为模型的评分结果。

参赛者请按测试代码标准将所有文件打包为一个.zip 文件：

选手在开发过程中，可以运行测试框架提供的评估脚本和旅行沙盒进行交互，验证策略实现是否存在问题。

资源和格式限制

模型单条数据样本的推理时间不超过 5 分钟，超时则失败。

十三、更新与答疑

赛题可能会进行更新，为了回复参赛者在参赛过程中遇到的问题，赛题将单独设立选手答疑群（选手正式报名后可看到）。

十四、比赛流程及奖项设置

（一）报名阶段

参赛者在比赛官方网站完成报名注册，提交个人或团队信息，获取初赛数据下载链接。

（二）初赛阶段

参赛者利用赛事方提供的训练数据集进行算法模型设计，并提交测试集结果，赛事方根据结果对比，计算模型得分。

（三）复赛（省赛）阶段

初赛结束后进入复赛阶段，开放复赛数据集下载链接。仅有初赛阶段提交有效结果的参赛团队可以进入复赛。复赛期间，参赛者根据提供的复赛阶段数据进行算法模型调试，并于复赛截止前提交算法模型与技术报告，赛事方验证算法运行结果，根据性能评估指标与提交材料完整性、质量等进行综合打分。

（四）复赛（省赛）成绩公布

在比赛官方网站上公布复赛成绩。以进入复赛参赛团队数量作为计奖基数，按照不超过大赛省赛设奖比例，评选出复赛一、二、三等奖（颁发省赛获奖证书）。评选复赛奖过程中，参赛者提交的算法性能低于赛事方提供的基线参考分数的判定为无效成绩，不予授奖。复赛一、二等奖晋级参加国赛总决赛。

（五）决赛（国赛）阶段

1. 决赛线上评选。基于决赛验证数据集对进入复赛的团队所提交的算法模型进行评测打分，以进入决赛参赛团队数量作为计奖基数，按照不超过大赛国赛设奖比例，评选出国赛一等奖候选名单及国赛二、三等奖获奖名单（颁发国赛二、三等奖证书）。

2. 决赛作品提交：国赛一等奖候选团队在规定时间内完善提交技术文档、算法

代码和模型文件、演示视频、补充材料等。提交截止后，不再接受任何形式的修改和补充。

3. 决赛评审阶段：由专业评审团队对国赛一等奖候选团队的参赛作品进行评审，根据性能评估指标和提交材料的完整性、质量等进行综合打分。评审过程中如有疑问，可要求参赛者进行解释说明。

4. 决赛线下评选：国赛一等奖候选团队在规定时间内提交完善后的技术文档、算法代码和模型文件、演示视频、补充材料，参加国赛线下总决赛复核答辩，最终确定国赛一等奖获奖名单及其排名（未参加线下复核答辩视同放弃奖项）。国赛一等奖颁发荣誉证书。

十五、奖金设置

为了鼓励参赛选手参赛积极性，本赛题根据总决赛成绩，对成绩排名前五的参赛团队设置奖金。

1. 冠军奖：第 1 名，奖金 5000 元/每团队
2. 亚军奖：第 2 名，奖金 2000 元/每团队
3. 季军奖：第 3 名，奖金 1000 元/每团队
4. 优秀奖：第 4-5 名，奖金 500 元/每团队

十六、其它说明

1. 公平性：严禁任何形式的作弊行为，包括但不限于数据泄露、模型预训练数据与测试数据重叠、抄袭他人代码等。一经发现，立即取消参赛资格，并追究相关责任。

2. 知识产权：参赛者提交的作品必须为原创，未在其他比赛中获奖或公开发表。比赛主办方有权对参赛作品进行展示、宣传等相关活动，但知识产权仍归参赛者所有。

十七、联系方式

赛项交流 QQ 群：768721120

邮 箱：guolz@nju.edu.cn

报名官网：www.aicomp.cn