

A Review on the Use of Large Language Models in the Context of Open Government Data

Daniel Staegemann¹^a, Christian Haertel¹^b, Matthias Pohl²^c and Klaus Turowski¹^d

¹Magdeburg Research and Competence Cluster VLBA, Otto-von-Guericke University Magdeburg, Magdeburg, Germany

²Institute of Data Science, German Aerospace Center (DLR), Jena, Germany

Keywords: Generative AI, GenAI, Large Language Model, LLM, Literature Review, Open Data, Open Government Data.

Abstract: Since ChatGPT was released to the public in 2022, large language models (LLM) have drawn enormous interest from academia and industry alike. Their ability to create complex texts based on provided inputs positions them to be a valuable tool in many domains. Moreover, since some time, many governments want to increase transparency and enable the offering of new services by making their data freely available. However, these efforts towards Open Government Data (OGD) face various challenges with many being related to the question how the data can be made easily findable and accessible. To address this issue, the use of LLMs appears to be a promising solution. To provide an overview of the corresponding research, in this work, the results of a structured literature review on the use of LLMs in the context of OGD are presented. Hereby, numerous application areas as well as challenges were identified and described, providing researchers and practitioners alike with a synoptic overview of the domain.

1 INTRODUCTION

Since ChatGPT was released to the public in 2022, large language models (LLM) and generative artificial intelligence (GenAI) have been in the centre of interest inside and outside of academia (Chang et al. 2024; Raiaan et al. 2024). Due to their ability to produce sophisticated outputs based on a provided prompt, they are widely seen as a promising tool to enhance the operations of organizations across numerous domains (Brynjolfsson et al. 2023; Filippucci et al. 2024; Simons et al. 2024).

While they are generally prone to occasionally making up information, which is referred to as *hallucinations* (Huang et al. 2024; Perković et al. 2024), this can be addressed through techniques such as specific training and fine-tuning or the utilization of retrieval augmented generation (RAG). Hereby, the model is given access to specific data and documents that it can then draw from to produce higher quality results that are based on the provided information (Fan et al. 2024).

A huge producer of data are governments. Whenever services are provided, decisions are made, or statistics are recorded, this adds to the body of related data. While this is generally positive, as this allows for their analysis, this also makes it harder to manage the resulting data deluge. Further, many governments have pivoted towards not only keeping the data and utilizing them themselves, but to also provide their citizens and the general public with access to many of these information (Attard et al. 2015; Bonina and Eaton 2020). This Open Government Data (OGD) movement, in turn, increases transparency and accountability, might lead to additional insights and services based on the data, and can help to increase trust (Janssen et al. 2012; Kucera and Chlapek 2014).

Yet, the effective use of these data is often rather challenging, because their sheer volume makes it hard to get an overview of the available information and the associated meta-information are often insufficient, which makes the discovery of potentially

^a <https://orcid.org/0000-0001-9957-1003>

^b <https://orcid.org/0009-0001-4904-5643>

^c <https://orcid.org/0000-0002-6241-7675>

^d <https://orcid.org/0000-0002-4388-8914>

useful data extremely cumbersome (Ahmed 2023; Ansari et al. 2022; Quarati 2023).

However, combining the high accessibility of LLMs with the wealth of information contained in OGD could be a valuable asset for shaping the society of the future, increasing the involvement of citizens, and offering a plethora of new services.

Thus, the goal of this work is to explore the current state of the scientific literature on this research stream, highlighting the potentials and challenges and outlining the most promising avenues for future work. To this end, a structured literature review (SLR) was conducted. Therefore, within this paper, the following research question (RQ) shall be answered:

RQ: What is the current state of the scientific literature on the use of large language models in the context of Open Government Data?

To answer the RQ, the remainder of this publication is structured as follows. Following the introduction, the SLR is outlined. Then the identified literature is examined. This is followed by a discussion of the findings. Finally, a conclusion is given, and avenues for future work are outlined.

2 THE REVIEW

To answer the RQ, a SLR was conducted. As frequently pointed out, the value of a SLR highly depends on its rigour and reproducibility (Kraus et al. 2022; vom Brocke et al. 2009). Therefore, adhering to common practices (Okoli 2015; vom Brocke et al. 2015), before starting the search, a protocol was developed to guide the process. In the following, the corresponding steps, as well as the underlying considerations and the obtained results are outlined.

To identify the relevant literature, *Scopus*¹ was chosen as the primary source, because it provides a comprehensive coverage across many scientific databases and publishers. However, since Scopus alone usually does not find all relevant papers, as will become visible in Table 3, multiple other databases and scientific search engines were used in addition to ensure a broader coverage. For instance, *IEEE Xplore*² (IEEE) was added because of IEEE's significance in the computer science domain. These two were complemented by the AIS electronic Library³ (AISEL), which, inter alia, contains the proceedings of some of the most renowned

conferences in the information systems domain and the ACM Digital Library⁴ (ACM), which is operated by the world's largest computing society (ACM History Committee 2025). Finally, Springer Nature Link⁵ (Springer) was added, since in the past it has shown to be a strong complement to the aforementioned sources.

While there are some differences in the design of the search masks, the search terms used in all of these were kept as similar as possible. In each case, the search term consisted of two components.

The first covers the realm of LLMs, while also including papers that refer to GenAI in general. Further, due to ChatGPT currently being the most popular LLM, it was also explicitly included in the term, whereas others were not. In each case, different spellings and abbreviations were covered to ensure comprehensiveness:

Part 1: *llm OR "large language model" OR "generative artificial intelligence" OR "generative ai" OR "gen ai" OR genai OR gpt OR chatgpt*

In the second part, the field of Open Data is addressed. While the focus of the work is on Open Government Data, this was done to ensure a broader coverage and include relevant papers that might have been missed otherwise. Therefore, the corresponding term was as follows:

Part 2: *"open data" OR "public data"*

To make sure that both parts are present in the found papers, these two parts were connected with an AND. Further, to increase comprehensiveness, the terms were not only searched in the title but more broadly, as shown in Table 1.3

Table 1: The utilization of the search terms.

Source	Part 1 used in	Part 2 used in
AISEL	All Fields	All Fields
ACM	Anywhere	Title
IEEE	All Metadata	All Metadata
Scopus	Article title, Abstract, Keyword	Article title, Abstract, Keyword
Springer	Keywords	Title

Thus, the final search term in Scopus was, for instance, as follows:

(TITLE-ABS-KEY (llm OR "large language model" OR "generative artificial intelligence" OR "generative ai" OR "gen ai" OR genai OR gpt OR

¹ <https://www.scopus.com>

² <https://ieeexplore.ieee.org>

³ <https://aisel.aisnet.org/>

⁴ <https://dl.acm.org/>

⁵ <https://link.springer.com/>

chatgpt) AND TITLE-ABS-KEY ("Open Data" OR "public data"))

The initial search in Scopus resulted in the identification of a total of 138 items. IEEE, in turn, yielded 36 papers and Springer Nature Link 10. Through ACM, 5 papers were found, and AISEL contributed 69 additional papers. Thus, overall, the keyword search brought 258 items. However, since multiple databases were used for the search, several duplicates occurred that were removed in the next step. After doing so, 238 items remained.

Naturally, not each of these fit the intended scope, which made additional filtering necessary. Aligned with common practices (vom Brocke et al. 2015), as shown in Figure 1, this was performed in multiple steps to assure a high degree of diligence while still maintaining efficiency.

For all of these phases, a joint set of inclusion and exclusion criteria, as depicted in Table 2, was defined in advance to serve as the foundation of the filter process. Hereby, for a paper to be deemed suitable, each of the inclusion criteria had to be met, whereas when at least one of the exclusion criteria applied, it was removed from the set.

To ensure the necessary quality, it was decided to only include *research articles* that were published as *conference papers* or *journal articles*. In turn, other items such as conference reviews, editorials, introductions to a minitrack, catch word articles, comments, or summarizations of panel discussions were not included. Further, book chapters were also not considered, since they are usually not peer-reviewed. This is also the reason why preprint services like arXiv⁶ were not included in the initial search, since there are "concerns about the research accuracy, quality, and credibility of preprints" (Adarkwah et al. 2024). Due to the provided meta-data of the publications regarding their type not always being correct and precise, this required manual checking. For this reason, the differentiation between the obtained publication types in the description of this search process is also only included after this step. After removing documents of the wrong type, 202 items remained, with 152 of them being conference papers and 50 journal articles.

Following this, the language was considered as an additional factor, to ensure that the authors of the publication at hand can comprehend the content without needing the help of translation services, which might involuntarily distort the content. For this reason, only papers that were written in English were kept. This resulted in the removal of three papers, which were all journal articles. Two of them were written in Chinese and one in Portuguese.

Table 2: The search's inclusion and exclusion criteria.

Inclusion Criteria	Exclusion Criteria
The paper is written in English	The paper is a duplicate
The paper is published in the proceedings of a scientific conference or in a scientific journal	The paper only briefly mentions Open Government Data without actually discussing it further
The paper focusses on the application of LLMs in the context of Open Government Data	The paper only briefly mentions LLMs without actually discussing them further
The paper discusses application scenarios or (potential) use cases	The paper is a short paper (here defined as not having a length of more than 5 pages)
	The found item is a conference review, an editorial, an introduction to a minitrack, a catch word article, a comment, presents the results of a panel or is a similar document type that is no research article in the narrower sense

Because the goal was to gain profound insights into the field, papers that had a length of not more than five pages were also excluded, since they were deemed too short to provide the necessary depth. After this step, 162 papers were left, 116 conference papers and 46 journal articles.

Once these formal aspects were considered and the papers were filtered accordingly, the actual content of the remaining ones was taken into account.

Due to this work's RQ, the papers had to provide insights into the use of large language models in the context of Open Government Data. Thus, items that focused one aspect and only briefly mentioned the other one, were not relevant.

In the first step of this phase, the remaining papers were filtered based on their title. When it was obvious that a publication did not fit the intended scope, it was excluded from the list, whereas, in case of uncertainty, the papers were kept. Once this was finished, 70 conference papers and 18 journal articles were left.

Afterwards, the abstracts and keywords were used to further filter the list. This, again, reduced the number of papers significantly, to 26, of which 20 were conference papers and 6 journal articles.

Finally, to conclusively appraise the suitability of the remaining papers to the RQ's scope, they were read in total, and those that did not fit were excluded,

⁶ <https://arxiv.org/>

resulting in a final literature set, as shown in Table 3, that comprises 16 papers, of which 12 are from conferences and 4 appeared in journals.

An overview of the identified relevant literature is given in Table 3. Besides the publication year, the reference, and the type of paper, it is also shown from

which database the items originate. As can be seen, for this SLR, ACM and IEEE were not crucial, whereas the other three each contributed at least one unique item. However, this could not have been known in advance and their inclusion still increased the search's comprehensiveness.

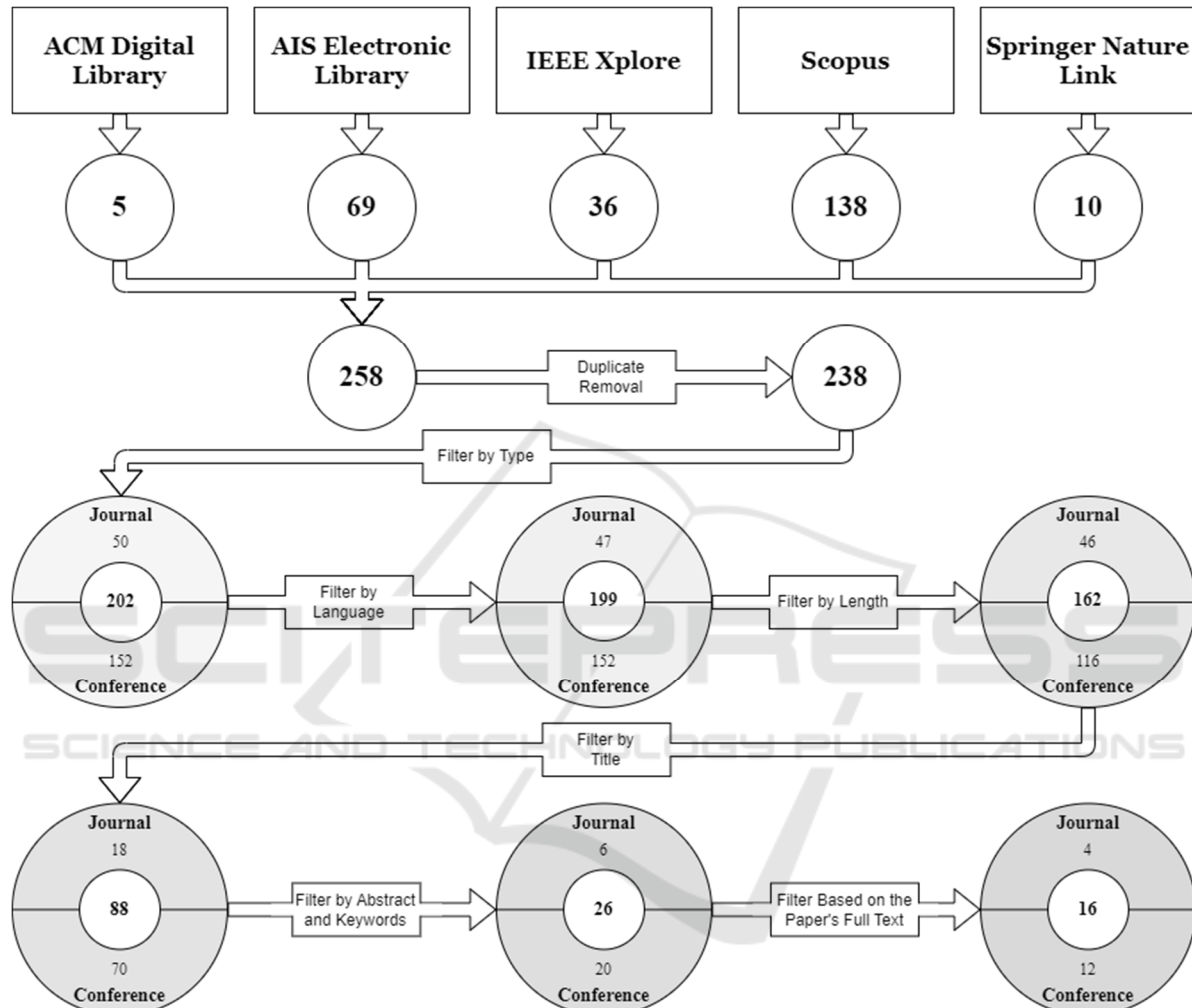


Figure 1: The search process.

Table 3: The identified papers.

ID	Title	Year	Type	Found in	Reference
1	Reimagining open data ecosystems: a practical approach using AI, CI, and Knowledge Graphs	2023	Conference Paper	Scopus	(Ahmed 2023)
2	ChatGPT Application vis-a-vis Open Government Data (OGD): Capabilities, Public Values, Issues and a Research Agenda	2023	Conference Paper	Scopus	(Loukis et al. 2023)
3	Can Large Language Models Revolutionize Open Government Data Portals? A Case of Using ChatGPT in statistics.gov.scot	2023	Conference Paper	Scopus	(Mamalis et al. 2023)
4	BRYT: Automated keyword extraction for open datasets	2024	Journal Article	Scopus	(Ahmed et al. 2024)

Table 3: The identified papers(cont).

ID	Title	Year	Type	Found in	Reference
5	Exploring Interpretability in Open Government Data with ChatGPT	2024	Conference Paper	Scopus	(Barcellos et al. 2024)
6	Instruct Large Language Models for Public Administration Document Information Extraction	2024	Conference Paper	Scopus	(Carta et al. 2024)
7	Aragón Open Data Assistant, Lesson Learned of an Intelligent Assistant for Open Data Access	2024	Conference Paper	Springer Nature Link	(Del Hoyo-Alonso et al. 2024)
8	TAGIFY: LLM-powered Tagging Interface for Improved Data Findability on OGD portals	2024	Conference Paper	IEEE; Scopus	(Klimask and Nikiforova 2024)
9	On Enabling Dynamic, Transparent, and Inclusive Government Consultative Processes with GenAIOpenGov DSS	2024	Conference Paper	AISel	(Marjanovic et al. 2024)
10	From the Evolution of Public Data Ecosystems to the Evolving Horizons of the Forward-Looking Intelligent Public Data Ecosystem Empowered by Emerging Technologies	2024	Conference Paper	Scopus; Springer Nature Link	(Nikiforova et al. 2024)
11	Designing a Large Language Model Based Open Data Assistant for Effective Use	2024	Conference Paper	Scopus; Springer Nature Link	(Schelhorn et al. 2024)
12	Unveiling inequality: A deep dive into racial and gender disparities in US court case closures	2024	Journal Article	Scopus	(Takefuji 2024)
13	Web Open Data to SDG Indicators: Towards an LLM-Augmented Knowledge Graph Solution	2025	Conference Paper	Springer Nature Link	(Benjira et al. 2025b)
	Automated mapping between SDG indicators and open data: An LLM-augmented knowledge graph approach	2025	Journal Article	Scopus	(Benjira et al. 2025a)
14	The Convergence of Open Data, Linked Data, Ontologies, and Large Language Models: Enabling Next-Generation Knowledge Systems	2025	Conference Paper	Springer Nature Link	(Cigliano and Fallucchi 2025)
15	AI for Good: History, Open Data and Some ESG-based Applications	2025	Journal Article	Scopus	(O'Leary 2025)

While in total 16 papers were found, one journal article (Benjira et al. 2025a) is an expanded version of a conference paper (Benjira et al. 2025b), which is why the two are grouped in the table. Further, (Ahmed 2023) and (Ahmed et al. 2024) are also somewhat related but without one being a direct expansion of the other one. Therefore, they are considered as independent publications.

3 FINDINGS

Based on the literature, numerous application avenues for LLMs in the context of OGD can be identified. Hereby, a plethora of stakeholders can participate

from their capabilities and the associated added convenience. These include but are not limited to citizens, journalists, scientists, companies, and government entities themselves (Loukis et al. 2023).

However, despite the versatility of tasks that can be facilitated through LLMs, these can be broadly divided into two categories, the provisioning and the utilization of OGD, which will also be reflected in the following two sub-sections.

3.1 OGD Provisioning

One major stream of research is focused on the capabilities that LLMs offer in the context of amending the OGD that are made available with additional information. As these are currently

oftentimes rather badly described and lacking in proper descriptions or keywords to support potential data consumers in finding datasets that are actually relevant to them (Ahmed 2023; Ansari et al. 2022; Quarati 2023), addressing this issue is a significant step to increase the value obtained from providing the OGD (Alexopoulos et al. 2024). Hereby, this becomes increasingly important, the more data are shared and the higher their complexity is, since maintaining an overview of them becomes more challenging.

To address this issue, the most commonly found utilization is the extraction of suitable keywords to describe the data and facilitate a more targeted search. While, theoretically, these could also be manually determined and provided by the creators themselves, their current limited availability highlights the limitations of this approach. Moreover, requiring the keywords to be submitted alongside the data would be an additional barrier that could lead to a lowered willingness to participate in the provisioning of OGD, which would be counterproductive when trying to increase participation (Barry and Bannister 2014). Further, this approach would not cover already published datasets, which is another downside. Additionally, maintaining consistency throughout the assignment of keywords is also not an easy task when many different government bodies, with potentially varying intrinsic interest in contributing (Alexopoulos et al. 2024), are involved. Thus, while on their own, the keywords could be contentually accurate, a diverse use of synonyms, spellings, or interpretations could still turn into a barrier.

Yet, LLMs like ChatGPT appear very promising for this task, which is why their utilization for keyword extraction was proposed in several of the identified papers (Ahmed 2023; Ahmed et al. 2024; Kliimask and Nikiforova 2024). Additionally, the use of LLMs to create annotations in a slightly more general sense was suggested in (Nikiforova et al. 2024). This would reduce the manual effort involved in the creation of keywords and, at the same time, allow to facilitate comparability. However, in the current state, LLMs are not performing flawless in this task, which might require a certain degree of human oversight depending on the requirements regarding the quality of the provided keywords and annotations (Kliimask and Nikiforova 2024).

Another way to increase the findability of OGD lies in the identification of themes the datasets belong to, to allow for a suitable categorization. Again, similar challenges as with the keyword generation apply, which could be addressed by utilizing LLMs

for this task (Ahmed 2023; Ahmed et al. 2024; Marjanovic et al. 2024).

To further enhance the descriptiveness of datasets, LLMs can also aide in summarizing their content. This can either be in the form of information triplets that are extracted as discussed in (Carta et al. 2024; Cigliano and Fallucchi 2025) or by providing actual summaries that allow potential users to understand the contents of datasets after they have performed a pre-selection (Marjanovic et al. 2024; Nikiforova et al. 2024). Depending on the use case, enhancing the datasets' interpretability and usability by adding additional context is another task that LLMs can be used for as highlighted in (Nikiforova et al. 2024).

Another potential use for LLMs could be the translation of contents to either support offering one's services in multiple languages or to be able to ingest data from different language sources and integrate them with one's own portfolio (Alexopoulos et al. 2024). However, in the example presented in (Kliimask and Nikiforova 2024), where translations were incorporated, instead of relying on the capabilities of ChatGPT, which was used for the generation of tags, DeepL⁷ was harnessed. Further, while not explicitly addressed in the identified papers, data completeness is also a serious issue (Alexopoulos et al. 2024) that could potentially be addressed through LLM-based control mechanisms.

Moving from the existing datasets to a more forward looking perspective, (Loukis et al. 2023) suggested that LLMs could also be helpful for analysing the OGD portals' user requests and usage data to detect needs as well as to identify data whose usefulness is limited, which is a somewhat common issue (Alexopoulos et al. 2024).

3.2 OGD Utilization

While generally the availability of vast OGD sets can be seen in a positive light, since it shows a desire for transparency and allows to cover many different domains and aspects, it also comes with its own challenges for the (potential) users. Besides the fact that some of the data might be incomplete or rather useless for value creation (Alexopoulos et al. 2024), one of the biggest challenges is the discovery of datasets or specific information that are relevant to the respective user's intentions. While the approaches mentioned in the previous section can help by increasing the descriptiveness of the datasets, huge potential also lies in directly supporting the users with this task.

For this reason, a lot of the current research is focused on doing so. However, the diversity of

⁷ <https://www.deepl.com>

proposed approaches is noteworthy, indicating on the one hand the versatility of LLMs in the context of OGD but on the other hand also a lack of maturity and best practices.

As could be expected, based on their strengths, the harnessing of LLMs to facilitate chatbots or data assistants that allow users to use natural language to inquire for information is highly present in the literature (Del Hoyo-Alonso et al. 2024; Loukis et al. 2023; Nikiforova et al. 2024; Schelhorn et al. 2024). This way, the data discovery could be considerably simplified, since users do not need to specifically search for suitable datasets, which they then have to explore and analyse, but instead, they just ask for the information they are interested in. Besides the improved convenience, this also greatly increases the accessibility, since the necessary capabilities to effectively interact with OGD (portals) are not always given (Alexopoulos et al. 2024; Schelhorn et al. 2024). Further, while the scope of the studies was limited, in (Barcellos et al. 2024) and (Schelhorn et al. 2024), the interest of potential users in such a solution was explored, showing that the incorporation of LLMs into the interaction with OGD is generally welcomed.

Besides chatbots that provide the desired answers by searching across the available datasets, the identified literature also contains several examples that focused on the provisioning of information once the relevant datasets are known. Showing how LLMs can help in the interaction with databases, (Mamalis et al. 2023) describes the creation of a LLM-based chatbot that interacts with the SPARQL endpoint of the statistical data portal of Scotland. Thus, instead of having to familiarize themselves with SPARQL queries, the users can use natural language, significantly decreasing the barrier for interaction. Another example can be found in (O'Leary 2025), where it was experimented with how ChatGPT can support the interaction with OGD that are available as Excel-files. Further, in (Cigliano and Fallucchi 2025), the use of LLMs as natural language interfaces to extract information from knowledge graphs is proposed. While the authors' suggestion is made without mentioning a specific application domain, this could help in making certain OGD more accessible to laypeople.

Moreover, as also mentioned in the section focused on the OGD provisioning, LLMs ability to provide summaries of information that would otherwise need a long time to be explored can be highly valuable, especially when it comes to developing an initial understanding of a topic (Loukis

et al. 2023; Marjanovic et al. 2024; Nikiforova et al. 2024).

Another aspect where LLMs might be of help is the linking of information across different datasets to provide more sophisticated insights (Benjira et al. 2025a, 2025b; Cigliano and Fallucchi 2025). Its utility can, for instance, be highlighted on the example of tracking the fulfilment of sustainable development goals, which are given in one document, based on corresponding indicators that are spread across a plethora of datasets (Benjira et al. 2025a, 2025b). To limit the manual work required for the matching based on the context, LLMs could be harnessed.

For concepts, where users can either submit requests for the provisioning of data or suggests certain government measures, LLMs can help in several ways. On the one hand, the aforementioned creation of summaries can allow users to quickly get an idea of already existing propositions to avoid duplicate request. On the other hand, LLMs can assist the users in writing their own proposals by helping them to adhere to guidelines regarding aspects such as language, writing style, or structure (Marjanovic et al. 2024). This reduces the required effort for users to provide suitable proposals, reducing the barrier for doing so, while also reducing the effort on the part of the operators of the portal because the (formal) quality of the obtained requests is increased and their structure is better aligned with internal standards and formatting specifications, simplifying their processing.

Another potential use case that is mentioned in (Loukis et al. 2023) and exemplarily demonstrated in (Takefuji 2024), is the harnessing of LLMs to provide assistance to programmers for creating OGD-based applications. This way, utilizing OGD to generate economic or social value from the available datasets becomes more accessible, allowing a bigger audience to contribute. Thereby, on a basic level, LLMs could be seen as an even more accessible alternative to the low-code development concept (Hintsch et al. 2021).

4 DISCUSSION

As the analysis of the identified literature showed, LLMs can be harnessed to perform or aid with a plethora of tasks related to OGD (Nikiforova et al. 2024). Hereby, their main potential is currently seen in making the available data more accessible and easier to handle. Instead of being mostly limited to users that possess the technical skills and the required OGD-literacy to effectively navigate the data

(Alexopoulos et al. 2024), LLMs can this way, a broader audience is enabled to interact with them, which, in turn, increases their provided benefit and strengthens the citizens' engagement and participation in the government.

Further, the use of LLMs can also facilitate the improvement of existing government-related services as well as the creation of new ones (Loukis et al. 2023). Hereby, these can be entirely based on OGD or by combining OGD and private data. An example for such a service is discussed in (Costa et al. 2024) where data that indicate bike-accidents are combined with map data to create an easily understandable visual map indicating different risk-zones for cyclists in the city of Porto. However, since in this example, the origin (are they derived from OGD or, for instance, from private hospitals or services) of the accident data is not entirely clear, this paper was not included in the literature review. Yet, it clearly shows the potential and is therefore mentioned in this place.

Nevertheless, despite the great potential, the use of LLMs in the context of OGD also comes with significant challenges that need to be thoroughly considered. This is also reflected in the literature, where multiple authors highlight that (at least current generation) LLMs are still sometimes producing erroneous responses (Cigliano and Fallucchi 2025; Del Hoyo-Alonso et al. 2024; Kliimask and Nikiforova 2024; Loukis et al. 2023; Schelhorn et al. 2024). Consequently, improving the models' accuracy and consistency or incorporating further control and correction measures is important for those cases, where a high quality of the results is essential. This might, for instance, be especially relevant in cases where the data are used as a basis for discussions on very controversial topics, when mistakes by the LLM could be misinterpreted as deliberate falsification of facts, eroding the trust in the data, the portal, or even the government itself.

Further, these models are usually black box in nature, making it hard to verify and understand the results, which could deter users from utilizing them (Loukis et al. 2023). This, however, could be addressed by integrating explainable AI principles to increase trust (Nikiforova et al. 2024; Schelhorn et al. 2024; Zoeten et al. 2024).

Other areas that were pointed out in the literature as potential barriers are related to data and user privacy, ethical concerns and regulatory frameworks (Loukis et al. 2023; Nikiforova et al. 2024), whereby especially the latter, however, also applies to OGD in general (Alexopoulos et al. 2024).

It is mentioned that linguistic differences across countries could result in algorithmic biases that might

compromise the performance of the used LLMs (Loukis et al. 2023).

Moreover, OGD consumers are sometimes faced with incomplete or erroneous data as well as with data formats that are highly heterogeneous, incompatible, or inappropriate (Alexopoulos et al. 2024). Yet, while LLMs can help in increasing the quality of data and also in integrating data from different sets and in different formats with each other, they also require high-quality, and if possible, well-curated input data in order to produce reliable and meaningful outputs (Nikiforova et al. 2024).

5 CONCLUSION

While the opening of government data to the public can bring enormous benefits through increased transparency and the potential creation of new services, accessing and analysing these data can be quite challenging. However, with the emergence of LLMs, a powerful new tool has become available that can facilitate the provisioning and utilization of OGD in many ways. In the publication at hand, a structured literature review is described that was conducted to compile the scientific literature on the use of LLMs in the context of OGD. Hereby, numerous application areas were identified and described. Further, in addition to the opportunities, also potential challenges were outlined. This way, researchers and practitioners alike are provided with an overview that allows them to develop an understanding of the domain, which they can then use to inform their own work. However, this work can only be seen as an introduction. To gain a deeper understanding, especially with regard to the technical details of the described implementations, the reader is advised to consult the identified papers. Further, the study's focus was only on OGD. Yet, expanding it to the use of LLMs in the context of open data in general may yield additional insights that might also be applicable in the realm of OGD. Therefore, this appears to be a worthwhile direction for future research. Overall, the conducted research showed that there is a lot of activity in the field, highlighting the perceived potential of LLMs to facilitate the use of OGD. Though, it also became clear that the field is still rather immature as indicated by the many different approaches and strategies being explored due to a lack of established best practices.

REFERENCES

- ACM History Committee. (2025). "ACM History," available at <https://www.acm.org/about-acm/acm-history>, accessed on Mar 14 2025.
- Adarkwah, M. A., Islam, A. Y. M. A., Schneider, K., Luckin, R., Thomas, M., and Spector, J. M. (2024). "Are Preprints a Threat to the Credibility and Quality of Artificial Intelligence Literature in the ChatGPT Era? A Scoping Review and Qualitative Study," *International Journal of Human-Computer Interaction*, pp. 1-14 (doi: 10.1080/10447318.2024.2364140).
- Ahmed, U. (2023). "Reimagining open data ecosystems: a practical approach using AI, CI, and knowledge graphs," in *BIR Workshops*, Ascoli Piceno, Italy. 13.09.2023 - 15.09.2023.
- Ahmed, U., Alexopoulos, C., Piangerelli, M., and Polini, A. (2024). "BRYT: Automated keyword extraction for open datasets," *Intelligent Systems with Applications* (23), p. 200421 (doi: 10.1016/j.iswa.2024.200421).
- Alexopoulos, C., Saxena, S., Janssen, M., Rizun, N., Lnenicka, M., and Matheus, R. (2024). "Why do Open Government Data initiatives fail in developing countries? A root cause analysis of the most prevalent barriers and problems," *The Electronic Journal of Information Systems in Developing Countries* (90:2) (doi: 10.1002/isd2.12297).
- Ansari, B., Barati, M., and Martin, E. G. (2022). "Enhancing the usability and usefulness of open government data: A comprehensive review of the state of open government data visualization research," *Government Information Quarterly* (39:1), p. 101657 (doi: 10.1016/j.giq.2021.101657).
- Attard, J., Orlandi, F., Scerri, S., and Auer, S. (2015). "A systematic review of open government data initiatives," *Government Information Quarterly* (32:4), pp. 399-418 (doi: 10.1016/j.giq.2015.07.006).
- Barcellos, R., Bernardini, F., Zuiderwijk, A., and Viterbo, J. (2024). "Exploring Interpretability in Open Government Data with ChatGPT," in *Proceedings of the 25th Annual International Conference on Digital Government Research*, H.-C. Liao, D. D. Cid, M. A. Macadar and F. Bernardini (eds.), Taipei Taiwan. 11.06.2024 - 14.06.2024, New York, NY, USA: ACM, pp. 186-195 (doi: 10.1145/3657054.3657079).
- Barry, E., and Bannister, F. (2014). "Barriers to open data release: A view from the top," *Information Polity* (19:1,2), pp. 129-152 (doi: 10.3233/IP-140327).
- Benjira, W., Atigui, F., Bucher, B., Grim-Yefsah, M., and Travers, N. (2025a). "Automated mapping between SDG indicators and open data: An LLM-augmented knowledge graph approach," *Data & Knowledge Engineering* (156), p. 102405 (doi: 10.1016/j.datak.2024.102405).
- Benjira, W., Atigui, F., Bucher, B., Grim-Yefsah, M., and Travers, N. (2025b). "Web Open Data to SDG Indicators: Towards an LLM-Augmented Knowledge Graph Solution," in *Web Information Systems Engineering – WISE 2024 PhD Symposium, Demos and Workshops*, M. Barhamgi, H. Wang, X. Wang, E. Aïmeur, M. Mrissa, B. Chikhaoui, K. Boukadi, R. Grati and Z. Maamar (eds.), Springer Nature Singapore, pp. 90-100 (doi: 10.1007/978-981-96-1483-7_7).
- Bonina, C., and Eaton, B. (2020). "Cultivating open government data platform ecosystems through governance: Lessons from Buenos Aires, Mexico City and Montevideo," *Government Information Quarterly* (37:3), p. 101479 (doi: 10.1016/j.giq.2020.101479).
- Brynjolfsson, E., Li, D., and Raymond, L. (2023). "Generative AI at Work," *NBER Working Paper Series* 31161, Cambridge, MA: National Bureau of Economic.
- Carta, S., Giuliani, A., Manca, M. M., Piano, L., Pisu, A., and Tiddia, S. G. (2024). "Instruct Large Language Models for Public Administration Document Information Extraction," in *Proceedings of the Ital-IA 2024 Thematic Workshops*, Naples, Italy. 29.05.2024 - 30.05.2024, pp. 424-429.
- Chang, Y., Wang, X., Wang, J., Wu, Y., Yang, L., Zhu, K., Chen, H., Yi, X., Wang, C., Wang, Y., Ye, W., Zhang, Y., Chang, Y., Yu, P. S., Yang, Q., and Xie, X. (2024). "A Survey on Evaluation of Large Language Models," *ACM Transactions on Intelligent Systems and Technology* (15:3), pp. 1-45 (doi: 10.1145/3641289).
- Cigliano, A., and Fallucchi, F. (2025). "The Convergence of Open Data, Linked Data, Ontologies, and Large Language Models: Enabling Next-Generation Knowledge Systems," in *Metadata and Semantic Research*, M. Sfakakis, E. Garoufallou, M. Damigos, A. Salaba and C. Papatheodorou (eds.), pp. 197-213 (doi: 10.1007/978-3-031-81974-2_17).
- Costa, D. G., Silva, I., Medeiros, M., Bittencourt, J. C. N., and Andrade, M. (2024). "A method to promote safe cycling powered by large language models and AI agents," *MethodsX* (13), p. 102880 (doi: 10.1016/j.mex.2024.102880).
- Del Hoyo-Alonso, R., Rodrigalvarez-Chamarro, V., Veas-Murgía, J., Zubizarreta, I., and Moyano-Collado, J. (2024). "Aragón Open Data Assistant, Lesson Learned of an Intelligent Assistant for Open Data Access," in *Chatbot Research and Design*, A. Folstad, T. Araujo, S. Papadopoulos, E. L.-C. Law, E. Luger, M. Goodwin, S. Hobert and P. B. Brandtzaeg (eds.), Cham: Springer Nature, pp. 42-57 (doi: 10.1007/978-3-031-54975-5_3).
- Fan, W., Ding, Y., Ning, L., Wang, S., Li, H., Yin, D., Chua, T.-S., and Li, Q. (2024). "A Survey on RAG Meeting LLMs: Towards Retrieval-Augmented Large Language Models," in *Proceedings of the 30th ACM SIGKDD Conference on Knowledge Discovery and Data Mining*, R. Baeza-Yates and F. Bonchi (eds.), Barcelona Spain. 25.08.2024 - 29.08.2024, New York, NY, USA: ACM, pp. 6491-6501 (doi: 10.1145/3637528.3671470).
- Filippucci, F., Gal, P., Jona-Lasinio, C., Leandro, A., and Nicoletti, G. (2024). "The impact of Artificial Intelligence on productivity, distribution and growth: Key mechanisms, initial evidence and policy challenges," *OECD Artificial Intelligence Papers*.
- Hintsch, J., Staegemann, D., Volk, M., and Turowski, K. (2021). "Low-code Development Platform Usage: Towards Bringing Citizen Development and Enterprise IT into Harmony," in *Proceedings of the 32nd*

- Australasian Conference on Information Systems (ACIS2021)*, Online. 06.12.2021 - 10.12.2021.
- Huang, L., Yu, W., Ma, W., Zhong, W., Feng, Z., Wang, H., Chen, Q., Peng, W., Feng, X., Qin, B., and Liu, T. (2024). "A Survey on Hallucination in Large Language Models: Principles, Taxonomy, Challenges, and Open Questions," *ACM Transactions on Information Systems* (doi: 10.1145/3703155).
- Janssen, M., Charalabidis, Y., and Zuiderwijk, A. (2012). "Benefits, Adoption Barriers and Myths of Open Data and Open Government," *Information Systems Management* (29:4), pp. 258-268 (doi: 10.1080/10580530.2012.716740).
- Kliimask, K., and Nikiforova, A. (2024). "TAGIFY: LLM-powered Tagging Interface for Improved Data Findability on OGD portals," in *2024 Fifth International Conference on Intelligent Data Science Technologies and Applications (IDSTA)*, Dubrovnik, Croatia. 24.09.2024 - 27.09.2024, IEEE, pp. 18-27 (doi: 10.1109/IDSTA62194.2024.10746941).
- Kraus, S., Breier, M., Lim, W. M., Dabić, M., Kumar, S., Kanbach, D., Mukherjee, D., Corvello, V., Piñeiro-Chousa, J., Liguori, E., Palacios-Marqués, D., Schiavone, F., Ferraris, A., Fernandes, C., and Ferreira, J. J. (2022). "Literature reviews as independent studies: guidelines for academic practice," *Review of Managerial Science* (16:8), pp. 2577-2595 (doi: 10.1007/s11846-022-00588-8).
- Kucera, J., and Chlapek, D. (2014). "Benefits and Risks of Open Government Data," *Journal of Systems Integration*, pp. 30-41 (doi: 10.20470/jsi.v5i1.185).
- Loukis, E., Saxena, S., Rizun, N., Maratsi, M. I., Ali, M., and Alexopoulos, C. (2023). "ChatGPT Application vis-a-vis Open Government Data (OGD): Capabilities, Public Values, Issues and a Research Agenda," in *Electronic Government*, I. Lindgren, C. Csáki, E. Kalampokis, M. Janssen, G. Viale Pereira, S. Virkar, E. Tambouris and A. Zuiderwijk (eds.), Springer Nature Switzerland, pp. 95-110 (doi: 10.1007/978-3-031-41138-0_7).
- Mamalis, M. E., Kalampokis, E., Karamanou, A., Brimos, P., and Tarabanis, K. (2023). "Can Large Language Models Revolutionize Open Government Data Portals? A Case of Using ChatGPT in statistics.gov.scot," in *Proceedings of the 27th Pan-Hellenic Conference on Progress in Computing and Informatics*, N. N. Karanikolas, M. G. Vassilakopoulos, C. Marinagi, A. Kakarountas and I. Voyiatzis (eds.), Lamia, Greece. 24.11.2023 - 26.11.2023, New York, NY, USA: ACM, pp. 53-59 (doi: 10.1145/3635059.3635068).
- Marjanovic, O., Hristova, D., Dash, S., and and Abedin, B. (2024). "On Enabling Dynamic, Transparent, and Inclusive Government Consultative Processes with GenAIOpenGov DSS," in *ACIS 2024 Proceedings*, Canberra, Australia. 04.12.2024 - 06.12.2024.
- Nikiforova, A., Lnenicka, M., Milić, P., Luterek, M., and Rodríguez Bolívar, M. P. (2024). "From the Evolution of Public Data Ecosystems to the Evolving Horizons of the Forward-Looking Intelligent Public Data Ecosystem Empowered by Emerging Technologies," in *Electronic Government*, M. Janssen, J. Cromptvoets, J. R. Gil-Garcia, H. Lee, I. Lindgren, A. Nikiforova and G. Viale Pereira (eds.), Springer Nature Switzerland, pp. 402-418 (doi: 10.1007/978-3-031-70274-7_25).
- O'Leary, D. (2025). "AI for Good: History, Open Data and Some ESG-based Applications," *Journal of Decision Systems* (34:1) (doi: 10.1080/12460125.2024.2443182).
- Okoli, C. (2015). "A Guide to Conducting a Standalone Systematic Literature Review," *Communications of the Association for Information Systems* (37), pp. 879-910 (doi: 10.17705/1CAIS.03743).
- Perković, G., Drobnjak, A., and Botički, I. (2024). "Hallucinations in LLMs: Understanding and Addressing Challenges," in *2024 47th MIPRO ICT and Electronics Convention (MIPRO)*, Opatija, Croatia. 20.05.2024 - 24.05.2024, IEEE, pp. 2084-2088 (doi: 10.1109/MIPRO60963.2024.10569238).
- Quarati, A. (2023). "Open Government Data: Usage trends and metadata quality," *Journal of Information Science* (49:4), pp. 887-910 (doi: 10.1177/01655515211027775).
- Raiaan, M. A. K., Mukta, M. S. H., Fatema, K., Fahad, N. M., Sakib, S., Mim, M. M. J., Ahmad, J., Ali, M. E., and Azam, S. (2024). "A Review on Large Language Models: Architectures, Applications, Taxonomies, Open Issues and Challenges," *IEEE Access* (12), pp. 26839-26874 (doi: 10.1109/ACCESS.2024.3365742).
- Schelhorn, T. C., Gnewuch, U., and Maedche, A. (2024). "Designing a Large Language Model Based Open Data Assistant for Effective Use," in *Design Science Research for a Resilient Future*, M. Mandviwalla, M. Söllner and T. Tuunanen (eds.), Springer Nature Switzerland, pp. 398-411 (doi: 10.1007/978-3-031-61175-9_27).
- Simons, W., Turrini, A., and Vivian, L. (2024). "Artificial Intelligence: Economic Impact, Opportunities, Challenges, Implications for Policy," *European Economy Discussion Papers* 210, European Union.
- Takefuji, Y. (2024). "Unveiling inequality: A deep dive into racial and gender disparities in US court case closures," *Cities* (154), p. 105398 (doi: 10.1016/j.cities.2024.105398).
- vom Brocke, J., Simons, A., Niehaves, B., Reimer, K., Plattfaut, R., and Cleven, A. (2009). "Reconstructing the Giant: On the Importance of Rigour in Documenting the Literature Search Process," in *Proceedings of the ECIS 2009*, Verona, Italy. 08.06.2009 - 10.06.2009.
- vom Brocke, J., Simons, A., Riemer, K., Niehaves, B., Plattfaut, R., and Cleven, A. (2015). "Standing on the Shoulders of Giants: Challenges and Recommendations of Literature Search in Information Systems Research," *Communications of the Association for Information Systems* (37) (doi: 10.17705/1CAIS.03709).
- Zoeten, M. C. de, Ernst, C.-P. H., and Rothlauf, F. (2024). "The Effect of Explainable AI on AI-Trust and Its Antecedents over the Course of an Interaction," in *ECIS 2024 Proceedings*, Paphos, Cyprus. 13.06.2024 - 19.06.2024.