



Let's Hide from LLMs: An Adaptive Contextual Privacy Preservation Method for Time Series Data

Ubaid Ur Rehman

Kyung Hee University (Global Campus), Yongin-si,
Republic of Korea
ubaid.rehman@khu.ac.kr

Tri D.T. Nguyen

Kyung Hee University (Global Campus), Yongin-si,
Republic of Korea
tringuyendt@khu.ac.kr

Musarrat Hussain

Kyung Hee University (Global Campus), Yongin-si,
Republic of Korea
musarrat.hussain@oslab.khu.ac.kr

Sungyoung Lee

Kyung Hee University (Global Campus), Yongin-si,
Republic of Korea
sylee@oslab.khu.ac.kr

ABSTRACT

The emergence of the Internet of Things (IoT) has evolved various application areas, such as healthcare, smart energy management, and autonomous vehicles. These devices continuously transmit time-series data that can be utilized by a variety of applications to provide personalized services. Recently, Large Language Models (LLMs) have been widely adopted in these application areas to input time-series data into prompts for in-context learning and to retrieve relevant responses accordingly. The time-series data contains sensitive information, and its processing can lead to privacy concerns. Several solutions have been proposed in the literature using differential privacy, which protects single data points or batch-wise privacy preservation through manual configuration of the privacy parameter (ϵ). In this paper, we propose an adaptive contextual privacy preservation method that analyzes the data attributes required for specific application services, acting as context. It then identifies sensitive attributes and adaptively selects the value of ϵ for each data attribute to maintain a balance between privacy and service requirements. The proposed approach was evaluated using power consumption and solar power generation datasets. The results show that the proposed approach dynamically selects the privacy parameter for each data attribute. Moreover, the original and anonymized data were fed into the prompt to assess the textual responses generated by LLM. The results show that our proposed approach achieved an average degree of semantic similarity score of 94.5% for power consumption data and 95.23% for solar power generation data.

CCS CONCEPTS

• Security and privacy; • Security services; • Pseudonymity, anonymity and untraceability;

KEYWORDS

Large Language Models, Privacy Preservation, Differential Privacy, Time Series Data



This work is licensed under a Creative Commons Attribution International 4.0 License.

AICCC 2023, December 16–18, 2023, Kyoto, Japan

© 2023 Copyright held by the owner/author(s).

ACM ISBN 979-8-4007-1622-5/23/12

<https://doi.org/10.1145/3639592.3639619>

ACM Reference Format:

Ubaid Ur Rehman, Musarrat Hussain, Tri D.T. Nguyen, and Sungyoung Lee. 2023. Let's Hide from LLMs: An Adaptive Contextual Privacy Preservation Method for Time Series Data. In *2023 6th Artificial Intelligence and Cloud Computing Conference (AICCC) (AICCC 2023), December 16–18, 2023, Kyoto, Japan*. ACM, New York, NY, USA, 8 pages. <https://doi.org/10.1145/3639592.3639619>

1 INTRODUCTION

The rapid advancement of IoT has raised privacy concerns because it continuously transmits time-series data that includes sensitive information. The acquired data are either stored in data repositories managed by the data curator for retrieval as needed or utilized by domain-specific applications to provide personalized services or decision-making. Usually, data curators are considered trustworthy and manage the data received from data sources in a secure and organized manner. However, with the emergence of the zero-trust security model, which relies on the principle of "never trust, always verify," the concern arises regarding what happens if the data curator is untrusted. This leads to privacy concerns because the data managed by the data curator includes sensitive and personally identifiable information, and any leakage or mishandling may lead to serious consequences. Therefore, some privacy mechanisms need to be deployed alongside authentication and authorization for data curators to ensure data protection, compliance with the General Data Protection Regulation (GDPR), and the Health Insurance Portability and Accountability Act (HIPAA), and to mitigate associated risks.

Another major concern is sharing the acquired time-series data with third-party services, such as data analysts, service providers, and Large Language Models (LLMs). These services utilize data-hungry machine learning systems, expert heuristics, and statistical methods to analyze the data, extract insights, and process the data according to the target requirements for making informed decisions. Therefore, the need for privacy and compliance with data governance is crucial when sharing data with these third-party services. Specifically, the widespread adoption of LLMs by different domain-specific applications has raised concerns regarding the potential leakage of sensitive information. LLMs interact in natural language, answer questions, and assist users with various tasks. To retrieve specific responses from LLMs, domain-specific applications feed the time-series data into prompts for in-context

learning and retrieve relevant responses, leading to privacy concerns. Carlini et al. in [1] demonstrated an attack on Generative Pre-trained Transformer (GPT-2) to extract verbatim text sequences and personally identifiable information used to train GPT-2. Therefore, a privacy preservation approach is needed to ensure a balance between privacy and the required service.

In this study, we have proposed an adaptive contextual privacy preservation method for time series data using differential privacy. The proposed approach addresses the two privacy concerns, including the untrusted data curator and data handling with third-party services. For each service, it identifies the context by grouping the required data, then identifies sensitive data attributes and adaptively allocates the privacy parameter for each attribute. It then computes and adds Laplace noise to the targeted time-series data to ensure differential privacy. We have evaluated our proposed approach with two datasets to assess contextual privacy preservation and dynamic allocation of privacy parameters. The results show that the proposed approach efficiently identifies the context and adaptively allocates the privacy parameter. To ensure the privacy preservation of time-series data in LLM, we have created prompt templates, inserted the original and noisy data, and acquired responses from LLM given the prompt. We then compared the responses received using original and noisy prompts, which showed a very low impact of changes on the generated responses. Therefore, we can utilize the noisy prompt to obtain responses from LLM instead of considering the original data due to privacy concerns.

The rest of the paper is organized as follows: Section 2 provides an overview of related work. Section 3 briefly describes our proposed adaptive contextual privacy preservation method. The evaluation of our proposed approach based on two datasets and its evaluation with LLM is presented in Section 4. Finally, Section 5 summarizes the proposed approach, its limitations, and sets future directions.

2 RELATED WORK

In this section, we describe the most recent and relevant literature related to our proposed methodology. We have identified eight studies, among which six studies have considered LLMs, specifically privacy preservation mechanisms for named entity types, but did not consider time series data with LLMs. However, two of the studies have considered privacy preservation techniques but not from the perspective of LLMs. In this study, we bridge this gap and consider privacy preservation for time-series data along with LLMs. The detailed descriptions of the existing approaches are as follows.

2.1 Privacy Preservation Techniques Considering LLMs

In this section, we describe the existing approaches developed or applied in the context of LLMs. Yermilov et al. in [2] analyzed the effectiveness of pseudonymization techniques using various datasets and models used for text classification and summarization. The authors emphasized identifying a suitable pseudonymization method for LLMs that can strike a balance between privacy and utility. The proposed approach focused on named entity types (person, location, and organization) but did not consider the assessment of time series data.

In [3], the authors analyzed the responses generated from LLMs for input copying and regurgitation. For this purpose, they created three prompts that comply with privacy regulations and protect sensitive information from leakage to retrieve responses from LLMs. The authors concluded that only prompts compliant with privacy regulations cannot guarantee privacy preservation. Therefore, an effective privacy preservation method can reduce the risk of information leakage to LLMs.

Kim et al. [4] proposed a probing tool to create awareness about the leakage of Personally Identifiable Information (PII) to LLMs. The approach used masked templates and acquired responses from LLMs. If the retrieved information resembled the ground truth, it indicated a high privacy risk to PII. However, the approach did not propose a privacy preservation mechanism to prevent the leakage of PII.

Chen et al. in [5] proposed an anonymization technique to protect PII information in a prompt by substituting it with filler or a mask. The proposed approach only considered named entity types and is not suitable for time series data.

In [6], the author proposed a data curator who collects data from different stakeholders and applies masking to ensure privacy. They then used the masked data to extract valuable insights from LLMs. The approach identified the named entity types from the aggregated data, masked them, and provided that data in prompts to retrieve a response. The approach by South et al. is not suitable for time series data.

Beyer et al. in [7] preserved the privacy of quasi-identifiers (age, marital status, or occupation) and proposed a prompting technique to engage LLMs in generating responses between a hypothetical user and a virtual assistant. The approach required an expert to annotate the quasi-identifiers dataset, which is a time-consuming process and requires specialized professionals.

2.2 Privacy Preservation Techniques Considering Time Series Data

In this section, we describe the methods or techniques proposed to preserve the privacy of time-series data. Katsomallos et al. in [8] proposed landmark privacy that quantifies the privacy loss under temporal correlation. The proposed approach preserves the actual timestamp of the landmarks but only considered the timestamp, not the actual value at the corresponding timestamp.

In [9], the author proposed a data curator who collects time-series data from IoTs and based on the service request, tailors it to meet the requirements, ensuring that PII is not shared. This approach did not consider the privacy concern of an untrusted data curator.

3 ADAPTIVE CONTEXTUAL PRIVACY PRESERVATION METHOD

We have identified from the existing literature that the data curator is usually considered a trusted resource, and data from physical resources can be directly stored and managed by the data curator. However, if the data curator itself is untrusted, then data privacy is at risk. In our proposed approach, we have considered two privacy concerns, as shown in Figure 1.

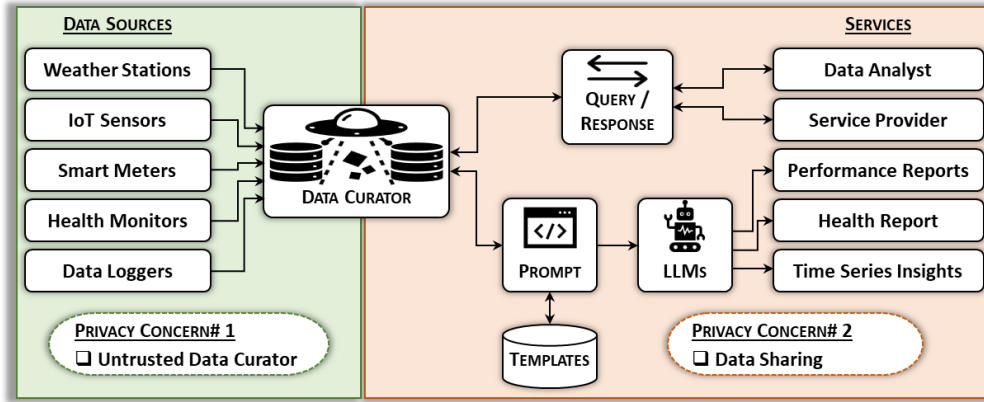


Figure 1: Privacy Concerns in terms of Untrusted Data Curator and Data Sharing

The first concern highlights that when data sources acquire data and share it with the data curator, no privacy preservation techniques are applied. The data curator can analyze and extract behavior patterns accordingly. Therefore, we need to eliminate the assumption of considering the data curator as a trusted party and take necessary measures accordingly.

The second concern is the sharing of data with targeted services. Usually, the acquired data can be utilized by data analysts, service providers, and, more recently, Large Language Models (LLMs). Data analysts can identify patterns, behaviors, insights, and correlations between different attributes. Service providers can utilize machine learning, expert heuristics, and statistical methods to provide relevant domain-specific personalized services. Additionally, with the popularity of LLMs and their adoption by different service providers, such data are now inserted into predefined prompt templates to provide human-interpreted text in the form of performance reports, health reports, and time series insights. Therefore, sharing data with these third-party services raises significant privacy concerns, especially in the domains of healthcare, energy management, and performance management.

Given the issues of an untrusted data curator and data sharing privacy, we have proposed an adaptive contextual privacy preservation method. This method is based on the concept of a zero-trust model, which eliminates the assumption of trust in any resource or communicating entity. Figure 2 illustrates the conceptual workflow of our proposed methodology. The details of each component are described as follows.

3.1 Data Sources

The data sources consist of various physical or virtual resources that can assess their surroundings and acquire data. Since the focus of our proposed approach is on privacy preservation for time series data, we have specified the types of resources that can acquire such data. These include weather stations, IoT sensors, smart meters, health monitors, and data loggers.

3.1.1 Weather Stations. Weather station data includes information on temperature, humidity, wind speed, and atmospheric pressure,

which is used for weather forecasting. Weather station data plays a crucial role in proactively managing power consumption.

3.1.2 IoT Sensors. IoT sensors encompass a variety of sensor types designed for different purposes, such as ambient sensors for environmental conditions, motion sensors for occupancy detection, proximity sensors for range identification, and biometric sensors for security purposes. Data acquired from IoT sensors is typically very sensitive and requires proper handling to avoid privacy risks.

3.1.3 Smart Meters. Smart meter data usually consists of consumer information related to electricity, gas, and water usage. Analyzing such data can help identify user consumption behavior and provide cost benefits for smart energy management. Therefore, the privacy preservation of such data is crucial in the energy management domain.

3.1.4 Health Monitors. Health monitors include data from wearables and medical devices, which are directly related to personal health information and may contain personally identifiable information. This data is used by different healthcare applications to provide various services, such as medication adherence and wellness-related decisions.

3.1.5 Data Logger. Data loggers store information over time that can be used for anomaly detection, security information, event management, and disaster management. They typically contain very sensitive records and require proper security and privacy measures.

3.2 Data and Service Mapper

The data and service mappers acquire two types of inputs: (i) from data sources and (ii) from services. As illustrated in Figure 1, the data sources include weather stations, IoT sensors, smart meters, health monitors, and data loggers. Similarly, the services include data analysts, service providers, and LLMs. The data and service mapper acquire this information and map it accordingly.

3.2.1 Data Analyst. Data analysts are interested in identifying patterns, behaviors, insights, and correlations between different

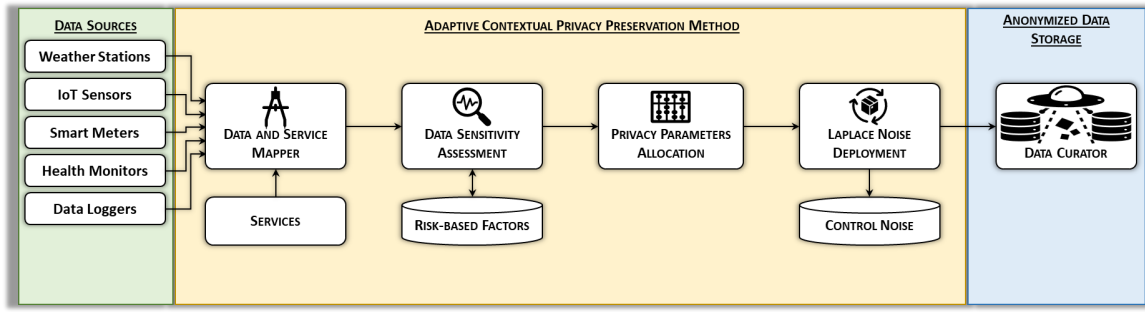


Figure 2: Adaptive Contextual Privacy Preservation Method

attributes. Therefore, the data and service mapper map the data analyst with each resource available in the data sources.

3.2.2 Service Provider. The services provided depend on the types of services, resources, and operational processes. It relies on the integration of our proposed approach with service provider applications. In the case of healthcare and medication adherence applications, data from IoT sensors and health monitors can be provided. For anomaly detection and prevention, the service provider can gain access to all types of data. Similarly, in the case of energy storage system scheduling and energy price estimation, the service is mapped with data from smart meters, IoT sensors, and weather station data.

3.2.3 Large Language Models (LLMs). LLMs understand natural language and generate responses based on given prompts. The prompt plays a crucial role because it helps LLMs understand specific queries and generate responses accordingly. Therefore, the prompt needs to contain specific information to provide a basic context to LLMs for generating relevant reports. The data and service mapper maps the data with the required information by the prompt to generate a specific textual report.

For example, in the case of a performance report, the prompt may require access to all types of data to provide trend analysis, narrative generation, and recommendations. In the case of a health report, the prompt may need access to vital signs, physical activities, and medical history. Similarly, for time series insights generated from LLMs, the prompt needs to provide historical data for effective response generation.

3.3 Data Sensitivity Assessment

The data sensitivity assessment involves acquiring data that has been mapped with services from the data and service mapper. It then identifies attributes that may have privacy constraints, such as vital signs, medical history, and smart meter data. We have adopted a risk-based approach for analyzing data sensitivity, which is a rule-based method that includes expert heuristics in the form of rules to categorize each attribute as sensitive (1) or non-sensitive (0).

After processing the data, each attribute is assigned a flag. For instance, vital signs assigned with 1 are categorized as sensitive, while wind speed assigned with 0 is categorized as non-sensitive. It is possible that one attribute can be considered very sensitive

for a specific service, but the same attribute may be considered non-sensitive for another type of service. Therefore, the risk-based approach considers all these constraints, including the targeted services, and assigns sensitivity tags accordingly.

3.4 Privacy Parameters Allocation

Privacy parameter allocation is an important step in our proposed solution. The context about a specific service and sensitive attributes has already been identified in the previous steps. Privacy parameter allocation adaptively selects parameters in such a way that it ensures a proper balance between privacy and service requirements. This ensures that the transformed privacy-preserved data does not degrade the accuracy of the targeted service.

Since our proposed approach is related to time series data, we have used a statistical privacy-preserving technique called differential privacy [10]. This technique adds controlled noise to each data point based on a Laplace or Gaussian distribution, ensuring that individual data points cannot be re-identified. We have used differential privacy with a Laplace distribution to account for the heavy-tailed behavior of time series data, which may have occasional extreme values. Differential privacy with Laplace distribution depends on the privacy parameter (ϵ), sensitivity (Δf), and location parameter (μ).

3.4.1 Privacy Parameter (ϵ). The privacy parameter (ϵ) quantifies the level of differential privacy. A small ϵ value means more noise is introduced, making it harder for adversaries to distinguish, leading to stronger privacy. A larger ϵ value means less noise and weaker privacy. Our proposed approach adaptively selects the value of ϵ from the set [1, 10, 100, 1000], which corresponds to [Strong Privacy, Medium-High Privacy, Medium Privacy, Low Privacy], respectively. The method analyzes each attribute with its sensitivity flag and selects the appropriate ϵ value for each attribute, ensuring that the cumulative privacy-preserved data can be used by each service for performing specific tasks. For example, for a dataset with three attributes, two of which are highly sensitive, and one with low sensitivity, the approach may allocate ϵ values of 10 and 100 to the highly sensitive attributes and 1000 to the low sensitivity attribute, resulting in an overall medium level of privacy for the dataset.

3.4.2 Sensitivity (Δf). Sensitivity (Δf) quantifies the change that occurs when adding or removing a data point from the dataset. It identifies the controlled noise to be added to the query result to

achieve differential privacy. Since our proposed approach deals with time series data, which contains numerical values, we define sensitivity as $\Delta f = \text{Maximum}_{\text{value}} - \text{Minimum}_{\text{value}}$. This equation calculates the difference between the maximum and minimum values in the selected dataset attribute. The difference quantifies the change, where a higher Δf value indicates greater sensitivity and vice versa.

3.4.3 Location Parameter (μ). The location parameter (μ) specifies the center of the distribution. It adjusts the center of the controlled noise added to the query results. For our proposed approach, we use the same scale of differential privacy, defined as $\mu = \Delta f / \epsilon$. A larger μ value means the noise distribution is centered farther away from the query result, resulting in less noisy data and weaker privacy. A smaller μ value centers the noise closer to the query result, leading to more noisy data and stronger privacy.

3.5 Laplace Noise Deployment

After privacy parameters allocation (ϵ , Δf , μ), Laplace noise is computed and added to the targeted time-series data to ensure differential privacy and maintain data utility. Our proposed approach stores the controlled noise added to each dataset in a repository with a unique identifier. Storing the controlled noise for each dataset ensures that the same privacy preservation mechanism is consistently applied to the specific data in the future.

3.6 Anonymized Data Storage

The data curator acquires the anonymized data and stores it in an organized manner. The transformed data protects individual identities while retaining the statistical properties and patterns of the original data. Therefore, it is suitable for identifying patterns, behaviors, insights, and correlations between different attributes. The service provider can utilize the anonymized data to provide utilities with the same level of decision-making accuracy as the original data. Additionally, the same data attributes can be used by the prompt to generate textual responses and prevent information leakage to the LLMs. In this way, our proposed approach addresses the privacy concerns related to an untrusted data curator and data sharing with third-party services.

4 EVALUATION

To evaluate the proposed adaptive contextual privacy preservation method, we selected a case study of an energy management system, where the smart meter collects power consumption data from consumers, and the solar energy meter measures the power generated from solar photovoltaic systems. We acquired two datasets related to power consumption and solar power generation and applied our proposed methodology independently to assess its impact. The details of each dataset are described as follows.

4.1 Power Consumption Data

Power consumption data can be used by adversaries to identify the living behavior of individuals and could potentially lead to burglary if leaked to malicious users. Therefore, a privacy preservation mechanism is required to prevent such information leakage. We acquired domestic power consumption data consisting of 473 data

points with four features, including equipment, host interface device, date-time, and power consumption. Our proposed approach identifies power consumption as highly sensitive and considers the other three features as low sensitivity within the domain of the energy management system.

The proposed approach first computes differential privacy by considering ϵ values of [1, 10, 100, 1000], and then identifies the correlation between the original and noisy data. Based on the results obtained from the correlation analysis, privacy requirements, and service requirements, the proposed approach adaptively selects the data that can be utilized by data analysts, service providers, and LLMs to provide specific services. Figure 3 illustrates the results obtained for each privacy parameter. The proposed approach selected the $\epsilon = 100$ data as the final anonymized power consumption data, ensuring a balance between privacy and utility requirements.

To assess the impact of anonymized data on service requirements, we used an LLM to generate textual reports using both the original and noisy data of power consumption.

4.1.1 Assessing the Impact of Anonymized Power Consumption Data. To evaluate the impact of anonymized power consumption data on textual report generation, we utilized the GPT-2 library from Hugging Face as an LLM [11]. We then created three prompt templates and provided both the original and noisy values to obtain responses from the LLM, as illustrated in Figure 4.

To compute the similarity between the responses obtained from utilizing the original and noisy power consumption data, we employed a Semantic Textual Similarity (STS) model. This model computes the contextual embedding of each response using a pre-trained Bidirectional Encoder Representations from Transformers (BERT), and subsequently calculates the cosine similarity between the contextual embeddings to obtain the similarity score. Figure 5 presents the similarity scores obtained using STS.

The results show that the value of ϵ was appropriately selected for the power consumption dataset, achieving an average degree of semantic similarity of 94.5% for these three prompts. Additionally, the generated textual description is dependent on the given prompt template, as illustrated in Figure 5, where the results obtained from prompts 1 and 2 are different, compared with those from prompt 3.

4.2 Solar Power Generation Data

With the emergence of green energy, consumers become prosumers by producing electricity from photovoltaic sources and then selling the excess energy to utility companies to either generate revenue or reduce utility costs. Energy companies have introduced smart bidding platforms to collect bids from prosumers, consumers, and distributors, using energy trading mechanisms to identify lists of buyers and sellers with cost benefits to meet energy utility requirements. If an adversary obtains information about an individual's generated power, they could potentially compute the financial benefits that the individual could receive from energy trading, which raises privacy concerns.

To address these concerns, we applied our proposed approach to the solar power generation dataset, which consists of 3,465 data points with four features: date, total (kW), maximum (kWh), and minimum (kWh) power generated per hour. Our proposed approach assessed the correlation between the original and noisy attributes

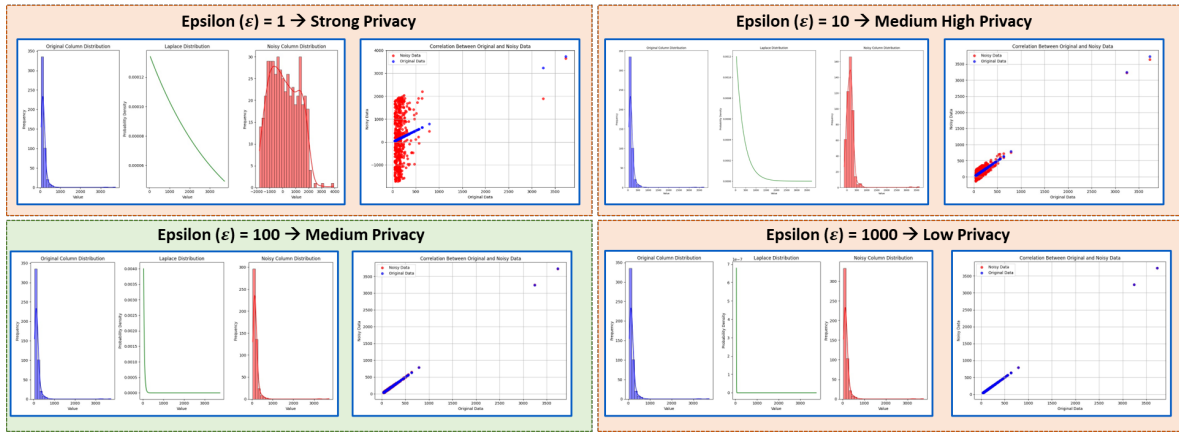
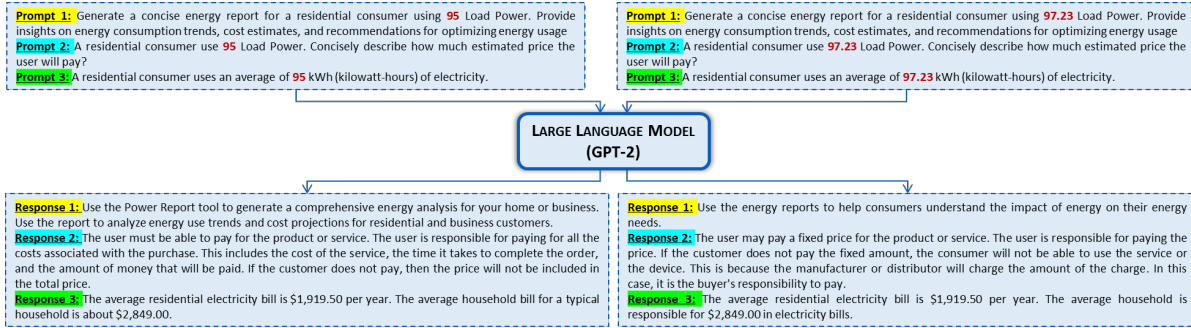
Figure 3: Adaptive Selection of Privacy Parameter ($\epsilon = 100$) on Power Consumption Data

Figure 4: Retrieving LLM Response Using Original and Noisy Power Consumption Data

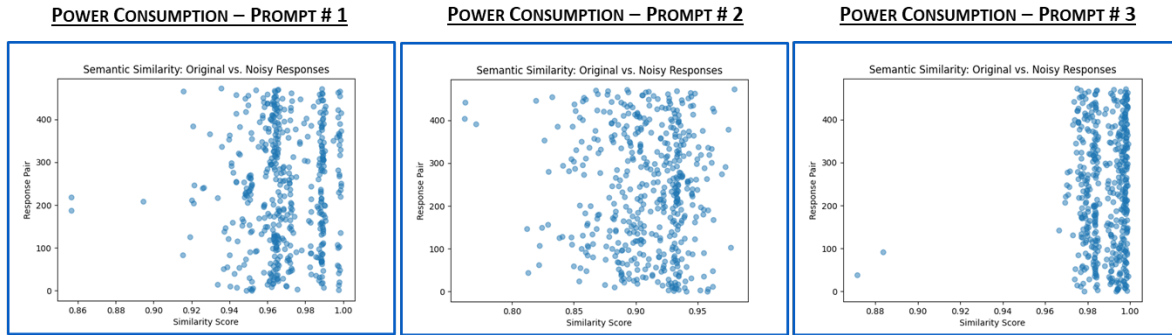


Figure 5: Semantic Textual Similarity Results on Power Consumption Dataset

after the privacy parameter allocation and then selected the most appropriate value of ϵ for each attribute, as shown in Figure 6. The proposed approach adaptively selected $\epsilon = 1000$ for total (kW), $\epsilon = 100$ for maximum (kWh), and $\epsilon = 10$ for minimum (kWh). The dynamic selection of ϵ values was due to the varying impact of each attribute.

Total (kW) was assigned a lower level of privacy because it has a high impact on decision-making, followed by maximum (kWh),

and then minimum (kWh). Therefore, our proposed approach considered these constraints to ensure privacy preservation. If the anonymized value significantly deviates from the original value, it can lead to inaccurate results. In this case, the overall privacy preservation of solar power generation is at a medium level.

4.2.1 Assessing the Impact of Anonymized Solar Power Generated Data. We also assessed the impact of anonymized solar power generation data on textual report generation. To do so, we formulated three prompt templates and inserted both the original and noisy

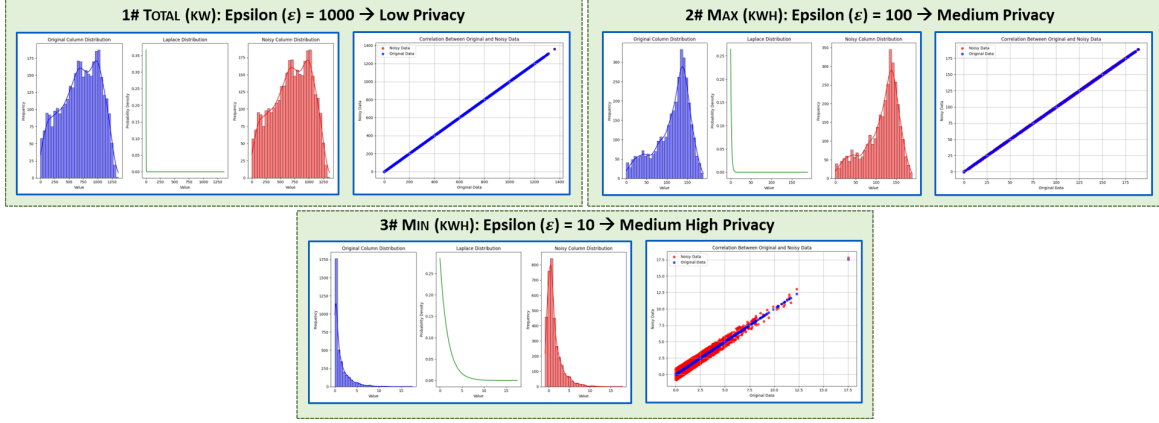
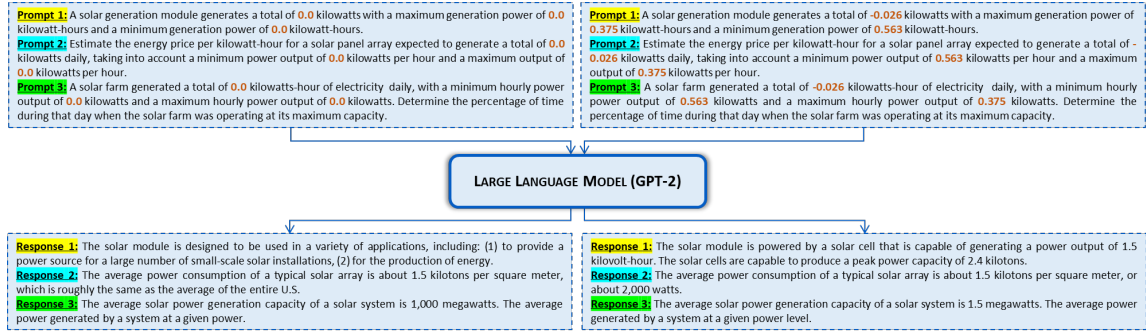
Figure 6: Adaptive Selection of Privacy Parameter ($\epsilon = [1000, 100, 10]$) on Solar Power Generation Data

Figure 7: Retrieving LLM Response Using Original and Noisy Solar Power Generation Data

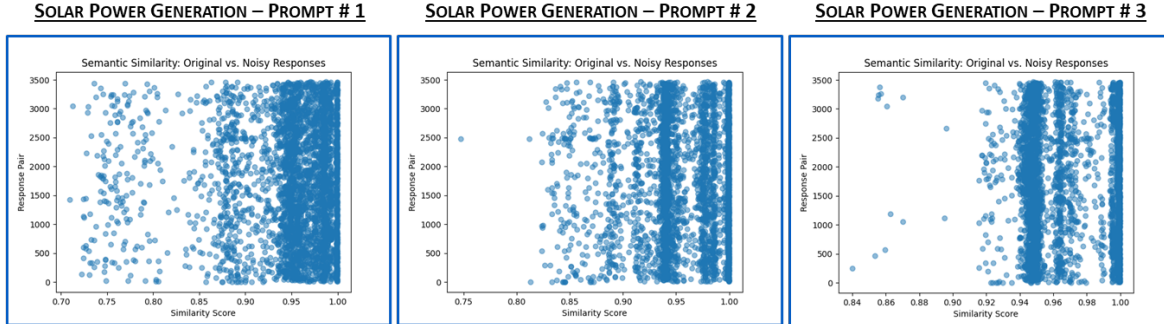


Figure 8: Semantic Textual Similarity Results on Solar Power Generation Dataset

values to generate responses from the LLM, as illustrated in Figure 7.

The similarity score was computed using STS between the responses obtained from the original and noisy prompts. Figure 8 illustrates the similarity scores obtained from these three prompts using solar power generation data. The results obtained from the dynamic selection of ϵ values achieved an average of 95.23% degree of semantic similarity based on these three prompts. Moreover, the prompts included more specific information about the desired

report. Therefore, the generated reports were very concise and specific, leading to higher similarity scores in this case.

5 CONCLUSION

In this study, we proposed an adaptive contextual privacy preservation method for time series data using differential privacy. Our approach groups data based on the targeted service, identifies sensitive attributes through risk-based factors, and adaptively selects the privacy parameter based on correlation, privacy levels, and service

requirements. After privacy parameter allocation, Laplace noise is computed and added to the targeted time series data to ensure differential privacy. Our objective was to eliminate the assumption of trust and ensure privacy preservation during data sharing with third-party services, including LLMs. We evaluated our approach using power consumption and solar power generation datasets, and the results showed that the proposed approach dynamically selects the privacy parameter for each data attribute. Furthermore, we assessed the impact of anonymized data on textual report generation using GPT-2 and computed the STS of responses obtained from original and noisy prompts. Our proposed approach achieved an average degree of semantic similarity score of 94.5% for power consumption data and 95.23% for solar power generation data. We also identified that the generation of textual descriptions is dependent on the given prompt template.

However, our proposed approach has limitations, as it may not be suitable for application areas that require specific values for decision-making. Additionally, the approach is resource-intensive, as it computes the Laplace distribution based on given ϵ values and selects the most appropriate privacy parameter for each data attribute based on correlation results. In the future, we will address these limitations to enhance our proposed approach and evaluate it with diverse data from different domain applications.

ACKNOWLEDGMENTS

This research was supported by the MSIT (Ministry of Science and ICT), Korea, under the Grand Information Technology Research Center support program (IITP-2022-2020-0-01489), the ITRC (Information Technology Research Center) support program (RS-2023-00259004) supervised by the IITP (Institute for Information & Communications Technology Planning & Evaluation) and by

the Institute of Information & Communications Technology Planning & Evaluation (IITP) grant funded by the Korean government (MSIT) (IITP-2022-0-00078, Explainable Logical Reasoning for Medical Knowledge Generation), (IITP-2017-0-00655, Lean UX core technology and platform for any digital artifacts UX evaluation).

REFERENCES

- [1] Carlini, Nicholas, Florian Tramer, Eric Wallace, Matthew Jagielski, Ariel Herbert-Voss, Katherine Lee, Adam Roberts *et al.* "Extracting training data from large language models." In 30th USENIX Security Symposium (USENIX Security 21), pp. 2633-2650. 2021.
- [2] Yermilov, Oleksandr, Vipul Raheja, and Artem Chernodub. "Privacy-and Utility-Preserving NLP with Anonymized Data: A case study of Pseudonymization." arXiv preprint arXiv:2306.05561 (2023).
- [3] Priyanshu, Aman, Supriti Vijay, Ayush Kumar, Rakshit Naidu, and Fatemehsadat Mirehshgallah. "Are Chatbots Ready for Privacy-Sensitive Applications? An Investigation into Input Regurgitation and Prompt-Induced Sanitization." arXiv preprint arXiv:2305.15008 (2023).
- [4] Kim, Siwon, Sangdoo Yun, Hwaran Lee, Martin Gubri, Sungroh Yoon, and Seong Joon Oh. "Propile: Probing privacy leakage in large language models." arXiv preprint arXiv:2307.01881 (2023).
- [5] Chen, Yu, Tingxin Li, Huiming Liu, and Yang Yu. "Hide and Seek (HaS): A Lightweight Framework for Prompt Privacy Protection." arXiv preprint arXiv:2309.03057 (2023).
- [6] South, Tobin, Guy Zuskind, Robert Mahari, and Thomas Hardjono. "Secure Community Transformers: Private Pooled Data for LLMs."
- [7] Beyer, Franka, Zahra Kolagar, and Darina Gold. "Let's Write Privacy: Survey and Experiment Design on Textual Privacy in Conversation." (2023).
- [8] Katsomallos, Manos, Katerina Tzompanaki, and Dimitris Kotzinos. "Landmark privacy: Configurable differential privacy protection for time series." In Proceedings of the Twelfth ACM Conference on Data and Application Security and Privacy, pp. 179-190. 2022.
- [9] Stach, Christoph. "VAULT: A Privacy Approach towards High-Utility Time Series Data." In Proceedings of the Thirteenth International Conference on Emerging Security Information, Systems and Technologies (SECURWARE), Nice, France, pp. 27-31. 2019.
- [10] Dwork, Cynthia. "Differential privacy." In International colloquium on automata, languages, and programming, pp. 1-12. Berlin, Heidelberg: Springer Berlin Heidelberg, 2006.
- [11] Radford, Alec, Jeffrey Wu, Rewon Child, David Luan, Dario Amodei, and Ilya Sutskever. "Language models are unsupervised multitask learners." OpenAI blog 1, no. 8 (2019): 9.