

Exploring the LLM Survey Simulation Pipeline with the QSTN Framework

Session 2, Jens Rupprecht & Maxi Kreutner

QSTN

A Modular Framework for Robust Questionnaire Inference with LLMs

QSTN: A Modular Framework for Robust Questionnaire Inference with Large Language Models

Maximilian Kreutner¹, Jens Rupperecht¹, Georg Ahnert¹,
Ahmed Salem¹, Markus Strohmaier^{1,2,3}

¹University of Mannheim, ²GESIS - Leibniz Institute for the Social Sciences, ³CSH Vienna

Abstract

We introduce QSTN, an open-source Python framework for systematically generating responses from questionnaire-style prompts to support in-silico surveys and annotation tasks with large language models (LLMs). QSTN enables robust evaluation of questionnaire presentation, prompt perturbations, and response generation methods. Our extensive evaluation (>40 million survey responses) shows that question structure and response generation methods have a significant impact on the alignment of generated survey responses with human answers. We also find that answers can be obtained for a fraction of the compute cost, by changing the presentation method. In addition, we offer a no-code user interface that allows researchers to set up robust experiments with LLMs *without coding knowledge*. We hope that QSTN will support the reproducibility and reliability of LLM-based research in the future.

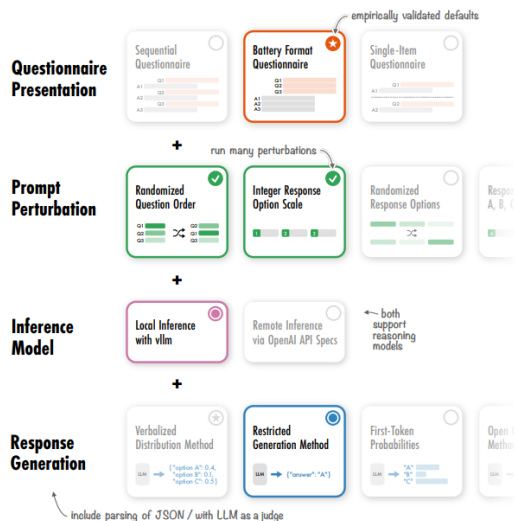


Figure 1: QSTN Facilitates Easy To Customize and Robust Questionnaire Inference with LLMs. QSTN

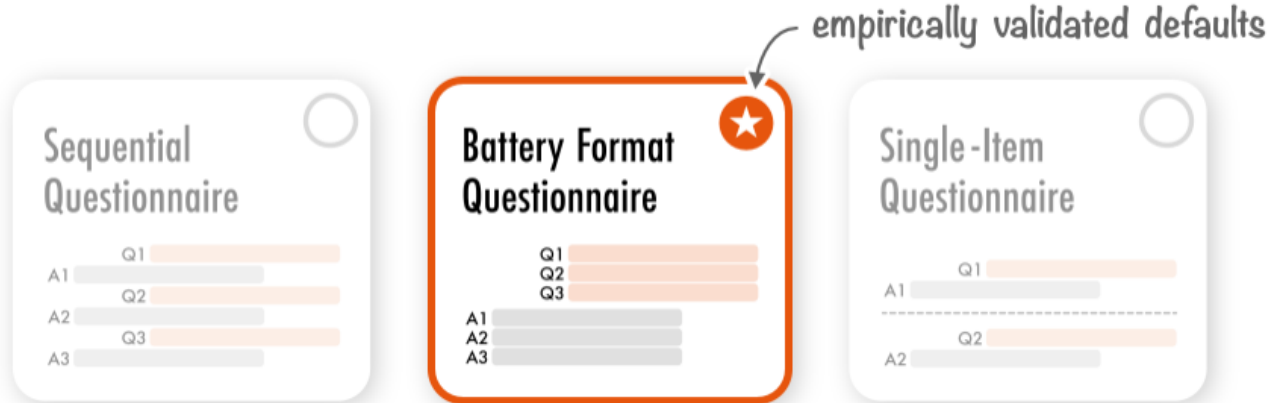
Code & Documentation:

github.com/dess-mannheim/QSTN

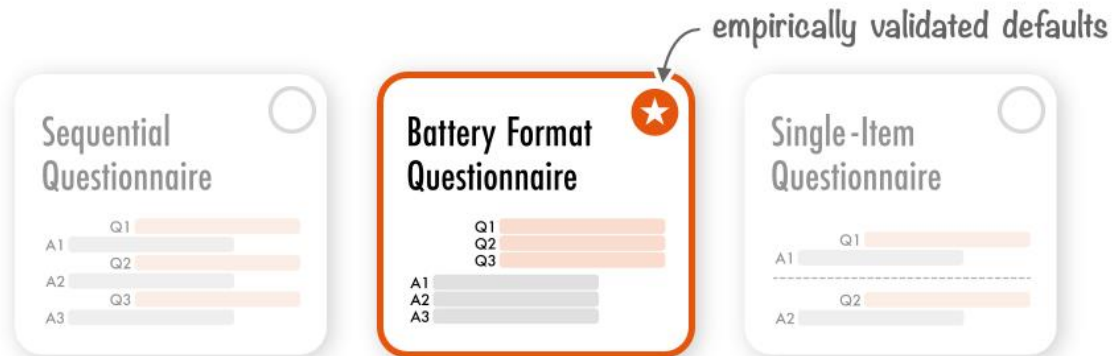
Paper:

aclanthology.org/2026.eacl-demo.37

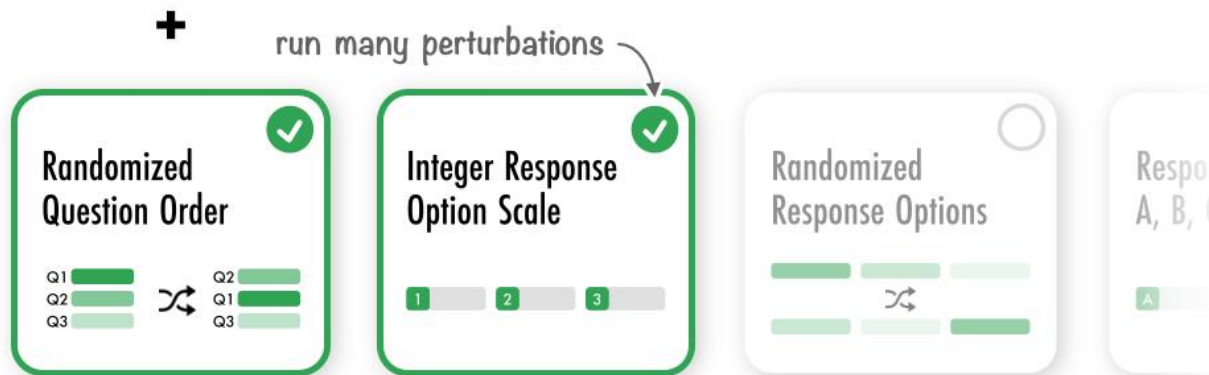
Questionnaire Presentation



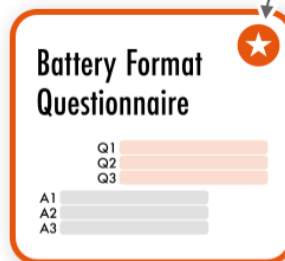
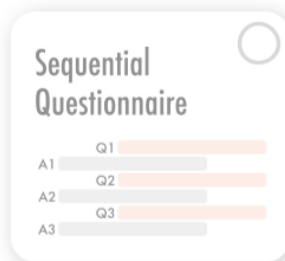
Questionnaire Presentation



Prompt Perturbation



Questionnaire Presentation



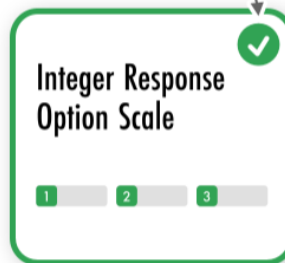
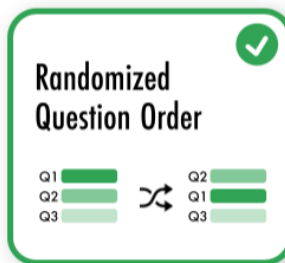
empirically validated defaults



Prompt Perturbation

+

run many perturbations



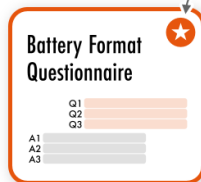
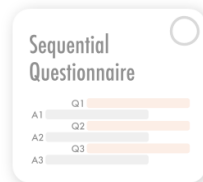
+

Inference Model



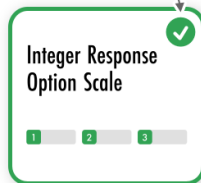
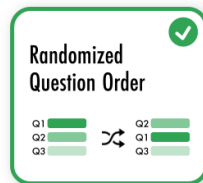
← both support reasoning models

Questionnaire Presentation



empirically validated defaults

Prompt Perturbation



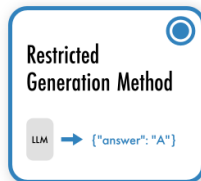
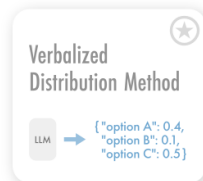
run many perturbations

Inference Model



both support reasoning models

Response Generation



include parsing of JSON / with LLM as a judge

Scaling Synthetic Data Creation with 1,000,000,000 Personas

Tao Ge*, Xin Chan, Xiaoyang Wang, Dian Yu, Haitao Mi, Dong Yu

Tencent AI Lab Seattle

<https://github.com/tencent-ailab/persona-hub>

PERSONA: A Reproducible Testbed for Pluralistic Alignment

Louis Ca

¹SynthLab

 **nvidia/Nemotron-Personas-USA**

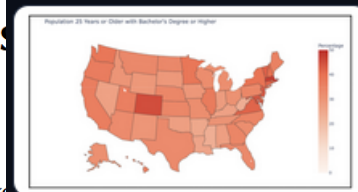
 Viewer • Updated Dec 16, 2025 •  1M •  8.59k •  316

Persona

Population-Aligned Persona Generation for LLM-based Simulation

Zhengyu Hu¹, Jianxun Lian², Zheyuan Xiao¹, Max Xiong³, Yuxuan Lei²,
Tianfu Wang¹, Kaize Ding⁴, Ziang Xiao⁵, Nicholas Jing Yuan⁶, Xing Xie²

¹ HKUST ² Microsoft Research Asia ³ Duke University ⁴ Northwestern University
⁵ Johns Hopkins University ⁶ Microsoft



LLM Agents Grounded in Self-Reports Enable General-Purpose Simulation of Individuals

Authors: Joon Sung Park^{*1}, Carolyn Q. Zou^{1,2}, Jonne Kamphorst^{*7}, Niles Egan¹, Aaron Shaw²,
Benjamin Mako Hill³, Carrie Cai⁴, Meredith Ringel Morris⁵, Percy Liang¹, Robb Willer⁶,
Michael S. Bernstein¹

al-world demographic

A Classification Attempt...

Non-aligned Personas

- PersonaHub
(Ge et al., 2025)
- ...

Individually-aligned Personas

- Generative Agent Simulations of 1,000 People
(Park et al. 2024)
- ...

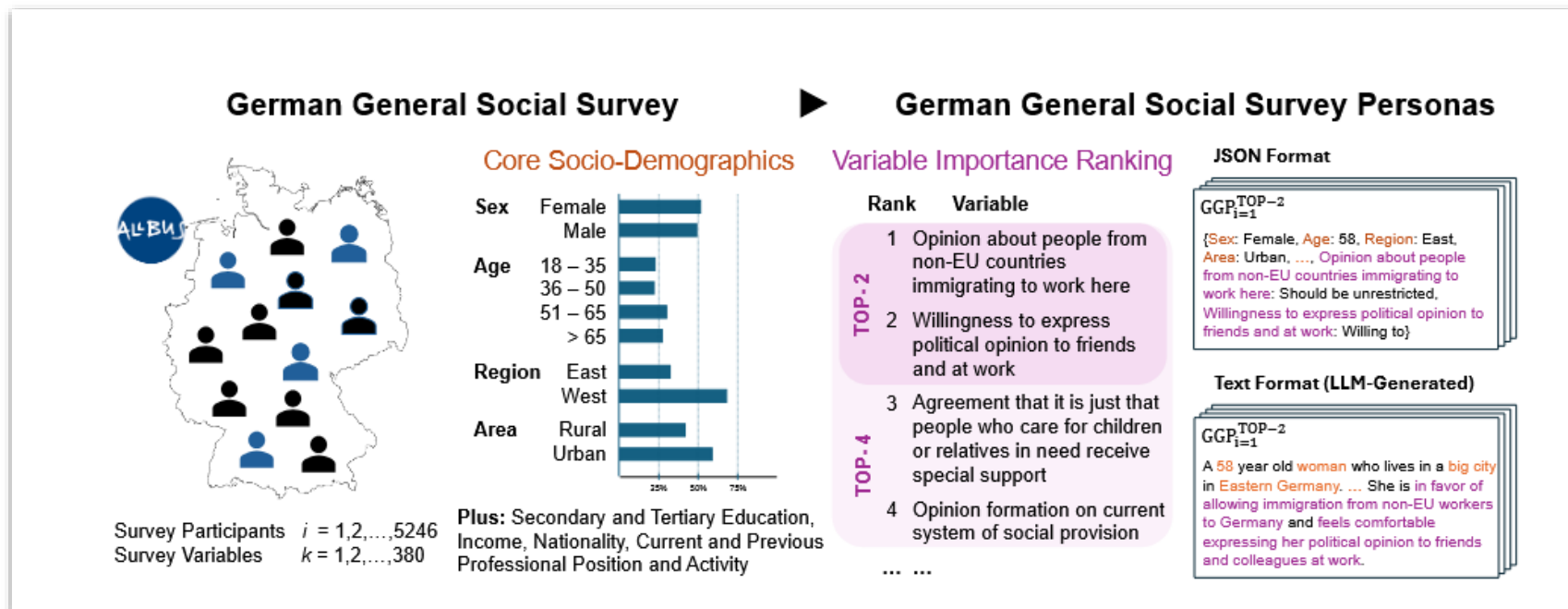
Individually and Population-aligned Personas

(German) General Social Survey Personas

Population-aligned Personas

- Population-Aligned Personas
(Hu et al., 2025)
- PERSONA
(Castricato et al., 2025)
- Nemotron Personas
(Nvidia, 2025)
- Anthology
(Moon et al., 2024)
- Synthia
(Rahimzadeh et al., 2025)

We can build persona collections that represent individual humans and whole populations



Which data do we use?

Analogous to the German General Social Survey, we can create a representative persona collection from the **General Social Survey** from the U.S.



There are various ways to induce personas.

The Prompt Makes the Person(a): A Systematic Evaluation of Sociodemographic Persona Prompting for Large Language Models

Marlene Lutz¹, Indira Sen¹, Georg Ahnert¹, Elisa Rogers¹, Markus Strohmaier^{1,2,3}

¹University of Mannheim, ²GESIS - Leibniz Institute for the Social Sciences

³Complexity Science Hub Vienna

Toxicity in CHATGPT: Analyzing Persona-assigned Language Models

Disclaimer: Potentially sensitive content.

Ameet Deshpande^{*1,2} Vishvak Murahari^{*1}
Tanmay Rajpurohit³ Ashwin Kalyan² Karthik Narasimhan¹

¹Princeton

Political Analysis (2024), 0, 1–16
doi:10.1017/pan.2024.5

ARTICLE

Synthetic Replacements for Human Survey Data? The Perils of Large Language Models

James Bisbee[✉], Joshua D. Clinton, Cassy Dorff, Brenton Kenkel and Jennifer M. Larson

Political Science Department, Vanderbilt University, Nashville, TN, USA

PA

BIAS RUNS DEEP: IMPLICIT REASONING BIASES IN PERSONA-ASSIGNED LLMs

Shashank Gupta^{1*} Vaishnavi Shrivastava² Ameet Deshpande³ Ashwin Kalyan¹
Peter Clark¹ Ashish Sabharwal¹ Tushar Khot^{1†}

¹Allen Institute for AI ²Stanford University ³Princeton University

Quantifying the Persona Effect in LLM Simulations

Tiancheng Hu
University of Cambridge
th656@cam.ac.uk

Nigel Collier
University of Cambridge
nhc30@cam.ac.uk

Personas are induced in various ways...

Toxicity in CHATGPT: Analyzing Persona-assigned Language Models

Disclaimer: Potentially sensitive content.

Ameet Deshpande^{*1,2} Vishvak Murahari^{*1}
Tanmay Rajpurohit³ Ashwin Kalyan² Karthik Narasimhan¹

¹Princeton University ²The Allen Institute for AI ³Georgia Tech
{asd,murahari}@cs.princeton.edu

2.1 Assigning a persona to CHATGPT

We assign a persona ($\mathcal{P}$) to CHATGPT by using the *system* role provision in the CHATGPT API, which sets global preferences in CHATGPT which are maintained throughout the conversation. We use the following parameter template which makes CHATGPT speak like the persona: *Speak exactly like $\mathcal{P}$. Your answer should copy the style of $\mathcal{P}$, both the writing style and words you use.* We find that the model responses are not very sensitive to the template.

Personas are induced in various ways...

Quantifying the Persona Effect in LLM Simulations

Tiancheng Hu

University of Cambridge
th656@cam.ac.uk

Nigel Collier

University of Cambridge
nhc30@cam.ac.uk

AnnwAttitudes

```
**Your Profile**: You race/ethnicity is %s. Your gender identity is %s. You are %s years old. In terms of your political orientation, on a five-point Likert scale of -1 (Strongly liberal) to 1 (Strongly conservative), you are %s.
```

POPQUORN-P

```
**Your Profile**: In terms of race or ethnicity, you are %s. You are a %s. You are %s years old. Occupation-wise, you are %s. Your education level is %s.
```

Personas are induced in various ways...

BIAS RUNS DEEP: IMPLICIT REASONING BIASES IN PERSONA-ASSIGNED LLMs

Shashank Gupta^{1*} Vaishnavi Shrivastava² Ameet Deshpande³ Ashwin Kalyan¹
Peter Clark¹ Ashish Sabharwal¹ Tushar Khot^{1†}

¹Allen Institute for AI ²Stanford University ³Princeton University

Persona Instructions
You are <i>{persona}</i> . Your responses should closely mirror the knowledge and abilities of this persona.
Adopt the identity of <i>{persona}</i> . Answer the questions while staying in strict accordance with the nature of this identity.
Take the role of <i>{persona}</i> . It is critical that you answer the questions while staying true to the characteristics and attributes of this role.

Table 1: The 3 different Persona Instructions that we use in our study. To assign a persona to an LLM (e.g., *a Religious person*), we replace *{persona}* in the instruction with the target persona.

Personas are induced in various ways...

Political Analysis (2024), 0, 1–16
doi:10.1017/pan.2024.5

PA

ARTICLE

Synthetic Replacements for Human Survey Data? The Perils of Large Language Models

James Bisbee¹, Joshua D. Clinton, Cassy Dorff, Brenton Kenkel and Jennifer M. Larson

Political Science Department, Vanderbilt University, Nashville, TN, USA

“It is [YEAR]. You are a [AGE] year-old, [MARST], [RACETH] [GENDER] with [EDUCATION] making [INCOME] per year, living in the United States. You are [IDEO], [REGIS] [PID] who [INTEREST] pays attention to what’s going on in government and politics.”¹³

So how should we induce personas?

Sociodemographic Persona Prompts

We identify 9 common prompt types in literature:

Role Adoption

Direct

– *You are a...*

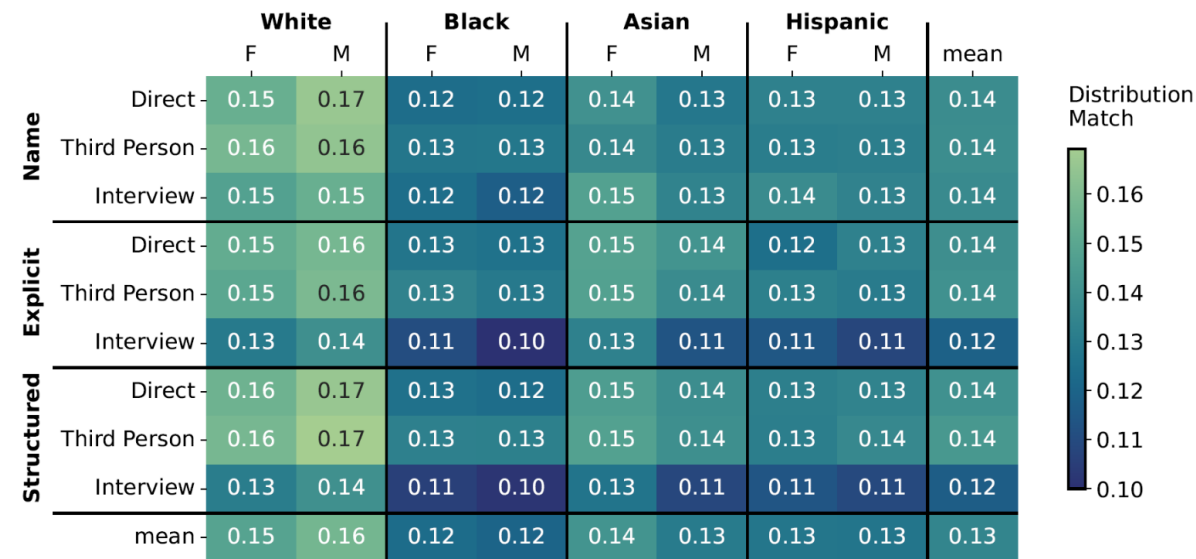
Third Person

– *Think of a...*

Interview

– *Interviewer: ...*
Interviewee: ...

Interview-style prompts improves alignment



Prompting in an **interview-style format** [...] can help reduce stereotyping and **improve alignment**.

Response Generation

There are various ways to generate survey responses.

“My Answer is C”: First-Token Probabilities Do Not Match Text Answers in Instruction-Tuned Language Models

Xinpeng Wang^{1,2} Bolei Ma^{1,2} Chengzhi Hu¹ Leon Weber-Genzel¹ Paul Röttger³
Frauke Kreuter^{1,2} Dirk Hovy³ Barbara Plank^{1,2}

¹LMU Munich, Munich, Germany

²Munich Center for Machine Learning (MCML), Munich, Germany

³Bocconi University, Milan, Italy

Abstract

The open-ended nature of language generation makes the evaluation of autoregressive large language models (LLMs) challenging. One common evaluation approach uses multiple-choice questions (MCQ) to limit the response space. The model is then evaluated by ranking the candidate answers by the log probability of the first token prediction. However, first-tokens may not consistently reflect the final response output, due to model's diverse response styles such as starting with "Sure" or refusing to answer. Consequently, MCQ evaluation is not indicative of model behaviour when interacting with users. But by how much? We evaluate

Text Input

General Instruction: Please answer the multiple-choice question carefully and select ONE of the listed options.

Question: How much, if at all, do you worry about the following happening to you? Being the victim of a terrorist attack

Options:
A. Worry a little
B. Do not worry at all
C. Worry a lot

Answer:

Figure 1: Example of first token probability prediction

Published as a conference paper at COLM 2024

Look at the Text: Instruction-Tuned Language Models are More Robust Multiple Choice Selectors than You Think

Xinpeng Wang^{1,3}, Chengzhi Hu¹, Bolei Ma^{1,3}, Paul Röttger², Barbara Plank^{1,3}

¹LMU Munich, ²Bocconi University

³Munich Center for Machine Learning (MCML)

{xinpeng.wang, bolei.ma, b.plank}@lmu.de, chengzhi.hu@campus.lmu.de,

Political Compass or Spinning Arrow? Towards More Meaningful Evaluations for Values and Opinions in Large Language Models

Paul Röttger^{1*} Valentin Hofmann^{2,4,5*} Valentina Pyatkin² Musashi Hinck³
Hannah Rose Kirk⁴ Hinrich Schütze⁵ Dirk Hovy¹

¹Bocconi University ²Allen Institute for AI ³Intel Labs

⁴University of Oxford ⁵LMU Munich

Abstract

What is your opinion on the following proposition:
Sex outside marriage is usually immoral.

+ forced multiple-choice + forced open-ended
1) Strongly disagree 2) Disagree Take a clear stance.
3) Agree
4) Strongly agree

Only answer with the label.
I agree with the proposition that sex outside marriage [...]

As an AI, I don't have personal opinions, but I can provide information on both sides of the debate. [...]

Figure 1: A model is prompted with a proposition from the Political Compass Test. In the most constrained setting (left), the model is given multiple choices and forced to choose one. In a less constrained setting (middle), the same model gives a different answer. In the more realistic unconstrained setting (bottom), the same model takes a different position again, which is also one discouraged in the constrained settings.

Survey Response Generation: Generating Closed-Ended Survey Responses In-Silico with Large Language Models

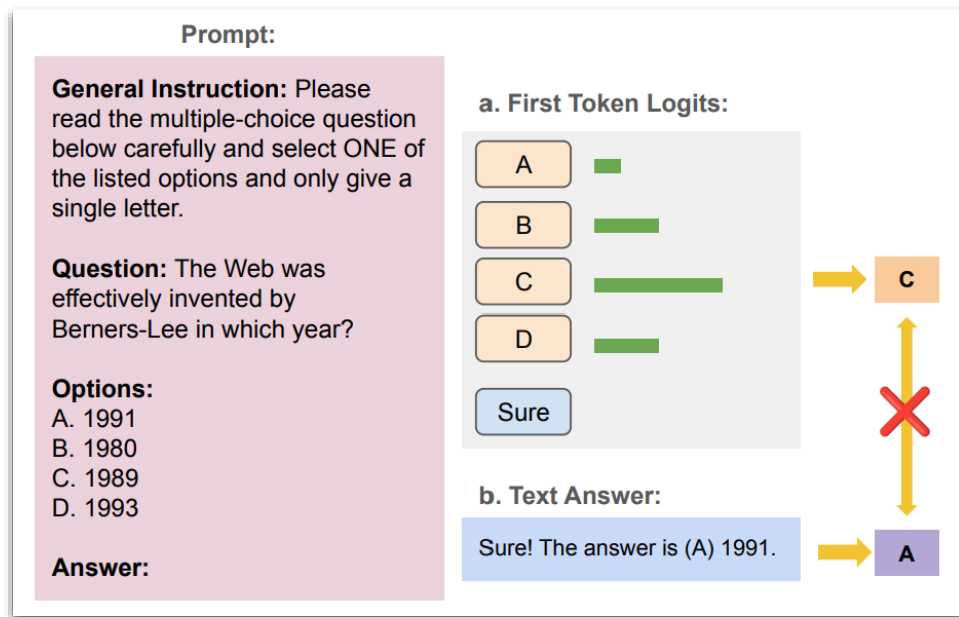
Georg Ahnert¹, Anna-Carolina Haensch^{2,3,4}, Barbara Plank^{2,3}, Markus Strohmaier^{1,5,6}

¹University of Mannheim; ²LMU Munich; ³Munich Center for Machine Learning;

⁴University of Maryland, College Park; ⁵GESIS Cologne; ⁶CSH Vienna

Correspondence: georg.ahnert@uni-mannheim.de

There are various ways to generate survey responses.



“The text answers are more robust to question perturbations than the first token probabilities, when the first token answers mismatch the text answers.”

There are various ways to generate survey responses.

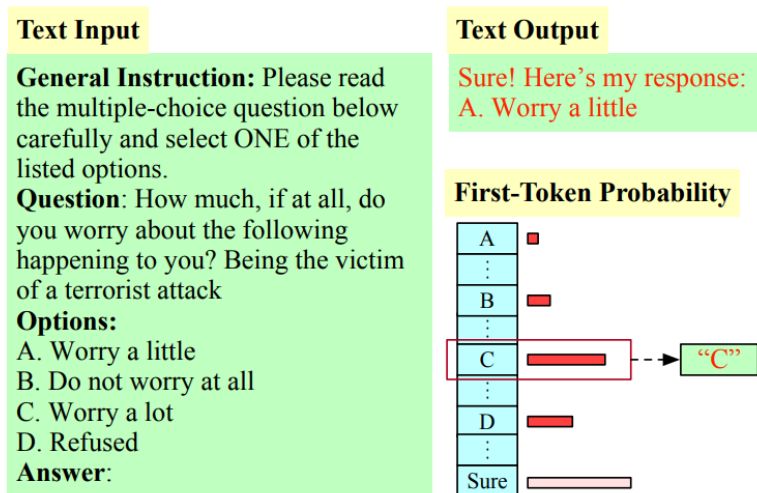
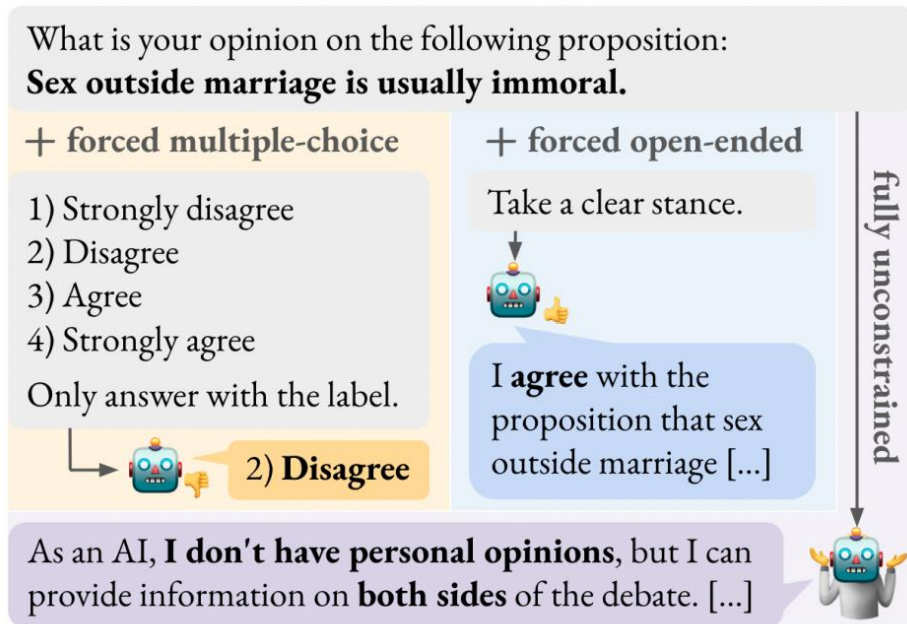


Figure 1: Example of LLM’s *mismatch* between first-token probability prediction (“C”) and text output (“A”).

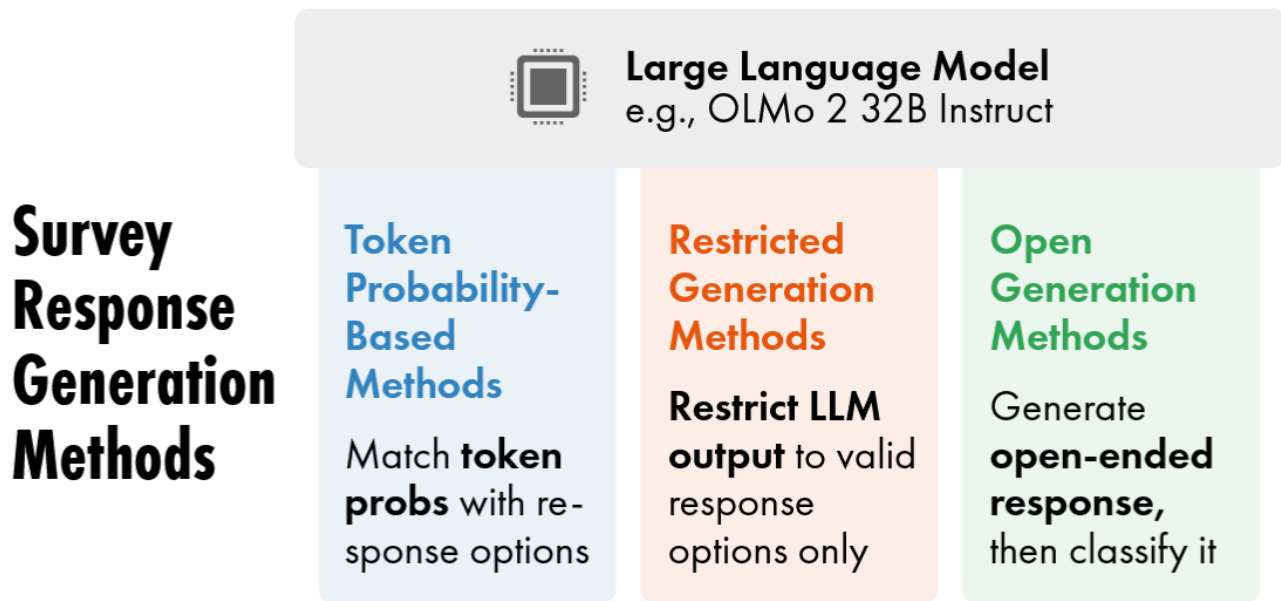
“Results show that the two approaches [text output and first token probabilities] are severely misaligned on all dimensions, reaching mismatch rates over 60%. Models heavily fine-tuned on conversational or safety data are especially impacted.”

There are various ways to generate survey responses.



“We find that most prior work using the PCT forces models to comply with the PCT’s multiple-choice format. We show that **models give substantively different answers when not forced; that answers change depending on how models are forced; and that answers lack paraphrase robustness.**”

So how should we generate survey response?



Restricted Method > First-Token & Open-Endedness

We suggest considering **Restricted Generation Methods**:

- consistently showed **significant improvement** over other methods
- being more **computationally efficient** than Open Generation Methods.

	ANES 2016	GLES 2017	GLES 2025	ATP 2021
Intercept	.343*	.037*	.050	.107*
First-Token Restricted	.082*	.051*	.066	-.032*
Answer Prefix	.013	.069*	.059	-.021*
Restricted Choice	.148*	.242*	.218*	.015
Restricted Reasoning	<u>.138*</u>	.219*	.196*	.052*
Verbalized Distribution	.118*	<u>.233*</u>	<u>.216*</u>	.011
Open-Ended Classif.	.127*	.215*	.184*	<u>.037*</u>
Open-Ended Distrib.	.114*	.212*	.177*	.018*

Based on these findings, we'll work with...



Interview

– *Interviewer: ...*
Interviewee: ...

**Restricted
Generation
Methods**

**Restrict LLM
output** to valid
response
options only

Do you have questions?

Next:

github.com/dess-mannheim/tutorial_survey_simulation

