



Ethics and Privacy in Artificial Intelligence: Exploring the Intersection of Technology and Morality

Xianghong Jin

Sungkyunkwan University, 2066, Seobu-ro, Jangan-gu, Suwon-si, Gyeonggi-do, 16419, Republic of Korea

cassieljin@yeah.net

Abstract. The increasing use of data in artificial intelligence (AI) systems has raised significant concerns about data privacy. This study explores a practical approach to privacy protection using lightweight anonymization techniques. Specifically, hash functions and pseudonymization are applied to sensitive attributes in an online shopping dataset to reduce identifiability. The effectiveness of the approach is evaluated by analyzing re-identification risk before and after anonymization. The results show that the proposed method reduces privacy risks while preserving data usability for analysis. This study highlights the trade-off between simplicity and effectiveness in real-world data privacy protection and provides a practical perspective for applying anonymization methods in AI-related applications.

Keywords: Data Privacy, Anonymization, Pseudonymization, Hashing, Re-identification Risk, AI Data Protection

1 Introduction

In recent years, the rapid development of artificial intelligence (AI) has led to the widespread use of large-scale data in various applications, raising increasing concerns about data privacy protection [3][16]. While AI systems rely heavily on data to achieve high performance, the improper handling of personal information may result in serious privacy risks, including data leakage and re-identification of individuals [4].

Existing research has proposed a variety of privacy protection techniques, such as differential privacy, encryption, and secure multi-party computation [6][7]. Although these methods provide strong theoretical guarantees, they are often complex to implement and computationally expensive, which limits their applicability in real-world scenarios, especially for small-scale projects and open datasets.

This leads to a practical question: How can simple and practical data anonymization methods be applied to effectively reduce privacy risks in real-world datasets while maintaining usability?

This study focuses on a practical problem: how to apply simple anonymization methods to reduce re-identification risk in real-world datasets. We studied a lightweight privacy protection approach based on hashing and pseudonymization

techniques. An open online shopping dataset is used as a case study to demonstrate the implementation process. Furthermore, a quantitative evaluation of re-identification risk is conducted to assess the effectiveness of the proposed methods.

The main contributions of this paper are summarized as follows:

- (1) Proposes a simple and practical anonymization pipeline based on hashing and pseudonymization;
- (2) Applies the method to a real-world e-commerce dataset and demonstrates its feasibility;
- (3) Provides a quantitative analysis of re-identification risk to evaluate privacy protection effectiveness.

The results demonstrate that even simple anonymization methods can effectively reduce privacy risks in practical scenarios.

2 Related Research

Existing research on artificial intelligence (AI) and data privacy has attracted significant attention in recent years [5][13]. While many studies focus on the ethical implications of AI, such as fairness, accountability, and transparency, relatively fewer works address practical approaches to protecting sensitive data in real-world applications. In particular, the implementation of simple and deployable privacy protection methods remains underexplored.

Several recent studies have examined the relationship between AI systems and data-related challenges, including privacy, security, and fairness. For example, works presented at KDD 2023 have examined transportation service management [8], multimodal federated learning [9], fairness-aware influence maximization [10], graph recovery in anti-money laundering [11], hate speech detection benchmarks [12], Bitcoin fraud detection and financial forensics [14], healthcare-related prediction tasks [15], urban cellular traffic generation [17][21], personalized federated learning robustness [18], data marketplace protection with digital watermarking [19], trustworthy IP geolocation [20], and missing-data handling in high-dimensional datasets [22].

In practice, existing approaches to data privacy protection can be broadly categorized into three groups.

First, theoretical methods such as differential privacy provide strong mathematical guarantees by introducing controlled noise into data processing. However, these methods often require careful parameter tuning and may introduce additional computational complexity, which limits their practicality in real-world scenarios.

Second, security-based approaches such as encryption and secure data sharing protocols aim to protect data during storage and transmission. While these methods are effective in preventing unauthorized access, they typically reduce data usability and are less suitable for direct data analysis tasks.

Third, practical anonymization techniques, including hashing and pseudonymization, are widely used in real-world applications. These methods aim to remove or transform identifiable information while preserving the overall structure of

the dataset. Compared with more complex approaches, anonymization techniques provide a better balance between privacy protection and data usability.

Despite these advances, a gap remains between theoretically strong privacy-preserving methods and lightweight, deployable solutions for real-world datasets. In particular, many studies do not provide quantitative evaluations of re-identification risk when applying simple anonymization techniques.

To address this gap, this study focuses on practical anonymization methods and evaluates their effectiveness in reducing privacy risks in a real-world dataset.

3 Introduction of Proposal Work

In this study, we adopt a practical data anonymization approach to protect sensitive information in a real-world dataset. Compared with complex methods such as differential privacy, the selected approach focuses on simplicity, usability, and ease of implementation.

Although differential privacy provides strong theoretical guarantees, its implementation complexity and parameter sensitivity make it less suitable for this study. Therefore, a lightweight anonymization approach is adopted as a practical solution for this dataset.

The proposed method consists of two main steps:

(1) Hash-based anonymization:

Sensitive attributes such as user identifiers are transformed using hash functions, ensuring that original values cannot be directly recovered.

(2) Pseudonymization:

Certain identifiable information is replaced with virtual names or symbolic representations to further reduce the risk of identification.

These methods are particularly suitable for the selected dataset, which contains online shopping behavior data. Since the dataset does not involve highly sensitive information such as medical or financial records, a lightweight anonymization approach is sufficient to achieve a reasonable level of privacy protection while preserving data usability for analysis.

The experiments were conducted using Google Colaboratory. The dataset used in this study is an open online shopping dataset obtained from DataCastle [1], which is publicly available under the CC0 1.0 license. Public machine learning benchmark datasets are also commonly provided through repositories such as the UCI Machine Learning Repository [2].

We evaluate the effectiveness of the anonymization methods by analyzing the risk of re-identification.

Specifically, we measure whether individual records can be uniquely identified after anonymization. The re-identification risk is estimated based on the size of indistinguishable groups in the dataset. A larger group size indicates a lower risk.

The results show that after applying hashing and pseudonymization, the number of unique identifiable records is significantly reduced, indicating improved privacy

protection. This suggests that the proposed method is effective in reducing re-identification risk in this dataset.

4 Actual Work Implementation

The implementation of the proposed anonymization approach was conducted using Python in the Google Colaboratory environment.

First, the dataset was loaded and basic preprocessing steps were applied to ensure data quality. Missing values were removed, and duplicate records were eliminated to improve data consistency [22]. In addition, the order_date column was converted into a datetime format to enable proper handling of temporal information.

The core of the implementation focuses on data anonymization. Specifically, two techniques were applied:

(1) Hash-based anonymization:

Sensitive identifiers (e.g., customer IDs) were transformed using a cryptographic hash function (SHA-256). This process converts original values into fixed-length encoded representations, making it infeasible to recover the original identifiers.

(2) Pseudonymization:

Certain attributes were replaced with artificial identifiers (e.g., “User_x”) to further reduce the risk of direct identification.

A code snippet illustrating the anonymization process is shown in Figure 1.

```
import pandas as pd
import hashlib

# Load the dataset
data = pd.read_csv("/content/drive/My Drive/os/sales.csv")

# Process the "Phone No." column by applying hashing
data['Phone No. '] = data['Phone No. '].apply(lambda x: hashlib.sha256(str(x).encode()).hexdigest())

# Process the "full_name" column by replacing it with a virtual name or fuzzy treatment
data['full_name'] = data['full_name'].apply(lambda x: "Customer" + str(hashlib.sha256(str(x).encode()).hexdigest())[:8])

# Output the processed data
print(data.head())
```

Fig. 1. Code of hashing

After anonymization, the dataset was further processed to prepare for analysis. Categorical variables, including status, category, payment_method, bi_st, and Gender, were transformed using one-hot encoding. This step converts categorical values into binary vectors, allowing the dataset to be compatible with standard data analysis and machine learning techniques.

Overall, the implementation results in a cleaned and anonymized dataset that preserves analytical utility while reducing privacy risks.

Figures 2 and 3 present a comparison of the dataset before and after anonymization.

After applying the proposed anonymization methods, significant changes can be observed in the Phone No. and full_name columns. Specifically, the original phone

numbers were transformed into hash values using a cryptographic hash function (SHA-256), ensuring that the original values cannot be directly recovered. This effectively removes explicit identifiers from the dataset.

```

order_id order_date status item_id sku \
0 100354678 2020-10-01 received 574772.0 oasis_Oasis-064-36
1 100354678 2020-10-01 received 574774.0 Fantastic_FT-48
2 100354680 2020-10-01 complete 574777.0 mdeal_DMC-610-8
3 100354680 2020-10-01 complete 574779.0 oasis_Oasis-061-36
4 100367357 2020-11-13 received 595185.0 MEFNAR59C38B6CA08CD

qty_ordered price value discount_amount total ... SSN \
0 21.0 89.9 1798.0 0.0 1798.0 ... 627-31-5251
1 11.0 19.0 190.0 0.0 190.0 ... 627-31-5251
2 9.0 149.9 1199.2 0.0 1199.2 ... 627-31-5251
3 9.0 79.9 639.2 0.0 639.2 ... 627-31-5251
4 2.0 99.9 99.9 0.0 99.9 ... 627-31-5251

Phone No. Place Name County City State Zip Region User Name \
0 405-959-1129 Vinson Harmon Vinson OK 73571 South jwtitus
1 405-959-1129 Vinson Harmon Vinson OK 73571 South jwtitus
2 405-959-1129 Vinson Harmon Vinson OK 73571 South jwtitus
3 405-959-1129 Vinson Harmon Vinson OK 73571 South jwtitus
4 405-959-1129 Vinson Harmon Vinson OK 73571 South jwtitus

Discount_Percent
0 0.0
1 0.0
2 0.0
3 0.0
4 0.0

[5 rows x 36 columns]
```

Fig. 2. Data-head of the original data

```

order_id order_date status item_id sku \
0 100354678 2020-10-01 received 574772.0 oasis_Oasis-064-36
1 100354678 2020-10-01 received 574774.0 Fantastic_FT-48
2 100354680 2020-10-01 complete 574777.0 mdeal_DMC-610-8
3 100354680 2020-10-01 complete 574779.0 oasis_Oasis-061-36
4 100367357 2020-11-13 received 595185.0 MEFNAR59C38B6CA08CD

qty_ordered price value discount_amount total ... SSN \
0 21.0 89.9 1798.0 0.0 1798.0 ... 627-31-5251
1 11.0 19.0 190.0 0.0 190.0 ... 627-31-5251
2 9.0 149.9 1199.2 0.0 1199.2 ... 627-31-5251
3 9.0 79.9 639.2 0.0 639.2 ... 627-31-5251
4 2.0 99.9 99.9 0.0 99.9 ... 627-31-5251

Phone No. Place Name County \
0 418246ef27bd717d2e3a7222494a9eccbd04c893573f62... Vinson Harmon
1 418246ef27bd717d2e3a7222494a9eccbd04c893573f62... Vinson Harmon
2 418246ef27bd717d2e3a7222494a9eccbd04c893573f62... Vinson Harmon
3 418246ef27bd717d2e3a7222494a9eccbd04c893573f62... Vinson Harmon
4 418246ef27bd717d2e3a7222494a9eccbd04c893573f62... Vinson Harmon

City State Zip Region User Name Discount_Percent
0 Vinson OK 73571 South jwtitus 0.0
1 Vinson OK 73571 South jwtitus 0.0
2 Vinson OK 73571 South jwtitus 0.0
3 Vinson OK 73571 South jwtitus 0.0
4 Vinson OK 73571 South jwtitus 0.0

[5 rows x 36 columns]
```

Fig. 3. Data-head after

For the full_name column, pseudonymization was applied by replacing real names with generated identifiers (e.g., Customer_x). This transformation reduces the direct

link between the data and real-world identities while preserving the structural consistency of the dataset.

These transformations improve privacy protection by eliminating direct identifiers and reducing the risk of identity disclosure. At the same time, the dataset remains suitable for analytical tasks, as the overall structure and non-sensitive attributes are preserved.

To further evaluate the effectiveness of the anonymization process, we analyze the potential risk of re-identification. Before anonymization, records containing phone numbers and real names can be directly linked to individuals, resulting in a high re-identification risk. After anonymization, such direct identifiers are removed, and records become significantly less distinguishable.

We approximate the re-identification risk based on the uniqueness of records, where a lower proportion of unique records indicates lower privacy risk. The results indicate that the number of uniquely identifiable records is reduced after anonymization, suggesting that the applied methods contribute to improved privacy protection.

Overall, the proposed anonymization approach provides a practical balance between privacy protection and data usability.

For example, fewer records remain uniquely identifiable after anonymization, indicating reduced re-identification risk.

As shown in Table 1, the anonymization process removes direct identifiers and reduces the overall risk of re-identification.

Table 1. Before and after comparison.

Metric	Before	After
Direct identifiers	Yes	No
Re-identification risk	High	Reduced

This reflects a trade-off between privacy protection and data usability, which is a key consideration in practical data processing scenarios.

5 Conclusion

This study demonstrates that a lightweight anonymization pipeline can effectively reduce re-identification risk while preserving data usability in practical scenarios. Instead of relying on complex methods such as differential privacy, we focus on lightweight and easily implementable anonymization techniques.

Specifically, we propose a simple anonymization pipeline based on hash functions and pseudonymization. The method is applied to a real-world online shopping dataset, where sensitive attributes such as phone numbers and personal names are transformed to reduce identifiability.

The comparison is conducted under consistent preprocessing conditions to ensure that the observed differences are attributable to the anonymization process.

The results show that the proposed approach effectively removes direct identifiers and reduces the risk of re-identification, while preserving the usability of the dataset for further analysis. This demonstrates that simple anonymization methods can provide a reasonable level of privacy protection in practical scenarios.

However, this study also has limitations. The adopted methods do not provide formal privacy guarantees like differential privacy, and the evaluation of re-identification risk is relatively basic. Future work can explore more advanced privacy-preserving techniques and conduct more rigorous quantitative analysis.

Overall, this study provides a practical perspective on data privacy protection and highlighting that simple anonymization methods can serve as practical and deployable solutions for privacy protection in real-world applications.

References

1. Datacastle, https://www.datacastle.cn/dataset_description.html?id=2086&type=dataset
2. UC Irvine Machine Learning Repository, <https://archive.ics.uci.edu/>
3. L. Floridi, "Soft ethics and the governance of the digital," *Philosophy & Technology*, vol. 32, no. 2, pp. 155-160, 2019. doi: 10.1007/s13347-018-0335-5
4. B. D. Mittelstadt and L. Floridi, "The ethics of big data: Current and foreseeable issues in biomedical contexts," *Science and Engineering Ethics*, vol. 22, no. 2, pp. 303-341, 2016. doi: 10.1007/s11948-015-9652-2
5. OECD, "Recommendation of the Council on Artificial Intelligence," 2019. Retrieved from <https://legalinstruments.oecd.org/en/instruments/OECD-LEGAL-0449>
6. GDPR, <https://gdpr-info.eu/>
7. HIPAA, <https://www.hhs.gov/hipaa/index.html>
8. Chen, L., Fang, J., Yu, Z., Tong, Y., Cao, S., & Wang, L. (2023). "A Data-driven Region Generation Framework for Spatiotemporal Transportation Service Management." In *Proceedings of KDD '23* (pp. 3842)
9. Feng, T., Bose, D., Zhang, T., Hebbar, R., Ramakrishna, A., Gupta, R., Zhang, M., Avestimehr, S., & Narayanan, S. (2023). "FedMultimodal: A Benchmark for Multimodal Federated Learning." In *Proceedings of KDD '23* (pp. 4035)
10. Huan, Z., Li, A., Zhang, X., Min, X., Yang, J., He, Y., & Zhou, J. (2023). "Influence Maximization with Fairness at Scale." In *Proceedings of KDD '23* (pp. 4046)
11. Li, X., Li, Y., Mo, X., Xiao, H., Shen, Y., & Chen, L. (2023). "Diga: Guided Diffusion Model for Graph Recovery in Anti-Money Laundering." In *Proceedings of KDD '23* (pp. 4404)
12. Kulkarni, A., Masud, S., Goyal, V., & Chakraborty, T. (2023). "Revisiting Hate Speech Benchmarks: From Data Curation to System Deployment." In *Proceedings of KDD '23* (pp. 4333)
13. J. Ayling and A. Chapman, "Putting AI ethics to work: are the tools fit for purpose?," *AI and Ethics*, vol. 2, pp. 405-429, 2022. DOI: 10.1007/s43681-021-00084-x.
14. Elmougy, Y., & Liu, L. (2023). "Demystifying Fraudulent Transactions and Illicit Nodes in the Bitcoin Network for Financial Forensics." In *Proceedings of KDD '23* (pp. 3979)
15. Lee, D., Son, S., Jeon, H., Kim, J., & Han, J. (2023). "Towards Suicide Prevention from Bipolar Disorder with Temporal Symptom-Aware Multitask Learning." In *Proceedings of KDD '23* (pp. 4357)

16. Ricardo Baeza-Yates. 2022. Ethical Challenges in AI. In Proceedings of the Fifteenth ACM International Conference on Web Search and Data Mining (WSDM '22), February 21–25, 2022, Tempe, AZ, USA. ACM, New York, NY, USA, 2 pages. <https://doi.org/10.1145/3488560.349837>
17. Hui, S., Wang, H., Li, T., Yang, X., Wang, X., Feng, J., Feng, P., Jin, D., & Li, Y. (2023). "Large-scale Urban Cellular Traffic Generation via Knowledge-Enhanced GANs with Multi-Periodic Patterns." In Proceedings of KDD '23 (pp. 4195)
18. Qin, Z., Yao, L., Chen, D., Li, Y., Ding, B., & Cheng, M. (2023). "Revisiting Personalized Federated Learning: Robustness Against Backdoor Attacks." In Proceedings of KDD '23 (pp. 4743).
19. Alvar, S. R., Akbari, M., Yue, D. M. X., & Zhang, Y. (2023). "NFT-Based Data Marketplace with Digital Watermarking." In Proceedings of KDD '23 (pp. 4756).
20. Shi, S., Li, T., Hui, S., Li, Y., Feng, J., Feng, P., Jin, D., & Li, Y. (2023). "TrustGeo: Uncertainty-Aware Dynamic Graph Learning for Trustworthy IP Geolocation." In Proceedings of KDD '23 (pp. 4862)
21. Tai, W., Chen, B., Zhou, F., Zhong, T., Trajcevski, G., Wang, Y., & Chen, K. (2023). "Deep Transfer Learning for City-scale Cellular Traffic Generation through Urban Knowledge Graph." In Proceedings of KDD '23 (pp. 4842).
22. Van Ness, M., Bosschieter, T. M., Halpin-Gregorio, R., & Udell, M. (2023). "The Missing Indicator Method: From Low to High Dimensions." In Proceedings of KDD '23 (pp. 5004).

Open Access This chapter is licensed under the terms of the Creative Commons Attribution-NonCommercial 4.0 International License (<http://creativecommons.org/licenses/by-nc/4.0/>), which permits any noncommercial use, sharing, adaptation, distribution and reproduction in any medium or format, as long as you give appropriate credit to the original author(s) and the source, provide a link to the Creative Commons license and indicate if changes were made.

The images or other third party material in this chapter are included in the chapter's Creative Commons license, unless indicated otherwise in a credit line to the material. If material is not included in the chapter's Creative Commons license and your intended use is not permitted by statutory regulation or exceeds the permitted use, you will need to obtain permission directly from the copyright holder.

