



Review

A Survey on Privacy Preserving Data Mining Techniques

Aziza Shamis Aldfaii^{1,*} and Rabie A. Ramadan¹

¹Department of Information Systems, College of Economics, Management, and Information Systems, University of Nizwa, Nizwa, Sultanate of Oman

*Correspondence: 28375243@uofn.edu.om

Received: January 1st, 2025; Accepted: March 3rd, 2025; Published: May 1st, 2025

Abstract: Privacy-preserving data mining (PPDM) has become a significant area of interest for researchers, facilitating the sharing and analysis of sensitive information while ensuring privacy protection. This paper investigates methods for maintaining data confidentiality while retaining the critical attributes necessary for analysis. The authors assess the efficacy of various PPDM techniques against criteria such as performance, data usability, and levels of uncertainty. The key findings and limitations of each approach are thoroughly reviewed and summarized. Various PPDM techniques present distinct advantages alongside certain limitations: Anonymization guarantees the anonymity of data owners but is vulnerable to linking attacks. Perturbation protects attributes independently but does not allow for the reconstruction of original values from the altered data. Randomization provides robust privacy protection but diminishes data utility due to the introduction of noise. Cryptographic methods offer strong security and utility but tend to be less efficient than other strategies. No single technique outperforms all criteria; rather, each is more effective under particular circumstances. This paper delivers a comparative analysis of PPDM techniques, emphasizing their strengths and weaknesses, and offers insights into their applicability across different scenarios.

Keywords: privacy preserving; data mining; anonymization; perturbation; cryptography; randomization

I. Introduction

Data Mining refers to the process of extracting valuable insights from large datasets [1]. It is defined as the method of uncovering significant knowledge from extensive volumes of data stored in databases or other data repositories [2]. Through the application of data mining techniques, it becomes feasible to identify patterns, noteworthy information, or sophisticated insights from the data examined from multiple perspectives. The knowledge obtained can subsequently be utilized for query processing, informed decision-making, data management, and process optimization. Recognized as a fundamental aspect of database systems, data mining stands out as one of the most dependable interdisciplinary innovations within the field of Information Technology.

The exploration of data mining focuses on extracting potentially beneficial information. The data extracted from a vast array of sources encompasses various application areas, including client relationship management and market basket analysis. The information derived from extensive databases may manifest as clusters, rules, patterns, or classification models. Throughout the data mining process, which begins with data collection and culminates in knowledge discovery, the datasets typically contain sensitive personal information about numerous individuals, which they prefer to keep confidential from dataset owners, collectors, users, and miners. There exists a significant risk of mishandling this sensitive information if it is inadvertently disclosed.

The abundance of data provides ample opportunities to glean extensive insights about individuals from publicly available information, thereby imposing a greater responsibility on those who handle such data. Concerning the privacy of an individual's personal data, it is important to recognize that comprehensive information about a person often includes sensitive details. The careless dissemination of such information constitutes an immediate breach of individual privacy. Privacy is characterized as the state of being shielded or isolated from the observation or presence of others. When data mining intersects with privacy concerns, it emphasizes the necessity of safeguarding an individual's information from unrestricted access by others.

Privacy is not deemed violated as long as an individual does not perceive that their personal information has been misused. However, once sensitive personal information is disclosed, the individual loses the ability to prevent its potential misuse. The preservation of privacy is deemed essential to prevent data leakage and to ensure the optimal utilization of extensive data sets. This preservation entails the secure storage of data in electronic formats without infringing upon the individual's rights. Therefore, it is imperative to maintain privacy before, during, and after the data mining process.

Privacy preservation has emerged as a significant concern regarding the success of the data mining process. Privacy-Preserving Data Mining (PPDM) is employed to safeguard the privacy of individuals' personal data or sensitive information without necessitating a

complete forfeiture of the required data's utility. Instances of privacy violations related to personal data are prevalent, leading to increased public awareness and a natural reluctance to share confidential information.

The significance of issues surrounding PPDM has become more pronounced in recent times [6], primarily due to the enhanced capability to store users' personal data and the growing sophistication of data mining algorithms that utilize this information [7]. Implementing privacy constraints cannot be achieved in a single step; rather, it is essential to apply PPDM techniques throughout the entire data mining process, from data collection to the generation of information and knowledge. The objective of PPDM is to develop methodologies that transform raw data in a manner that preserves the confidentiality of private knowledge and personal data, even after the data mining process has been completed.

II. Techniques

The primary objective of Privacy-Preserving Data Mining (PPDM) is to establish data mining practices that mitigate the risk of individual data misuse while maintaining the integrity of information usage. Most of these techniques involve altering the raw data in various ways to ensure privacy preservation [4]. Consequently, the modified data prepared for mining must adhere to privacy requirements without sacrificing the advantages of the mining process [5].

Specific personal information is organized in a tabular format, consisting of rows (or records) and columns (or attributes). Numerous techniques aimed at preserving privacy in the context of data publication, including cell suppression, randomization, data swapping, sampling, and perturbation, have been developed specifically for the dissemination of microdata. The field of privacy preservation has progressed through various developmental phases. Given the inherent complexity of existing methods, privacy preservation is increasingly recognized as a distinct area of research.

Typically, personal identifiers are eliminated before the release of data to facilitate data mining. The safeguarding of privacy, regarded as a critical issue, is achieved through the application of diverse techniques. Privacy-preserving data mining methods can be categorized into three primary groups: Perturbation, Anonymization, and Cryptography, as illustrated in Figure 1.

II.1 Perturbation

Data perturbation is a technique employed to alter data through the application of random processes [8]. This approach effectively modifies sensitive data values by employing mathematical operations such as addition, subtraction, or other similar methods. It is



Figure 1. Classification of Privacy Preserving Data Mining Techniques.

Figure 1: Privacy-Preserving Data Mining Techniques

capable of handling a variety of data types, including Boolean, character, integer, and classification types. Prior to applying the perturbation method, it is imperative to pre-process the raw dataset.

Data perturbation is also referred to as data noise or data distortion. The protection of sensitive information is of utmost importance, and the data perturbation process plays a vital role in safeguarding such information. Distortion can be achieved through various techniques, including the use of data rearrangement matrices, the introduction of noise, and the addition of unfamiliar values, among others.

II.2 Anonymization

Information is often disseminated by removing essential identity markers such as social security numbers and names from personal records. However, the combination of various attributes from different datasets, known as quasi-identifiers, can still enable accurate identification of individual records. For example, attributes such as date of birth, race, zip code, and gender are present in voter registration lists. When these indicators are found in sensitive databases, such as medical records, quasi-identifiers can be utilized to ascertain the identity of the individual by linking the two datasets.

To safeguard privacy, the k -anonymity model, which employs suppression and general-

ization techniques, has been proposed [10]. This model asserts that an individual becomes indistinguishable if there are at least $k-1$ other individuals sharing the same quasi-identifier details. The generalization process involves recoding a specific entry to a less precise yet meaningfully relevant entry. For instance, to mitigate the risk of identification, a birth date may be generalized to a broader range, such as the year of birth.

Suppression refers to the complete concealment of a value. It is evident that, although such strategies mitigate the risk of identification when utilizing public records, they simultaneously diminish the accuracy of operations performed on the altered data. Additionally, this approach may be susceptible to two further types of attacks known as the homogeneity attack and the background knowledge attack.

II.2.1 L-diversity

The two primary threats known as the Homogeneity attack and the Background Knowledge attack have resulted in the development of a novel technique referred to as l-diversity. This approach represents an enhancement of the k-anonymity model, designed to safeguard privacy even when the data owner is unaware of the information possessed by the intruder [11]. l-Diversity is based on the k-anonymity model, wherein k records within the dataset correspond to $k-1$ other records, thereby reducing the granularity or detail of the dataset to create an l-diverse dataset.

II.2.2 T-closeness

To attain l-diversity, each group of records within the dataset that meets k-anonymity must contain l adequately represented values for every sensitive attribute. Furthermore, the aforementioned technique does not protect the dataset from the potential disclosure of these attributes. To address this issue, the t-closeness method was developed, which effectively resolves the limitations of k-anonymity and l-diversity.

The concept of t-closeness, as introduced in this section, is defined such that a dataset is considered to satisfy t-closeness if every equivalence class exhibits t-closeness [12]. This model serves as an enhancement to the l-diversity framework. A significant characteristic of the l-diversity model is its uniform treatment of all attribute values, regardless of their distribution within the dataset.

II.3 Randomization

Randomization is regarded as a commonly employed technique in the field of Privacy-Preserving Data Mining (PPDM) research [13]. This approach entails the introduction of noise to the original data in order to generate values for each record. The combination of

perturbation with genuine data is substantial enough to ensure privacy, thereby preventing the recovery of the original data [9].

The Randomized Response scheme and random-noise-based perturbation facilitate Randomization methods in accomplishing both knowledge discovery and the preservation of privacy [14]. This technique, despite resulting in significant information loss, is regarded as a more effective and efficient approach. Randomization demonstrates the capability to maintain certain semantic elements while anonymizing the entire dataset.

Among the various privacy-preserving data mining techniques currently in use, randomization is considered a fundamental method [13]. It strikes a balance between utility and privacy, as well as facilitating knowledge discovery [15]. Once the data has been appropriately randomized, it is sent to the intended recipient. The recipient utilizes a distribution reconstruction algorithm to access the data. This method provides a straightforward and effective means of safeguarding individual privacy while also preserving data utility to a certain degree.

II.4 Cryptography

Cryptography serves as a crucial technique for safeguarding sensitive information [16]. This method is highly regarded due to its ability to ensure the safety and security of confidential data, as noted by various authors [17]. The privacy of an individual's records can be compromised through the process of data mining.

For instance, consider a scenario where multiple medical institutions aim to conduct a collaborative study utilizing combined datasets for shared benefits, while striving to protect sensitive information. However, when a data mining algorithm is applied to a dataset created from the amalgamation of two separate datasets, there exists a risk that the outcomes may inadvertently reveal private details about individuals. Unfortunately, such data leakage is often unavoidable.

II.4.1 The Two-Party Case

A protocol known as constant-round was introduced for the computation of any probabilistic polynomial time function, accommodating scenarios where the adversary may be either malicious or partially honest [18]. To illustrate, consider two parties possessing inputs a and b . These parties are keen to collaboratively execute a function based on their inputs for mutual advantage. Let the function be defined as $g(x,y) = (g1(x,y), g2(x,y))$. In this arrangement, $g1(x,y)$ is provided to the first party, while $g2(x,y)$ is allocated to the second party. The security aspect ensures that only the output is disclosed to the parties; beyond this output, they are unable to glean any additional information regarding the protocol.

II.4.2 The Multiparty Case

The protocols designed for multiparty cases enable participants to compute their inputs using a collaborative approach while ensuring that sensitive data related to those inputs remains confidential [21]. This allows the involved parties to assess the function while safeguarding the privacy of their inputs, similar to previous models. Numerous researchers have successfully demonstrated this concept in various scenarios.

In multiparty protocols, it is essential for each pair of parties to exchange messages to facilitate the effective computation of functions at each gate of the circuit. However, this requirement poses challenges in certain contexts, such as web applications, where the interaction between the server and client does not facilitate efficient communication among all parties. Additionally, in this context, the relationship between communication and computation is linear concerning the size of the circuit.

II.4.3 Oblivious Transfer

This protocol is regarded as a fundamental component for secure computation [22]. The concept of the 1-out-2 oblivious transfer protocol was introduced earlier in the field [23]. In this oblivious transfer protocol, two parties are involved: a sender and a receiver. The sender's input consists of a pair (X_0, X_1) , while the receiver's input is Q , which can be either 0 or 1. Upon completion of the protocol, the sender gains no information, and the receiver is only able to learn X_Q .

While cryptographic methods ensure the accuracy and security of modified data [24], they often fall short when multiple participants are involved [25]. Additionally, the confidentiality of individual records may be compromised by the outcomes of data mining. Despite the existence of numerous solutions utilizing semi-honest models, there has been a limited amount of research focused on malicious models [26].

II.5 Generalization

Generalization is recognized as one of the conventional methods of anonymization [27]. This technique transforms a dataset that is highly specific to individuals into one that is less so [28]. It is widely employed to replace quasi-identifier (QI) values with broader, yet semantically relevant, alternatives. However, generalization can lead to significant information loss due to the extensive range of QI values.

To mitigate this loss, it is essential to maintain records within the same conformity class in close proximity [29]. Another drawback of generalization is that it can render the data ineffective [30-32]. The nuanced analysis of attribute interrelationships may also be compromised due to the discrete generalization applied to each attribute. The

benefits and drawbacks of all Privacy-Preserving Data Mining (PPDM) techniques are summarized in Table 1.

Table 1: Advantages and Limitations of PPDM Techniques

Technique	Advantages	Limitations
Anonymization technique of PPDM	Data owner's sensitive or private data are to be secreted.	More information loss, Linking attack
Perturbation technique of PPDM	Preserves various attributes independently.	Information loss and Cannot regenerate original data values.
Randomized Response technique of PPDM	It provides good efficiency. Simple and useful for keeping the individual information secretly.	Loss in individual's information. Not much good for database containing several attributes.
Cryptography technique of PPDM	Data transformation is accurate and protected. Provides better privacy and data utility.	It is particularly hard to scale if multiple parties are involved.

III. Conclusion

The main aim of Privacy-Preserving Data Mining (PPDM) is to develop algorithms that protect sensitive information and ensure privacy. Such sensitive data remains undisclosed to unauthorized individuals or intruders. In the realm of data mining, there exists a balance between the utility and privacy of data; achieving one often results in a negative effect on the other. The paper reviews various existing PPDM techniques. Ultimately, it concludes that no single PPDM technique excels over all others across every conceivable criterion, including data usage, performance, complexity, and compatibility with data mining processes. A specific algorithm may outperform another based on a particular criterion, and different algorithms may excel in various aspects. Researchers are actively engaged in efforts to safeguard individuals' sensitive data while maintaining its utility for diverse applications.

Author Contributions

Conceptualization, A.S.A. and R.A.R.; methodology, A.S.A.; validation, R.A.R.; formal analysis, A.S.A.; investigation, A.S.A.; resources, R.A.R.; data curation, A.S.A.; writing—original draft preparation, A.S.A.; writing—review and editing, R.A.R.; visualization, A.S.A.; supervision, R.A.R.; project administration, A.S.A. All authors have read and agreed to the published version of the manuscript.

Funding

The APC was funded by University of Nizwa.

Acknowledgments

The authors thank University of Nizwa for their support.

Conflicts of Interest

The authors declare no conflict of interest.

References

- [1] Agrawal, R., & Srikant, R. (2000). Privacy-preserving data mining. In Proceedings of the 2000 ACM SIGMOD International Conference on Management of Data (pp. 439-450).
- [2] Agrawal, D., & Aggarwal, C. C. (2001). On the design and quantification of privacy preserving data mining algorithms. In Proceedings of the twentieth ACM SIGMOD-SIGACT-SIGART symposium on Principles of Database Systems (pp. 247-255).
- [3] Sweeney, L. (2002). k-anonymity: A model for protecting privacy. *International Journal of Uncertainty, Fuzziness and Knowledge-Based Systems*, 10(05), 557-570.
- [4] Evfimievski, A., Srikant, R., Agrawal, R., & Gehrke, J. (2002). Privacy preserving mining of association rules. In Proceedings of the eighth ACM SIGKDD international conference on Knowledge discovery and data mining (pp. 217-228).
- [5] Lindell, Y., & Pinkas, B. (2000). Privacy preserving data mining. In Annual International Cryptology Conference (pp. 36-54). Springer, Berlin, Heidelberg.

- [6] Kargupta, H., Datta, S., Wang, Q., & Sivakumar, K. (2003). On the privacy preserving properties of random data perturbation techniques. In Third IEEE International Conference on Data Mining (pp. 99-106).
- [7] Evfimievski, A. (2003). Randomization in privacy-preserving data mining. ACM SIGKDD Explorations Newsletter, 4(2), 43-48.
- [8] Pinkas, B. (2002). Cryptographic techniques for privacy-preserving data mining. ACM SIGKDD Explorations Newsletter, 4(2), 12-19.
- [9] Du, W., & Zhan, Z. (2004). Building decision tree classifier on private data. In Proceedings of the IEEE international conference on Privacy, security and data mining (pp. 1-8).
- [10] Verykios, V. S., Bertino, E., Fovino, I. N., Provenza, L. P., Saygin, Y., & Theodoridis, Y. (2004). State-of-the-art in privacy preserving data mining. ACM SIGMOD Record, 33(1), 50-57.
- [11] Goldreich, O., Micali, S., & Wigderson, A. (1987). How to play any mental game. In Proceedings of the nineteenth annual ACM symposium on Theory of computing (pp. 218-229).
- [12] Aggarwal, C. C., & Yu, P. S. (2005). On variable constraints in privacy preserving data mining. In Proceedings of the 2005 SIAM International Conference on Data Mining (pp. 115-125).
- [13] Fung, B. C., Wang, K., & Yu, P. S. (2005). Top-down specialization for information and privacy preservation. In Proceedings of the 21st International Conference on Data Engineering (pp. 205-216).
- [14] Bayardo, R. J., & Agrawal, R. (2005). Data privacy through optimal k-anonymization. In Proceedings of the 21st International Conference on Data Engineering (pp. 217-228).
- [15] Machanavajjhala, A., Gehrke, J., Kifer, D., & Venkatasubramanian, M. (2006). l-diversity: Privacy beyond k-anonymity. In 22nd International Conference on Data Engineering (pp. 24-24).
- [16] Machanavajjhala, A., Kifer, D., Gehrke, J., & Venkatasubramanian, M. (2007). l-diversity: Privacy beyond k-anonymity. ACM Transactions on Knowledge Discovery from Data (TKDD), 1(1), 3-es.
- [17] Li, N., Li, T., & Venkatasubramanian, S. (2007). t-closeness: Privacy beyond k-anonymity and l-diversity. In 2007 IEEE 23rd International Conference on Data Engineering (pp. 106-115).

- [18] Wang, H., & Jia, X. (2007). Preserving Privacy in Association Rule Mining: A Randomization Approach. In 2007 International Conference on Computational Intelligence and Security (pp. 676-680).
- [19] Yao, A. C. (1986). How to generate and exchange secrets. In 27th Annual Symposium on Foundations of Computer Science (sfcs 1986) (pp. 162-167).
- [20] Rabin, M. O. (1981). How To Exchange Secrets with Oblivious Transfer. Technical Report TR-81, Aiken Computation Lab, Harvard University.
- [21] Vaidya, J., Clifton, C. W., & Zhu, Y. M. (2008). Privacy preserving data mining (Vol. 19). New York: Springer.
- [22] Aggarwal, C. C., & Yu, P. S. (Eds.). (2008). Privacy-preserving data mining: models and algorithms (Vol. 34). Springer Science & Business Media.
- [23] Xiao, X., & Tao, Y. (2006). Personalized privacy preservation. In Proceedings of the 2006 ACM SIGMOD international conference on Management of data (pp. 229-240).
- [24] Rizvi, S. J., & Haritsa, J. R. (2002). Maintaining data privacy in association rule mining. In Proceedings of the 28th international conference on Very Large Data Bases (pp. 682-693).
- [25] Kantarcioglu, M., & Clifton, C. (2004). Privacy-preserving distributed mining of association rules on horizontally partitioned data. *IEEE transactions on knowledge and data engineering*, 16(9), 1026-1037.
- [26] Ciriani, V., di Vimercati, S. D. C., Foresti, S., & Samarati, P. (2008). k-anonymous data mining: A survey. In *Privacy-preserving data mining* (pp. 105-136). Springer, Boston, MA.
- [27] Dwork, C. (2008). Differential privacy: A survey of results. In *International conference on theory and applications of models of computation* (pp. 1-19). Springer, Berlin, Heidelberg.
- [28] Han, J., Kamber, M., & Pei, J. (2006). *Data mining: concepts and techniques*. Morgan kaufmann.
- [29] Witten, I. H., & Frank, E. (2005). *Data Mining: Practical machine learning tools and techniques*. Morgan Kaufmann.
- [30] Malina, L., & Hajny, J. (2013). Efficient security solution for privacy-preserving cloud services. In 36th International Conference on Telecommunications and Signal Processing (pp. 23-27).

- [31] Sachan, A., Roy, D., & Agrawal, P. V. (2013). An efficient intrusion detection system using CUDA enabled GPU. *International Journal of Advanced Research in Computer Science and Software Engineering*, 3(4), 156-165.
- [32] Ramadan, R. A., & Yadav, K. (2020). A Novel Hybrid Intrusion Detection System (IDS) for the Detection of Internet of Things (IoT) Network Attacks. *Annals of Emerging Technologies in Computing (AETiC)*, 4(5), 61-74.
- [33] Ramadan, R. A., Aboshosha, B. W., Alshudukhi, J. S., Alzahrani, A. J., El-Sayed, A., & Dessouky, M. M. (2021). Cybersecurity and Countermeasures at the Time of Pandemic. *Journal of Advanced Transportation*, 2021, Article ID 6627264.