

RESEARCH ARTICLE



Using ChatGPT with Prompt Engineering for Personalized Travel Destination Recommendations

Sarah Castratori¹, Domenico Dongiovanni¹, David Grossenbacher¹ and Thomas Hanne^{2,*}

¹*School of Business, University of Applied Sciences and Arts Northwestern Switzerland, Switzerland*

²*Institute for Information Systems, University of Applied Sciences and Arts Northwestern Switzerland, Switzerland*

Abstract: Artificial intelligence (AI) has become widely accessible and can support a wide range of use cases such as personalized recommendations. Our study investigates the capabilities of a pre-trained large language model for generating personalized travel destination recommendations. Based on a design science research approach, we evaluate the usability and effectiveness of suggestions. In particular, we explore how additional task-specific or user-specific input can enhance the response effectiveness. User preferences gathered from a survey are used for computational experiments. The results are presented in a second survey, and the participants' feedback is collected to assess the perceived quality of personalized recommendations. This feedback demonstrates that the model recommendations are positively evaluated with regard to user satisfaction, preference consideration, and timesaving. Yet, the participants also see potential for more detailed and specific recommendations, along with enhanced granularity for the budget breakdown and better transparency through shared sources. We also identify limitations, such as the availability and quality of gathered data, leading to potentially inadequate recommendations, as well as subjective evaluation measures requiring larger and more diverse sample data to confirm the generalizability of the approach. The findings indicate the potential benefits of personalized AI recommendations in the travel industry and suggest areas for improvement.

Keywords: large language models, ChatGPT, recommender systems, travel recommendations

1. Introduction

Artificial intelligence (AI) has achieved unprecedented influence with its latest developments and media coverage. With the arrival of new, advanced, and easily accessible AI tools such as large language models (LLMs), it can be argued that there is an excess of sophisticated use cases and research areas [1–4]. These tools may improve the efficiency of human activities, provide humans with new ideas, and help unleash more creativity. However, various practical and ethical issues that need to be addressed such as generated content that is inaccurate, inappropriate, or biased [5, 6]. Moreover, not all approaches can be realized easily as they may require further development of use case-specific solutions. However, in various cases, pre-trained transformer networks, such as ChatGPT [7], do not require further time-consuming training efforts but can provide beneficial results based on additional text input by the user.

Considering this, the purpose of our study is to explore a use case for additional input regarding recommendation requirements for ChatGPT as a pre-trained transformer network. With the help of experiments, the goal is to verify whether such approaches

equipped with additional personalized information are better in mundane decision-making tasks compared to the bare usage of the pre-trained model.

1.1. Problem statement

In recent months, LLMs such as ChatGPT have been the center of discussion in politics, social media, entertainment, and dinner tables of ordinary households. Reuters reported that in January 2023, over 100 million people actively engaged with the service a mere two months after its November 2022 release to the public, which sets an unprecedented number of users in the history of consumer applications [8]. Despite it being a widely used tool, it is still unclear how much mundane human decision-making can benefit from it. Such tasks can be the decision of what to cook for lunch, which movie to pick on movie night, where to go out for dinner, and which music to listen to on a road trip. In addition to such low-cost decisions, grander but still common decisions like where to go for a weekend trip with friends also consume significant time for planning and preparation. What if AI could reduce this time drastically by offering an easy-to-use solution, for example, by presenting a few distinct ideas or concrete proposals? According to Dwivedi et al. [9], there is, for instance, potential in the tourism industry because travelers may require extensive information from different sources to evaluate

*Corresponding author: Thomas Hanne, Institute for Information Systems, University of Applied Sciences and Arts Northwestern Switzerland, Switzerland. Email: thomas.hanne@fhnw.ch

services and to create meaningful itineraries. Recently, they have relied on search engines and had to consult multiple websites to gather such information. Specifically, LLMs such as ChatGPT could enable tourists to easily obtain comprehensive answers by combining various resources. It may provide quick and accurate information in natural language, helping them plan and enhance their travel experiences.

1.2. Thesis statement and research questions

This research paper aims to investigate the usability of large pre-trained transformer networks by testing the following thesis statement: “Additional input of task-specific or user-specific texts will enhance the adaptability and effectiveness of LLMs for personalized travel destination recommendation generation.” Adaptability in this sense means that the model with user- or task-specific input is expected to significantly improve the quality of recommendations.

Despite the success of pre-trained transformer networks in natural language processing, there is still a need to investigate their potential for personalized recommendations, as illustrated before. While pre-trained models can be readily accessible, additional input of task-specific or user-specific texts is necessary to adapt them to individual purposes. Thus, the research aims to explore the possibilities of using pre-trained transformer networks for generating personalized travel destination recommendations. For our study, we will focus on ChatGPT as a tool as it appears to be the most well-known and frequently used model by private users although other models have become competitive. Thus, we assume that results similar to those obtained with ChatGPT could be achieved with other LLMs as well.

The effectiveness of additional input for the pre-trained models in order to receive suitable travel recommendations shall be evaluated through computational experiments, with a focus on assessing the quality of personalized recommendations.

Ultimately, the study seeks to contribute to an adaptable and effective application of natural language processing models that can be used by everyone, regardless of their expertise in information systems, for personalized recommendations and complex decision-making (no code experience). In the context of this work, adaptability refers to using a pre-defined model, connected to a user's preferences and personal information, to enable an enhanced personalization of the output.

The following research questions are proposed to validate the thesis statement. The research is anticipated to contribute to the development of personalized recommendation systems using pre-trained transformer-based models. Our study evaluates the effectiveness of using ChatGPT for personalized recommendations of travel destinations and provides insights into the performance. However, we assume that results could be generalized for other LLMs as well. The research also demonstrates the importance of additional task-specific or user-specific input for adapting pre-trained transformer-based models to individual purposes.

The main research question to explore and evaluate the thesis statement is: “How can additional input of task-specific or user-specific texts improve the adaptability and effectiveness of pre-trained models for generating personalized travel destination recommendations?”

This question is divided into sub-research questions to guide the research phases. We follow a design science research (DSR) approach, which comprises the phases awareness, suggestion, implementation, evaluation, and conclusion. However, the DSR phases are not followed in the traditional manner due to the

strong emphasis on experimentation and evaluation versus creating an artifact. The awareness phase mainly considers the investigation of related literature and the identification of research gaps and addresses no specific research question. The same holds for the suggestion phase and the conclusions phase, which include a discussion of potential practical applications and limitations. The considered research questions are as follows:

1) Development/experimentation phase

RQ1: How do the manual inputs and prompts of the user need to be engineered to generate effective, suitable, and desirable travel destinations? Concretely, how to design standardized prompts to increase the quality of the responses?

RQ2: Can LLMs be used to design a customizable recommendation system, where input and prompts are standardized and fine-tuned such that the user does not need to test several approaches but simply fills in the required data for valid proposals?

2) Evaluation phase

RQ3: Which evaluation method and metrics shall be used to assess a satisfactory result?

RQ4: How well does ChatGPT learn individual preferences from case-by-case manual input? How do the experiment participants perceive the generated recommendations?

2. Literature Review

In the following, we illustrate specific findings in the field of the research problem and demonstrate the identified gap for this study.

Zhang and Chen [10] discuss the ongoing efforts to generate natural language explanations, primarily training sentence generation models using user reviews to generate “review-like” explanations. However, it is claimed that generating natural language explanations is still in its infancy, requiring further work to enable machines to provide them. Therefore, language models need to be adapted to generate personalized explanations that require external information. Duan et al. [11] note that while most LLMs excel in handling semantic problems, they exhibit limitations in inferential tasks. Hence, they suggest increasing the amount of pre-training data to enhance model performance. Li et al. [12] additionally emphasize that natural language generation tasks, such as explainable recommendation and review generation, require the use of user and item IDs to distinguish each of them from the other. The item IDs serve to specify each user's preferences “(…) about different item features (e.g., style vs quality), and different items may have different characteristics (e.g., fashionable vs comfortable).” These IDs capture individual preferences and item characteristics, crucial for personalized recommendations. Standard transformer models face challenges in handling personalized natural language generation due to the semantic distinction between IDs and words. To address this, Li et al. [12] propose a modified transformer attention mechanism that offers a simple and efficient approach for personalized generation, leveraging the transformer's language modeling capability for generating explanations in recommender systems. For further research, they note that there are numerous applications for personalized generation, many of which are yet to be fully explored with these models.

The usage of an LLM in a scenario of retrieval-augmented generation (RAG) is explored in Reference [13], where the approach focuses on sustainability evaluation and utilizes specific travel information from documents. In Reference [14], a solution

(called TravelAgent) is suggested involving an LLM-based front end and including specific planning functionalities using a database besides recommendation aspects. Singh et al. [15] suggested a tool called TravelPlanner+ to consider user-related information, preferences, and personal concepts in travel planning, making use of an LLM. Using an LLM also for judging the obtained recommendation, they conclude that personal plans show clear improvement in relevance and suitability, with preference rates up to 74.4% on validation and 87.3% on the test set in comparison to generic plans. Aribas and Daglarli [16] investigated the possibilities of integrating personality models with travel recommendations in an LLM-based framework to enrich the understanding of user preferences and behavior. Using an RAG approach, they achieved promising results for user satisfaction, system accuracy, and the performance rate for different personality traits by using their personalized travel recommendation system.

In Reference [17], a multi-agent-based solution is suggested for providing day-by-day itineraries based on the budget, interests, and travel styles of users. Another recent approach for recommending personalized, context-aware city tourist itineraries, called GPT-TP, is suggested in Reference [18]. In another study by Andreev et al. [19], the problem of biases in LLM-generated travel recommendations is investigated.

Based on the previous explanation [12], it is apparent that there are still further applications available for research to explore personalized recommender systems. In particular, further insights into the alignment of recommendations with user preferences should be addressed by respective studies. In addition, the specific information to be provided for personalized recommendations is not sufficiently discussed in previous studies. Our study addresses the suitability of a transformer model, namely, ChatGPT, for deriving personalized travel recommendations with a focus on evaluating their perceived performance.

3. Research Design

Our study is conducted with a mixed-method approach. The guiding methodology is DSR for suggesting an artifact (an existing LLM with use case-specific prompting) for solving a practical problem. The hypothetical-deductive method [20] is followed, where one starts with the existing knowledge, identifies a problem and/or a gap, and creates, in this case, a hypothesis, which is tested subsequently. Recommendations are evaluated in the testing phase with participant feedback. Thereby, the typical DSR phases have been followed according to Hevner and Chatterjee [21]. The specific application in our study is illustrated in Table 1 based on the DSR framework in [22]. Let us note that there are other concepts and variations of DSR phases [22].

The research was then conducted via the following steps:

- 1) Review the existing knowledge.
- 2) Identify a real-world problem and find its gap in the existing body of knowledge.
- 3) Create (in this case) a hypothesis and test the hypothesis.

After identification of the stated problem, it is scoped via the problem statement. The problem awareness phase is met with the literature review. The problem is broken down into research questions that shall help to test the hypothesis of the thesis statement.

The hypothesis is tested with experiments. First of all, an artifact is needed in the sense of how to use ChatGPT as a support for mundane decision-making for regular persons. In our case, this requires a suitable prompt design. The experiments are then conducted iteratively and refined throughout an experiment cycle. With the information learned and gained through the experiments, an artifact will be further developed.

This study combines experimental, survey, quantitative, and qualitative research techniques to evaluate the effectiveness of personalizing pre-trained models for travel recommendation generation.

The specific artifact to test the hypothesis is intended to include the following:

- 1) Prompt engineering to determine a well-suited prompt and data input for satisfying recommendation generation by ChatGPT. This is done with iterative experiments to find the right input and sequence. According to Short and Short [23]: “Prompt engineering, or ‘prompting’ for short, refers to the process by which a user can fine-tune their inputs to achieve a desired output, often incorporating task descriptions to further guide a specific model. Thus, at its core, prompt engineering is about “how to talk to AI to get it to do what you want.”
- 2) Design a standardized template for data collection as manual input for ChatGPT to learn how to create personalized travel destination recommendations.
- 3) Data collection for manual input is done via the standardized survey tool MS Forms.
- 4) Scale up experiments to $n = 10\text{--}15$ users. Use the data gathered by the survey from the experiment participants and feed it as prompts and inputs to ChatGPT. This is done manually by us.
- 5) For evaluation, a 5-point Likert scale is chosen. According to Robinson [24], it is a scientifically accepted and validated method to measure human attitude. The original Likert scale foresees a list of statements, called items, to be studied. The participants are asked to indicate for each item their level of agreement on a metric scale. Dawes [25] states that a 10-point

Table 1
Design science research framework

DSR phases	Application in this research paper
1. Awareness of problem	Use ChatGPT with prompt engineering to get personalized recommendations
2. Suggestion	Consult state-of-the-art literature to identify research gaps in recommender systems
3. Development	Gather data via a survey for personal travel preferences and develop prompts with experiments to get responses. Based on the data, develop and conduct a second survey to gather feedback from respondents on personalized recommendations
4. Evaluation	Evaluation of survey results
5. Conclusion	Discussion of the results, answering the research questions and conclusion, including limitations, and potential for further research

scale produces marginally lower relative means compared to a 5- or 7-point scale. However, regarding other characteristics, “(...) no scale format produced data with markedly lower variances about the mean.” In addition, the scale should be driven by the applicability of the topic concerned. So, the respondents’ understanding and the expertise of the response item creators are relevant. Given the limitations with regard to both aspects in the context of this paper, a 5-point scale is therefore regarded as the best option.

- 6) Evaluate results via the chosen metric. Also, qualitative feedback from experiment participants is considered.

As specific methodological limitations and risks of our study, we should note that the research focus is to interact with ChatGPT like an ordinary user. Thus, the same interface as publicly available is used, which is <https://chat.openai.com>. In case the model behind the chat changes during the experimentation period, this might affect research and experiments.

4. Experiments and Results

This section describes the actual experiments, the considerations taken, and the steps performed to reach the evaluation stage. The experiments are split into two phases, each with its own activities. The process of experimentation can be separated into different categories according to Table 2.

Table 2
Own experimentation process categories

Category	Description
Ideation	Brainstorming, experiments to find a suitable approach
Preparation	Any step that leads toward experimentation or evaluation
Execution	Executing experiments with the support of ChatGPT and collected data
Evaluation	Studying collected feedback data
Conclusion	Insights, learning, and conclusions from feedback

The experiments are conducted as outlined in the overviews as illustrated in Tables 3 and 4. The following sections explain how they are conducted in detail.

4.1. Creation of personalized travel recommendations by ChatGPT

As part of this stage of the experimentation, several key steps were identified, which are outlined in the following subsections.

4.1.1. Naïve experiments with prompt engineering

We individually explored ChatGPT and conducted experiments to understand how the tool reacts to the input given in the context of generating personalized travel recommendations. The experiments included back-and-forth conversations, where with each exchange more preferences and information were shared, and ChatGPT further tailored its response to increase the personalization. By “chatting” with ChatGPT, relevant information could be gained to prepare the data collection and prompt engineering. Some of these insights include:

Table 3
Experimentation steps part one

Step	Description
1. Ideation	Naïve experiments with prompt engineering
2. Preparation	Design a survey to collect data
3. Preparation	Recruit experiment participants
4. Preparation	Collect participants’ personal and preference information
5. Preparation	Edit and clean data
6. Execution	Create individual prompts for each participant
7. Execution	Generate personalized ChatGPT travel recommendations for each participant
8. Preparation	Prepare personalized prompt and recommendation files for each participant

Table 4
Experimentation steps part two

Step	Description
9. Preparation	Determine the evaluation metric to assess the experiment
10. Preparation	Second survey for evaluation
11. Preparation	Survey evaluation by experiment participants
12. Evaluation	Analyze collected data and evaluate it
13. Conclusion	Draw conclusions

- 1) The relevant attributes that need to be collected with the survey from the participants and that should be part of the prompts.
- 2) An understanding of some of the assumptions ChatGPT autonomously takes, like budgets given to be consumed by the whole party of travelers, for example, when prompting to create a recommendation for a couple. The budget input is then assumed to cover the costs of the two traveling persons. Gender is probably not a relevant input unless someone would like to travel to a country where one gender is not at liberty to travel by itself. Furthermore, to avoid gender bias, this attribute was not considered further.

If concrete parameters are sought, concrete prompt engineering must be used. Prompt engineering in this context is understood as:

- 1) Giving ChatGPT a persona to act as, like “You are a travel agent, responsible for proposing a personalized travel destination recommendation.”
- 2) Provide clear input about expectations. For instance, if one only starts with the persona, ChatGPT will show a list of questions, which is not the most effective approach. Instead, anticipating the input needed will speed up the conversation.
- 3) Clear input is also needed to achieve the details and parameters wished for. In this context, parameters refer to details such as the number of destinations to be recommended.
- 4) The output varies in detail and description style with just minor differences in inputs. Especially if one wants to make

sure that an input was considered, it is suggested to prompt getting proof for it in the response and in what detail. An example of this is the budget. Not asking about it will not reveal any details on how it was considered. If a request for a budget breakdown is included in the prompt, one will be provided.

- 5) Recommendations are often but not always padded, with that lively description of a location, like the following ChatGPT-generated output: “Zanzibar is an island paradise with some of the most beautiful beaches in the world. You can relax on the white sand beaches, go snorkeling or diving in the crystal-clear waters, or explore the island’s history and culture. Zanzibar is also a great place to try local spices and seafood, and to experience the warm hospitality of the locals. You can find affordable accommodations in guesthouses and homestays, which will allow you to immerse yourself in the local culture.”

4.1.2. Survey for data collection

Based on the insights gained from the first experiments, a survey was designed using Microsoft Forms as the tool to gather data. The survey asked for 16 inputs including personal information and preference information regarding travel activities. Some of the information was not used for the experiments, but only to identify the participants, to later get their feedback, like their full name and e-mail address. Other input was collected to inquire about personal preferences, like the number of travelers, starting location, budget, number of travelers, age, preferred and not preferred destinations, activities, scenery to include, and mode of transportation preferred. The survey invitations were shared within the private and professional ecosystem of the authors.

The survey link is shared via e-mail, and the participants are given seven days to participate. The Excel file that is automatically generated by Microsoft Forms is used for the continuation of the research. The survey is closed on day eight, with 31 participants sharing their input to conduct the experiments.

Figure 1 shows the age distribution; 45%, most of the participants were in the age group 30–34, followed by 29% aged 25–29 and 19% aged 35–39. 61% of the participants were male and 39% female.

Figure 2 shows the traveler type, where 30% plan to travel as a couple, followed by 19% with their family, and solo and friends’ travelers with 13%. Figure 3 shows the chosen month for the travel destination recommendation. The budget per trip is, on average, 3’306 CHF (see Table 5).

For further processing, some preparation (editing and cleaning) of the travel preference information of participants is

Figure 1
Age distribution survey

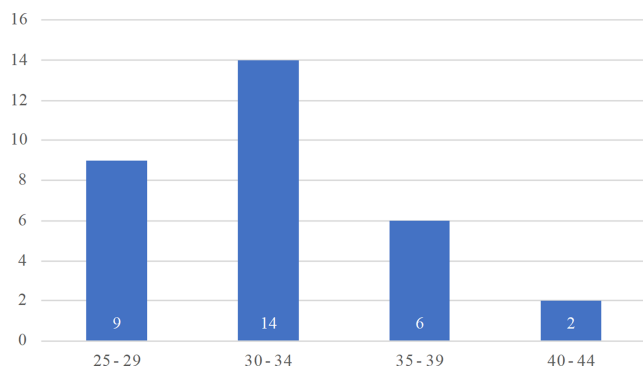


Figure 2
Survey of traveler types

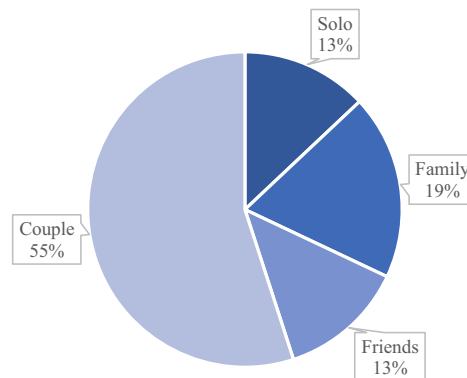


Figure 3
Survey of chosen travel periods

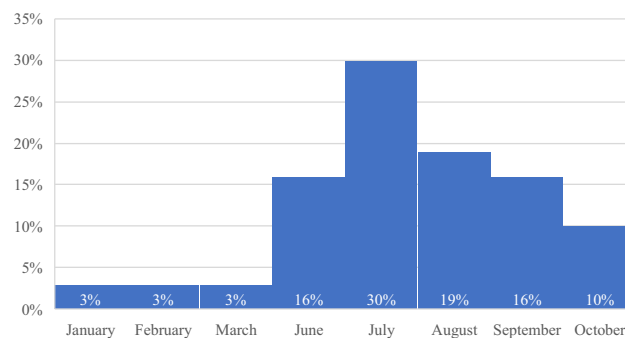


Table 5
Survey of budget preferences

Budget	CHF
Average	3306
Min	600
Max	10000
Per day average	290
Per day min	71
Per day max	679

required. New attributes are created to be used in the experiment. For example, only the first name is extracted to ensure that no detailed personal data, like a full name, is fed as input to ChatGPT.

The currency is extracted from the answer, but as the exchange rate of the local currency CHF and USD was in the range 1–1.3 during recent years, we considered the impact as small enough to neglect it so that the attribute was dropped. Participants who indicate traveling with children are flagged to personalize their prompts clearly. Attributes without input are flagged with “no data.”

All relevant inputs are labeled with an attribute number, which supports the generation of individual prompts for each participant, based on their dataset. The attribute is the data (answer) collected from the participants via the survey. An overview of the attributes, survey questions, answers, and data types, as well as a sample set, is shown in Table 6.

Table 6
Attributes for personalization and overview of inputs from the survey

Attributes for personalization	Attribute name	Answer type	Data type	Sample dataset
Not used	ID	Auto generated	Integer	<i>1</i>
Not used	Name Full	Open	String	<i>Sarah Castratori</i>
Attribute1	First Name Only	Open	String	<i>Sarah</i>
Attribute2	E-mail	Open	String	castratori@123.ch
Attribute3	Age (if you prefer not to say, skip this question)	Open	Integer	<i>39</i>
Attribute4	Where will your journey start?	Multiple-choice and open (Switzerland, France, Germany, other)	String	<i>Switzerland</i>
Attribute5	What kind of trip do you plan?	Multiple-choice (solo, couple, friends, family)	String	<i>Solo</i>
Attribute6	Number of travelers	Open	Integer and string (adult, children info)	<i>1</i>
Attribute7	Children Y/N	Inferred from Attribute 6	Integer	<i>n</i>
Attribute8	Travel duration in days	Open	Integer	<i>12</i>
Attribute9	Travel budget	Open	Integer and string (currency)	<i>2200</i>
not used	Currency	Inferred from Attribute 9	String	<i>CHF</i>
Attribute11	Travel period	Multiple-choice (list of all calendar months)	String	<i>August</i>
Attribute12	Destination preferences	Multiple-choice (surprise vs open (individual user preference))	String	<i>No</i>
Attribute13	Do you want any particular scenery to be part of your travel experience?	Multiple-choice (none, nature, beach, mountains, urban, open/individual user preference)	String	<i>Nature, beach, urban (at least partially explore city life)</i>
Attribute14	Any specific activities?	Multiple-choice (water-sport, nightlife, hiking, open (individual user preference))	String	<i>Pandas</i>
Attribute15	If you plan to stay on your continent of residency (likely Europe), please indicate if you have any transport mode preferences to consider.	Multiple-choice (public transport, car, mixed mode, by foot/bicycle, no preference, open (individual user preference))	String	<i>No data</i>
Attribute16	Any previous travel destinations and experiences you loved and want to be considered as good experiences for the generated recommendation?	Open	String	<i>No data</i>
Attribute17	Destinations to avoid	Open	String	<i>No data</i>
Attribute18	Anything else you want the travel recommendation to consider?	Open	String	<i>No data</i>

4.1.3. Prompt engineering

Based on the generated attributes that are aimed to be used as additional input for personalization, a template for the prompts is designed. The template uses natural language to express a prompt the way a user would. In place of the personal information and user preferences, the attribute is references. Where attributes are optional, conditional statements are set so that the sentence is only printed if data are available. This is illustrated in Figure 4.

For each participant, an individual prompt needs to be created. Instead of writing the prompts manually, ChatGPT is used for this task. It is given the prompt template, the datasets, flagged with ID and Name, so they can be connected throughout the experiments, as well as this general prompt to be considered for all datasets (see Figure 5).

This generic prompt serves to ensure that exactly two recommendations are generated and that the budget is detailed. With the sample dataset shown in Table 6 and the prompts shown in Figures 4 and 5, ChatGPT created the following prompt (Figure 6). To note, this dataset does not consist of the full array of possible attributes, but for personal data protection reasons, it is the chosen example.

The prompts were created for all 31 participants. Due to the capabilities of the model, the datasets had to be fed in several batches; all 31 at once would not have been feasible due to the input size limitations of the used model. Some apparent memory loss of ChatGPT had to be overcome by reminding the general prompt instructions (Figure 5) as well. All prompts were generated within a single chat, which reflects a typical private usage scenario with potentially unrelated content entered into the LLM before.

4.2. Evaluation of experiments

To evaluate the experiments, it is considered important not to simply check whether ChatGPT answered the question of the

individual prompt with all requested details. This was already secured by thorough prompt engineering. It is considered relevant that the experiment is evaluated by quantitative and qualitative feedback from the participants themselves. Hence, the experiment continues with part two, which is described in the following sections.

4.2.1. Assessment of experimental results

As indicated before, we used a 5-point Likert scale for the evaluation in most questions. Two questions offer multiple-choice selection. In addition, the questions have an optional text box to give additional qualitative feedback, which shall help to understand the rating and give greater insights into the participants' perception and feelings about the experiment and recommendations they receive.

4.2.2. Second survey for evaluation

The insights that are intended to be gained are categorized into five dimensions, and the questions are designed accordingly (see Table 7).

The list of questions and answer types is illustrated in Table 8. It is a lengthy survey that bears the risk of survey fatigue and dropouts. The risk is taken to secure gaining valuable insight instead of shallow feedback. MS Forms was again chosen as the tool.

Each participant received the Word file with their prompt, their travel destination recommendation, and the link to the survey by e-mail. Seven days were granted to provide feedback. A reminder was sent after 5 days. Seventeen participants provided their feedback, which ensured that our goal of 10–15 was reached.

The details of the participation rates in the evaluation survey are as follows: Out of 44 people who received the call to participate in the experiment, 31 answered the survey, resulting in a participation rate of 70.5%. Out of the 26 survey participants

Figure 4
Prompt template

```
Hey, I am "Attribute1", and I am "Attribute3" years old. Please create
personalized holiday destination recommendations based on my individual
input. I will depart from "Attribute4" as a "Attribute5" Traveler.

[If Attribute8 exists: ]My trip should last "Attribute8" days.

[If Attribute9 exists: ]I want to spend a maximum of "Attribute9" CHF.

The trip should take place in the month of "Attribute11", and the
destination should feature "Attribute12".

[If Attribute13 exists: ]I am interested in experiencing "Attribute13".

[If Attribute14 exists: ]The following activities should be considered for
my trip: "Attribute14".

When it comes to travel mode, I would like to use "Attribute15".

[If Attribute16 exists: ]I'd like to share with you what I enjoyed during
previous trips, which was "Attribute16".

[If Attribute17 exists: ]Please avoid the following places: "Attribute17".

[If Attribute18 exists: ]And also, please consider "Attribute18".

Thanks!
```

Figure 5
Generic add-on prompt for individual travel recommendations

```
For each of these ID's generate individual travel destination
recommendations. include at least two options for the travel destination.
for the budget be specific for how it is broken down. indicate the ID in
your answer. And always print the input given by me followed by your answer
together.
```

Figure 6
Individual prompt example

ID 1 Sarah

Hey, I am Sarah, and I am 39 years old. Please create personalized holiday destination recommendations based on my individual input. I will depart from Switzerland as a Solo Traveler.

My trip should last 12 days.
I want to spend a maximum of 2'200 CHF.
The trip should take place in the month of August, and the destination should feature Nature, Beach, and Urban |(at least partially explore city life).

I am interested in experiencing Pandas.

Thanks!

Table 7
Dimensions of feedback evaluation

Dimension	Question	Answer type	Answer options
Name	Participant's name, to tie the feedback to their recommendation	Open	
S	How satisfied are you overall with the personalized recommendation?	Likert 1–5	5—Very satisfied, 4—Satisfied, 3—Neither satisfied nor dissatisfied, 2—Somewhat dissatisfied, 1—Very dissatisfied
S	Feel free to share feedback on why you are satisfied/dissatisfied with the recommendation	Open	Free text
S	How satisfied are you with the personalized recommendation for the following aspects:		
S	Preferred destination?	Likert 1–5	5—Very satisfied, 4—Satisfied, 3—Neither satisfied nor dissatisfied, 2—Somewhat dissatisfied, 1—Very dissatisfied
S	Preferred scenery?		
S	Desired activities?		
S	Preferred transport mode?		
S	How satisfied are you with the provided budget breakdown for each destination?		
S	If you are not satisfied or just partially, feel free to share what would improve your satisfaction with the budget insights provided.	Open	Free text
T	How much do you trust the recommendation(s) regarding the following aspects:		
T	Consideration of personal preferences?	Likert 1–5	5—Strongly trust, 4—Trust, 3—Neither trust nor distrust, 2—Distrust, 1—Strongly distrust
T	Authenticity and correctness?		
T	Budget adherence?		
T	What would further increase your trust in the recommendation?	Open	Free text
PC	Do you consider that the generated recommendation(s) will save you time in your decision-making process to pick a travel destination?	Likert 1–5	5—Strongly agree, 4—Agree, 3—Neither agree nor disagree, 2—Disagree, 1—Strongly disagree
PC	If not, what extra information would you wish to find in the recommendation that would help you save time?	Open	Free text
A	How likely is it that you would choose one of the recommendations as one of your next travel destinations?	5-point Likert scale	5—Very likely, 4—Likely, 3—Neither likely nor unlikely, 2—Unlikely, 1—Very unlikely

(Continued)

Table 7
(Continued)

Dimension	Question	Answer type	Answer options
A	If not or unlikely, why?	Open	Free text
PC	Would you consider ChatGPT or similar tools in the future to help you provide personalized travel recommendations?	5-point Likert scale	5—Strongly agree, 4—Agree, 3—Neither agree nor disagree, 2—Disagree, 1—Strongly disagree
PC	If not, why?	Open	Free text
PC	When it comes to creating your own recommendation with the support of ChatGPT and similar models, would you rather:	Multiple-choice	Prefer to interact with a chatbot directly (like OpenAI ChatGPT web interface); prefer to use a tailor-made “travel recommendation” website/mobile app
PC	Would you consider ChatGPT or similar tools in the future to generate personalized recommendations other than for travel purposes? (use the “other” box to tell us what you would or already use it for, in the context of recommendations)	Multiple-choice	No Yes, for other personalized service recommendations Yes, for personalized product recommendations free text
AR	Do you have any additional remarks regarding this experiment? Especially the recommendation you received?	Open	Free text

Table 8
Dimensions of feedback evaluation and expected insights

Dimension	Expected insights	Reviewed with
Satisfaction (S)	How satisfied are the participants with the recommendation they receive?	5-point Likert scale and open text
Trust (T)	How much do the participants trust the recommendation?	
Agreement (A)	How likely are the participants to follow/choose the recommendation?	
Preference/_ Consideration (PC)	Will the participants consider ChatGPT as a tool for personalized recommendations in the future?	Yes/no
Additional Remarks (AR)	Any other feedback the participants want to share	Open text

who were selected to support the evaluation of the experiment by reviewing, validating, and rating their personalized recommendation, 17 reported back, resulting in a participation rate of 65.5%.

4.2.3. Analysis of collected data

The analysis is first presented for the quantitative data collected. In the second part, a qualitative analysis is presented, based on the open-text feedback that is gathered from the participants. The collected data are analyzed with statistical means. Table 9 outlines the chosen metrics, and Table 10 shows the results of the analysis.

The personalized travel recommendations are perceived with satisfaction, with an overall average of 81.86% and an average score of 4.08 out of 5. Table 11 provides insights into the percentage and distribution of replies by summarizing the frequency of terms in a dataset, identifying the highest distribution toward the rating 4 at 50.5% and rating 5 at 33.0%. The areas where the recommendations scored highest are personal preferences understanding, overall satisfaction, and potential for timesaving, whereas the areas to be improved are transport mode options, budget breakdown details, and budget adherence.

Personalized Recommendation Satisfaction: With an average score of 4.35 and 94.12% satisfaction, participants show positive

satisfaction. The standard deviation indicates acceptance of the recommendation. No dissatisfaction is expressed.

Preferred Destination/Scenery/Activities: The highest average scores are found here, with a range from 4.24 to 4.47, together with an average satisfaction score across the three questions of over 88.24%, while dissatisfaction did not occur.

Preferred Transport Mode: An average score of 3.88 accompanied by a satisfaction level of 70.59%. The standard deviation and mode suggest a wider range of replies, indicating some users were not entirely satisfied with the suggested transport options, with a dissatisfaction level of 11.76%.

Satisfaction with Budget Breakdown: With 3.94, the score is lower than the total average with a higher standard deviation, indicating that the participants' opinions diverge. With 76.47% satisfaction and 11.76% dissatisfaction, there is potential for improving the budget planning aspect.

Consideration of Personal Preferences and Authenticity and Correctness: These trust metrics scored averages of 4.12 and 3.88, with relatively high satisfaction of 82.35% and 76.47% and low dissatisfaction of 0% and 5.88%.

Budget Adherence: Lowest-scoring category with a score of 3.41 and the highest dissatisfaction rate with 23.53%. With 1.06, the highest standard deviation is observed, suggesting a wide range of responses, as well as the lowest mode and median with 3.0.

Table 9
Metric explanation for quantitative analysis

Column name	Explanation
Average score	Mean of the scores provided by the participants on a scale of 1–5
Standard deviation	Indicates the dispersion or variability in the scores. Lower values suggest scores close to the mean; higher values suggest a wider range
% of satisfaction	Percentage (sum) of participants who rated their satisfaction level as either 4 or 5
% of dissatisfaction	Percentage (sum) of participants who rated their satisfaction level as either 1 or 2
Mode	The most frequently occurring value in the satisfaction scores
Median	Middle value in a sorted list of numbers. Divides the dataset into two equal halves

Table 10
Participant feedback for quantitative analysis

Dimension	Question	Average score	Standard deviation	Mode	Median	Satisfaction (score 4–5)	Dissatisfaction (score 1–2)
S	How satisfied are you overall with the personalized recommendation?	4.35	0.61	4	4	94.12%	0.00%
	How satisfied are you with the personalized recommendation for the following aspects:						
S	Preferred destination?	4.47	0.62	5	5	94.12%	0.00%
S	Preferred scenery?	4.47	0.72	5	5	88.24%	0.00%
S	Desired activities?	4.24	0.56	4	4	94.12%	0.00%
S	Preferred transport mode?	3.88	0.99	4	4	70.59%	11.76%
S	How satisfied are you with the provided budget breakdown for each destination?	3.94	0.97	4	4	76.47%	11.76%
	How much do you trust the recommendation(s) regarding the following aspects:						
T	Consideration of personal preferences?	4.12	0.70	4	4	82.35%	0.00%
T	Authenticity and correctness?	3.88	0.78	4	4	76.47%	5.88%
T	Budget adherence?	3.41	1.06	3	3	47.06%	23.53%
PC	Do you consider that the generated recommendation(s) save you time in your decision-making process to pick a travel destination?	4.06	0.43	4	4	94.12%	0.00%
A	How likely is it that you would choose one of the recommendations as one of your next travel destinations?	3.94	1.03	5	4	70.59%	11.76%
PC	Would you consider Chat-GPT or similar tools in the future to help you provide personalized travel recommendations?	4.18	0.95	4	4	94.12%	5.88%
	Overall average per metric	4.08	0.79	4.2	4.1	81.86%	5.88%

Timesaving Potential of Recommendations: Average score of 4.06 and 94.12% satisfaction rate. The system's ability to accelerate decision-making is perceived as positive.

Likelihood of Choosing Recommendations: The average score is 3.94 with a 70.59% satisfaction rate and 11.76% dissatisfaction. The standard deviation of 1.03 indicates a wider spread of acceptance.

Consideration of AI-generated Personalized Recommendations in the future: Based on an average rating of 4.18 and 94.12% satisfaction, supported by a 5.88% dissatisfaction rate, the participants express a high likelihood of using such a similar solution in the future.

Qualitative feedback provided us with deeper insights into why the participants were scoring as they chose; open questions

Table 11
Rating frequencies

Rating	Frequency	Percentage
4	101	50.5
5	66	33.0
3	25	12.5
2	11	5.5
1	1	0.5

are offered. Seven such questions are shared for the participants to give free-text feedback. They are listed in Table 12. The summary of the feedback is paraphrased to convey the general input received. If input was received for a different subject, the response numbers were corrected. Where not stated otherwise, the sentiment was expressed only once.

Satisfaction: For the general satisfaction with their recommendation, all 11 feedbacks are positive, with some additional input on how to improve the satisfaction further.

Two participants were requesting more detailed recommendations, highlighting a wish to find clearer activity, hotel, and budget information. One was confirming that they had already done the suggested trip or would go for one of the recommendations.

Budget: All seven feedbacks express negative sentiments. The budget insights received are assessed as vague, doubtful, or the accuracy or coverage is questioned:

Trust: When prompted to share how trust can be increased, 13 participants share their needs. Five of them confirm a need for more clarity about the budget. Four would like direct links to offers, and three want to know the sources of the details they received. Improving areas to increase trust can be associated with specifying sources for information, including user habits, and offering better estimation.

Timesaving: Most respondents agree that the recommendation saves time in decision-making. Three people suggest improving this further by providing weather forecasts, more detailed budgets, and direct links to offers and comparisons.

Not choosing one of the provided recommendations: The three responses claim that the destination does not fit the season, that they prefer to base their decision on others' experience or travel agents, or that it is too generic.

Future consideration of the solution: One person rated the likelihood of choosing any of the recommendations with 1, stating that there is no value in it, due to the marketing content on the internet.

Additional feedback on experiments and recommendations: Eight participants share additional reflections in support of the

experiments and further constructive feedback. The positive sentiments are praise for the experiments, the great idea, the big potential, and the inspiration received. The critical input once again emphasizes the need for a direct link to offers and information, clearer insights into the budget needed, more details, and a wish to chat with a person to take the recommendations further.

5. Discussion

The following discussion follows the research questions outlined in Section 1.2.

RQ1 (development of prompts and generation of personalized recommendations): Through various iterations of experiments and trials, attributes for personalization were identified as valuable to generate effective travel destination recommendations. ChatGPT was employed for this process. These attributes were then used for the survey to collect the data from the participants. While ChatGPT would not need the prompt in the human conversational style that was used in our study, it was considered useful as users can better assess the whole chat conversation. When scaling up such an experiment, this could be simplified by using a persona-specific prompt like "act as a travel agent," and the list of individual attributes could be provided in addition. Since no participant gave negative feedback about the style or understanding of the prompts and recommendations, it can be implied that they were constructive.

Looking at the feedback received regarding the dimensions, where the users lack trust and satisfaction, it is clear that the generic prompt that steers the creation of a budget overview and two recommendations could be enriched further. This would ensure that the recommendations would receive a higher score in these dimensions and in their perceived quality.

RQ2 (usability of ChatGPT for a customizable recommendation system): Once the attributes and prompts for personalized recommendations are specified, the suggested solution can be regarded as a standardized system that is fine-tuned to the users' individual preferences. It enables upscaling the recommendations easily. User-specific engineering of the prompt in terms of how and what to ask becomes redundant, and the template can be followed.

RQ3 (evaluation method and metric): Choosing a survey tool enabled an efficient and easy way to collect data and query the participants about their opinions. The Likert scale supported the acquisition of quantitative insights and clearly drew attention to the lack of trust and the budget. The many open feedback options enabled the gathering of quantitative input. This clearly revealed the wishes of the participants to enrich the recommendations with desired details, like sources and budget insights. If the experiments

Table 12
Topics of open questions from feedback survey

Dimension	Question topic (paraphrased)	Responses
S	Feedback regarding satisfaction/dissatisfaction	11
S	Improvement wishes for the budget insights	7
T	How to further increase trust in recommendations	13
PC	If no timesaving is seen, how could it be achieved	3
A	Reason to not choose one of the recommendations	1
PC	Reason not to consider ChatGPT or a similar tool for future travel recommendations	2
AR	Additional remarks for the experiment and the recommendation	8

were scaled up, a shorter evaluation survey should be considered to ensure more participants' feedback.

RQ4 (evaluation of the experiment regarding personalized travel destination recommendations): The effectiveness of the ChatGPT-generated personalized travel recommendations did well in terms of user satisfaction, desire consideration, and timesaving, as shown in the analysis section of the participant feedback. However, the scores for authenticity, correctness, and budget adherence were less satisfying. To increase the degrees of trust and user satisfaction, more detailed and specific recommendations are directly requested by the survey participants, as well as enhanced granularity for the budget breakdown, and more transparency by sharing sources. The expressed wishes to incorporate links to resources and booking offers suggest a business potential that directly connects recommendations to booking platforms or travel agents, for example. Hence, a future method of such a travel destination recommendation should consider enriching the information that is shared with users with the discussed sources and details.

When thinking about implementing our approach as a business model, with a web or mobile interface, it is seen as key that the user could customize how important certain dimensions are for them and, by doing so, steer the depth of details received on budget or activities. The needs of the details may vary for different aspects of the dimensions, which is concluded by the fact that some of the participants are content with their recommendations, while other participants were seeking more details. Having the possibility for manual intervention about the depth of information would allow different degrees of detail and support the respective user-specific needs. Regardless of these important factors, the participants gave positive feedback regarding the experiments, praising it as a great idea and something with big potential.

Let us finally discuss some limitations of the conducted experiments and possibilities to address them: Since the participants did not have the chance to directly conduct the experiments, they could not ask follow-up questions. If they had, they would have gained clarity about what is covered in the budget and could have directly managed to receive more details. A simple prompt asking "is this budget covering all travelers" would have provided clarity and increased trust. However, this would have come with a payoff of more time spent on the task. Further prompt engineering could potentially increase user satisfaction and trust in this regard. If the idea were to be put into production, with an application or web interface, it would be an apparent need of the user to have very clear insights into the reasoning of the answers with references and sources, the scope of the information provided, in terms of budget, traveling party, and details considered. This was made clear with several of the open feedback suggestions.

The chosen transformer model, with its cutoff date, cannot provide the latest information about political stability, environmental factors, and safety concerns (training for ChatGPT 3.5 only includes data up until September 2021). This is seen as a limitation, as there is a risk of recommending someone to go to an unsuitable destination. ChatGPT is mitigating this, however, by expressing that its own research must be conducted. With the service offering web browsing for paying ChatGPT Plus customers, this is further alleviated.

Another limitation was the usage of ChatGPT via the web interface and not with a more technically advanced method like connecting bigger datasets via Application Programming Interface (API). This resulted in the need to repeatedly feed data

and re-trigger a continuation of the responses. This, however, was accepted as part of the experiment since an everyday user's perspective was taken and not that of more advanced information system experts.

Moreover, we would like to mention the limitation due to the final sample size of 17 participants and the selection from the authors' personal and professional networks. We would recommend applying a larger sample size and considering random or stratified samples for future research to improve external validity.

6. Conclusions

Our study focused on AI opportunities to support human decision-making by reducing time and presenting clear and specific suggestions, helping people overcome decision paralysis. In the context of travel recommendations, a solution like ChatGPT has considerable potential. It can simplify travel planning by providing comprehensive and relevant information in natural language, reducing users' need to consult multiple sources.

Through the experiments conducted with real user preferences and data, as well as their personal evaluations, this research paper can confirm the thesis statement: "Additional input of task-specific or user-specific texts will enhance the adaptability and effectiveness of pre-trained models for personalized travel destination recommendation generation." Although the results of this research contribute to the field, they only represent a first attempt to apply additional input for personalized recommendations in the travel area. Consequently, further empirical research would be required to gain more insights and detailed results to increase the satisfaction, trustworthiness, and reliability of personalized travel recommendations generated by ChatGPT and similar recommender systems. This should be done by addressing the mentioned limitations, such as gathering more data for experiments to improve the quality and diversity (which should also allow for statistical validation) or focusing on the sentiment analysis aspect, which was briefly mentioned in the literature review part, to incorporate more data in the form of reviews.

In addition, we recommend a more refined analysis of prompting approaches and provided personal information (e.g., in the form of an ablation study), which could also benefit from the insights gained from our current study. In particular, in future studies, LLMs could be asked for more detailed recommendations, especially regarding the budget needed, and a higher transparency of provided information, for example, by asking for used sources.

Ethical Statement

Data were taken from individuals in a non-clinical context and without any inclusion of private data for further processing; that is, only anonymized data were processed in the preparation of the article. According to Swiss law (Art. 2 of the Federal Act on Research Involving Human Beings [Human Research Act]) of September 30, 2011 [Status as of September 1, 2023], cf. <https://www.fedlex.admin.ch/eli/cc/2013/617/en>, research projects involving anonymized datasets do not require approval (<https://www.bag.admin.ch/en/human-research-approval-of-research-projects>).

In addition, all participants were informed about the purpose of data collection and the related study and provided verbal consent to provide the data used in the form of their responses to the survey questions.

Conflicts of Interest

The authors declare that they have no conflicts of interest to this work.

Data Availability Statement

Data are available from the corresponding author upon reasonable request.

Author Contribution Statement

Sarah Castratori: Conceptualization, Methodology, Formal analysis, Investigation, Data curation, Writing – original draft, Project administration. **Domenico Dongiovanni:** Conceptualization, Methodology, Formal analysis, Investigation, Data curation, Writing – original draft, Project administration. **David Grossenbacher:** Conceptualization, Methodology, Formal analysis, Investigation, Data curation, Writing – original draft, Project administration. **Thomas Hanne:** Conceptualization, Validation, Writing – review & editing, Supervision, Project administration.

References

- [1] Annepaka, Y., & Pakray, P. (2025). Large language models: A survey of their development, capabilities, and applications. *Knowledge and Information Systems*, 67(3), 2967–3022. <https://doi.org/10.1007/s10115-024-02310-4>
- [2] Cheng, J., Ghate, K., Hua, W., Wang, W. Y., Shen, H., & Fang, F. (2025). REALM: A dataset of real-world LLM use cases. In *Findings of the Association for Computational Linguistics: ACL 2025*, 8331–8341. <https://doi.org/10.18653/v1/2025.findings-acl.437>
- [3] McGinness, L., Baumgartner, P., Onyango, E., & Lema, Z. (2025). Highlighting case studies in LLM literature review of interdisciplinary system science. In *AI 2024: Advances in Artificial Intelligence: 37th Australasian Joint Conference on Artificial Intelligence*, 29–43. https://doi.org/10.1007/978-981-96-0348-0_3
- [4] Chkirbene, Z., Hamila, R., Gouisse, A., & Devrim, U. (2024). Large language models (LLM) in industry: A survey of applications, challenges, and trends. In *2024 IEEE 21st International Conference on Smart Communities: Improving Quality of Life Using AI, Robotics and IoT*, 229–234. <https://doi.org/10.1109/HONET63146.2024.10822885>
- [5] Patil, R., & Gudivada, V. (2024). A review of current trends, techniques, and challenges in large language models (LLMs). *Applied Sciences*, 14(5), 2074. <https://doi.org/10.3390/app14052074>
- [6] Huang, L., Yu, W., Ma, W., Zhong, W., Feng, Z., Wang, H., . . . , & Liu, T. (2025). A survey on hallucination in large language models: Principles, taxonomy, challenges, and open questions. *ACM Transactions on Information Systems*, 43(2), 42. <https://doi.org/10.1145/3703155>
- [7] Shafik, W. (2024). Introduction to ChatGPT. In A. J. Obaid, B. Bhushan, M. S. & S. S. Rajest (Eds.), *Advanced applications of generative AI and natural language processing models* (pp. 1–25). IGI Global Scientific Publishing. <https://doi.org/10.4018/979-8-3693-0502-7.ch001>
- [8] Hu, K. (2023). ChatGPT sets record for fastest-growing user base—Analyst note. *Reuters*. <https://www.reuters.com/technology/chatgpt-sets-record-fastest-growing-user-base-analyst-note-2023-02-01/>
- [9] Dwivedi, Y. K., Kshetri, N., Hughes, L., Slade, E. L., Jeyaraj, A., Kar, A. K., . . . , & Wright, R. (2023). Opinion paper: “So what if ChatGPT wrote it?” Multidisciplinary perspectives on opportunities, challenges and implications of generative conversational AI for research, practice and policy. *International Journal of Information Management*, 71, 102642. <https://doi.org/10.1016/j.ijinfomgt.2023.102642>
- [10] Zhang, Y., & Chen, X. (2020). Explainable recommendation: A survey and new perspectives. *Foundations and Trends® in Information Retrieval*, 14(1), 1–101. <https://doi.org/10.1561/15000000066>
- [11] Duan, J., Zhao, H., Zhou, Q., Qiu, M., & Liu, M. (2020). A study of pre-trained language models in natural language processing. In *2020 IEEE International Conference on Smart Cloud*, 116–121. <https://doi.org/10.1109/SmartCloud49737.2020.00030>
- [12] Li, L., Zhang, Y., & Chen, L. (2021). Personalized transformer for explainable recommendation. 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)*, 4947–4957. <https://doi.org/10.18653/v1/2021.acl-long.383>
- [13] Banerjee, A., Satish, A., & Wörndl, W. (2025). Enhancing tourism recommender systems for sustainable city trips using retrieval-augmented generation. In *Recommender Systems for Sustainability and Social Good: First International Workshop*, 19–34. https://doi.org/10.1007/978-3-031-87654-7_3
- [14] Chen, A., Ge, X., Fu, Z., Xiao, Y., & Chen, J. (2024). TravelAgent: An AI assistant for personalized travel planning. *arXiv*. <https://doi.org/10.48550/ARXIV.2409.08069>
- [15] Singh, H., Verma, N., Wang, Y., Bharadwaj, M., Fashandi, H., Ferreira, K., & Lee, C. (2024). Personal large language model agents: A case study on tailored travel planning. In *Proceedings of the 2024 Conference on Empirical Methods in Natural Language Processing: Industry Track*, 486–514. <https://doi.org/10.18653/v1/2024.emnlp-industry.37>
- [16] Aribas, E., & Daglarli, E. (2024). Transforming personalized travel recommendations: Integrating generative AI with personality models. *Electronics*, 13(23), 4751. <https://doi.org/10.3390/electronics13234751>
- [17] Banerjee, A., Satish, A., Aisyah, F. N., Wörndl, W., & Deldjoo, Y. (2026). Collab-REC: An LLM-based agentic framework for balancing recommendations in tourism. *ACM Transactions on Recommender Systems*. Advance online publication. <https://doi.org/10.1145/3837862>
- [18] González-Sanz, E. L., Cantador, I., & Bellogín, A. (2025). LLM-based generation of personalized, context-aware city tourist itineraries: A user study with GPT trip planner. In *Workshop on Recommenders in Tourism*.
- [19] Andreev, H., Kosmas, P., Livieratos, A. D., Theocharous, A., & Zopiatas, A. (2025). Destination (un)known: Auditing bias and fairness in LLM-based travel recommendations. *AI*, 6(9), 236. <https://doi.org/10.3390/ai6090236>
- [20] Orbik, Z. (2025). Hypothetico-deductive method in management sciences: Its development and application. *Scientific Papers of Silesian University of Technology, Organization & Management*, 2025(229), 379–396. <http://dx.doi.org/10.29119/1641-3466.2025.229.21>

- [21] Hevner, A., & Chatterjee, S. (2010). *Design research in information systems: Theory and practice*. USA: Springer. <https://doi.org/10.1007/978-1-4419-5653-8>
- [22] Nagle, T., Doyle, C., Alhassan, I. M., & Sammon, D. (2022). The research method we need or deserve? A literature review of the design science research landscape. *Communications of the Association for Information Systems*, 50(1), 358–395. <https://doi.org/10.17705/1CAIS.05015>
- [23] Short, C. E., & Short, J. C. (2023). The artificially intelligent entrepreneur: ChatGPT, prompt engineering, and entrepreneurial rhetoric creation. *Journal of Business Venturing Insights*, 19, e00388. <https://doi.org/10.1016/j.jbvi.2023.e00388>
- [24] Robinson, J. (2023). Likert scale. In F. Maggino (Ed.), *Encyclopedia of quality of life and well-being research* (pp. 3917–3918). Springer. https://doi.org/10.1007/978-3-031-17299-1_1654
- [25] Dawes, J. (2008). Do data characteristics change according to the number of scale points used? An experiment using 5-point, 7-point and 10-point scales. *International Journal of Market Research*, 50(1), 61–104. <https://doi.org/10.1177/147078530805000106>

How to Cite: Castratori, S., Dongiovanni, D., Grossenbacher, D., & Hanne, T. (2026). Using ChatGPT with Prompt Engineering for Personalized Travel Destination Recommendations. *Journal of Data Science and Intelligent Systems*. <https://doi.org/10.47852/bonviewJDSIS62028449>