



Recent Developments in Privacy-preserving Mining of Clinical Data

CHANCE DESMET and DIANE J. COOK, Washington State University

28

With the dramatic improvements in both the capability to collect personal data and the capability to analyze large amounts of data, increasingly sophisticated and personal insights are being drawn. These insights are valuable for clinical applications but also open up possibilities for identification and abuse of personal information. In this article, we survey recent research on classical methods of privacy-preserving data mining. Looking at dominant techniques and recent innovations to them, we examine the applicability of these methods to the privacy-preserving analysis of clinical data. We also discuss promising directions for future research in this area.

CCS Concepts: • **Security and privacy** → **Data anonymization and sanitization**; • **Information systems** → *Data mining*;

Additional Key Words and Phrases: Privacy, privacy preserving data mining, PPDM, clinical PPDM

ACM Reference format:

Chance DeSmet and Diane J. Cook. 2021. Recent Developments in Privacy-preserving Mining of Clinical Data. *ACM/IMS Trans. Data Sci.* 2, 4, Article 28 (November 2021), 32 pages.
<https://doi.org/10.1145/3447774>

1 INTRODUCTION

The acquisition and analysis of data form the backbone of the Industrial Revolution 4.0 and fuels much of current clinical research. At the same time, the **Health Insurance Portability and Accountability Act (HIPAA)** is a “privacy rule” that demands that individuals’ health information be protected. Data mining offers essential insights in medical, industrial, and governmental fields, thus prevention of the abuse of mined data is a critical yet often difficult task [36, 122]. Maintaining anonymity has typically consisted of merely removing key attributes such as a person’s name, address, social security number, and other unique identifiers. However, the recent proliferation of high-dimensional data sets introduces the possibility of piecing together a person’s complete profile from seemingly disparate and anonymized pieces of information [83, 157]. This danger is heightened when collected information is linked to ubiquitous, location-tracking mobile devices [36, 44, 90, 171].

This increased awareness of digital exposure has sparked a similar rise in research to maintain the privacy of sensitive information in the face of data mining. New **privacy-preserving data-mining (PPDM)** methods are being continuously proposed to combat the corresponding

Authors’ address: C. DeSmet and D. J. Cook, Washington State University, P.O. Box 642752, Pullman, Washington, 99164-2752; emails: {chance.desmet, djcook}@wsu.edu.

Permission to make digital or hard copies of all or part of this work for personal or classroom use is granted without fee provided that copies are not made or distributed for profit or commercial advantage and that copies bear this notice and the full citation on the first page. Copyrights for components of this work owned by others than ACM must be honored. Abstracting with credit is permitted. To copy otherwise, or republish, to post on servers or to redistribute to lists, requires prior specific permission and/or a fee. Request permissions from permissions@acm.org.

© 2021 Association for Computing Machinery.

2577-3224/2021/11-ART28 \$15.00

<https://doi.org/10.1145/3447774>

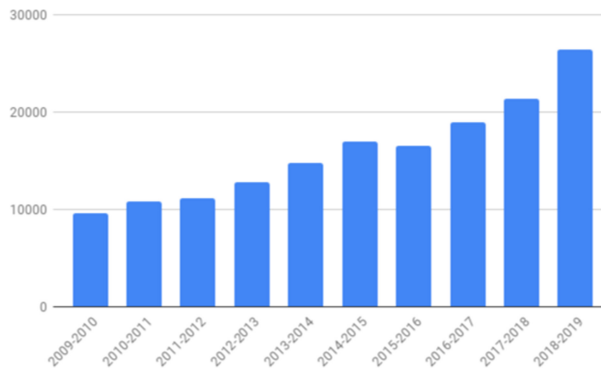


Fig. 1. This rate of growth has been steadily increasing.

expansion of data exploitation methods. Figure 1 illustrates how the number of PPDM articles has grown over the past decade, with no indication of slowing down. This coincides with a rise in clinical vulnerability to data compromise, as in recent years there has been a marked increase in the use of online, open access data sharing services [110].

A factor in this surge of interest might be attributed to the desire for commercial entities to protect themselves from the loss of their customers' data. According to the General Data Protection Regulation set in effect in the European Union, organizations are responsible for the misuse of information that is processed on their systems [90]. Thus, it is not just the individual person that is interested in the security of their data [6, 140, 148, 176], but many commercial enterprises who process these data are motivated to ensure that they are not subject to unintended disclosure through neglect or otherwise.

Another factor in the growing desire for patient privacy preservation stems from the United States government's HIPAA act, which regulates how a health care center may use its client's data. Medical electronic data use increases led to a surge in accidental disclosures, costing medical centers time and resources [72, 161]. However, it is frequently desirable for clinical data to be shared with other organizations, including other medical institutions, public health organizations, law enforcement, and even military inquirers [39, 119, 160]. Therefore, it is in the best interest of medical centers to ensure that the data they provide to external sources cannot be traced back to their clients.

In this article, we survey the development of PPDM approaches and their current clinical usage. Because of the increasing importance and influence of privacy-preservation on the data-mining field, this has become a popular area of research. Aggarwal and Yu [4] provide an early survey of the topic. In recent years, authors focus on specific aspects of PPDM, while others provide a longitudinal look at the field [176]. When examining the field as a whole, some authors focus on particular methods such as random noise addition, mapping, or learned models [118, 185, 187]. Others, such as Wagner and Eckhoff [185], review a range of alternative privacy metrics. Still, others concentrate on a domain of application, such as transactional medical data or big data analytics [141, 146, 180].

The goal of this article is to provide a comprehensive look at PPDM methods and their value for clinical application. Easily accessible data creates more opportunities for the exposure of personal information [3]. We, therefore, focus on clinical applications of privacy-preserving data mining. The rest of this article is organized as follows: We first define PPDM terms in Section 2, then review and compare classes of PPDM methods and metrics in Section 4. Section 5 discusses adversarial strategies to combat PPDM methods. Because location information is valuable for monitoring

and assessing health, Section 6 presents the unique challenges preserving the privacy of location information. Finally, we close with a summary of the surveyed topics in Section 7 and examine directions for future clinically relevant PPDM research in Section 8.

2 DEFINITIONS

Designing privacy-preserving data-mining techniques poses a challenge for researchers and practitioners because of the multiple, sometimes conflicting, goals associated with this endeavor. While PPDM methods should obscure the identity of human subjects and other sensitive information to the greatest degree possible, the integrity of the shared data and resulting models also needs to be ensured. Similarly, researchers need to balance the thoroughness of any PPDM technique with the additional computational expense. Considering these varied and conflicting goals, several metrics are used to evaluate PPDM algorithms. Here, we introduce and discuss the relative merits of these common performance measures.

Clinical data: This survey focuses on PPDM techniques that process clinical data. Based on a definition by Iavindrassan et al. [67], we restrict clinical data to be those that relate to the behavior or medical condition of a person. Thus, we discuss mitigation strategies that address attacks on the gathering (i.e., clinical pathways, discussed in Section 4.6.3) or application of such data.

Quasi-identifier: A quasi-identifier is a piece of information that on its own may not identify an individual in data, but a cohort of these quasi-identifiers may have enough strength together to divulge an individual's identity.

Sensitive attribute: A sensitive attribute is one that, if divulged, violates the privacy of the referenced individual.

Data composition vulnerability: Data are not always inherently vulnerable to re-identification; data that bear no relation to the individual they came from introduce no threat to that individual's privacy. Data containing quasi-identifiers, however, can be vulnerable to disclosure. A quasi-identifier is an attribute that, while on its own cannot positively identify an individual, can be used in conjunction with other quasi-identifiers to identify that individual. Ensuring that quasi-identifiers are suppressed or altered in such a way that they do not reveal a user's identity is a primary goal of PPDM. In this way, the inherent data composition can be seen as part of the vulnerability. Defining metrics to evaluate data composition vulnerability is largely an untapped problem, particularly since the amount and specificity of quasi-identifiers vary widely between datasets.

Performance/privacy trade-off: A critical decision PPDM researchers make is how to balance the desire for privacy with the goal of maintaining usable data [54, 82], because these two goals are inversely related. Increasing the privacy of a data point generally involves distorting the point in some way, which damages its usefulness as a representative of real-world phenomena [62, 82, 120]. To increase utility of privacy-preserved data, many of the methods we survey exhibit varying levels of privacy protection. Because PPDM methods reach peak performance at different privacy settings, comparisons between the methods sometimes prove difficult.

K-Anonymity: K-Anonymity is a property that can be used to describe the security of a data set. A data set with this property ensures that every point is indistinguishable from k other data points [170]. Formally, if X represents a data set and Q represents the set of all non-sensitive quasi-identifiers in X , then X satisfies K-Anonymity if, for all combinations of Q_i in X , there are at least K examples of each Q_i [170]. Equation (1) states this description mathematically:

$$\begin{aligned} \forall Q_i \in X, \\ ||Q_i|| \geq K. \end{aligned} \tag{1}$$

Table 1. Example Clinical Data Illustrating K-Anonymity Measure

Sex	Age	Country of origin
Male	[20–40]	United States
Male	[20–40]	United States
Male	[40–60]	Australia
Male	[40–60]	Australia
Male	[40–60]	Australia
Female	[60–80]	Montenegro
Female	[60–80]	Montenegro

This table displays a K-anonymity of 2, as each entry has at least one identical record.

This property is very helpful in ensuring that data outliers are not immediately identifiable. Table 1 shows an example of a data set with K-Anonymity where $k = 2$, because each unique combination of attributes is exhibited by at least two data points.

L-Diversity: The idea of K-Anonymity can be extended to L-Diversity, which requires that each sensitive attribute in the data set also contains at least l examples with the same value for that attribute [107, 123]. Through this process, L-Diversity improves K-Anonymity by ensuring that not only are samples well represented, but there are enough varied examples to prevent easy identification of data points. Formally, let D be a data set, Q the non-sensitive quasi-identifiers in this data set, Q_i the combinations of Q that exist in D , and S_i the set of sensitive attributes associated with each Q_i . Equation (2) shows how L-Diversity can be defined in this context [107]. L-Diversity has been additionally extended for increased utility, resulting in such measures such a c-diversity (a categorically minded extension of L-Diversity) [78] and t-closeness (ensuring that the distribution of provided sensitive attributes is no more than t distance away from the true sensitive attributes) [94]:

$$\begin{aligned}
 &\forall Q_i \in X, \\
 &\forall S_i \in Q_i, \\
 &||S_i|| \geq l.
 \end{aligned} \tag{2}$$

Table 1 is not L-Diverse for $l = 2$, because if one saw this data set and knew that their queried person was female or that their age was 76, they would be able to determine that the corresponding country of origin was Montenegro, even though there are more than one of these examples in the data set.

Differential Privacy: Measures such as K-Anonymity and L-Diversity attempt to define the privacy of an individual point within a data set. Similarly, differential privacy is used to measure if the omission of a member's data from a set would have a greater loss of privacy than ϵ when an operation T is performed on the data set before and after removal [24, 40, 62, 80]. The amount of disclosure risk afforded by ϵ varies based on the properties of the data [43, 99, 116, 186]. As a result, it is difficult to set a standard ϵ that signifies confidence in a user's privacy. However, among similar data sets, differential privacy can be used to determine how sensitive a data set is to small changes in its composition.

Disclosure: Disclosure is the discovery of one's private information in a data set by an unauthorized actor. Disclosure has many causes, including accidental disclosure and disclosure due to the re-identification of a person in an anonymized data set [33]. The amount of information that must be leaked to signify a disclosure is often disputed, but in some cases, discovering a single

feature about a person would constitute a disclosure [92]. The likelihood that a data set can have elements disclosed is referred to as *disclosure risk* and is quantified in several ways. One measure of this disclosure risk is the proportion of elements in a data set that are unique [164]. Similar in concept to K-Anonymity, this metric could be used to compare different sets of data to see which ones are the most susceptible to disclosure.

Utility: One criterion by which PPDM algorithms differentiate themselves is the extent to which they impact the utility of the resultant data. Some PPDM methods can have a variable impact on performance, allowing users to choose an acceptable balance of utility and privacy to fit their needs. The utility of a specific PPDM technique is measured for a specific application. On the one hand, PPDM methods may calculate utility loss in terms of the deviation of the new data from the old. This deviation can be quantified using metrics such as Wasserstein's distance or Kullback–Leibler divergence. On the other hand, the impact can be measured as a loss in the predictive performance of a model that is trained on the manipulated data rather than the original. A number of measures have been introduced to quantify such predictive performance, such as accuracy, sensitivity/specificity, f1 measure, and area under the ROC curve. We will refer to accuracy throughout the article as a representative predictive metric. In a clinical setting, the clinical utility of data represents the amount that the data may be used to facilitate treatment [61, 205]. Manipulating the data to retain privacy may decrease the effectiveness of treatment that emanates from the new data.

3 PPDM CASE STUDIES

Following the introduction of the HIPAA privacy rule and guidance from the European General Data Protection Regulation [12], the common-practice method of anonymization was to remove obviously identifiable information from collected data, including names, birth dates, and social security numbers. However, recent investigations into the security of public data sets revealed that in many instances, data thought to be anonymized contained flaws that led to the identification of members within the data set [171, 173]. As the disclosure of these data can be disastrous for those involved, researchers have not only investigated known events of privacy loss but have also taken a closer at data sets that could be vulnerable to compromise.

One well-publicized case of a compromised data set with far-reaching consequences is the Facebook-to-Cambridge Analytica data leak, resulting in unauthorized actors gaining access to private information of over 83 million individuals [90]. Through inadequate access control, Facebook was also found to be inadvertently providing third parties with the ability to view user's birth dates, widely considered a private attribute [37]. Similarly, the AccuWeather application transmitted location data for its iOS users to a third party that used this data for targeted advertisements, a severe invasion of user privacy [90].

While medical and government data are often viewed as most at-risk, other data sources are also vulnerable to exposure. Power grid information such as resource usage or consumption rates is considered private as it may lead to an adversary obtaining knowledge of the consumers' lifestyles, or even an absence from their house, resulting in burglary [29]. Automated safety messages sent out by automobiles are also a privacy concern, as they can reveal location data of the occupants to unintended recipients [47].

On a clinical note, in the state of Washington, researchers accessing medical data that had been de-identified were able to find newspaper stories on injuries that led to the identification of 43% of the patient medical records [171]. This de-identification was accomplished by crosschecking newspaper print dates with hospital admission/injury reports.

In another instance of clinical data vulnerability, based on South Korean government-issued identity numbers, researchers were able to manipulate publicly available check-sum and encoded member data to positively identify every person in a 23,163-person list of weakly encoded

Table 2. Sample Patient Data

Sex	Age	Blood Pressure
Male	21	57
Male	39	76
Female	45	67
Female	47	78
Mean	38.00	69.50
Std Dev	11.83	9.61

Table 3. Abstracted Data

Sex	Age	Blood Pressure	Weight
Male	[20–39]	[50–79]	2
Female	[40–59]	[50–79]	2

Here sample data are aggregated into two different weighted groups.

prescription data [173]. This breach was possible because each prescription contained demographic information about the recipient, including date of birth, gender, and place of birth. While the data was assumed to be secure, because numbers were substituted for letters in the identifiers, this process was reversed using logical reasoning from known patterns in the data [173].

Privacy-preserving data mining can take many forms, and there is a correspondingly diverse set of metrics to evaluate its success. In this article, we review recent methods that address privacy preservation with an eye toward a clinical environment. We categorize historic approaches as well as recent privacy-preserving data-mining techniques into four groups: abstraction methods, random methods, mapping methods, machine learning methods, and synthetic methods.

4 METHODS

At the same time that growing evidence supports the necessity of privacy preservation, researchers have introduced new strategies to ensure data privacy. For this article, we will focus on surveying approaches for data anonymization and privacy preservation. We categorize these as random, mapping, abstraction, learned-model, and synthetic generation methods. Here, we review these popular methods, highlight recent innovations, and contrast their approaches to data and inference security, particularly for clinical applications. To illustrate the alternative ways these PPDM methods modify the data, we utilize an example set of patient data provided in Table 2.

4.1 Abstraction Methods

Many privacy-preserving data-mining methods alter the form of a data point in some way, such as adding noise to distort the value, mapping it to a new point in the space, or swapping some attributes with another data point. However, there are other methods that create new points using combinations of the original data points. Also known as substituting or abstracting the data, these methods group data points into increasingly larger sets, until all identifiable data points have been subsumed by an aggregation of the larger set [23, 52, 95, 132]. Abstraction methods often merge points into a combined group until a pre-determined privacy threshold has been reached. A pre-defined measure of privacy such as K-Anonymization may provide such a threshold. In the case of a K-Anonymization threshold, points will be combined into larger groups until each original data point in the set is not distinguishable from $k - 1$ other points (Section 2).

As an example, we modify Table 3 from Table 2 by abstracting attributes in several ways. Sex cannot be abstracted without combining all feature values into one category, so that remains

unchanged. However, age is discretized into ranges [0–19, 20–39, 40–59, 60–79, 80–), and blood pressure is discretized into ranges [0–49, 50–79]. In both of these cases, the abstracted ranges were derived from K-Anonymization with $k = 1$. As can be seen from this example, one of the chief concerns in using the abstraction method is the loss of information that occurs when over-abstracting the data. More so than some other methods, the accuracy/privacy trade-off is prevalent for abstraction methods. Thus, abstracted data may result in generally poor modeling performance if privacy demands are great. This can be shown by observing that as the groups grow, the corresponding features correspond to the entire possible value ranges, removing the possibility of distinguishing between population subgroups.

There are many ways to abstract data. Individual data points can be iteratively subsumed into greater approximations until the desired privacy level is reached [170]. These privacy levels can be based on K-Anonymity thresholds or more stringent privacy requirements such as variants of L-Diversity. One L-Diversity variant was introduced by Gong et al. [53]. Using their proposed (K, L)-diversity method, data are abstracted until a desired privacy level is reached. As before, this algorithm abstracts feature values ranges. Additionally, this method also handles overlaps between multiple datasets. Specifically, one datum may appear in more than one dataset (with overlapping features). When this occurs, abstraction is applied to both entries to ensure that the privacy metric (e.g., K-Anonymization) is met for both entries in both datasets. This method, called 1:M generalisation, offers an important capability, as standard PPDM methods suffer when duplicates exist [175].

Another abstraction approach was proposed by Lin et al. [101]. These researchers cluster data for similar patients to relay significant adverse medication reactions without divulging user identities. Similarly, Abidi et al. cluster data and then define the sensitive attributes of each data point to be the cluster mean [3]. As highlighted by these methods, data abstraction performs a similar role as the random methods discussed in Section 4.2. Specifically, abstraction loosens precision on individual data points just enough that privacy is maintained. As Savi et al. observe [155], the degree of abstraction will have a direct impact on the resulting classification accuracy and thus should be chosen carefully.

While many types of abstraction PPDM methods aggregate precise feature values into value ranges, data can also be abstracted into a new, synthetic version that bear similarity to the original data, but do not contain any actual entries that may be used to identify an individual person [184, 197]. Typically, synthetic data are generated by combining observed values to create new data points, or by utilizing statistical information about a data set such as the distribution of features to create data points that exhibit the same statistical properties [21, 42, 89]. Synthetic data are often then employed for purposes such as testing software or validating models. To ensure that user privacy is being preserved throughout the data generation process, Vreeken et al. [184] define a criteria to ensure that a sample from the original data set is unlikely to appear in the generated set unless it is very common in the original data. This is an important criterion for generative methods, because if the generator randomly combines data feature values, then there is a possibility that a unique, real example could be included in the generated set. This is discussed further in Section 4.5.

4.1.1 Clinical Usage. Abstraction-based methods offer a useful approach for many clinical goals due to their ability to easily handle both categorical and text data. These data types are commonly found in clinical data and represent limitations for many other PPDM methods.

An abstraction method designed to cluster and sanitize candidates from data was introduced by Wu et al. [196]. The authors demonstrate that generating sanitized data with minimum deviance from the original data is an NP-hard problem. To approximate the optimal privacy abstraction trade-off, the authors propose a greedy approach that, each iteration, marks individual data points

for sanitation or subsumption based on their customized privacy metric. The greedy iterations continue until the desired trade-off is reached between privacy and classification accuracy. As testing of this model indicated a high level of privacy protection as well as minimal data loss, this represents a useful method for securing clinical data. Abstraction was also adopted by Khan et al. Like Lin et al. [101], these researchers hypothesized that a clustering and minimal-abstraction approach could be successful in protecting HIPAA-compliant health data [79]. Khan et al. used differential privacy risk (described further in Section 4.6.4) to cluster sensitive attributes into separate “buckets.” The design disallows linkage attacks (Section 5.1) between members of different buckets. These methodologies exemplify the power of abstraction methods that are capable of removing data specificity until individual members of the data are no longer at risk while still maintaining much of the original data information content.

4.2 Random Methods

Random PPDM methods exploit the original data distribution to randomly inject “noise” into each data entry [149, 162]. This noise can be generated using a variety of statistical manipulations that make it difficult for an adversary to discern the original data point [154].

To explain the general framework for injecting noise into data, let x represent an original data point, c represent noise that is added to the data, and $\hat{x}$ represent the resulting perturbed data point that will be added to the data set. Here, x and $\hat{x}$ each contain n features [5, 44, 149]. Equation (3) formalizes the process of adding random noise to a sample. In this equation, x_n represents a feature of x , and $\hat{x}_n$ represents the perturbed version of that feature. The value c_n represents a unique amount of noise that is added to the corresponding feature, influenced by the distribution of each feature within the data:

$$\begin{aligned} \forall x_n \in x, \\ y_n = x_n + c_n. \end{aligned} \tag{3}$$

Random PPMD methods often distinguish themselves by adopting unique approaches to generating values for c [44, 139, 162]. Traditionally, c is a random term with a mean set at 0, drawn from a distribution that is dependent on the feature it is perturbing. In one of the seminal papers on this method, Agrawal and Srikant experimented with both uniform and Gaussian distributions [5]. Using a decision tree classifier, they evaluated the classification accuracy of the data modified by noise drawn from these distributions [5, 75]. The change in classification accuracy was most apparent when choosing to modify the data more dramatically with the goal of heightened privacy, pointing to a need for random methods to be able to provide enough noise for a specific feature to not reveal sensitive information, but still retain usability. Both Gaussian and uniform noise addition were adept at preserving the classification, with accuracy staying between 5% and 15% of the original classification margin [5]. This accuracy was consistent throughout several different privacy levels, which dictated the breadth of the distribution that was used to generate the noise [5, 75].

We illustrate the process of perturbing data in Table 4. For this example, c is drawn from a normal distribution, and the sex of the person is not considered a private attribute. It can be seen from this table that the perturbation process does affect both the mean and the standard distribution of both blood pressure and age, thus the perturbation caused these to shift considerably. In a data set containing a larger sample, it is likely the mean and standard deviation would exhibit less variance once perturbed.

While random noise addition works well at obfuscating data, adding a noise value to each feature independently of the others can damage relationships between features that contain dependencies [70]. Age and blood pressure are considered to be independent in this example, so the noise factor

Table 4. Perturbed Data Where the Noise Value c Is Drawn from a Gaussian Distribution and Is Used to Modify Values from Table 2

Sex	Age	Blood Pressure	Perturbed Age	Perturbed Pressure
Male	21	57	15.52	65.99
Male	39	76	49.85	73.37
Female	45	67	41.60	47.89
Female	47	78	39.07	72.78
Mean	38.00	69.50	36.50	65.01
Std Dev	11.83	9.61	14.73	11.89

c was calculated independently for each feature. To perturb data sets with dependent variables, a method was introduced in which matrices perform the noise addition, shown by Equation (4), where X represents a set of data points, $\hat{X}$ represents the new perturbed set of points, E is a co-variance matrix representing relationships between the features of X , and α is a random variable used to permute E [125, 149]:

$$\begin{aligned}\hat{X} &= X + E, \\ E_{i,j} &= \alpha * E_{i,j}.\end{aligned}\tag{4}$$

As seen in Equation (4), this new variation of random noise addition relies on a matrix drawn from a random distribution with the same co-variance as the original data. This equation creates new data that possess the same relationship between features as is exhibited in the original data [125, 149].

While the previously discussed approaches employ standard distributions such as Laplace and Gaussian, some authors explored methods that create noise based on characteristics of each individual dataset. As an example, Eyupoglu et al. [44] introduce a data perturbation algorithm that is based on chaos theory. In this method, data points are selected as shown in Equation (5), based on the number of unique features. These points are then modified by the logistic mapping function, which is a chaotic function:

$$\begin{aligned}x_{n+1} &= \lambda * x_n * (1 - x_n), \\ \lambda &\in (3.99, 4).\end{aligned}\tag{5}$$

A chaotic function is one where small changes to the input values have a large effect on the behavior of the series [44]. In Equation (5), the initial value of x is specified *a priori* [44]. Here, values close to 4.0 are used for λ as they generate the maximum variance and unpredictability for the mapped values. This chaotic function makes it nearly impossible for an adversary to determine the initial conditions and therefore determine the specifics of the noise that was added to the data.

Though random noise strategies can be effective tools at tailoring the amount of data privacy, they are applicable primarily to continuous-valued data. Often, clinical usage may necessitate the use of data that are described by categorical attributes as well [9]. Adding noise is difficult for such data, and many attempts to do so operate on associations between different categorical terms, rather than gaining an understanding of what the terms signify [149]. To combat this tendency, Rodriguez-Garcia et al. [149] integrate ontological relationships to advance the data obfuscation principle of noise addition. In the text mining applications that they consider, they examine the meaning of an expressed sentiment and find replacement terms that are taxonomically similar to the word or object. For example, the word “Headache,” or an instance of a specific type of headache, might be replaced with Concussion, Fracture, or Migraine to generalize the phrase but still convey

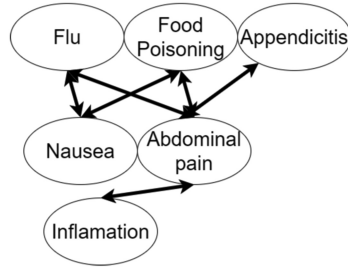


Fig. 2. A natural language taxonomy allows words to be replaced with similar nominal term values. Using this graph, words at one level of the taxonomy can be grouped with terms at a higher level.

a meaning that is similar to the original word. Figure 2 illustrates one example taxonomy to abstract words contained within a clinical document.

4.2.1 Clinical Usage. While on the surface it may appear as a less-sophisticated privacy-preserving method, random noise addition remains useful for clinical PPDM, both as a standalone method and as an augmentation to other strategies. For example, the Priward algorithm [152], introduced by R  th et al., added noise by allowing two parties to calculate likelihoods from hidden Markov models without disclosing either the model or the observation sequences to the other party. By using cryptographic techniques and secure operators, each party can input their portion of the data and obtain a result without discovering or being able to deduce contributions from other parties. This algorithm offers a unique benefit, because a relatively simple addition of random noise provides enough abstraction to contribute to an otherwise-unrelated algorithm’s ability to provide privacy protection.

Another recent example of random noise addition was offered by Ni et al. [127]. In their MCDB-SCAN clustering algorithm, the goal is to ensure that differential privacy is not violated for data points within each cluster. To achieve this goal, they inject Laplacian noise to individual data points, adding uncertainty to the individual points within the cluster. By adding variance to the data points within the clusters, differential privacy is ensured for each added data point while the resulting clusters will preserve privacy as well. This approach actually combines elements of noise injection, data abstraction (Section 4.1), and differential privacy (Section 4.6.4) guarantees, approaching the problem of privacy preservation from several angles. Finally, Aaronson and Rothburn showed how a differential privacy-satisfying level of added noise can be abstracted to the gentle measurement of quantum states [1]. This investigation was performed by observing the relationship between varying the amount of data privacy change due to inclusion or omission and finding quantum states that cause as little disruption to individual states as possible. The investigators noted how well differential privacy concepts can be extended to other disciplines where information about individual members of a set should be considered in the context of the group. From these recent endeavors, the utility of random noise injection can be seen. Not only does this strategy privatize data with minimal impact on the number of samples or the form of the samples, but it pairs well with differential privacy guarantees. This is because random injection supports an easy-to-calculate differential privacy score from the noise parameters.

4.3 Mapping Methods

Many times while mining data, the relationships between different data elements offer critical insights. As a result, privacy preservation needs to retain the relationships to the extent they exist

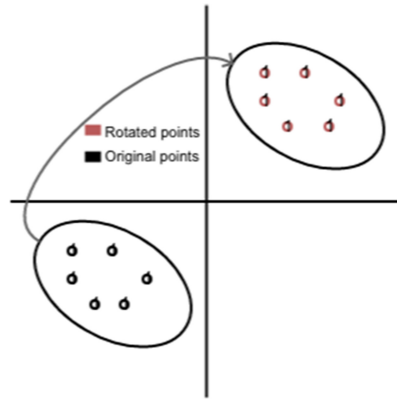


Fig. 3. An example of data rotated 180° around the origin.

in the original data [145]. For example, if a patient has a medical condition that always requires attention when their blood pressure is double their heart rate, then an algorithm that attempts to predict either of these values will need this relationship to be retained when the data are privatized to maintain the same predictive accuracy. To meet this need, PPDM methods have been developed that transform the data into a new form, while still preserving if not replicating the internal relationships. This may be accomplished by mapping the data into a new space where the individual's traits are unrecognizable. Alternatively, the PPDM algorithm can perform internal *rotations*. These rotations are typically performed by selecting two or three random features in a data set and rotating them around a given axis. The resulting data bear minimal similarity to their original form. If done properly, however, then they retain the distribution and relational dependencies from the original data set. While the rotations can damage the predictive ability of some classifiers, other methods, such as SVM and k-nearest neighbor classifiers, are often rotation-invariant. As a result, the classification error of these methods is not affected by such rotations [26]. Figure 3 shows an example of how a two-dimensional data set may be rotated. In this diagram, two features are rotated 180° around the origin, changing the data substantially while still preserving the distances between the individual points.

Sometimes, mapping or rotation may occur within previously defined clusters, generated using methods from Section 4.1 [23, 60, 75, 121, 133, 145]. In these cases, rotation PPDM methods are constrained to occur within clusters, thus ensuring that the rotation keeps similar points together while differentiating distinct clusters [23]. This process ensures that swapping only occurs between similar values, to preserve as much structure in the data as possible.

In an influential paper on rotational methods, Olivera and Zaiane proposed several different mapping methods [130]. In one such method, TDP, each feature in the data is offset by the same amount, perturbing the data, but possibly having an adverse effect on the proportions between data points, and therefore utility [130]. They then proposed another method where two features at a time are selected and rotated simultaneously within an $\mathbb{R}^2$ space, repeating the process until every feature has been rotated at least once [130]. This method, called RDP, was found effective at preserving both privacy and classification accuracy [76].

Once points are clustered, mapping methods can be applied. One unique design by Upadhyay et al. [182] extended the RDP method by selecting three features at a time to rotate in an $\mathbb{R}^3$ space, repeating this process until all features had been rotated at least once. This method further improved data privacy while still supporting machine learning-based classification. The method yielded predictive performance within 1% of the original data using K nearest neighbors, J48

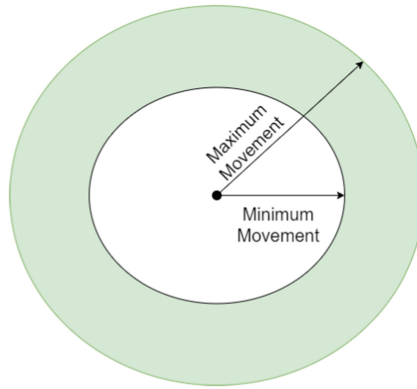


Fig. 4. The donut method rotates a selected point to the gray area, forcing a minimum distance from the original feature value.

decision trees, and naive Bayes classification methods, outperforming two-dimensional rotations such as RDP [182].

Another style of mapping, known as the “donut method,” maps each data point a distance between a minimum and maximum value, creating a torus, or donut shape when applied to increasingly high dimensions [60]. This algorithm is called the donut method, because the inclusion of a minimum distance means that the possible area for the new point is bounded between two concentric circles, as opposed to other methods with no minimum threshold, which creates a “circle.” This method was developed for the anonymization of patient location data by moving the location in a random direction within a specified range. Figure 4 shows how the authors improved their method over a standard rotation. When mapping location coordinates, a maximum translation distance may be specified. In the standard approach, the translation value, r , may vary anywhere from the original point (the central point in Figure 4) to the maximum value. However, using the donut method, a minimum distance is also enforced, forcing r to be selected from values in the gray area between the minimum and maximum. This method is adept at preserving privacy in cases where individual data points are easily distinguishable, possibly due to the minimum distance threshold ensuring that each data point is sufficiently rotated [60]. The donut method outperformed standard aggregation measures in both sensitivity and specificity while preserving the privacy of users’ locations.

As with random methods, researchers have investigated privacy-preserving mapping methods for nominal data sets. As an example, Rodriguez-Garcia et al. [150] extended their work on nominal data to taxonomically classify ailments and use these classifications to employ swapping. By identifying words and phrases that are close in meaning, terms can be transformed with categorical synonyms, resulting in privacy preservation of data that can still be useful for research, as it contains information about a very similar class of problems.

Another style of mapping involves observing the relationship between different data points. In one instance, distributed medical data was able to be mined for information between different parties by observing the relationship and distances between different clusters of data [156]. This PPDM topic is particularly relevant for clinicians, as it would support learning from data without introducing the risk of compromising actual data [34]. Learning from distributed data with mapping methodologies was also investigated by Teo et al., where secure operators were introduced that allowed each party to use information from the others without knowing the actual information contained within [178].

4.3.1 Clinical Usage. With the ability to allow high utility as well as easily shareable permutations, mapping methods are a flexible, albeit computationally expensive option for mining clinical data. A primary challenge with these methods is that they are difficult to apply to streaming data, as they generally process the entire set of data at once.

Recent research involving mapping methods includes the work done by Chamikara et al. [23], who use a covariance matrix generated by points within a cluster to perform intra-cluster rotation. Once the rotation is complete, the clusters are merged and the data points are randomly ordered, yielding the new data. This method was tested on several data sets using the k-nearest-neighbors classifier and exhibited generally superior accuracy when compared to basic rotation and abstraction methods. The results indicate that rotation of data-defined clusters can be used to generate new, private, data samples that provide predictive accuracy comparable to the original data. Additionally, this method was shown to preserve the proportional relationships between the original and mapped data, further improving the overall data utility.

Mapping methodologies apply to multiple types of clinical data. In a new work by Aloufi et al., transforms of collected data, including waveforms of voice recordings, were used to privatize the recordings. Mapping these clinical data to an unrecognizable dimension ensures privacy of the unique information [8]. This example illustrates a potential advantage of mapping methods. Both the addition of random noise and abstraction of waveform data may run the risk of tending this data toward the mean, greatly degrading its quality and usability. For example, perturbing waveform data may degrade the corresponding voice recording to gibberish. However, by mapping the data to an unrecognizable dimension, the component of the audio that is considered sensitive, the emotion, remained private while the speaker and the speech were still recognizable.

Finally, mapping methods were combined with machine learning and cryptography in work by Ping et al. [137]. This work introduces a model that facilitates private support vector clustering between clients and a server, with data undergoing a mapping transform to maintain privacy. This work illustrates how mapping methods may complement many different styles of data protection, such as encryption. They provide an easy-to-enact way of obscuring real distances and relationships between sensitive data, while still allowing the underlying correlations to be maintained.

4.4 Learned Models

To this point, our discussion has centered on PPDM techniques that are designed to safeguard the whole or part of a dataset. In some cases, data privacy can be maintained by sharing a learned model of the data (or inferences derived from the data) rather than sharing the data themselves [55, 57, 88, 105]. There are many methods that, when used correctly, generate models that do not reveal individual-specific information. As an example, Mao et al. [115] demonstrated how facial recognition-based deep learners could preserve individuals' privacy. This result benefits medical applications that deal with the imaging of specific disease patterns, as they could use deep learning to detect these diseases without compromising privacy. Such models have been learned via random forests, perceptrons, and deep learning methods [13, 20, 25, 68, 98, 115, 134, 167, 194, 206, 207]. These learned models differentiate themselves from other learning methods that leave members of the data vulnerable to re-identification. Models that run the risk of re-identification include support vector machines and naive Bayes models trained on small data sets [49, 100]. It is possible, however, to utilize some of these normally insecure methods in such a way that they still ensure a level of privacy. As an example, Lin and Chen [100] modified the typical support vector machine classifier in such a way that the support vectors were not made up of individual data points, and therefore the classifier produced a result that was privacy-preserving. This was done by modifying the support vectors to include ones that provided the same decision boundary but were not drawn from the original data, similar to a mapping method mentioned in Section 4.3.

4.4.1 Clinical Usage. Some research has introduced PPDM solutions for specific clinical use cases. Recently, Alabdulkarim et al. employed a random forest to protect privacy by presenting the most likely maladies an individual might possess without giving specific details on the patient, helping physicians perform differential diagnoses [7]. Another recent method supporting the private usage of clinical data was a federated deep learning model for the segmentation of brain tumors by Li et al. [97]. In this context, “federated” refers to the fact that there are multiple collaborating deep networks, allowing researchers to use information from the trained model without requiring access to the original training samples. The deep networks shared information only after the gradients of each network had been modified by Laplacian noise, guaranteeing a degree of differential privacy.

One downside to using these privacy-preserved machine learning models is that they provide insight into only the target concept. While the learned model may address the original analysis question, methods that retain as much of the original data as possible offer insights for a broader range of clinical analyses. Therefore, it may often be better to use PPDM methods that provide as much raw data as possible, allowing the end-user to design their own machine learning method for answering additional questions about the data.

4.5 Synthetic Data

In general, the previously discussed PPDM archetypes modify existing data to make the individuals safe from re-identification. Synthetic PPDM approaches instead attempt to provide privacy through the generation of synthetic data and have shown to be a useful tool in the acquisition of knowledge in a clinical setting [91, 200]. Additionally, the proliferation of big data for clinical use has resulted in concerns over the applicability of the data, and whether it can wholly encompass the population being measured [17, 59]. Synthetic data can help to alleviate this issue by ensuring that the output data is of a realistic form characteristic of the entire populace.

As mentioned in Section 4.1, synthetic data generation offers an effective method for providing privacy while maintaining model utility [69]. Just as abstraction methods attempt to group part or whole of the data to protect vulnerable elements, synthetic data generation augments part or whole of the data with additional artificial samples that do not need privacy. The resulting infusion lends privacy to the original members of the data. Developing new methods of synthetic data generation that are more adept at recognizing patterns in original data could yield superior artificial data aimed at privacy-preserving. These generative techniques are designed to use many measures of the data such as distribution, clustering cosine similarity, outlier analysis. In this way, data generation could mirror the original information as closely as possible, providing more data to researchers. Along with the privacy protection provided by synthetic data generation, this strategy can bring “new life” to historic data that has been shown to be less representative of how the current populace [151]. Representative synthetic data generation could greatly increase the quality and quantity of available data in terms of both privacy to users and utility to clinicians.

Deep learning is revolutionizing many aspects of machine learning and has begun to affect PPDM processes as well [131]. Along with other deep learning systems, **generative adversarial networks (GANs)** can be used to maximize privacy preservation while ensuring the accuracy remains as high as possible, balancing these two “adversarial” goals [103, 181, 190].

4.5.1 Clinical Usage. In recent work, Abay et al. [2, 13], used a deep learner to generate synthetic data, yielding promising results for both accuracy and privacy. While GANs can generate high-quality synthetic data, the results are not always both sufficiently private and accurate. Yale et al. [200] attempted to address this through the introduction of medGAN, a GAN

optimized for clinical synthetic data generation. Dash et al. also successfully applied medGAN to generate private time series data [35]. Demonstrating application to time series data is important for process mining, the analysis of how an entire patient event log can help determine the efficacy of the treatments [124, 135].

4.6 Ancillary Approaches

In addition to methods that protect privacy through data manipulation, other methods may be used to augment privacy by limiting the accessibility of the data, changing the form of the data, or assess the privacy of the data. While these often support the previously discussed methods, they can still be employed on their own.

4.6.1 Cryptography. An important contribution of recent ancillary methods is cryptographic techniques. These are used to secure data and grant access only to authorized users [96, 138, 153, 159], making it extremely difficult for an adversary to gain access to the data. As discussed in Section 5.3.2, cryptographic methods are often less efficient than PPDM methods for clinical data sharing and distribution. This is largely due to the difficulty of ensuring that only authorized recipients have access to the data. A further contribution is the computational expense of the cryptographic methods themselves [73].

4.6.2 Sanitation. A harsher approach to PPDM is to sanitize attributes from the data [112, 129]. Sanitation refers to removing all items viewed as “sensitive” from the data, rendering the resulting data devoid of any similarities to the original sensitive attributes [48]. This data can still offer some utility, but much value may be lost in this sanitation process. This is a different method than discussed in Section 4.1. Instead of grouping the data into non-uniquely identifiable sections, sensitive data is strictly removed, additionally carrying the risk that some sensitive data may remain.

Focused on removing access to the data rather than making them confidential, sanitation methods are sometimes included in the literature as PPDM methods. However, in isolation, cryptography and sanitation are often unsuited for clinical data, as they severely limit the cohort that may be able to use the data or diminish the utility of data.

4.6.3 Clinical Data Variations. While much of this article focuses on the relationship between PPDM methods and clinical data related to patient health, health data may take many different forms, including images (in the form of x-rays or other diagnostic visual aids) and processes (in the form of a clinical pathway, also known as a care map). Privacy preservation of images used in a clinical setting often takes the form of cryptographic methods, designed to ensure that only trusted individuals gain access as well as establishing control of the image [66, 77, 202]. Despite this, there exist some contemporary clinical image PPDM methods designed to facilitate this sharing of sensitive images. Li et al. demonstrated that through the addition of noise to a deep learner’s weights, information learned from medical images may be shared with outside observers with a differential privacy guarantee on the data [97]. In a similar vein, Kim et al. constructed an encoder to obfuscate medical images presented to it, while still preserving enough utility in the images to be useful in “task-specific” analysis [81].

Kinsman et al. [87] define a clinical pathway as a recorded log or series of medical interventions that are performed for a patient [128]. This record makes clinical pathways available for process mining, facilitating the improvement of treatment protocols. These plans of care may be vulnerable, however, to exposing the patient’s treatment regimen or even their condition to outside observers [114, 142]. To mitigate such attacks, privacy-preserving methods will suppress or generalize logs to include only abstract information [136], or sanitize logs to meet K-anonymity and T-closeness

requirements [45]. Recently, a clinical pathway PPDM method was proposed by Mannhardt et al. [113]. This method added noise to log queries from non-trusted entities. We note that while specific clinical tasks, such as supporting clinical pathways, spark the creation of new algorithms, the underlying PPDM methods remain consistent with those introduced in the rest of this survey.

4.6.4 Differential Privacy. As discussed in Section 2, differential privacy is often used as a guarantee of the desired privacy level for a given purpose [168, 203, 209]. We elaborate on differential privacy as an ancillary method due to its increasing usage in PPDM work as well as its great ability to augment and validate other PPDM methods. This can be seen in several examples. Cheu et al. [28] proposed a shuffling methodology evaluated by differential privacy to verify the sensitivity of messages sent between two parties. Differential privacy was also used to clarify the level of different protections given to defenses against attacks on machine learning models [64, 93]. Finally, Xu et al. used differentially private guarantees to address multi-party learning and ensure that all members in this collaborative environment retained a suitable amount of privacy [198].

Recently, differential privacy was enhanced by a method called “integral privacy,” which is a strengthening of differential privacy to include not just a member of data, but subsets of the data. This refinement is useful to many clinical and pharmaceutical endeavors as they often look at data sub-components [65]. Using this measure, the privacy of “niche” data subsets can be evaluated in addition to the privacy of the entire dataset [65]. Differential privacy has also been adapted to suit the type of privacy it is guaranteeing. Additionally, differential privacy can take the form of *central differential privacy* or *local differential privacy* [10, 46, 56, 117, 199]. Central differential privacy ensures the privacy of data once they have all been collected, whereas, in local differential privacy, each submitting contributor ensures the privacy of their data before they are included [10, 117].

In addition to strengthening differential privacy requirements when needed, differential privacy requirements can also be relaxed to address situations when such stringent privacy specifications are not needed [24]. This can be seen in the work by Asi et al., where differential privacy may be relaxed to allow users of differing involvement to be segmented by their differing privacy needs. For example, if a hospital employee was treated at that same hospital, the person may not be harmed by being listed as having visited that hospital, while another person that only visited this hospital once for a specific health concern may be harmed [11, 35]. The concept of differential privacy relaxation was further extended by Kim et al. [84] in their presentation of MPPDS, a privacy-preserving sharing system. This system used personalized differential privacy to facilitate different levels of privacy depending on trust between users.

5 RE-IDENTIFICATION

When designing and comparing PPDM methods, it is wise to also consider possible attack avenues. Awareness of attack techniques can motivate a choice of PPDM method and a desired privacy level [193].

5.1 Attack Vectors

As reported by case studies in Section 3, many parties attempt to identify private features from supposedly secure data sets. These parties may be malicious, or they may simply be curious researchers or journalists. No matter the intent, it is still up to the data collectors to ensure that the sensitive features are not exposed [14]. Studies on re-identification attempts have shown that the success rate for these re-identification attacks is typically between 26% and 34% [14]. While these findings do not take into account the low degree of confidence in the results, they still demonstrate how often an attack can yield at least some information about supposedly secure data.

A popular re-identification method links two different sets of data [14, 16, 109, 169, 191, 208]. Many linkage strategies are based on the work of Sunteb and Fellegi [169], who compare two data sets by examining the probability that a point from each of the sets reference the same point. This method has been used by several re-identification strategies [173] and has been extended to big data [191]. In another case study analysts successfully linked newspaper-recorded deaths to stored family structures, allowing the analysts to discover detailed genealogical information for over half of the individuals [109]. Links can be discovered in numerous public sources, revealing private information [12, 41, 109, 163, 172]. Linkages aid in identifying individuals from sparse information even when supposedly private information has been removed. For example, 86% of the United States population is identifiable using only their birth date, sex, and 5 digit zip code [14, 109].

In clinical data, Reisaro et al. [144] found that adversaries could link different parts of genomic data together to identify participants. While a common attack strategy, linking is also practiced within clinical research as a way of discovering additional information in data, using association rule mining [30, 86, 189]. Recently, work has been done to secure data against these forms of linking attacks. Telikani et al. [177] used evolutionary computation to keep the data impenetrable. This evolutionary process employed swarm-based optimization to make the data increasingly impervious to association rule mining invasions.

5.2 Potential Vulnerabilities

Of the methods surveyed in Section 4, potentially most vulnerable to linkage-based attack is abstraction aggregation. As mentioned, linking attacks attempt to identify common elements from multiple different data, using similarities between shared elements to attempt to discover relations between these different elements. Aggregation creates opportunities for data to be linked with other data sets, even when aggregated [192]. Applying linking methods, attackers can determine with a variable degree of certainty to which records a person belongs. Aggregation is also vulnerable to data outliers as well as attackers' knowledge of real constraints on data types such as realistic age ranges [14, 193].

Mapping and random methods are somewhat more secure than abstraction, but both do have inherent vulnerabilities. Mapping methods may reveal a weak point around the axis of movement, as points there experience the least rotation. Because these points move less compared to others, an attacker may use the smaller movements to determine the overall mapping of some or all of the set [27]. Similarly, simple swapping methods exchange feature values within small clusters, allowing an adversary to determine what the possible original values might be for the points within that neighborhood [27, 193].

Random methods are further vulnerable to discovering the degree of added noise, allowing attackers to determine the range of possible initial values [27, 74, 193]. If an attacker can discover the distribution of added random noise, then they can infer a likely range of initial values [75, 154]. This sort of discovery is also possible if the adversary can find a sample of unperturbed examples and their corresponding perturbed permutation. An adversary may also use spectral graph or primary components analysis filtering to determine with a high degree of accuracy the original data [75]. This represents a difficult challenge for data perturbation methods as increasing the amount of perturbation can weaken the utility of the data set [154, 155, 162].

5.3 Mitigation Strategies

To address the vulnerabilities outlined in Section 5.2, methods of mitigating attacks have been developed. In this article, we survey two methods for combating privacy attacks: a blending of multiple PPDM approaches and a merging of PPDM practices with those in the cryptography field.

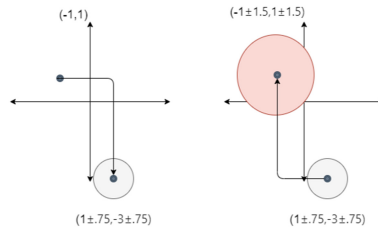


Fig. 5. Applying random noise to the mapping function creates uncertainty as to the origin of the point.

5.3.1 Combining Methods. An effective way to combat attack vulnerabilities is to combine different PPDM methods, as this can leverage multiple security designs, potentially thwarting the attempts of an attacker to learn the original data [27, 127, 193]. As seen throughout Section 4, many current clinical methods combine different styles of PPDM. The combinations range from differential privacy and clustering to abstraction and learned models. Some methods combine very well, for example, mapping and noise addition complement each other, because together they incapacitate re-identification techniques that are targeted for only one method: A popular approach to re-identify mapped data is to utilize known unperturbed examples and their subsequent transformations to discover how the data are mapped. Similarly, for random methods, the goal is to discover the distribution of added noise to intuit the likely original data. Combining mapping and random methods render these strategies ineffective, because the addition of random noise means that possessing previous samples does not give away the mapping. The original point could have been mapped to a variety of regions, with the noise influencing the final location. Mapping the data to new positions before adding noise also thwarts attempts to discover the distribution, because even if the distribution of the mapped data is discovered, this does not necessarily describe the original, unmapped data. Figure 5 shows how the addition of noise to a mapping method makes the original location ambiguous, due to the unknown noise value. Additionally, the combination of differential privacy and random noise injection provides a privacy guarantee, allowing clinical users to determine the degree of safety that they wish to impart on mined data. Combining these PPDM methods and measures can improve the effectiveness of privacy preservation over traditional or novel approaches used in isolation.

5.3.2 Multi-party Computation. As the PPDM field matures, researchers incorporate more diverse computer science ideas to enhance both the privacy and utility of the privacy-preserved data. As discussed in Section 4.5, the introduction of neural networks such as GANs exemplifies how using external techniques can yield promising results for the private generation of synthetic data.

Multi-party Computation is a modern security technique that allows multiple groups to perform an analysis on data without fear of that data leaking [106]. This area of cryptography is quite similar to the goal of many PPDM methods, attempting to facilitate wide access to sensitive information. A subset of multi-party communication is known as homomorphic techniques [15, 85, 174, 176]. Homomorphism stems from encryption and is used to denote a process whereby results can be gathered on encrypted data that mirror the results that would have been gathered on non-encrypted data [201]. Applying homomorphic and other cryptographic concepts to PPDM is a novel way to increase security without having to deal with difficulties using encrypted data, such as ensuring trust between parties, efficiently sharing keys, and facing expensive decryption costs [32, 71]. The parallels between homomorphic encryption and PPDM are clear; homomorphism may be seen as an extreme application of a mapping method. Both methods provide users with new data that are representative of the original, protect the privacy of the individuals, and may be widely disseminated without concerns about end-user “trust.” Recently, PPDM researchers have

explored new strategies that exhibit this feature to provide strong security and privacy. Song et al. [165] used homomorphism to develop a privacy algorithm based on cryptographic models. They combined homomorphic encryption with learned models, merging these disciplines with PPDM.

These homomorphic PPDM strategies reflect a trend for these methods to not only be robust against adversarial attacks but, in a similar vein as cryptographic methods, to integrate these methods into their design and operation [183]. As homomorphic methods may bridge the gap between PPDM and cryptography, they may become increasingly popular, particularly for widely shared data.

6 PPDM FOR LOCATION INFORMATION

With the rapidly increasing ubiquity of mobile devices, as well as clinical applications for IOT devices location has become an increasingly common data feature whose privacy must be maintained. Many smartphone applications rely upon enabling location services. Doing this opens the door for the network provider and device provider as well as the app designer to collect (and disseminate) location information. An attacker also uses these locations to learn intimate details about a person's life [31, 50, 51, 126]. Location data are also providing increasingly critical insights for clinicians. Knowledge of a user's location offers context when examining the influences and symptoms of an individual's health. Such contexts include knowledge of frequented locations, activity level, interruptions in daily routines, alerts of possible wandering behavior, social interactions, and symptoms of specific diseases [19, 22, 38, 108, 111, 143, 166]. Therefore, privacy protection of location data is an important component of ensuring private, applicable clinical data. While location-based privacy preservation is similar to traditional PPDM methods, unique challenges arise due to both the comparatively few features and the known value constraints. These challenges mean that while many ideas and practices can be transferred to location-based problems, they must often be altered to adequately protect privacy while conveying useful information.

Location data can be difficult to keep private, because some mobile operating systems store this information when location services are enabled. However, too severe of a privacy threshold greatly degrades the usability of the location data [51, 179]. Due to these unique constraints on privacy-preserving location mining, standard PPDM methods must adapt both their PPDM goals and their strategies.

User location data often appear as a series of (latitude, longitude, altitude) coordinates indicating the movement of a user over a time period [63]. Therefore, methods that attempt to preserve the privacy of user locations typically modify the reported location values, location time stamps, or both [50, 51, 179, 204]. Location-based privacy should be addressed separately because of the unique nature of these data. Location trajectories are time series, containing spatio-temporal relationships between individual readings. As a result, making changes to individual data points can easily distort the underlying, valuable information.

Many of the methods can be considered analogous to common PPDM methods discussed in Section 4. Moving the locations by incremental amounts is very similar to random methods [51, 188]. Also, clustering/partitioning location is very similar to mapping and abstraction methods. One general methodology groups location points and abstracts them to a broader neighborhood within which multiple clusters can fit [50, 126, 179, 188]. To further ensure privacy, noise can be added to these broad locations, subsequently increasing the difficulty of determining the cluster locations [188], similarly to the combined strategies discussed in Section 5. Finally, some versions of mobile privacy introduce the concept of "trusted nodes," to which the mobile element will only connect, decreasing the risk of a malicious entity gaining unauthorized information [195]. Location privacy remains an open challenge that requires additional research to retain both the value of location data and the privacy of the individuals being tracked. The increased frequency of patient

location data being collected from a variety of sensors presents a unique challenge to the PPDM conscious researcher and is an increasingly relevant vulnerability that often must be addressed to safeguard the security of data members. Due to the potential inclusion of location data into a clinical record, safeguards to ensure the security of this data are necessary. As such, while the methods and motivations of privatizing location data may not primarily be focused on clinical usage, the inclusion of location into clinical data necessitates the investigation of this PPDM area.

7 DISCUSSION

Throughout this article, we survey recent methods for privacy-preserving data mining, assess the vulnerability of the methods to re-identification, and discuss how to adapt such methods to location-based clinical data. As discussed in Section 5, accessing sensitive data remains a clear and present threat. Because safeguarding patient personal information is a high priority, this threat motivates us to find ways to ensure data privacy, while maintaining data utility. Here, we shed light on the strengths and weaknesses of PPDM techniques as well as highlight directions that warrant continued research. Table 5 summarizes many of the surveyed approaches. From this table, we can view differences between strategies.

As Table 5 indicates, Abstraction methods typically lose data fidelity when privacy is increased, making them appropriate only when details of the original data are not required. They may be an effective approach for preserving the privacy of data that possesses small margins between classes. An example of this could be detecting the volume of hard-to-locate tumors [18].

Random methods can be very effective in that they change the individual data points slightly while still keeping the information as similar to the original as possible. This approach can be difficult to enact correctly, however, as additions of too little noise can result in an adversary being able to “see past” the noise and discover the original data. At the same time, trying to fix this problem by aggressively adding noise may jeopardize the integrity of the original data. Random methods are also weak at providing privacy to non-continuous data; the minimum amount of noise to be added in such cases is an integer value. Despite these flaws, noise addition is still popular with applications using differential privacy [58, 62, 102]. The addition of noise adds a quantifiable amount of uncertainty, clearly defining a trade-off between privacy and utility. Because random methods offer flexibility in the amount of data manipulation that is performed, they allow practitioners to increase privacy for vulnerable populations while opting to retain data purity for less sensitive cases [156].

Mapping methods are adept at preserving the relationships between different groups of data. This is a useful trait when the goal of the project is classification. Another use of mapping is feature swapping, which creates semi-new data points out of the ones in the original data. Mappings are useful when the original data distribution needs to be preserved. One example in a medical setting is identifying outliers, as will occur when searching for medication errors [158].

Rather than operate on some data to privatize them as the previous three methods do, machine learning methods refine the knowledge into a model (e.g., decision tree, deep network). While this may decrease the overall utility of the information, secure learned models can provide precisely the needed information while not including details that may leave individuals vulnerable to exposure [57].

The generation of synthetic data represents another shift in approach to the privatization of data. Through a learned approximation of the original data—by means of statistics or deep learning—a generator creates data that ostensibly could have come from the original distribution, and exemplifies all the characteristics of the original data. This approach is difficult to successfully perform with data possessing extreme outliers or abnormal distributions, as accounting for these

Table 5. Comparison between Different Presented PPDM Methods

Method	Algorithm	Performance	Vulnerability	Data Type	Runtime
Abstraction	Condensation [206]	ΔAcc : within 5%	sparse; redundant features	categorical; continuous	$\frac{N * F}{K}$
	HM:PFSOM [3]	IL : ~ 0.35	unclustered data	categorical; continuous	$F * N^2$
	1:M-Generalization [53]	IL : ~ 0.15	sparse	categorical; continuous	$N \log(N)$
	MS(k, θ)-anonymity [101]	IL : ~ 0.1	sparse	categorical reports	N^2
Random Noise	GADP [125]	$Dist$: no change from original	topological irregularities	continuous	N
	Chaos Method [44]	ΔAcc : within 1%	loss of correlation	categorical; continuous	N
Mapping	Donut Method [60]	CS : within 10%	topological irregularities	location	N
	Translation Data Perturbation [130]	$\Delta Error$: within 7%	highly correlated data	categorical; continuous	$N * F$
	Geometric Data Perturbation [182]	ΔAcc : within 5%	mapping method compromise	categorical; continuous	$N * F$
	P2RoCAI [23]	ΔAcc : within 2%	large compute time	categorical; continuous	N^3
Generation	Privacy Data Generator [184]	$Diss$: ~ 0.06	representative input	categorical; continuous	$N * F$

ΔAcc = difference in accuracy after privacy preservation, IL = information loss after privacy preservation, $Dist$ = change in distribution of data after privacy preservation, $Diss$ = dissimilarity measure between data before and after privacy preservation, CS = cluster similarity before and after privacy preservation, N = number of data points, F = number of features, K = number of clusters.

may leave those individuals prone to exposure, but ignoring them may severely degrade the quality of the data. Conversely, when a model can successfully approximate the distribution of the original data, the generation of synthetic data is a powerful way to effectively sidestep the issue of privacy. When faced with an extreme desire for privacy, such as when dealing with a novel or uncommon affliction, synthetic data generation provides an avenue to share data that are similar to the original but reference no real participants [147].

Clinical research and practice impose their own constraints on the choice of optimal PPDM methods. Clinical pathways, medical imaging, and location information all offer a unique challenge for the researcher. Much like more standard clinical data, effectively safeguarding these data types requires consideration of both the form of the data and the desired use of that data. When attempting to safeguard sensitive data, the intended use for the data plays as large of a role in the choice of method as the data. To better illustrate some potential uses for each of the discussed methods, Table 6 ranks each of the discussed categories across several use cases.

Finally, the computational cost of PPDM methods may impact their value. As seen in Table 5, random noise addition and mapping methods are not computationally costly, while abstraction and generation both require additional, potentially costly, steps. Abstraction methods must often use a metric such as K-anonymity to decide which entries to modify. Data generation methods, especially those using deep learning, must be trained on the data before providing useful results. As a result, such methods may impose significant computational constraints.

Table 6. Sample Clinical Tasks and Ranked Suitability of Alternative PPDM Methods to the Task
(1 = Most Suitable . . . 5 = Least Suitable)

PPDM Method	Hypothesis test	Clustering	Base stats	Classification	Detect anomaly
Abstraction	3	4	4	5	4
Random	1	3	1	4	2
Mapping	2	2	3	1	1
Learned	5	1	5	2	5
Generation	4	5	2	3	3

8 SUGGESTED DIRECTIONS FOR FUTURE PPDM RESEARCH

While we highlight novel and robust methodologies in this survey, there are several avenues of research that are needed to extend and strengthen PPDM. As seen in Section 7, there is no general consensus on best practices to use for evaluating the efficacy of PPDM methods. While differential privacy has become an increasingly common method of validating the privacy of privatized data, developing a measure that combines the privacy given to data along with the preserved utility would be a good method of providing insight into the overall utility of the proposed method. In particular, developing an evaluation criterion that works across multiple domains, types of data, and classes of PPDM models would be of great benefit to the community as a whole. Similarly, standardizing the data and testing methods used for newly proposed PPDM methods would facilitate comparisons between these different methods as well as the selection of an appropriate approach to a particular data set. These measures could use many different aspects of the input data, such as composition (described in Section 2). An example of a criterion that could perform well is representing the overall utility of a PPDM method as the area under a curve, where the X-axis represents varying degrees of privacy and the Y-axis represents the utility of the data. Methods that exhibit a large area under the curve would be ones that retain high levels of utility as PPDM parameters vary. Another possibly more focused avenue for PPDM metrics is a general-purpose privacy metric for synthetic data. As it currently stands, it is often quite difficult to quantify the privacy provided through the use of synthetic data. If synthetic data is generated correctly, then there exist no correct ties to the original data, making it difficult to establish a link between a subset synthetic data and any possible originating record in the original data. However, as synthetic data is not always generated completely free of relation to the original data, proposing a metric aimed at grading synthetic data quality would be extremely useful.

Along with standardizing the evaluation criteria of proposed PPDM models, an effective further direction for this field is the integration of a re-identification agent within a PPDM framework. As seen in Sections 4.4 and 4.5, deep learning models, especially GANs, have shown to be an effective way of augmenting or generating data that protects the privacy of the members contained within. Creating a GAN that not only evaluates the synthetic data for realism but also attempts to re-identify the generated data, could result in a mechanism that produces synthetic data that is representative of the original, but is also robust against adversarial attacks on the data.

A rewarding avenue for PPDM research may also be the introduction of a class of methods that attempt to provide private data through a transformation of related data. Transfer learning and domain adaptation are popular research areas and may be repurposed to facilitate taking secure information and translating it into an insecure domain. This proposed style of transfer PPDM methods would exhibit the privacy characteristics of synthetic data, but the relevance to real data of mapping methods. Finally, PPDM practices may be used in the field of adversarial learning [104]. Adversarial learning is characterized by the interplay between a learning model and an agent who attempts to poison the performance of that model. Modification of the model's training data using

PPDM methods may increase the robustness of the learned models involved, due to the decreased similarity between data used for the model and data used by an adversary.

REFERENCES

- [1] Scott Aaronson and Guy N. Rothblum. 2019. Gentle measurement of quantum states and differential privacy. In *Proceedings of the Annual ACM Symposium on Theory of Computing*. 322–333. <https://doi.org/10.1145/3313276.3316378>
- [2] Nazmiye Ceren Abay, Yan Zhou, and Bhavani Thuraisingham. 2018. Privacy preserving synthetic data release using deep learning. In *Proceedings of the Joint European Conference on Machine Learning and Knowledge Discovery in Databases*. 510–526. <https://doi.org/10.1007/978-3-662-44851-9>
- [3] Karim Abouelmehdi, Abderrahim Beni-Hessane, and Hayat Khaloufi. 2018. Big healthcare data: Preserving security and privacy. *J. Big Data* 5, 1 (2018), 1–18. <https://doi.org/10.1186/s40537-017-0110-7>
- [4] Charu C. Aggarwal and Philip S. Yu. 2008. *A General Survey of Privacy-Preserving Data Mining Models and Algorithms*. 11–52. https://doi.org/10.1007/978-0-387-70992-5_2
- [5] R. Agrawal and R. Srikant. 2000. Privacy preserving data mining. *SIGMOD Rec.* 29, 2 (2000), 439–450. <https://doi.org/10.19026/rjaset.9.1445>
- [6] Stanley C. Ahalt, Christopher G. Chute, Karamarie Fecho, Gustavo Glusman, Jennifer Hadlock, Casey Overby Taylor, Emily R. Pfaff, Peter N. Robinson, Harold Solbrig, Casey Ta, Nicholas Tatonetti, and Chunhua Weng. 2019. Clinical data: Sources and types, regulatory constraints, applications. *Clin. Translation. Sci.* 12, 4 (2019), 329–333. <https://doi.org/10.1111/cts.12638>
- [7] Alia Alabdulkarim, Mznah Al-Rodhaan, Tian, and Yuan Abdullah Al-Dhelaan. 2019. A privacy-preserving algorithm for clinical decision-support systems using random forest. *Comput. Mater. Cont.* 58, 3 (2019), 585–601. <https://doi.org/10.32604/cmc.2019.05637>
- [8] Ranya Aloufi, Hamed Haddadi, and David Boyle. 2019. Emotionless: Privacy-preserving speech analysis for voice assistants. Retrieved from <http://arxiv.org/abs/1908.03632>.
- [9] Kaiomars P. Anklesaria. 1986. Estimating the future state of a system through time-series nominal data analysis. *J. Operation. Res. Soc.* 37, 12 (1986), 1105–1112.
- [10] Pathum Chamikara Mahawaga Arachchige, Peter Bertok, Ibrahim Khalil, Dongxi Liu, Seyit Camtepe, and Mohammed Atiquzzaman. 2019. Local differential privacy for deep learning. *IEEE Internet Things J.* 7, 7 (2019), 1–16. <https://doi.org/10.1109/jiot.2019.2952146>
- [11] Hilal Asi, John Duchi, and Omid Javidi. 2019. Element level differential privacy: The right granularity of privacy. Retrieved from <http://arxiv.org/abs/1912.04042>.
- [12] Dixie Baker, Bartha M. Knoppers, Mark Phillips, David van Enckevort, Petra Kaufmann, Hanns Lochmuller, and Domenica Taruscio. 2018. Privacy-preserving linkage of genomic and clinical data sets. *IEEE/ACM Transactions on Computational Biology and Bioinformatics* 16, 4 (2018), 1–7. <https://doi.org/10.1109/TCBB.2018.2855125>
- [13] Brett K. Beaulieu-Jones, William Yuan, Samuel G. Finlayson, and Zhiwei Steven Wu. 2018. Privacy-preserving distributed deep learning for clinical data. In *Proceedings of the Machine Learning for Health Workshop (ML4H'18)*. Retrieved from <http://arxiv.org/abs/1812.01484>.
- [14] Kathleen Benitez and Bradley Malin. 2010. Evaluating re-identification risks with respect to the HIPAA privacy rule. *J. Amer. Med. Inform. Assoc.* 17, 2 (2010), 169–177. <https://doi.org/10.1136/jamia.2009.000026>
- [15] Bonnie Berger and Hyunghoon Cho. 2019. Emerging technologies towards enhancing privacy in genomic data sharing. *Genome Biol.* 20, 1 (2019), 19–21. <https://doi.org/10.1186/s13059-019-1741-0>
- [16] Jiang Bian, Alexander Loiacono, Andrei Sura, Tonatiuh Mendoza Viramontes, Gloria Lipori, Yi Guo, Elizabeth Shenkman, and William Hogan. 2019. Implementing a hash-based privacy-preserving record linkage tool in the OneFlorida clinical research network. *JAMIA Open* 2, 4 (2019), 562–569. <https://doi.org/10.1093/jamiaopen/ooz050>
- [17] Alessandro Blasimme, Effy Vayena, and Ine Van Hoyweghen. 2019. Big data, precision medicine and private insurance: A delicate balancing act. *Big Data Soc.* 6, 1 (2019), 1–6. <https://doi.org/10.1177/2053951719830111>
- [18] Ian Boon, Tracy Au Yong, and Cheng Boon. 2018. Assessing the role of artificial intelligence (AI) in clinical oncology: Utility of machine learning in radiotherapy target volume delineation. *Medicine* 5, 4 (2018), 131. <https://doi.org/10.3390/medicines5040131>
- [19] Mehdi Boukhechba, Yu Huang, Philip Chow, Karl Fua, Bethany A. Teachman, and Laura E. Barnes. 2017. Monitoring social anxiety from mobility and communication patterns. In *Proceedings of the ACM International Joint Conference on Pervasive and Ubiquitous Computing and 2020 ACM International Symposium on Wearable Computers (UBICOMP/ISWC'17)*. 749–753. <https://doi.org/10.1145/3123024.3125607>
- [20] Tianxi Cai, Molei Liu, and Yin Xia. 2019. Individual data protected integrative regression analysis of high-dimensional heterogeneous data. Retrieved from <http://arxiv.org/abs/1902.06115>.

- [21] Gregory Caiola and Jerome P. Reiter. 2010. Random forests for generating partially synthetic, categorical data. *Trans. Data Priv.* 3, 1 (2010), 27–42.
- [22] Matteo Cella, Łukasz Okruszek, Megan Lawrence, Valerio Zarlenga, Zhimin He, and Til Wykes. 2018. Using wearable technology to detect the autonomic signature of illness severity in schizophrenia. *Schizophrenia Res.* 195 (2018), 537–542. <https://doi.org/10.1016/j.schres.2017.09.028>
- [23] M. A. P. Chamikara, P. Bertok, D. Liu, S. Camtepe, and I. Khalil. 2018. Efficient data perturbation for privacy preserving and accurate data stream mining. *Pervas. Mobile Comput.* 48 (2018), 1–19. <https://doi.org/10.1016/j.pmcj.2018.05.003>
- [24] Kamalika Chaudhuri, Jacob Imola, and Ashwin Machanavajjhala. 2019. Capacity bounded differential privacy. In *Proceedings of the 33rd Conference on Neural Information Processing Systems*. Retrieved from <http://arxiv.org/abs/1907.02159>.
- [25] Kamalika Chaudhuri and Claire Monteleoni. 2009. Privacy-preserving logistic regression. In *Proceedings of the 21st Conference on Advances in Neural Information Processing Systems*. 289–296. <https://doi.org/10.12720/jait.6.3.88-95>
- [26] Keke Chen and Ling Liu. 2005. Privacy preserving data classification with rotation perturbation. In *Proceedings of the IEEE International Conference on Data Mining (ICDM'05)*. 589–592. <https://doi.org/10.1109/ICDM.2005.121>
- [27] Keke Chen, Gordon Sun, and L. Liu. 2007. Towards attack-resilient geometric data perturbation. In *Proceedings of the 7th SIAM International Conference on Data Mining*. 78–89. <https://doi.org/doi:10.1137/1.9781611972771.8>
- [28] Albert Cheu, Adam Smith, Jonathan Ullman, David Zeber, and Maxim Zhilyaev. 2019. Distributed differential privacy via shuffling. *Lecture Notes in Computer Science (including subseries Lecture Notes in Artificial Intelligence and Lecture Notes in Bioinformatics)* 11476 LNCS. 375–403. https://doi.org/10.1007/978-3-030-17653-2_13
- [29] Jun Xing Chin, Giulio Giaconi, Tomas Tinoco De Rubira, Deniz Gimduz, and Gabriela Hug. 2018. Considering time correlation in the estimation of privacy loss for consumers with smart meters. In *Proceedings of the 20th Power Systems Computation Conference (PSCC'18)*. <https://doi.org/10.23919/PSCC.2018.8442899>
- [30] Hyunghoon Cho, Sean Simmons, Ryan Kim, and Bonnie Berger. 2020. Privacy-preserving biomedical database queries with optimal privacy-utility trade-offs. Retrieved from <https://www.biorxiv.org/content/10.1101/2020.01.16.909010v1>.
- [31] Jeppe H. Christensen, Niels H. Pontoppidan, Rikke Rossing, Marco Anisetti, Doris Eva Bamiou, George Spanoudakis, Louisa Murdin, Thanos Bibas, Dimitris Kikidiks, Nikos Dimakopoulos, Giorgos Giotis, and Apostolos Ecomomou. 2019. Fully synthetic longitudinal real-world data from hearing aid wearers for public health policy modeling. *Front. Neurosci.* 13 (Aug. 2019), 1–5. <https://doi.org/10.3389/fnins.2019.00850>
- [32] Elenora Ciceri, Marco Mosconi, Melek Önen, and Orhan Ermiş. 2019. PAPAYA: A platform for privacy preserving data analytics. Retrieved from <https://www.papaya-project.eu/>.
- [33] Kimberly Claffy and Erin E. Kenneally. 2010. Dialing privacy and utility: A proposed data-sharing framework to advance internet research. *IEEE Secur. Priv.* 8, 4 (2010), 31–39. <https://doi.org/10.1109/MSP.2010.57>
- [34] Andrea Damiani, Carlotta Masciocchi, Luca Boldrini, Roberto Gatta, Nicola Dinapoli, Jacopo Lenkowicz, Giuditta Chiloire, Maria Antonietta Gambacorta, Luca Tagliaferri, Rosa Autorino, Monica Maria Pagliara, Maria Antonietta Blasi, Johan Van Soest, Andre Dekker, and Vincenzo Valentini. 2018. Preliminary data analysis in healthcare multicentric data mining: A privacy-preserving distributed approach. *J. E-Learn. Knowl. Soc.* 14, 1 (2018), 71–81. <https://doi.org/10.20368/1971-8829/1454>
- [35] Saloni Dash, Ritik Dutta, Isabelle Guyon, Adrien Pavao, Andrew Yale, and Kristin P. Bennett. 2019. Synthetic event time series health data generation. Retrieved from <http://arxiv.org/abs/1911.06411>.
- [36] Yves-Alexandre De Montjoye, Sebastien Gambs, Vincent Blondel, Geoffrey Canright, Nicolas De Cordes, Sébastien Deletaille, Kenth Engo-Monsen, Manuel Garcia-Herranz, Jake Kendall, Cameron Kerry, Gautier Krings, Emmanuel Letouze, Miguel Luengo, Nuria Oliver, Luc Rocher, Alex Rutherford, Zbigniew Smoreda, Jessica Steele, Erik Wetter, Alex Pentland, and Linus Bengtsson. 2018. On the privacy-conscious use of mobile phone data. *Nature Publish. Group* 5 (2018), 1–6. <https://doi.org/10.1038/sdata.2018.286>
- [37] Ratan Dey, Cong Tang, Keith Ross, and Nitesh Saxena. 2012. Estimating age privacy leakage in online social networks. *Proceedings of the IEEE International Conference on Computer Communications (INFOCOM'12)*. 2836–2840. <https://doi.org/10.1109/INFCOM.2012.6195711>
- [38] Sonia Difrancesco, Paolo Fraccaro, Sabine N. Van Der Veer, Bader Alshoumr, John Ainsworth, Riccardo Bellazzi, and Niels Peek. 2016. Out-of-home activity recognition from GPS data in schizophrenic patients. *Proceedings of the IEEE Symposium on Computer-Based Medical Systems*. 324–328. <https://doi.org/10.1109/CBMS.2016.54>
- [39] Nikunj Domadiya and Udai Pratap Rao. 2019. Privacy preserving distributed association rule mining approach on vertically partitioned healthcare data. *Procedia Comput. Sci.* 148 (2019), 303–312. <https://doi.org/10.1016/j.procs.2019.01.023>
- [40] Cynthia Dwork. 2006. Differential privacy. In *Proceedings of the 33rd International Colloquium on Automata, Languages and Programming, Part II (ICALP'06)*. https://doi.org/10.1007/11787006_1

- [41] Khaled El Emam, Elizabeth Jonker, Luk Arbuckle, and Bradley Malin. 2011. A systematic review of re-identification attacks on health data. *PLoS ONE* 6, 12 (2011). <https://doi.org/10.1371/journal.pone.0028071>
- [42] Josh Eno and Craig W. Thompson. 2008. Generating synthetic data to match data mining patterns. *IEEE Internet Comput.* 12, 3 (2008), 78–82. <https://doi.org/10.1109/MIC.2008.55>
- [43] Úlfar Erlingsson, Vitaly Feldman, Ilya Mironov, Ananth Raghunathan, Kunal Talwar, and Abhradeep Thakurta. 2019. Amplification by shuffling: From local to central differential privacy via anonymity. In *Proceedings of the Annual ACM-SIAM Symposium on Discrete Algorithms*. 2468–2479. <https://doi.org/10.1137/1.9781611975482.151>
- [44] Can Eyupoglu, Muhammed Aydin, Abdul Zaim, and Ahmet Sertbas. 2018. An efficient big data anonymization algorithm based on chaos and perturbation techniques. *Entropy* 20, 5 (2018), 373. <https://doi.org/10.3390/e20050373>
- [45] Stephan A. Fahrenkrog-Petersen, Han Van Der Aa, and Matthias Weidlich. 2019. PRETSA: Event log sanitization for privacy-aware process discovery. *Proceedings of the International Conference on Process Mining (ICPM'19)*. 1–8. <https://doi.org/10.1109/ICPM.2019.00012>
- [46] Andrew David Foote, Ashwin Machanavajjhala, and Kevin McKinney. 2019. Releasing earnings distributions using differential privacy. *J. Priv. Confidential.* 9, 2 (2019). <https://doi.org/10.29012/jpc.722>
- [47] Julien Freudiger, Maxim Raya, Márk Félegyházi, Panos Papadimitratos, and Jean-Pierre Hubaux. 2007. Mix-zones for location privacy in vehicular networks. In *Proceedings of the ACM Workshop on Wireless Networking for Intelligent Transportation Systems (WiN-ITS'07)*.
- [48] Max Friedrich, Arne Köhn, Gregor Wiedemann, and Chris Biemann. 2019. Adversarial learning of privacy-preserving text representations for de-identification of medical records. In *Proceedings of the 57th Annual Meeting of the Association for Computational Linguistics*. 5829–5839. <https://doi.org/10.18653/v1/p19-1584>
- [49] Chong zhi Gao, Qiong Cheng, Pei He, Willy Susilo, and Jin Li. 2018. Privacy-preserving Naive Bayes classifiers secure against the substitution-then-comparison attack. *Info. Sci.* 444 (2018), 72–88. <https://doi.org/10.1016/j.ins.2018.02.058>
- [50] B. Gedik and Ling Liu. 2005. Location privacy in mobile systems: A personalized anonymization model. In *Proceedings of the 25th IEEE International Conference on Distributed Computing Systems*. 620–629. <https://doi.org/10.1109/icdc.2005.48>
- [51] Buğra Gedik and Ling Liu. 2008. Protecting location privacy with personalized k-anonymity: Architecture and algorithms. *IEEE Trans. Mobile Comput.* 7, 1 (2008), 1–18. <https://doi.org/10.1109/TMC.2007.1062>
- [52] Gabriel Ghinita, Yufei Tao, and Panos Kalnis. 2008. On the anonymization of sparse high-dimensional data. In *Proceedings of the International Conference on Data Engineering*. 715–724.
- [53] Qi Yuan Gong, Junzhou Luo, Ming Yang, Weiwei Ni, and Xiao Bai Li. 2017. Anonymizing 1:M microdata with high utility. *Knowl.-Based Syst.* 115 (2017), 15–26. <https://doi.org/10.1016/j.knosys.2016.10.012>
- [54] Daniel L. Goroff. 2015. Balancing privacy versus accuracy in research protocols. *Science* 347, 6221 (2015), 479–480. <https://doi.org/10.1126/science.aaa3483>
- [55] Julian Gruendner, Thorsten Schwachhofer, Phillip Sippl, Nicolas Wolf, Marcel Erpenbeck, Christian Gulden, Lorenz A. Kapsner, Jakob Zierk, Sebastian Mate, Michael Stürzl, Roland Croner, Hans Ulrich Prokosch, and Dennis Toddenroth. 2019. Ketos: Clinical decision support and machine learning as a service — A training and deployment platform based on Docker, OMOP-CDM, and FHIR Web Services. *PLoS ONE* 14, 10 (2019), 1–16. <https://doi.org/10.1371/journal.pone.0223010>
- [56] Xiaolan Gu, Ming Li, Yang Cao, and Li Xiong. 2019. Supporting both range queries and frequency estimation with local differential privacy. In *Proceedings of the IEEE Conference on Communications and Network Security (CNS'19)*. 124–132. <https://doi.org/10.1109/CNS.2019.8802778>
- [57] Zhitao Guan, Zefang Lv, Xiaojian Du, Longfei Wu, and Mohsen Guizani. 2019. Achieving data utility-privacy tradeoff in internet of medical things: A machine learning approach. *Future Gen. Comput. Syst.* 98 (2019), 60–68. <https://doi.org/10.1016/j.future.2019.01.058>
- [58] Mehmet Emre Gursoy, Ali Inan, Mehmet Ercan Nergiz, and Yucel Saygin. 2017. *Differentially Private Nearest Neighbor Classification*. Vol. 31. Springer U.S., 1544–1575. <https://doi.org/10.1007/s10618-017-0532-z>
- [59] Nina Hallowell, Michael Parker, and Christoffer Nellåker. 2019. Big data phenotyping in rare diseases: some ethical issues. *Genet. Med.* 21, 2 (2019), 272–274. <https://doi.org/10.1038/s41436-018-0067-8>
- [60] Kristen H. Hampton, Molly K. Fitch, William B. Allshouse, Irene A. Doherty, Dionne C. Gesink, Peter A. Leone, Marc L. Serre, and William C. Miller. 2010. Mapping health data: Improved privacy protection with donut method geomasking. *Amer. J. Epidemiol.* 172, 9 (2010), 1062–1069. <https://doi.org/10.1093/aje/kwq248>
- [61] Daniel F. Hayes, Robert C. Bast, Christopher E. Desch, Herbert Fritsche, Nancy E. Kemeny, J. Milburn Jessup, Gershon Y. Locker, John S. Macdonald, Robert G. Mennel, Larry Norton, Peter Ravdin, Sheila Taube, and Rodger J. Winn. 1996. Tumor marker utility grading system: A framework to evaluate clinical utility of tumor markers. *J. Natl. Cancer Inst.* 88, 20 (1996), 1456–1466. <https://doi.org/10.1093/jnci/88.20.1456>
- [62] Michael Hilton. 2018. Differential privacy: A historical survey. Retrieved from <https://www.scopus.com/inward/record.uri?eid=2-s2.0-0021010509&partnerID=40&md5=9daf5f8b395159093ea1259e8291aeb6>.

- [63] Shen Shyang Ho and Shuhua Ruan. 2013. Preserving privacy for interesting location pattern mining from trajectory data. *Trans. Data Priv.* 6, 1 (2013), 87–106.
- [64] Zonghao Huang, Rui Hu, Yuanxiong Guo, Eric Chan-Tin, and Yanmin Gong. 2020. DP-ADMM: ADMM-based distributed learning with differential privacy. *IEEE Trans. Info. Forens. Secur.* 15 (2020), 1002–1012. <https://doi.org/10.1109/TIFS.2019.2931068>
- [65] Hisham Husain, Zac Cranko, and Richard Nock. 2018. Integral privacy for sampling from mollifier densities with approximation guarantees. Retrieved from <http://arxiv.org/abs/1806.04819>.
- [66] J. Hyma, G. Lakshmeeswari, D. S. Sampath Kumar, and Ayush Anand. 2016. An efficient privacy preserving medical image retrieval using ROI enabled searchable encryption. *Int. J. Appl. Eng. Res.* 11, 11 (2016), 7509–7516. <https://doi.org/10.37622/IJAER/11.11.2016.7509-7516>
- [67] J. Iavindrasana, G. Cohen, A. Depeursinge, H. Müller, R. Meyer, and A. Geissbuhler. 2009. Clinical data mining: A review. *Yearbook Med. Inform.* (2009), 121–133. <https://doi.org/10.1055/s-0038-1638651>
- [68] Selim Ickin, Konstantinos Vandikas, and Markus Fiedler. 2019. Privacy preserving QoE modeling using collaborative learning. In *Proceedings of the Annual International Conference on Mobile Computing and Networking (MOBICOM'19)*. 13–18. <https://doi.org/10.1145/3349611.3355548>
- [69] Joonas Jälkö, Eemil Lagerspetz, Jari Haukka, Sasu Tarkoma, Samuel Kaski, and Antti Honkela. 2019. Privacy-preserving data sharing via probabilistic modelling. Retrieved from <http://arxiv.org/abs/1912.04439>.
- [70] Jay Kim. 1986. A method for limiting disclosure in microdata based on random noise and transformation. *Sect. Survey Res. Methods* 3 (1986), 303–308.
- [71] Hao Jin, Yan Luo, Peilong Li, and Jomol Mathew. 2019. A review of secure and privacy-preserving medical data sharing. *IEEE Access* 7 (2019), 61656–61669. <https://doi.org/10.1109/ACCESS.2019.2916503>
- [72] M. Eric Johnson. 2009. Data hemorrhages in the health-care sector. In *Financial Cryptography and Data Security*, Roger Dingledine and Philippe Golle (Eds.). Springer, Berlin, 71–89.
- [73] Michael Jones, Matthew Johnson, Mark Shervy, Joel T. Dudley, and Noah Zimmerman. 2019. Privacy-preserving methods for feature engineering using blockchain: Review, evaluation, and proof-of-concept. *J. Med. Internet Res.* 21, 8 (2019), 1–18. <https://doi.org/10.2196/13600>
- [74] H. Kargupta, S. Datta, Q. Wang, and Krishnamoorthy Sivakumar. 2004. On the privacy preserving properties of random data perturbation techniques. In *Third IEEE International Conference on Data Mining*. 99–106. <https://doi.org/10.1109/icdm.2003.1250908>
- [75] Hillol Kargupta, Souptik Datta, Qi Wang, and Krishnamoorthy Sivakumar. 2005. Random-data perturbation techniques and privacy-preserving data mining. *Knowl. Info. Syst.* 7, 4 (2005), 387–414. <https://doi.org/10.1007/s10115-004-0173-6>
- [76] Mohammed Ketel and Abdollah Homaifar. 2007. Privacy-preserving mining by rotational data transformation. In *Proceedings of the 43rd ACM Southeast Conference*. 233. <https://doi.org/10.1145/1167350.1167419>
- [77] Y. I. Khamlichi, M. Machkour, K. Afdel, and A. Moudden. 2006. Medical image watermarked by simultaneous moment invariants and content-based for privacy and tamper detection. In *Proceedings of the 6th WSEAS International Conference on Multimedia Systems and Signal Processing*. 16–18.
- [78] Razaullah Khan, Xiaofeng Tao, Adeel Anjum, Haider Sajjad, Rehman Malik, Abid Khan, and Fatemeh Amiri. 2020. Privacy preserving for multiple sensitive attributes against fingerprint correlation attack satisfying c-diversity. *Wireless Commun. Mobile Comput.* 2020, 8416823 (2020), 18.
- [79] Saira Khan, Khalid Iqbal, Safi Faizullah, Muhammad Fahad, Jawad Ali, and Waqas Ahmed. 2019. Clustering based privacy preserving of big data using fuzzification and anonymization operation. *IJACSA* 10, 12 (2019), 282–289.
- [80] Daniel Kifer and Ashwin Machanavajjhala. 2011. No free lunch in data privacy. In *Proceedings of the 2011 ACM SIGMOD International Conference on Management of Data*. 193. <https://doi.org/10.1145/1989323.1989345>
- [81] Bach Ngoc Kim, Jose Dolz, Pierre-Marc Jodoin, and Christian Desrosiers. 2019. Privacy-net: An adversarial approach for identity-obfuscated segmentation of medical images. Retrieved from <http://arxiv.org/abs/1909.04087>.
- [82] Dongjin Kim, Zhiyuan Chen, and Aryya Gangopadhyay. 2012. Optimizing privacy-accuracy tradeoff for privacy preserving distance-based classification. *Int. J. Info. Secur. Privacy* 6, 2 (2012), 16–33. <https://doi.org/10.4018/jisp.2012040102>
- [83] Hyunsung Kim. 2019. Research issues on data centric security and privacy model for intelligent internet of things based healthcare. *Biomed. J. Sci. Tech. Res.* 16, 3 (2019), 12050–12052. <https://doi.org/10.26717/bjstr.2019.16.002856>
- [84] Jong Wook Kim, Kennedy Edemacu, and Beakcheol Jang. 2019. MPPDS: Multilevel privacy-preserving data sharing in a collaborative eHealth system. *IEEE Access* 7 (2019), 109910–109923. <https://doi.org/10.1109/access.2019.2933542>
- [85] Miran Kim, Junghye Lee, Lucila Ohno-Machado, and Xiaoqian Jiang. 2020. Secure and differentially private logistic regression for horizontally distributed data. *IEEE Trans. Info. Forens. Secur.* 15 (2020), 695–710. <https://doi.org/10.1109/TIFS.2019.2925496>

- [86] Youngjun Kim and Stéphane M. Meystre. 2020. Ensemble method-based extraction of medication and related information from clinical texts. *J. Amer. Med. Inform. Assoc.* 27, 1 (2020), 31–38. <https://doi.org/10.1093/jamia/ocz100>
- [87] Leigh Kinsman, Thomas Rotter, Erica James, Pamela Snow, and Jon Willis. 2010. What is a clinical pathway? Development of a definition to inform the debate. *BMC Med.* 8 (2010), 8–10. <https://doi.org/10.1186/1741-7015-8-31>
- [88] William A. Knaus and Richard D. Marks. 2019. New phenotypes for sepsis: The promise and problem of applying machine learning and artificial intelligence in clinical research. *J. Amer. Med. Assoc.* 321, 20 (2019), 1981–1982. <https://doi.org/10.1001/jama.2019.5794>
- [89] Bogdan Korel. 1990. Automated software test data generation. *IEEE Trans. Softw. Eng.* 16, 8 (1990), 870–879. <https://doi.org/10.1109/32.57624>
- [90] Christian Kurtz, Martin Semmann, and Wolfgang Schulz. 2018. Towards a framework for information privacy in complex service ecosystems. In *Thirty Ninth International Conference on Information Systems*. 1–9.
- [91] Christoph F. Kurz, Martin Rehm, Rolf Holle, Christina Teuner, Michael Laxy, and Larissa Schwarzkopf. 2019. The effect of bariatric surgery on health care costs: A synthetic control approach using Bayesian structural time series. *Health Econ. (UK)* 28, 11 (2019), 1293–1307. <https://doi.org/10.1002/hec.3941>
- [92] Diane Lambert. 1993. Measures of disclosure risk and harm. *J. Offic. Stat. Stockholm* 9 (1993), 313–313. Retrieved from <http://www.jos.nu/Articles/abstract.asp?article=92313>.
- [93] Mathias Lecuyer, Vaggelis Atlidakis, Roxana Geambasu, Daniel Hsu, and Suman Jana. 2019. Certified robustness to adversarial examples with differential privacy. *Proceedings of the IEEE Symposium on Security and Privacy*. 656–672. <https://doi.org/10.1109/SP.2019.00044>
- [94] Ninghui Li, Tiancheng Li, and Suresh Venkatasubramanian. 2007. *t-Closeness: Privacy Beyond k-anonymity and New-line: l-Diversity*.
- [95] Ninghui Li, Wahbeh Qardaji, and Dong Su. 2010. Provably private data anonymization: Or, k-anonymity meets differential privacy. Retrieved from <https://arXiv:1101.2604>. <https://doi.org/10.1007/s40279-014-0145-2>
- [96] Tong Li, Zhengan Huang, Ping Li, Zheli Liu, and Chunfu Jia. 2018. Outsourced privacy-preserving classification service over encrypted data. *J. Netw. Comput. Appl.* 106 (2018), 100–110. <https://doi.org/10.1016/j.jnca.2017.12.021>
- [97] Wenqi Li, Fausto Milletari, Daguang Xu, Nicola Rieke, Jonny Hancox, Wentao Zhu, Maximilian Baust, Yan Cheng, Sébastien Ourselin, M. Jorge Cardoso, and Andrew Feng. 2019. Privacy-preserving federated brain tumour segmentation. *Lecture Notes in Computer Science (Including Subseries Lecture Notes in Artificial Intelligence and Lecture Notes in Bioinformatics)* 11861 LNCS. 133–141. https://doi.org/10.1007/978-3-030-32692-0_16
- [98] Xiaoxiao Li, Yufeng Gu, Nicha Dvornek, Lawrence Staib, Pamela Ventola, and James S. Duncan. 2020. Multi-site fMRI analysis using privacy-preserving federated learning and domain adaptation: ABIDE results. *Abide I* (2020), 1–13. Retrieved from <http://arxiv.org/abs/2001.05647>.
- [99] Katrina Ligett, Seth Neel, Aaron Roth, Bo Waggoner, and Zhiwei Steven Wu. 2019. Accuracy first: Selecting a differential privacy level for accuracy-constrained ERM. *Adv. Neural Info. Process. Syst.* 9, 2 (2019), 2567–2577. <https://doi.org/10.29012/jpc.682>
- [100] Keng-pei Lin and Ming-syan Chen. 2011. On the design and analysis of the privacy-preserving SVM classifier. *IEEE Trans. Knowl. Data Eng.* 23, 11 (2011), 1704–1717. <https://doi.org/10.1109/TKDE.2010.193>
- [101] Wen Yang Lin, Duen Chuan Yang, and Jie Teng Wang. 2016. Privacy preserving data anonymization of spontaneous ADE reporting system dataset. *BMC Med. Inform. Decis. Mak.* 16, Suppl 1 (2016). <https://doi.org/10.1186/s12911-016-0293-4>
- [102] Ao Liu, Lirong Xia, Andrew Duchowski, Reynold Bailey, Kenneth Holmqvist, and Eakta Jain. 2019. Differential privacy for eye-tracking data. In *Proceedings of the Eye Tracking Research and Applications Symposium (ETRA'19)*. <https://doi.org/10.1145/3314111.3319823>
- [103] Yi Liu, Jialiang Peng, James J. Q. Yu, and Yi Wu. 2020. PPGAN: Privacy-preserving generative adversarial network. In *Proceedings of the IEEE International Conference on Parallel and Distributed Systems (ICPADS'20)*. 985–989. <https://doi.org/10.1109/icpads47876.2019.00150>
- [104] Daniel Lowd and Christopher Meek. 2005. Adversarial learning. In *the Eleventh ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*. 641–647.
- [105] Jing Ma, Joyce C. Ho, Qiuchen Zhang, Li Xiong, Jian Lou, and Xiaoqian Jiang. 2019. Privacy-preserving tensor factorization for collaborative health data analysis. In *Proceedings of the International Conference on Information and Knowledge Management, Proceedings*. 1291–1300. <https://doi.org/10.1145/3357384.3357878>
- [106] Rong Ma, Yi Li, Chenxing Li, Fangping Wan, Hailin Hu, Wei Xu, and Jianyang Zeng. 2020. Secure multiparty computation for privacy-preserving drug discovery. *Bioinformatics* 36, 9 (2020), 2872–2880.
- [107] Ashwin Machanavajjhala, Daniel Kifer, Johannes Gehrke, and Muthuramakrishnan Venkitasubramaniam. 2007. L-diversity. *ACM Trans. Knowl. Discov. Data* 1, 1 (2007), 3–es. <https://doi.org/10.1145/1217299.1217302>
- [108] Sumit Majumder, Tapas Mondal, and M. Jamal Deen. 2017. Wearable sensors for remote health monitoring. *Sensors (Switzerland)* 17, 1 (2017). <https://doi.org/10.3390/s17010130>

- [109] Bradley Malin. 2006. Re-identification of familial database records. In *Proceedings of the AMIA Annual Symposium*. 524–528. [https://doi.org/86122\[pii\]](https://doi.org/86122[pii])
- [110] Bradley Malin, David Karp, and Richard H. Scheuermann. 2010. Technical and policy approaches to balancing patient privacy and data sharing in clinical and translational research. *J. Investigat. Med.* 58, 1 (2010), 11–8. <https://doi.org/10.2310/JIM.0b013e3181c9b2ea>
- [111] Martina Mancini, Heather Schlueter, Mahmoud El-Gohary, Nora Mattek, Colette Duncan, Jeffrey Kaye, and Fay B. Horak. 2016. Continuous monitoring of turning mobility and its association to falls and cognitive function: A pilot study. *J. Gerontol. Ser. A Biol. Sci. Med. Sci.* 71, 8 (2016), 1102–1108. <https://doi.org/10.1093/gerona/glw019>
- [112] Jyothi Mandala and M. V. P. Chandra Sekhara Rao. 2019. Privacy preservation of data using crow search with adaptive awareness probability. *J. Inform. Secur. Appl.* 44 (2019), 157–169. <https://doi.org/10.1016/j.jisa.2018.12.005>
- [113] Felix Mannhardt, Agnes Koschmider, Nathalie Baracaldo, Matthias Weidlich, and Judith Michael. 2019. Privacy-preserving process mining: Differential privacy for event logs. *Bus. Inform. Syst. Eng.* 61, 5 (2019), 595–614. <https://doi.org/10.1007/s12599-019-00613-3>
- [114] Felix Mannhardt, Sobah Abbas Petersen, and Manuel Fradinho Oliveira. 2018. Privacy challenges for process mining in human-centered industrial environments. *Proceedings of the International Conference on Intelligent Environments (IE'18)*. 64–71. <https://doi.org/10.1109/IE.2018.00017>
- [115] Yunlong Mao, Shanhe Yi, Qun Li, Jinghao Feng, Fengyuan Xu, and Sheng Zhong. 2018. A privacy-preserving deep learning approach for face recognition with edge computing. Retrieved from <https://www.usenix.org/system/files/conference/hotedge18/hotedge18-papers-mao.pdf%0Ahttps://www.usenix.org/conference/hotedge18/presentation/mao>.
- [116] David McClure and Jerome P. Reiter. 2012. Differential privacy and statistical disclosure risk measures: An investigation with binary synthetic data. *Trans. Data Priv.* 5, 3 (2012), 535–552.
- [117] Ryan McKenna, Raj Kumar Maity, Arya Mazumdar, and Gerome Miklau. 2020. A workload-adaptive mechanism for linear queries under local differential privacy. Retrieved from <http://arxiv.org/abs/2002.01582>.
- [118] Maheyyah Md Siraj, Nurul Adibah Rahmat, and Mazura Mat Din. 2019. A survey on privacy preserving data mining approaches and techniques. *ACM Int. Conf. Proc. Ser.* F1479 (2019), 65–69. <https://doi.org/10.1145/3316615.3316632>
- [119] Rebecca T. Mercuri. 2004. The HIPAA-potamus in health care data security. *Commun. ACM* 47, 7 (2004), 25. <https://doi.org/10.1145/1005817.1005840>
- [120] Jennifer Miller. 2016. How full disclosure of clinical trial data will benefit the pharmaceutical industry. *The Pharmaceutical Journal* 296, 7890 (2016), 1–8. <https://doi.org/10.1211/pj.2016.20201274>
- [121] Jimmy Ming-Tai Wu, Jerry Chun-Wei Lin, Philippe Fournier-Viger, Youcef Djenouri, Chun-Hao Chen, and Zhongcui Li. 2019. The density-based clustering method for privacy-preserving data mining. *Math. Biosci. Eng.* 16, 3 (2019), 1718–1728. <https://doi.org/10.3934/mbe.2019082>
- [122] Brent Mittelstadt. 2019. The ethics of biomedical “big data” analytics. *Philos. Technol.* 32, 1 (2019), 17–21. <https://doi.org/10.1007/s13347-019-00344-z>
- [123] Mona Mohamed, Sahar Ghanem, and Magdy Nagi. 2020. Privacy-preserving for distributed data streams: Towards l-diversity. *Int. Arab J. Info. Technol.* 17, 1 (2020), 52–64. <https://doi.org/10.34028/iajit/17/1/7>
- [124] James J. Morrison. 2019. Evolution in private practice interventional radiology: Data mining trends in procedure volumes. *Sem. Intervent. Radiol.* 36, 1 (2019), 17–22. <https://doi.org/10.1055/s-0039-1683358>
- [125] Krishnamurty Muralidhar, Rahul Parsa, and Rathindra Sarathy. 1999. A general additive data perturbation method for database security. *Manage. Sci.* 45, 10 (1999), 1399–1415. <https://doi.org/10.1287/mnsc.45.10.1399>
- [126] Hoa Ngo and Jong Kim. 2015. Location privacy via differential private perturbation of cloaking area. In *Proceedings of the Computer Security Foundations Workshop*. 63–74. <https://doi.org/10.1109/CSF.2015.12>
- [127] Lina Ni, Chao Li, Xiao Wang, Honglu Jiang, and Jiguo Yu. 2018. DP-MCDBSCAN: Differential privacy preserving multi-core DBSCAN clustering for network user data. *IEEE Access* 6 (2018), 21053–21063. <https://doi.org/10.1109/ACCESS.2018.2824798>
- [128] Saskia Nuñez von Voigt, Stephan A. Fahrenkrog-Petersen, Dominik Janssen, Agnes Koschmider, Florian Tschorsch, Felix Mannhardt, Olaf Landsiedel, and Matthias Weidlich. 2020. Quantifying the re-identification risk of event logs for process mining: Empirical evaluation paper. *Lecture Notes in Computer Science (Including Subseries Lecture Notes in Artificial Intelligence and Lecture Notes in Bioinformatics)* 12127 LNCS. 252–267. https://doi.org/10.1007/978-3-030-49435-3_16
- [129] S. R. M. Oliveira and O. R. Zaiane. 2004. Protecting sensitive knowledge by data sanitization. In *Third IEEE International Conference on Data Mining*. 613–616. <https://doi.org/10.1109/icdm.2003.1250990>
- [130] Stanley R. M. Oliveira and Osmar R. Zaiane. 2010. Privacy preserving clustering by data transformation. *J. Info. Data Manage.* 1, 1 (2010), 37.
- [131] Seyed Ali Osia, Ali Shahin Shamsabadi, Sina Sajadmanesh, Ali Taheri, Kleomenis Katevas, Hamid R. Rabiee, Nicholas D. Lane, and Hamed Haddadi. 2017. A hybrid deep learning architecture for privacy-preserving mobile analytics. Retrieved from <http://arxiv.org/abs/1703.02952>.

- [132] Jisha Jose Panackal and Anitha S. Pillai. 2015. Adaptive utility-based anonymization model: Performance evaluation on big data sets. *Procedia Comput. Sci.* 50 (2015), 347–352. <https://doi.org/10.1016/j.procs.2015.04.037>
- [133] Rupa Parameswaran and Douglas M. Blough. 2005. A robust data-obfuscation approach for privacy preservation of clustered data. In *Proceedings of the Workshop on Privacy and Security Aspects of Data Mining*. 18–25. <https://doi.org/10.1021/om500167r>
- [134] Le Trieu Phong and Tran Thi Phuong. 2019. Privacy-preserving deep learning via weight transmission. *IEEE Trans. Info. Forens. Secur.* 14, 11 (2019), 3003–3015. <https://doi.org/10.1109/TIFS.2019.2911169>
- [135] Anastasiia Pika, Moe T. Wynn, Stephanus Budiono, Arthur H. M. ter Hofstede, Wil M. P. van der Aalst, and Hajo A. Reijers. 2019. Towards privacy-preserving process mining in healthcare. In *Lecture Notes in Business Information Processing*, Vol. 362 LNBP. 483–495. https://doi.org/10.1007/978-3-030-37453-2_39
- [136] Anastasiia Pika, Moe T. Wynn, Stephanus Budiono, Arthur H. M. Ter Hofstede, Wil M. P. van der Aalst, and Hajo A. Reijers. 2020. Privacy-preserving process mining in healthcare. *Int. J. Environ. Res. Public Health* 17, 5 (2020). <https://doi.org/10.3390/ijerph17051612>
- [137] Yuan Ping, Bin Hao, Xiali Hei, Jie Wu, and Baocang Wang. 2020. Maximized privacy-preserving outsourcing on support vector clustering. *Electronics* 9, 1 (2020), 178. <https://doi.org/10.3390/electronics9010178>
- [138] Benny Pinkas. 2007. Cryptographic techniques for privacy-preserving data mining. *ACM SIGKDD Explor. Newslett.* 4, 2 (2007), 12–19. <https://doi.org/10.1145/772862.772865>
- [139] H. Polat and Wenliang Du. 2004. Privacy-preserving collaborative filtering using randomized perturbation techniques. In *Proceedings of the IEEE International Conference on Data Mining*. 625–628. <https://doi.org/10.1109/icdm.2003.1250993>
- [140] W. Nicholson Price and I. Glenn Cohen. 2019. Privacy in the age of medical big data. *Nature Med.* 25, 1 (2019), 37–43. <https://doi.org/10.1038/s41591-018-0272-7>
- [141] Vartika Puri, Shelly Sachdeva, and Parmeet Kaur. 2019. Privacy preserving publication of relational and transaction data: Survey on the anonymization of patient data. *Computer Science Review* 32, 1 (5 2019), 45–61. <https://doi.org/10.1016/j.cosrev.2019.02.001>
- [142] Majid Rafiei, Leopold von Waldthausen, and Wil M. P. van der Aalst. 2020. Supporting confidentiality in process mining using abstraction and encryption. *Lect. Notes Bus. Info. Process.* 379, 2 (2020), 101–123. https://doi.org/10.1007/978-3-030-46633-6_6
- [143] Daniel Rainham, Ian McDowell, Daniel Krewski, and Mike Sawada. 2010. Conceptualizing the healthscape: Contributions of time geography, location technologies and spatial ecology to place and health research. *Soc. Sci. Med.* 70, 5 (2010), 668–676. <https://doi.org/10.1016/j.socscimed.2009.10.035>
- [144] Jean Louis Raisaro, Florian Tramér, Zhanglong Ji, Diyu Bu, Yongan Zhao, Knox Carey, David Lloyd, Heidi Sofia, Dixie Baker, Paul Flicek, Suyash Shringarpure, Carlos Bustamante, Shuang Wang, Xiaoqian Jiang, Lucila Ohno-Machado, Haixu Tang, Xiao Feng Wang, and Jean Pierre Hubaux. 2017. Addressing Beacon re-identification attacks: Quantification and mitigation of privacy risks. *J. Amer. Med. Inform. Assoc.* 24, 4 (2017), 799–805. <https://doi.org/10.1093/jamia/ocw167>
- [145] V. Rajalakshmi and G. S. Anandha Mala. 2014. Anonymization by data relocation using sub-clustering for privacy preserving data mining. *Indian J. Sci. Technol.* 7, 7 (2014), 975–980.
- [146] Priya Ranjan and Raj Kumar Paul. 2019. A survey on privacy preserving mining and limitations. *SHODH SANGAM* 2, 1 (2019), 63–68.
- [147] Debbie Rankin, Michaela Black, Raymond Bond, Jonathan Wallace, Maurice Mulvenna, and Gorka Epelde. 2020. Reliability of supervised machine learning using synthetic data in health care: Model to preserve privacy for data sharing. *JMIR Med. Inform.* 8, 7 (2020), e18910. <https://doi.org/10.2196/18910>
- [148] Haroon Ur Rashid, Fatma Hussain, and Khalid Masood. 2019. Patient privacy: Challenges and opportunities in the age of big data. *Current Science Perspectives* 5, 1 (2019), 1–5.
- [149] Mercedes Rodriguez-Garcia, Montserrat Batet, and David Sánchez. 2017. A semantic framework for noise addition with nominal data. *Knowl.-Based Syst.* 122 (2017), 103–118. <https://doi.org/10.1016/j.knosys.2017.01.032>
- [150] Mercedes Rodriguez-Garcia, Montserrat Batet, and David Sánchez. 2019. Utility-preserving privacy protection of nominal data sets via semantic rank swapping. *Info. Fusion* 45 (2019), 282–295. <https://doi.org/10.1016/j.inffus.2018.02.008>
- [151] M. G. Ruano, G. P. Almeida, F. Palma, J. F. Raposo, and R. T. Ribeiro. 2018. Reliability of medical databases for the use of real word data and data mining techniques for cardiovascular diseases progression in diabetic patients. In *Proceedings of the Global Medical Engineering Physics Exchanges/Pan American Health Care Exchanges, (GMEPE/PAHCE'18)*. 1–6. <https://doi.org/10.1109/GMEPE-PAHCE.2018.8400769>
- [152] Jan Henrik Ziegeldorf, Jan Metzke, Jan R  th, Martin Henze, and Klaus Wehrle. 2017. Privacy-preserving HMM forward computation. In *Proceedings of the Seventh ACM Conference on Data and Application Security and Privacy*. ACM, New York, NY, USA, 83–94. <https://doi.org/10.1145/3029806.3029816>

- [153] Rahul Saha, Gulshan Kumar, Mritunjay Kumar Rai, Reji Thomas, and Se Jung Lim. 2019. Privacy ensured e-healthcare for fog-enhanced IoT based applications. *IEEE Access* 7 (2019), 44536–44543. <https://doi.org/10.1109/ACCESS.2019.2908664>
- [154] Rathindra Sarathy and Krishnamurthy Muralidhar. 2011. Evaluating Laplace noise addition to satisfy differential privacy for numeric data. *Transactions on Data Privacy* 4, 1 (2011), 1–17.
- [155] Marco Savi, Cristina Rottondi, and Giacomo Verticale. 2015. Evaluation of the precision-privacy tradeoff of data perturbation for smart metering. *IEEE Trans. Smart Grid* 6, 5 (2015), 2409–2416. <https://doi.org/10.1109/TSG.2014.2387848>
- [156] Simone Scardapane, Rosa Altילו, Valentina Ciccarelli, Aurelio Uncini, and Massimo Panella. 2018. Privacy-preserving data mining for distributed medical scenarios. In *Smart Innovation, Systems and Technologies*, Anna Esposito, Marcos Faudez-Zanuy, Francesco Carlo Morabito, and Eros Pasero (Eds.). Smart Innovation, Systems and Technologies, Vol. 69. Springer International Publishing, Cham, 119–128. https://doi.org/10.1007/978-3-319-56904-8_12
- [157] Eric E. Schadt. 2012. The changing privacy landscape in the era of big data. *Mol. Syst. Biol.* 8, 612 (2012), 1–3. <https://doi.org/10.1038/msb.2012.47>
- [158] Gordon D. Schiff, Lynn A. Volk, Mayya Volodarskaya, Deborah H. Williams, Lake Walsh, Sara G. Myers, David W. Bates, and Ronen Rozenblum. 2017. Screening for medication errors using an outlier detection system. *J. Amer. Med. Inform. Assoc.* 24, 2 (2017), 281–287. <https://doi.org/10.1093/jamia/ocw171>
- [159] Thomas Schneider and Amos Treiber. 2020. A comment on privacy-preserving scalar product protocols as proposed in “SPOC.” *IEEE Transactions on Parallel and Distributed Systems* 31, 3 (3 2020), 543–546. <https://doi.org/10.1109/TPDS.2019.2939313>
- [160] Mahsa Shabani, Stephanie O. M. Dyke, Luca Marelli, and Pascal Borry. 2019. Variant data sharing by clinical laboratories through public databases: Consent, privacy and further contact for research policies. *Genet. Med.* 21, 5 (2019), 1031–1037. <https://doi.org/10.1038/s41436-018-0316-x>
- [161] S. Sharma, K. Chen, and, A. Sheth. 2018. Towards practical privacy-preserving analytics for IoT and cloud based healthcare systems. *IEEE Internet Computing* 22, 2 (3 2018), 42–51. <https://doi.org/10.1109/MIC.2018.112102519>
- [162] Desmond Ko Khang Siang, Siti Hajar Othman, and Raja Zahilah Raja Mohd Radzi. 2018. A comparative study on perturbation techniques in privacy preserving data mining. *Int. J. Innovat. Comput.* 8, 1 (2018), 27–32.
- [163] Siddharth Singh. 2019. Big dreams with big data! Use of clinical informatics to inform biomarker discovery. *Clin. Translat. Gastroenterol.* 10, 3 (2019), 1–6. <https://doi.org/10.14309/ctg.0000000000000018>
- [164] C. J. Skinner and M. J. Elliot. 2002. A measure of disclosure risk for microdata. *J. Roy. Stat. Soc. Ser. B: Stat. Methodol.* 64, 4 (2002), 855–867. <https://doi.org/10.1111/1467-9868.00365>
- [165] Baek Kyung Song, Joon Soo Yoo, Miyeon Hong, and Ji Won Yoon. 2019. A bitwise design and implementation for privacy-preserving data mining: From atomic operations to advanced algorithms. *Security and Communication Networks* 2019, 1 (10 2019), 1–14. <https://doi.org/10.1155/2019/3648671>
- [166] Gina Sprint, Diane J. Cook, Douglas L. Weeks, and Vladimir Borisov. 2015. Predicting functional independence measure scores during rehabilitation with wearable inertial sensors. *IEEE Access* 3 (2015), 1350–1366. <https://doi.org/10.1109/ACCESS.2015.2468213>
- [167] Likitha Sravya and Rajya Lakshmi. 2017. Privacy-preserving data mining with random decision tree framework. *IOSR J. Comput. Eng.* 19, 4 (2017), 43–49. <https://doi.org/10.9790/0661-1904034349>
- [168] Julian Steil, Inken Hagedstedt, Michael Xuelin Huang, and Andreas Bulling. 2019. Privacy-aware eye tracking using differential privacy. In *Proceedings of the Eye Tracking Research and Applications Symposium (ETRA’19)*. <https://doi.org/10.1145/3314111.3319915>
- [169] Alan B. Sunteb and Ivan P. Fellegi. 1969. A theory for record linkage. *J. Amer. Statist. Assoc.* 64, 328 (1969), 1183–1210.
- [170] Latanya Sweeney. 2002. k-anonymity: A model for protecting privacy. *Int. J. Uncertain. Fuzz. Knowl.-Based Syst.* 10, 05 (2002), 557–570. <https://doi.org/10.1142/s0218488502001648>
- [171] Latanya Sweeney. 2015. Only you, your doctor, and many others may know. *Technol. Sci.* Retrieved from <https://techscience.org/a/2015092903/>.
- [172] Latanya Sweeney, Akua Abu, and Julia Winn. 2013. Identifying participants in the personal genome project by name (a re-identification experiment). *SSRN Electr. J.* (2013), 1–4. <https://doi.org/10.2139/ssrn.2257732>
- [173] Latanya Sweeney and Ji Su Yoo. 2015. De-anonymizing South Korean resident registration numbers shared in prescription data. *Technol. Sci.* Retrieved from <https://techscience.org/a/2015092901/>. <https://doi.org/10.1007/s10964-009-9456-2>
- [174] Fengyi Tang, Wei Wu, Jian Liu, Huimei Wang, and Ming Xian. 2019. Privacy-preserving distributed deep learning via homomorphic re-encryption. *Electronics (Switzerland)* 8, 4 (2019). <https://doi.org/10.3390/electronics8040411>
- [175] Youdong Tao, Yunhai Tong, Shaohua Tan, Shiwei Tang, and Dongqing Yang. 2008. Protecting the publishing identity in multiple tuples. *Lecture Notes in Computer Science (Including Subseries Lecture Notes in Artificial Intelligence and Lecture Notes in Bioinformatics)* 5094 LNCS. 205–218. https://doi.org/10.1007/978-3-540-70567-3_16

- [176] G. Jelin Taric and E. Poovammal. 2017. A survey on privacy preserving data mining techniques. *Indian J. Sci. Technol.* 10, 5 (2017), 1–5. <https://doi.org/10.17485/ijst/2017/v10i5/111138>
- [177] Akbar Telikani, Amir H. Gandomi, Asadollah Shahbahrami, and Mohammad Naderi Dehkordi. 2019. Privacy-preserving in association rule mining using an improved discrete binary artificial bee colony. *Expert Syst. Appl.* (2019).
- [178] Sin Gee Teo, Jianneng Cao, and Vincent C. S. Lee. 2020. DAG: A general model for privacy-preserving data mining. *IEEE Trans. Knowl. Data Eng.* 32, 1 (1 2020), 40–53 <https://doi.org/10.1109/TKDE.2018.2880743>
- [179] Manolis Terrovitis. 2011. Privacy preservation in the dissemination of location data. *SIGKDD Explor.* 13, 1 (2011), 6–18.
- [180] Hong Yen Tran and Jiankun Hu. 2019. Privacy-preserving big data analytics a comprehensive survey. *J. Parallel Distrib. Comput.* 134 (2019), 207–218. <https://doi.org/10.1016/j.jpdc.2019.08.007>
- [181] Aleksei Triastcyn and Boi Faltings. 2019. Federated generative privacy. Retrieved from <http://arxiv.org/abs/1910.08385>.
- [182] Somya Upadhyay, Chetana Sharma, Pravishti Sharma, Prachi Bharadwaj, and K. R. Seeja. 2018. Privacy preserving data mining with 3-D rotation transformation. *J. King Saud Univ. Comput. Info. Sci.* 30, 4 (2018), 524–530. <https://doi.org/10.1016/j.jksuci.2016.11.009>
- [183] Anamaria Vizitiu, Cosmin Ioan Niță, Andrei Puiu, Constantin Suciu, and Lucian Mihai Itu. 2019. Towards privacy-preserving deep learning based medical imaging applications. In *Proceedings of the Symposium on Medical Measurements and Applications (MeMeA'19)*. <https://doi.org/10.1109/MeMeA.2019.8802193>
- [184] Jilles Vreeken, Matthijs Van Leeuwen, and Arno Siebes. 2007. Preserving privacy through data generation. In *Proceedings of the IEEE International Conference on Data Mining (ICDM'07)*. 685–690. <https://doi.org/10.1109/ICDM.2007.25>
- [185] Isabel Wagner and David Eckhoff. 2018. Technical privacy metrics. *Comput. Surveys* 51, 3 (2018), 1–38. <https://doi.org/10.1145/3168389>
- [186] Ning Wang, Xiaokui Xiao, Yin Yang, Jun Zhao, Siu Cheung Hui, Hyejin Shin, Junbum Shin, and Ge Yu. 2019. Collecting and analyzing multidimensional data with local differential privacy. In *Proceedings of the International Conference on Data Engineering*. 638–649. <https://doi.org/10.1109/ICDE.2019.00063>
- [187] Pingshui Wang, Tao Chen, and Zecheng Wang. 2019. Research on privacy preserving data mining. *J. Info. Hid. Priv. Protect.* 1, 2 (2019), 61–68. <https://doi.org/10.32604/jihpp.2019.05943>
- [188] Shuo Wang, Richard Sinnott, and Surya Nepal. 2018. Privacy-protected place of activity mining on big location data. *Proceedings of the IEEE International Conference on Big Data (BigData'17)*. 1101–1108. <https://doi.org/10.1109/BigData.2017.8258035>
- [189] Zhen Wang, Xiang Yue, Soheil Moosavinasab, Yungui Huang, Simon Lin, and Huan Sun. 2019. SurfCon: Synonym discovery on privacy-aware clinical data. In *Proceedings of the ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*. 1578–1586. <https://doi.org/10.1145/3292500.3330894>
- [190] Garrett Wilson and Diane J. Cook. 2020. A survey of unsupervised deep domain adaptation. *ACM Transactions on Intelligent Systems and Technology* 11, 5 (12 2020) <https://doi.org/10.1145/351812.02849v1>
- [191] William E. Winkler. 2014. Matching and record linkage. *Wiley Interdisciplinary Reviews: Computational Statistics* 6, 5 (2014), 313–325. <https://doi.org/10.1002/wics.1317>
- [192] William E. Winkler and D. C. Washington. 2004. *Re-identification Methods for Masked Microdata*. Technical Report.
- [193] William E. Winkler and D. C. Washington. 2005. Re-identification methods for evaluating the confidentiality of analytically valid microdata. *U.S. Census Research Report Series*.
- [194] Bingzhe Wu, Shiwang Zhao, Guangyu Sun, Xiaolu Zhang, Zhong Su, Caihong Zeng, and Zhihong Liu. 2019. P3SGD: Patient privacy preserving SGD for regularizing deep cnns in pathological image classification. *Proceedings of the IEEE Computer Society Conference on Computer Vision and Pattern Recognition*. 2094–2103. <https://doi.org/10.1109/CVPR.2019.00220>
- [195] Dapeng Wu, Lei Fan, Chenlu Zhang, Honggang Wang, and Ruyan Wang. 2018. Dynamical credibility assessment of privacy-preserving strategy for opportunistic mobile crowd sensing. *IEEE Access* 6 (2018), 37430–37443. <https://doi.org/10.1109/ACCESS.2018.2847251>
- [196] Tsu Yang Wu, Jerry Chun Wei Lin, Yuyu Zhang, and Chun Hao Chen. 2019. A grid-based swarm intelligence algorithm for privacy-preserving data mining. *Appl. Sci. (Switzerland)* 9, 4 (2019). <https://doi.org/10.3390/app9040774>
- [197] Xintao Wu, Chintan Sanghvi, Yongge Wang, and Yuliang Zheng. 2005. Privacy aware data generation for testing database applications. In *Proceedings of the International Database Engineering and Applications Symposium (IDEAS'05)*. 317–326. <https://doi.org/10.1109/IDEAS.2005.45>
- [198] Depeng Xu, Shuhan Yuan, and Xintao Wu. 2019. Achieving differential privacy in vertically partitioned multiparty learning. Retrieved from <http://arxiv.org/abs/1911.04587>.
- [199] Min Xu, Tianhao Wang, Bolin Ding, Jingren Zhou, Cheng Hong, and Zhicong Huang. 2018. DPSAAs: Multidimensional data sharing and analytics as services under local differential privacy. *Proc. VLDB Endow.* 12, 12 (2018), 1862–1865. <https://doi.org/10.14778/3352063.3352085>

- [200] Andrew Yale, Saloni Dash, Ritik Dutta, Isabelle Guyon, Adrien Pavao, Andrew Yale, Saloni Dash, Ritik Dutta, Isabelle Guyon, and Adrien Pavao. 2019. Privacy preserving synthetic health data. In *European Symposium on Artificial Neural Networks, Computational Intelligence and Machine Learning (ESANN'19)*. 1–10.
- [201] Yuki Yamada, Kurt Rohloff, and Masato Oguchi. 2019. Homomorphic encryption for privacy-preserving genome sequences search. *Proceedings of the IEEE International Conference on Smart Computing (SMARTCOMP'19)*. 7–12. <https://doi.org/10.1109/SMARTCOMP.2019.00021>
- [202] Yang Yang, Xingxing Xiao, Xue Cai, and Weiming Zhang. 2020. A secure and privacy-preserving technique based on contrast-enhancement reversible data hiding and plaintext encryption for medical images. *IEEE Signal Process. Lett.* 27 (2020), 256–260. <https://doi.org/10.1109/LSP.2020.2965826>
- [203] Qingqing Ye, Haibo Hu, Xiaofeng Meng, and Huadi Zheng. 2019. PrivKV: Key-value data collection with local differential privacy. *Proceedings of the IEEE Symposium on Security and Privacy*. 317–331. <https://doi.org/10.1109/SP.2019.00018>
- [204] Ling Yin, Qian Wang, Shih Lung Shaw, Zhixiang Fang, Jinxing Hu, Ye Tao, and Wei Wang. 2015. Re-identification risk versus data utility for aggregated mobility research using mobile phone location data. *PLoS ONE* 10, 10 (2015). <https://doi.org/10.1371/journal.pone.0140589>
- [205] Yuan Yuan, Eliezer M. Van Allen, Larsson Omberg, Nikhil Wagle, Ali Amin-Mansour, Artem Sokolov, Lauren A. Byers, Yanxun Xu, Kenneth R. Hess, Lixia Diao, Leng Han, Xuelin Huang, Michael S. Lawrence, John N. Weinstein, Josh M. Stuart, Gordon B. Mills, Levi A. Garraway, Adam A. Margolin, Gad Getz, and Han Liang. 2014. Assessing the clinical utility of cancer genomic and proteomic data across tumor types. *Nature Biotechnol.* 32, 7 (2014), 644–652. <https://doi.org/10.1038/nbt.2940>
- [206] Jinquan Zhang, Bowen Zhao, Guochao Song, Lina Ni, and Jiguo Yu. 2019. Maximum delay anonymous clustering feature tree based privacy-preserving data publishing in social networks. *Procedia Comput. Sci.* 147 (2019), 643–646. <https://doi.org/10.1016/j.procs.2019.01.190>
- [207] Qingchen Zhang, Laurence T. Yang, and Zhikui Chen. 2016. Privacy preserving deep computation model on cloud for big data feature learning. *IEEE Trans. Comput.* 65, 5 (2016), 1351–1362. <https://doi.org/10.1109/TC.2015.2470255>
- [208] Han Zhao, Jianfeng Chi, Yuan Tian, and Geoffrey J. Gordon. 2019. Adversarial privacy preservation under attribute inference attack. Retrieved from <http://arxiv.org/abs/1906.07902>.
- [209] Fengyu Zhou, James Anderson, and Steven H. Low. 2019. Differential privacy of aggregated dc optimal power flow data. *Proceedings of the American Control Conference*. 1307–1314. <https://doi.org/10.23919/acc.2019.8815257>

Received April 2020; revised December 2020; accepted January 2021