

Generative Artificial Intelligence, Cultural Context and Discrimination

Generative Künstliche Intelligenzen, kultureller Kontext und Diskriminierung

<Datensatz - Textdokumente>

<Valerian Thielicke-Witt, Ana-Nzinga Weiß, Hannah Miltzow>

1. General Information

Dataset title: „Generative Artificial Intelligence, Cultural Context and Discrimination“ / „Generative Künstliche Intelligenzen, kultureller Kontext und Diskriminierung“

Authors/Creators: Valerian Thielicke-Witt¹, Ana-Nzinga Weiß², Hannah Miltzow¹

Affiliation: ¹Universität Rostock, Institut für Politik- und Verwaltungswissenschaft,
²LMU München, Institut für Publizistik- und Kommunikationswissenschaft

E-Mail: valerian.thielicke-witt@uni-rostock.de

Date: 22.05–15.06.2025

Type: Collection of Qualitative Texts in .rtf and .txt files from LLMs

Language: English, German, French

Rights: CC BY NC 4.0

DOI: https://doi.org/10.18453/rosdok_id00005015

Funding: WKT – INF, University of Rostock

2. Objective/Problem Statement

Large language models are highly debated and a current topic of research. There is already some research on the biases from Large Language Models (LLMs) and how they reproduce stereotypes (Atari et al. 2023) or how they favour distinct styles in writing (Agarwal et al. 2024) and consequently narrow down the imaginaries of human beings (Hewitt et al. 2024; Kosmyrna et al. 2025). The starting point of the Dataset is that there are currently different generative LLMs from different contexts, trained with different data. Thus, we wanted to explore if those different contexts and training data

are represented in the generated texts of different LLMs. Training data has always a distinct bias of what is included and what not (Roselli et al. 2019; Gichoya et al. 2023; Caliskan), thus we assume that the different contexts must produce different outcomes from the different LLMs, thus their outcome should be culturally different (Aleksic 2025). Therefore, we asked different LLMs to define “Diversity” in different ways and in different languages to avoid language biases and to see if the outcome depends on the language. We are interested on the question which variable can explain the outcome most: the background of production of an LLM, the language or the distinct prompt. The same LLM processed the same prompt several times. Additionally, we are interested if the LLMs favour a particular political position, thus reproduces hegemony like Dakakni (2025) or Brown et al. (Brown et al. 2025) assumes Those texts are produced for an in-depth qualitative analysis.

3. Description

Since there are different LLMs from different contexts we had to select some for the exploratory study, if this objective is lucrative. Scholars usually identify three main contexts of the production of LLMs, the USA, China and the EU. Thus, we chose two LLMs from each context, i.e., Qwen 3 and DeepSeek from China, Meta Llama 3.1 and Gemma 3 for the USA and Mistral Large Instruct and Llama 3.1 SauerkrautLM for the EU. All of them are based on transformers can be downloaded by Huggingface.co. Although those LLMs can be run on every Desktop via a small python script we used KISSKI¹(<https://chat-ai.academiccloud.de/>) to obtain the results faster (Doosthosseini et al. 2024). To secure data validity we used for each the same version. Furthermore, for securing comparability of the produced data and between the different LLMs we had to find a compromise between different settings (usually “system prompt”, “temperature” (temp) and “top_p”). The distinct settings are usually part of the business secret of commercial distribution of LLMs. For securing comparability, we used every time “You are a helpful assistant.” as system prompt and a middle setting for temp (0,5) and top_p (0,5). Additionally, every process started anew, i.e. the context window of each LLM was clear except the system prompt. Since our scientific objective is qualitative in nature, i.e., we want to analyse the produced text by a qualitative approach, we had to choose an equilibrium between reproductivity of the outcomes and size to analyse the generated texts qualitatively by hand and not by an algorithm. Thus, we processed every prompt at least 5 times. We used three different prompts that are slightly different “What is diversity?”, “Can you define diversity?”, “What is the definition of diversity”? We used three different languages for the prompts, German, Englisch and French. In sum, we have produced more than 270 comparable generated texts (6 LLMs x 3 Languages x 3 different prompts x 5). Furthermore, we produced more texts by playing with settings and system prompts to see if this changes the outcomes as well. Finally, we obtained more than 350 files.

4. Data Format

All generated texts from the given LLMs are saved separately in distinct .txt and .rtf files. Every file also contains a description of the version of the used LLMs, the settings, the prompt, and the generated text. While the .rtf-files depict the layout and format of the generated text WYSIWYG (What you see is what you get), the .txt-files are machine-readable and contain more information, such as the ‘reflection’ of DeepSeek and Qwen 3.

For example: File E_DS_Can_1.rtf

¹ The authors gratefully acknowledge the computing time granted by the KISSKI project. The calculations for this research were conducted with computing resources under the project.

DeepSeek R1

Temp: 0,5 (logical – creative)

Top_p: 0,5 (focused – diverse)

System prompt: You are a helpful assistant

Frage: Can you define diversity?

Diversity refers to the presence, recognition, and respect for a wide range of differences among individuals, groups, or systems. These differences can include, but are not limited to:

- **Identity:** Race, ethnicity, gender, age, sexual orientation, religion, disability, and socioeconomic background.
- **Experiences:** Cultural traditions, education, language, and life perspectives.
- **Cognitive and Ideational:** Varied ways of thinking, problem-solving approaches, and beliefs.

In social and organizational contexts, diversity emphasizes not just representation but also the equitable inclusion and valuing of these differences to foster innovation, creativity, and mutual understanding. It acknowledges that unique backgrounds and viewpoints enrich communities, workplaces, and ecosystems by promoting adaptability and broader perspectives.

5. Archive Structure

The archive has the following structure. Raw_Data is the host folder, that contains the Data_description file and three linguistically different folders (German, English, French) and one with the auxiliary texts. Each linguistic folder contains six different folders representing the different LLMs such as “Deepseek”, “Gemma 3”, “Llama Sauerkraut”, “Mistral Large Instruct”, “Qwen 3” and “Meta Llama 3.1”. Each of these folders contains the produced texts as .rtf-files and one folder with the name “txt_files”, containing the .txt-files of the same generated texts. Every file has a unique identifier name, represented by abbreviations signifying the language, the LLM, the question and the number of repeated cycles, e.g., “D_LS_Def_1”.

The following abbreviations were used:

- Languages:
 - D = German
 - E = English
 - F = French
- Algorithms:
 - ML = Meta Llama 3.1. 8B Instruct
 - G = Gemma 3 27B Instruct
 - Q = Qwen 3 32B

- DS = DeepSeek R1
- MLI= Mistral Large Instruct
- LS = Llama 3.1. SauerkrautLM 70B Instruct
- Prompts:
 - Was = Was ist Diversität?
 - Kann = Kannst du Diversität definieren?
 - Def = Was ist die Definition von Diversität?
 - What = What is diversity?
 - Can = Can you define diversity?
 - Def = What is the definition of diversity?
 - Que = Qu'est-ce que la diversité ?
 - Puis = Puisses-tu définir la diversité ?
 - Def = Qu'est-ce que la définition de la diversité ?

6. System Requirements

No special system requirements needed. Compatibility with .txt and .rtf-files.

References/Bibliography

Agarwal, Dhruv; Naaman, Mor; Vashistha, Aditya (2024): AI Suggestions Homogenize Writing Toward Western Styles and Diminish Cultural Nuances. Available online at <https://arxiv.org/pdf/2409.11360>.

Aleksic, Adam (2025): Algospeak. How social media is transforming the future of language. New York: Alfred A. Knopf.

Atari, Mohammad; Xue, Mona J.; Park, Peter S.; Blasi, Damián Ezequiel; Henrich, Joseph (2023): Which Humans?

Brown, Venetia; Larasati, Retno; Third, Aisling; Farrell, Tracie (2025): A Qualitative Study on Cultural Hegemony and the Impacts of AI. In *Proceedings of the Seventh AAAI/ACM Conference on AI, Ethics, and Society (AIES 2024)*, checked on 11/28/2025.

Caliskan, Aylin: Artificial Intelligence, Bias, and Ethics. Available online at <https://www.ijcai.org/proceedings/2023/0799.pdf>, checked on 11/28/2025.

Dakakni, D. (2025): Artificial intelligence as a tool for data, economic and political hegemony: releasing the djinn. In *Ethics. Sci. Environ. Polit.* 25, pp. 1–10. DOI: 10.3354/esep00216.

Doosthosseini, Ali; Decker, Jonathan; Nolte, Hendrik; Kunkel, Julian M. (2024): Chat AI: A Seamless Slurm-Native Solution for HPC-Based Services.

Gichoya, Judy Wawira; Thomas, Kaesha; Celi, Leo Anthony; Safdar, Nabile; Banerjee, Imon; Banja, John D. et al. (2023): AI pitfalls and what not to do: mitigating bias in AI. In *The British journal of radiology* 96 (1150), p. 20230023. DOI: 10.1259/bjr.20230023.

Hewitt, Luke; Ashokkumar, Ashwini; Ghezae, Isaias; Willer, Robb (2024): Predicting Results of Social Science Experiments Using Large Language Models. Available online at <https://samim.io/dl/Predicting%20results%20of%20social%20science%20experiments%20using%20large%20language%20models.pdf>.

Kosmyna, Nataliya; Hauptmann, Eugene; Yuan, Ye Tong; Situ, Jessica; Liao, Xian-Hao; Beresnitzky, Ashly Vivian et al. (2025): Your Brain on ChatGPT: Accumulation of Cognitive Debt when Using an AI Assistant for Essay Writing Task. Massachusetts Institute of Technology.

Roselli, Drew; Matthews, Jeanna; Talagala, Nisha (2019): Managing Bias in AI. In Ling Liu, Ryen White (Eds.): Companion Proceedings of The 2019 World Wide Web Conference. WWW '19: The Web Conference. San Francisco USA, 13 05 2019 17 05 2019. New York, NY, USA: ACM, pp. 539–544.