

一种面向直方图发布的均衡差分隐私保护方法

杨旭东¹⁾ 高 岭^{1),2)} 王 海¹⁾ 郭红波¹⁾ 郑 杰¹⁾

¹⁾(西北大学信息科学与技术学院 西安 710127)

²⁾(西安工程大学计算机科学学院 西安 710048)

摘 要 作为一种常用的数据发布方法,直方图数据发布因其直观便捷的特点得到了广泛关注.直方图数据发布在带来方便的同时也面临着隐私泄露的风险.当前基于差分隐私的保护方法虽然提高了一定发布安全性,但仍然存在以下问题:(1)现有的差分隐私保护方法往往忽略直方图发布数据之间的关联特性;(2)同时,现有的方法缺乏有效评估直方图间接隐私泄露风险的方法;(3)现有的方法难以实现全面均衡的直方图隐私保护.本文针对上述问题,通过引入关联隐私泄露评估量化机制,设计了一种面向直方图数据发布的均衡差分隐私保护方法.首先,结合作用域的马尔科夫模型定义了直方图关联隐私;然后,基于隐私泄露损失因素,提出一种多指标决策的隐私泄露损失评估方法;最后,借鉴 Nash 博弈与 Stackberg 博弈思想,设计一种均衡差分隐私保护直方图发布方法.通过在不同数据集上的实验,验证了本文所提均衡隐私保护方法的有效性与鲁棒性,并证明了本文所提出的均衡差分隐私保护方法均衡地保护了直接与间接隐私泄露,且优于 AHP/GS 直方图隐私保护发布方法.

关键词 直方图发布;差分隐私;关联隐私;马尔科夫;Nash 均衡;博弈论

中图法分类号 TP311 **DOI 号** 10.11897/SP.J.1016.2020.01414

Balanced Correlation Differential Privacy Protection Method for Histogram Publishing

YANG Xu-Dong¹⁾ GAO Ling^{1),2)} WANG Hai¹⁾ GUO Hong-Bo¹⁾ ZHENG Jie¹⁾

¹⁾(Department of Information Science and Technology, Northwest University, Xi'an 710127)

²⁾(Department of Computer Science, Xi'an Polytechnic University, Xi'an 710048)

Abstract The histogram data as its feature of concise, clear and intuitive form of data expression is widely used in statistics, analysis, and other fields. The sanitary administrative organs use histogram data to issue epidemic warnings and the family-planning bureaucracy use histogram data to display and analyze population distribution. As a commonly used data publishing method, histogram data publishing has received extensive attention due to its intuitive and convenient features. Although histogram data publishing increases the efficiency of users' acquisition of data knowledge, privacy disclosure in published data is still the main factor hindering its development. To cope with this challenge, histogram data publish under privacy protection has attracted academic attention. Differential privacy as a strictly provable approach to privacy protection has gained more attention than other approaches. However, the existing differential privacy protection methods often ignore the association characteristics between the histogram published data and there are three main problems: (1) The effect of data association privacy in histogram sequence on privacy protection; (2) It is difficult to give full consideration to the direct and indirect privacy

收稿日期:2018-12-18;在线发布日期:2019-12-08. 本课题得到国家重点研发计划(2019YFC1521400)、国家自然科学基金(61672426, 61572401)资助. 杨旭东,博士研究生,中国计算机学会(CCF)会员,主要研究方向为隐私保护、边缘计算. E-mail: lfszyxd@qq.com. 高 岭(通信作者),博士,教授,中国计算机学会(CCF)会员,主要研究领域为计算机网络、边缘计算、隐私保护. E-mail: gl@nwu.edu.cn. 王 海,博士,副教授,中国计算机学会(CCF)会员,主要研究方向为隐私保护、边缘计算. 郭红波,博士,中国计算机学会(CCF)会员,主要研究方向为隐私保护、医学图像处理. 郑 杰,博士,中国计算机学会(CCF)会员,主要研究方向为异构网络、边缘计算、隐私保护.

of data; (3) The difficulty of equalizing the direct and indirect privacy of histogram data. Aiming at the protection privacy in publishing histogram data, we design a balanced differential privacy protection method for histogram data publishing by introducing the associated privacy leakage assessment and quantization mechanism. First, we define the histogram spatial correlation and privacy according to the scope-based Markov mode. The Markov transfer probability between spatial adjacent histograms is used to represent the correlation privacy between histograms, and it is shown that correlation privacy can cause the leakage of comprehensive privacy. Second, we propose a comprehensive quantitative evaluation method for multi-factor privacy leakage, which by combining various loss factors caused by privacy leakage. We determine the cognitive scope based on people's cognitive psychology and establish the evaluation index of indirect privacy leakage through local and global privacy leakage loss within the cognitive scope. Through the degree to which the index value belongs to the evaluation class, the single-index analysis is realized. Then, the multi-index analysis result of the weighted sum of the single index and high value is achieved. Through the comprehensive evaluation of the above results, the risk of leakage of indirect privacy is realized. Thirdly, based on the Nash game and Stackelberg game theory, a histogram distribution method for direct and indirect differential privacy balance protection based on game theory was designed. First, based on the above different leakage risk, the Nash game is used to realize the balanced distribution of the privacy budget. According to the direct privacy budget, noise is added to the histogram based on the lap mechanism. According to the indirect privacy budget, the histogram after adding noise is selected by the exponential mechanism for clustering publishing. Finally, Verify our method by running our method on two common datasets include the WakeTakere datasets and the Search log datasets in the following different situations: (1) By setting different parameters to verify the effectiveness of the method under different Settings; (2) By comparing with different algorithms, the superiority of our method is demonstrated. Based on the Analysis in above two different situations, the robustness and the applicability of the proposed method in this paper is proved.

Keywords histogram data publish; differential privacy; correlation privacy; Markov model; Nash equilibria game theory

1 引言

后摩尔时代的来临,造成数据爆发式的增长,伴随着“大数据时代”的来临,数据所蕴含的信息包括人们的生活轨迹、学习路径等,数据的发布也时刻影响着人们的生活.例如,电视台发布的近期疾病高爆发状况可以引起人们对于疾病卫生的注意;交通局发布的道路拥挤预警情况会驱使人们向着人群稀疏的地方行驶.这些数据信息的发布对于人们的生产生活有着至关重要的影响.

直方图数据作为一种数据发布的重要形式,以其准确直观的特点深受人们喜爱.但是,关于用户隐私的直方图数据发布会泄露人们的隐私信息,从而导致生活的困扰.例如,某医院发布的最近疾病

情况统计,其中癌症病人 10 例,20~30 岁有 2 人,30~40 岁有 2 人,40 岁以上有 6 人.当攻击者了解其中除 20~30 岁其他所有人的信息即可以推断出 20~30 岁中是否是自己认识的小张.这样的攻击会对被攻击者的隐私造成泄露,影响其生活.

为了应对直方图数据发布过程中伴随的隐私泄露风险,针对直方图数据发布的差分隐私保护研究受到了科研工作者的关注.直方图数据发布的隐私保护在提高数据发布可用性的同时也注重用户的隐私性.隐私保护数据发布的关键在于提供安全的发布方法的同时确保有足够的信息支撑数据分析.这就要求:(1)直方图发布的数据必须是经过噪声扰动后的数据,确保用户的个人信息安全;(2)为了保证数据是可分析的,发布的数据尽可能与真实数据接近,以具有较高可用性.

传统的数据发布隐私保护方法主要基于 k -匿名模型^[1-4],这类方法存在以下不足:(1)隐私保护效果受数据分布影响较大,稀疏数据的保护效果不够理想;(2)隐私保护方法需要获取攻击者的背景知识,获取难度比较大.基于差分隐私的数据发布方法因为具有如下的特点,很大程度上缓解了上述问题.一方面,差分隐私保护具有不受背景攻击影响的优点;另一方面,差分隐私添加噪声具有独立分布的特点.

现有差分隐私保护下的数据发布研究主要集中在直方图划分后添加噪声与添加噪声后进行分组两个主要研究方向.加噪声后划分的研究具有查询范围大,查询精度灵活的优点^[8-10];先分组后加噪声的研究具有响应查询度高的优点^[11,25].虽然上述研究一定程度上解决了差分隐私相关问题,但是现有差分隐私保护的数据发布方法仍然存在以下难题:

(1)如何保护直方图数据的间接关联隐私.现有的研究主要针对直方图数据直接信息进行保护^[6-10],忽略了攻击者可以通过直方图之间的关联性推断出其余直方图数据值,从而导致隐私泄露.

(2)如何根据直方图数据的发布特点进行隐私

泄露损失量化评估.现有的研究主要针对位置数据的时序关联隐私信息的转移概率来量化^[12-14],直方图数据发布缺乏有效隐私泄露保护.

(3)如何均衡保护直方图隐私是令研究人员头疼的难题.现有的研究仅考虑直接隐私忽略关联隐私,而关联推断攻击同样会导致隐私的泄露;或仅考虑关联隐私缺乏直接隐私也会造成隐私得不到有效保护^[17-18].

基于上述分析,本文针对直方图数据的关联隐私的保护问题,通过引入关联隐私泄露评估量化机制,设计了一种面向直方图数据发布的均衡差分隐私保护方法,并通过两个模拟数据集运行本文方法,验证了本文所提方法的有效性 & 鲁棒性.其主要贡献如图 1,具体如下:

(1)针对缺乏直方图关联隐私的关注,本文提出了基于直方图认知及关联隐私相关概念;

(2)结合直方图认知相关概念,构建关联隐私泄露评价指标,并据此实现一种多指标关联隐私风险泄露量化方法;

(3)基于隐私泄露量化结果,通过博弈思想构建一种全面均衡的差分隐私保护方法,实现隐私保护的直方图数据发布.

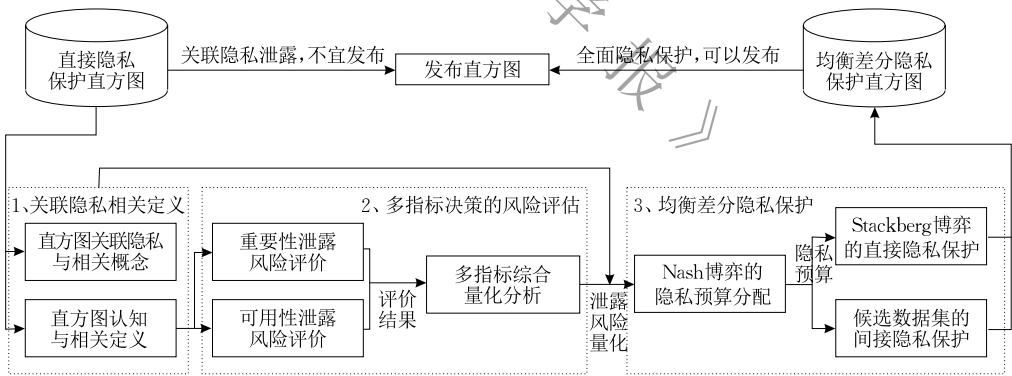


图 1 本文主要工作

本文第 2 节给出相关研究工作;第 3 节主要介绍差分隐私及直方图认知相关理论,认知相关理论是第 4 节量化指标计算的基础,差分隐私及相关理论为第 5 节实现隐私保护的基础;第 4 节介绍多指标综合隐私泄露量化评估方法,是第 5 节直方图均衡保护隐私预算分配的主要依据;第 5 节给出一种均衡差分隐私直方图保护方法,基于第 4 节隐私泄露风险量化结果,建立 Nash 博弈模型对于直接隐私与间接隐私进行保护;第 6 节给出实验设计及结论;第 7 节给出总结与展望.

2 相关工作

差分隐私大多集中在差分隐私保护下的数据发布、数据挖掘、位置隐私保护等.因其具有高度隐私保护能力和依赖严谨的数学推导已经得到国内外众多学者的关注.数据发布的差分隐私保护因为其在数据统计、资源管理等方面的应用价值受到了广泛研究.现有的研究主要基于 Dwork^[5]提出的拉普拉斯机制与指数机制,可以分为直接

隐私差分保护与间接隐私差分保护。

在数据发布的直接差分隐私保护方面,主要分为基于层次树划分的差分隐私保护与基于聚类分组的差分隐私保护。基于层次树划分的差分隐私保护研究包括:Boost2^[6]利用 M-ary 树对等宽直方图进行重新划分组织。Xiao 等人^[7]设计了基于小波变换的 Privelet 的差分隐私直方图发布保护方法。基于层次树划分的优点在于可以精确地响应长范围的查询,而它们的不足在于查询的敏感性较高。为了降低查询敏感性,提出了基于聚类划分的差分隐私保护。基于聚类划分的差分隐私保护研究主要集中在直方图先分组后添加噪声与先添加噪声后进行分组两个方向。其中:(1)先加噪后分组包括:LP^[8]采用等宽直方图进行聚类划分,一定程度上对隐私信息进行了保护,但是容易造成噪声过大的问题。Boost1^[9]考虑了结合一致性约束条件与最小二乘法的优化方法实现对直方图桶划分的优化,进而保证了直方图发布后的可用性;Xu 等人研究了直方图添加噪声后直方图分组的方法,命名为 NoiseFirst^[10],该方法仅支持直方图的等宽分组,缺乏启发式规则;(2)先分组后加噪包括:Zhang 等人^[11]研究了一种全局化聚类方法 AHP 并取得了不错的效果,但该方法忽略了直方图数据的有序性,导致可用性降低。Zhang 等人^[25]结合蒙特卡洛与指数机制提出一种贪心聚类的直方图发布方法,在一定程度上解决了有序性的问题。

在数据发布的间接差分隐私保护的数据发布方面,研究者主要集中在位置时空关联隐私信息方面。Xiao 等人^[12]基于高斯关联模型提出一种移动位置关联关系的移动群体感知的差分隐私保护模型,首先提出了隐私保护的相关定义,并分别从攻击者与保护者的角度提出隐私评估模型及相应保护模型以实现隐私保护,并且验证了所提出的方法具有不错的保护效果。Ou 等人^[13]对多用户移动位置之间的关联性隐私保护进行了研究,提出了基于隐马尔科夫模型的私有候选集的位置信息发布机制,实验结果表明该方法具有较好的差分隐私保护效果。Cao^[18]结合差分隐私保护技术,对时空序列中的具有关联关系数据的发布进行了研究,构造了一种基于马尔科夫转移模型的关联隐私模型及其量化方法,并提供了多种隐私保护方案。Hai 等人^[19]提出一种识别第三方服务器的隐私窃取攻击的方法并提供了具有时空关联的隐私保护方案。

从上述分析可以看出,现有的直方图隐私保护单纯依赖于直接隐私保护方法,现有间接隐私保护研究主要集中在位置信息方面;同时,并未应用于直方图数据发布领域。因此,缺乏一种有效的直方图数据的关联隐私保护方法。本研究提出了直方图数据的关联隐私的定义及其量化评估方法,同时设计了一种均衡直接与间接差分隐私保护的直方图发布方法。

3 理论基础与相关定义

定义 1. 直方图序列。

设 $H = \{H_1, H_2, \dots, H_n\}$ 代表直方图原始桶计数序列, H_i 为第 i 个桶计数, $H' = \{H'_1, \dots, H'_i, \dots, H'_n\}$ 代表相邻直方图桶计数序列,且 H_i 与 H'_i 桶计数不相同。

定义 2. ϵ -差分隐私^[15]。

给定直方图序列 H 发布方法 A , $\Pr(A)$ 为 A 的输出,若方法 A 在 H 与 H' 上任意输出结果 $\tilde{H}$ 满足以下不等式,则 A 满足 ϵ -差分隐私

$$\Pr[A(H) = \tilde{H}] \leq \exp(\epsilon) \times \Pr[A(H') = \tilde{H}] \quad (1)$$

对于 ϵ ,其值越小则发布的算法的隐私保护能力越强。差分隐私的实现方式包括噪声机制与指数机制,其都依赖方法 A 的全局敏感度,接下来介绍全局敏感度。

定义 3. 全局敏感度^[16]。

对于 H 与 H' 上的查询函数 f ,其全局敏感度为

$$\Delta f = \max \|f(H) - f(H')\| \quad (2)$$

其中, $\|\cdot\|$ 代表一阶范数距离,差分隐私的主要实现方法包括拉普拉斯噪声机制与指数机制,二者的实现都依赖于全局敏感度。

定义 4. 拉普拉斯机制^[16]。

对于 H 上任意查询函数 f ,算法输出结果为数值型函数,若 f 满足

$$A(H) = f(H) + \text{lap}\left(\frac{\Delta f}{\epsilon}\right) \quad (3)$$

则隐私保护算法 A 满足 ϵ 差分隐私,其中 $\text{lap}\left(\frac{\Delta f}{\epsilon}\right)$ 代表相互独立的拉普拉斯噪声变量, Δf 为查询全局敏感度。当敏感度越大时,所需噪音越多,一般为了限制噪声大小,通过划分的方法降低敏感度,从而降低噪声大小。

定义 5. 指数机制^[16]。

对于 H 上隐私保护算法 A 的输出函数 f , 算法输出结果为非数值型数据, 若 f 满足

$$A(H) = f[H_i \in H] \propto \exp \frac{\epsilon u(H, H_i)}{2\Delta u} \quad (4)$$

则 A 满足 ϵ 差分隐私. 其中, Δu 为打分函数 u 的全局敏感度.

上述指数机制与拉普拉斯机制作为实现差分隐私的两种重要手段受到了广泛应用. 可以看出, 这两种手段的使用很大程度受到隐私预算的影响.

定义 6. 组合定理^[16].

假设直方图序列 H 中的每个隐私算法 $A_i, A_i \in A$ 满足 ϵ_i 的差分隐私, 则其组合序列 $A = \sum_{A_i \in A} A_i$ 满足 $\epsilon = \sum_i \epsilon_i$ 的差分隐私.

直方图数据发布是为了提高人们对于数据蕴含知识的有效认知, 因此定义了直方图认知能力与认知作用域相关概念.

定义 7. 直方图认知能力.

直方图认知能力 (Cognitive Ability) 代表接收直方图后人们对于其信息的获取程度, 在一定程度上影响关联隐私保护程度. 结合统计理论与认知理论^[19]可以看出, 直方图认知依赖于认知作用域与用户背景知识. 假设用户对于发布背景知识 B 固定, 则认知的作用范围过大会造成认知需要背景知识越多, 只有降低认知的范围才能达到同样认知能力. 设直方图 H_i 关联关系作用域 R_i , 则认知能力 Cl_i 等于作用范围与用户背景的反比关系, 即

$$Cl_{ability, i} = \frac{B}{R_i} \quad (5)$$

其中, 由于直方图的认知基于序列才有用, 因此 B 可以利用用户知识覆盖认知直方图 H_i 邻域的覆盖率进行计算, 同时为了计算方便将其分为 3 个等级

$$B = \begin{cases} 0.1, & 0.1 \leq \frac{H_{ui}}{H'_i} < 0.3 \\ 0.45, & 0.3 \leq \frac{H_{ui}}{H'_i} < 0.6 \\ 0.75, & 0.6 \leq \frac{H_{ui}}{H'_i} < 0.9 \end{cases} \quad (6)$$

其中, H_{ui} 代表用户了解直方图 H'_i 邻域内直方图个数. 从公式中可以看出, 作用域范围越大, 认知有效性越低, 其中 R_i 初始化为 5.

定义 8. 认知作用域

直方图认知能力严重依赖于认知作用域. 直方

图认知作用域代表直方图认知过程中认知范围的上限, 即超出作用域后, 认知会有大幅度衰减. 初始直方图没有认知能力, 进而通过式 (5) 无法确定作用范围. 直方图相邻的直方图会对当前直方图作用范围进行影响, 因此结合 H_i 正相邻直方图认知作用域构建, 即在 H_i 前面的直方图 H_{i-1} 构建 H_i 的作用域, 主要通过基于认知范围与质量构建认知影响因子实现. 假设已知 H_{i-1} 的作用域 R_{i-1} 则

$$R_i = (H_i - H_{i-1}) \times \left[1 - \frac{CI_{ability, i-1}}{CI_{ability}(H)} \right] \times R_{i-1} \quad (7)$$

其中, $\overline{CI_{ability}(H)}$ 代表对于直方图发布序列的历史平均认知能力, $CI_{ability, i-1}$ 代表对于 H_{i-1} 的认知能力.

定义 9. 关联隐私.

假设发布序列 H 中单独桶计数 H_i 为马尔科夫状态变量, 则空间相邻直方图 H_i 与 H_{i+1} 的关联关系可以使用马尔科夫转移概率来表示, 并且具有马尔科夫特性, 即

$$\Pr(H_{i+1} | H_i, H_2, \dots, H_i, H_n) = \Pr(H_{i+1} | H_i) \quad (8)$$

则推断攻击利用直方图之间转移概率 $\Pr(H_i | H_{i+1})$ 推断出直方图敏感信息 H^* , 因此直方图序列中转移概率可以看做直方图序列 H 中的关联隐私, 相对于统计数隐私 (直接隐私) 称为间接隐私 (Indirect Correlation Privacy).

定义 10. 推断攻击.

当攻击者具有某些直方图的背景知识 H_B , 以及转移概率 $\Pr(H_B | H^*)$ 便推断出经过保护的直方图序列中敏感直方图 H^* 信息的过程, 称为直方图序列 H 推断攻击. 可以表示为

$$H^* = H_B \times \Pr(H^* | H_B) \quad (9)$$

其中, H_B 代表攻击者的背景知识, $\Pr(H^* | H_b)$ 代表获取的转移概率, $\times$ 为关联操作, 即将背景知识与转移概率结合起来推断出敏感数据的过程.

同时根据攻击目的可以分为通过关联隐私获取敏感直方图 H^* 的直接攻击与获取关联隐私 $\Pr(H_B | H^*)$ 的间接攻击. 可以看出两种攻击都可以获取关联隐私中的直方图信息, 因此本文所提的保护都基于关联隐私的直方图进行聚类保护实现.

4 多指标综合决策隐私泄露评估

结合直方图认知相关知识, 本节介绍了关联隐私泄露风险评价指标及综合量化方法, 如图 2.

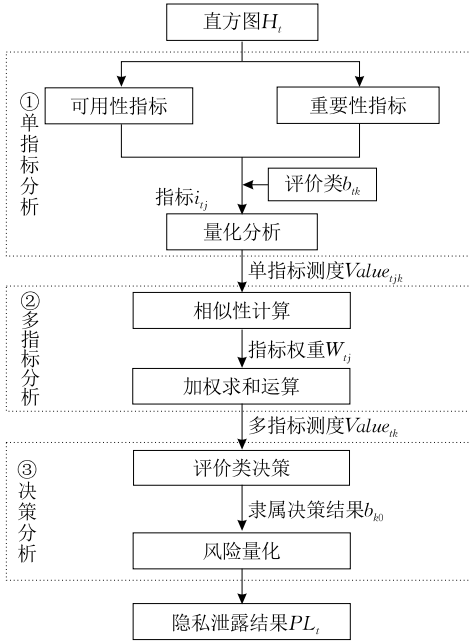


图2 多指标隐私量化综合评估过程

4.1 直方图关联隐私泄露评价指标

隐私泄露风险越大则应该采取更强大的保护手段进行保护. 基于上述思路, 我们利用隐私保护支出进行隐私泄露风险评价. 根据上节分析, 对于发布序列中的任意关联隐私 $\Pr(H_i | H_{i+1})$, 其中 $H_i, H_{i+1} \in H$ 隐藏 H_i 可以使得攻击者不能获取准确的 H_i 与 H_{i+1} 的关联隐私. 在此保护下, 隐私保护支出等价于 H_i 缺失发布造成的直方图损失, 即直方图的数值的统计以及基于直方图的认知作用的损失. 通过结合重要性与可用性损失结果对其进行量化评估, 包括缺失 H_i 对于发布直方图序列缺失数据导致的统计重要度、认知重要度与认知可用性及关联可用性损失.

统计重要度 (Statistical Important Factor) 损失主要代表直方图隐藏 H_i 发布直方图序列, 导致的直方图统计质量损失. 通过分布熵的局部与全局变化率结果的乘积进行衡量, 即通过不同桶计数在直方图中出现情况进行衡量. 同时考虑在不同认知范围情况, 结合 H_i 缺失造成局部统计重要度 (即在直方图作用域内的重要程度) 与全局统计重要度 (即在直方图发布序列内的重要程度) 的损失进行评价, 得到统计重要度损失评价如下:

$$SIF = \frac{G_i}{G'_i} \cdot \frac{P_i}{P'_i} \quad (10)$$

其中, P_i 代表在包含 H_i 认知域内的分布情况, G_i 代表在包含 H_i 认知域内缺乏直方图分布熵, G'_i 代表在全局序列的分布情况, G_i 代表在全局序列的认知域

内缺乏直方图分布熵, $0 < SIF \leq 1$, 其中,

$$G_i = \sum_{p_i \in G} \frac{1}{p_i} \log_2 \frac{1}{p_i} \quad (11)$$

$$P_i = \sum_{p_i \in G} \frac{1}{p_i} \log_2 \frac{1}{p_i} \quad (12)$$

可以从式(10)看出, 直方图缺失后的分布差异越大, 则其重要性越大.

认知重要度 (Cognitive Importance Factor) 损失主要代表隐藏 H_i 发布直方图序列, 导致用户对于直方图数据认知水平的影响. 通过包含直方图认知作用域与认知偏差因子的认知水平变化来表示, 认知重要度损失可以表示为

$$CIF = \left| \frac{R_i}{R'_i} \right| \cdot \left| \frac{CI_i}{\phi + \overline{CI_i}} \right| \quad (13)$$

其中, CI_i 代表缺乏 H_i 的认知能力, R'_i 代表缺乏直方图 H_i 后的认知区域, ϕ 为均衡因子且

$$\phi = \begin{cases} 1, & \text{若 } \overline{CI} < 0.5 \\ 0, & \text{其他} \end{cases} \quad (14)$$

通过上述分析可以看出, 认知重要程度变化不仅考虑对于自身认知能力的影响, 还考虑对于其他桶的影响, 这符合直方图序列的特点. 同时, 对于直方图缺失对于认知的影响力不区分正负程度, 更能体现直方图对于认知的重要性.

直方图发布的目的是使得用户对于直方所代表的信息能够有效认知. 因此, 我们结合直方图发布后的认知质量变化进行可用性评价, 体现在认知可用性, 其是决定用户是否有效接受所提供直方图信息的关键性因素.

认知可用性 (Cognitive Availability Factor) 代表隐藏 H_i 发布直方图序列后对于人们的认知质量 (Cognitive Quality) 的影响, 认知质量通过认知能力偏差 (Cognitive Bias) 进行评估, 则对于 H_i 的认知有用性损失具体如下:

$$CAF_i = CQ'_i - CQ_i \quad (15)$$

其中, CQ'_i 代表包含直方图 H_i 的认知质量, CQ_i 代表缺乏 H_i 序列的认知质量, 使用平均认知能力偏差表示, 具体如下:

$$CQ_i = \frac{\sum_{i' \in R_i} CI_{i'}}{|H_i|} + \sum_{i' \in R_i} \frac{CI_{i'} - \overline{CI_{i'}}}{|H_{i'}|} \quad (16)$$

$|H_{i'}|$ 代表当前邻域内的直方图数量, $CI_{i'}$ 代表有过记录认知能力, $\overline{CI_{i'}}$ 代表认知记录中直方图邻域内

的直方图数量. 这样设定, 通过平均认知偏差可以反映当前直方图与已发布直方图认知能力的影响的差异. 当包含平均认知与缺乏直方图 H_i 认知偏差越大, 说明 H_i 对于认知质量影响越重要.

在认知过程中, 直方图之间的关联认知也十分重要, 主要代表通过直方图之间的潜在关联关系可以推断出相关知识. 例如: 在具有递增关系直方图序列中, 可以通过近邻直方图之间的差值推断出直方图蕴含的知识的增幅等信息, 提高认知的有效性. 关联可用性 (Background-Correlation Availability Factor) 损失评价指标, 代表 H_i 缺失直方图发布对于直方图之间关联认知情况的影响进行可用性损失评价:

$$BAF_i = GCL_i + PCL_i \quad (17)$$

其中, GCL_i 代表全局关联损失, PCL_i 代表局部关联损失. 全局关联损失 (Global Correlation Loss) 代表直方图 H_i 缺失后对于用户获取全局关联认知知识的影响, 缺乏直方图会导致关联推理距离增大, 关联关系降低, 从而使得推理精度降低. 使用加权路径变化评价. 其中, 关联关系 (转移概率) 作为权重 $w_{i,t+1}$, 则直方图 H_i 影响的区域作为节点的加权路径长度表示

$$GCL_i = \sum_{H_j \in H} w_{i,t+1} \times H_j - \sum_{H'_j \in H'} w_{i,t+1} \times H'_j \quad (18)$$

其中, H 与 H' 分别代表包含 H_i 与缺乏 H_i 的全部直方图序列.

局部关联损失 (Part Correlation Loss) 代表直方图 H_i 缺失后对于用户获取局部认知关联知识的影响, 分析如下:

$$PCL_i = \sum_{H_{i'} \in R_i} w_{i,t+1} \times H_{i'} - \sum_{H_{i'} \in R'_i} w_{i,t+1} \times H_{i'} \quad (19)$$

其中, R_i 代表包含 H_i 的认知作用域, R'_i 代表不包含 H_i 的认知作用域.

4.2 隐私泄露损失多指标综合量化方法

对于直方图序列 H , 其中任一直方图 H_i 关联隐私泄露风险对应 4 个评价指标, 即 $I = \{I_1, I_2, I_3, I_4\}$, 对应上节 4 个指标 SIF, CIF, BAF, CAF ; 同时假设每个评价指标对应 k 个评价类, $E = \{e_1, e_2, e_3, e_4, \dots, e_k\}$ 构建对于直方图的评价矩阵为

$$\begin{matrix} & I_1 & I_2 & I_3 & I_4 \\ \begin{matrix} e_1 \\ e_2 \\ e_3 \\ \vdots \\ e_k \end{matrix} & \begin{vmatrix} b_{11} & b_{12} & b_{13} & b_{14} \\ b_{21} & b_{22} & b_{23} & b_{24} \\ b_{31} & b_{32} & b_{33} & b_{34} \\ \vdots & \vdots & \vdots & \vdots \\ b_{k1} & b_{k1} & b_{k1} & b_{k1} \end{vmatrix} \end{matrix},$$

其中 $b_{k1}, b_{k2}, b_{k3}, b_{k4}$ 代表样本属于第 j ($j=1, 2, 3, 4$) 个指标下第 k 个评价类的程度, 表示评价分析的量化结果, 满足 $b_{11} < b_{21} < \dots < b_{k1}$.

这时要判断 H_i 的隐私泄露风险, 可以根据其对应指标 I 属于评价类的程度判断, 具体包括根据单指标分析、多指标分析与决策系统 3 个过程.

(1) 单指标分析. 对于直方图 H_i 的第 j 个风险指标 i_{ij} , $0 \leq j \leq 4, 1 \leq k \leq K$ 表示其属于第 k 个评价类的程度, 我们称指标测度, 代表评价结果, 表示为 $Value_{ijk}$.

假设对于直方图中任意 H_i 样本第 j 个评价指标 $b_{i1} \leq b_{i2} \leq \dots \leq b_{ik}$, 则其单个指标分析结果可以计算如下:

当 $i_{ij} \leq b_{j1}$ 时,

$$Value_{ij1} = 1, Value_{ij2} = Value_{ij3} \dots Value_{ijK} = 0;$$

当 $i_{ij} > b_{iK}$ 时,

$$Value_{iK} = 1, Value_{ij1} = Value_{ij2} \dots Value_{ij3} = 0;$$

当 $b_{il} \leq i_{ij} \leq b_{il+1}$ 时,

$$Value_{ijl} = \frac{|i_{ij} - b_{il}|}{|b_{il} - b_{il+1}|}, Value_{ijl+1} = \frac{|i_{ij} - b_{il+1}|}{|b_{il} - b_{il+1}|},$$

进而得到 H_i 的损失评价矩阵 S_i

$$S = \{S_{ij} | S_{ij} = Value_{ijk}, 0 \leq Value_{ijk} < 1\} \quad (20)$$

(2) 多指标分析. 通过损失评价矩阵得到 H_i 在指标 I_j 的评价结果, 则多指标分析结果可以通过加权的方法求得

$$Value_{ik} = \sum_{j=1}^m W_{ij} \times Value_{ijk} \quad (21)$$

同时, 当难以确定多指标权重时, 可以通过计算平均测度获得

$$Value_{ik} = \frac{1}{m} \sum_{j=1}^m Value_{ijk} \quad (22)$$

为了客观赋予指标权重, 根据已经获得的单指标测度 Mea_{ijk} 与多指标评估 Mea_{ik} 的相似性确定指标的重要程度, 即相似程度越大代表指标的重要性越高, 使用皮尔逊相似系数评价相似性, 计算如下:

$$r_{ij} =$$

$$\frac{\sum_{k=1}^m (Value_{ijk} - \overline{Value_{ijk}}) \sum_{k=1}^m (Value_{ik} - \overline{Value_{ik}})}{\sqrt{\sum_{k=1}^m (Value_{ijk} - \overline{Value_{ijk}})^2} \sqrt{\sum_{k=1}^m (Value_{ik} - \overline{Value_{ik}})^2}} \quad (23)$$

权重计算公式为

$$w_{ij} = \frac{r_{ij}}{\sum_{i=1}^n r_{ij}} \quad (24)$$

最终使用式(17)获得多指标分析结果.

(3) 决策分析. 决策代表选择直方图多指标分析结果属于哪一个评价类的过程^[21]. 假设 H_t 的综合评价类的隶属程度 $b'_1 > \dots > b'_k$, 若 $Ame_{tk0} = \max(b'_k \times Value_{tk})$, 则认为 H_t 属于综合评价 k_0 类, $1 \leq k'_0 \leq K'$ 据此隐私泄露量化结果为

$$PL_t = \frac{k_0}{Ame_{tk0} + \alpha},$$

其中, α 防止分母为 0 的情况, 为一个很小的正数.

多因素综合决策隐私泄露风险量化算法 QPLMDM(Quantization Privacy Leak by Multifactor Decision Making)如下:

算法 1. QPLMDM 算法.

输入: 直方图原始序列: H

直方图认知范围序列: R

直方图认知能力序列: CI

输出: 直方图间接隐私泄露风险:

$$PL = \{PL_1, PL_2, \dots, PL_N\}$$

1. FOR EACH H_t IN H DO
2. //利用 R_t, CI_t 计算 H_t 可用性损失指标 SIF_t, CIF_t
3. 利用式(10)、(13)计算 SIF_t, CIF_t
4. //利用 R_t, CI_t 计算 H_t 重要性损失评价
5. 利用式(15)、(17)计算 BAF_t, CAF_t
6. //构建 H_t 的损失标准矩阵
7. $I = \text{getMatrix}(SIF_t, CIF_t, BAF_t, CAF_t)$
8. $B_t = \text{getMatrix}(I)$
9. //构建 H_t 的单指标评价测度
10. FOR EACH i_j IN I DO
11. IF $i_j \leq b_{j1}$
12. $Value_{j1} = 1, Value_{j2} = Value_{j3} \dots Value_{jk} = 0$
13. ELSE IF $i_j > b_{jk}$
14. $Value_{jk} = 1, Value_{j1} = Value_{j2} \dots Value_{jk-1} = 0$
15. ELSE
16. $Value_{jk} = \frac{|i_j - b_{k+1}|}{|b_k - b_{k+1}|}, 0 < k < m$
17. //构建 H_t 的多指标综合决策
18. (1) 计算多指标平均值
19. $Value_{tk} = \frac{1}{m} \sum_{j=1}^m Value_{tjk}$
20. (2) 计算多指标向量与单指标向量相似度
21. $r_{ij} =$

$$\frac{\sum_{k=1}^m (Value_{tjk} - \overline{Value_{tjk}}) \sum_{k=1}^m (Value_{tk} - \overline{Value_{tk}})}{\sqrt{\sum_{k=1}^m (Value_{tjk} - \overline{Value_{tjk}})^2} \sqrt{\sum_{k=1}^m (Value_{tk} - \overline{Value_{tk}})^2}}$$

22. (3) 获得指标权重 $w_{ij} = \frac{r_{ij}}{\sum_{i=1}^n r_{ij}}$

$$23. (4) \text{多指标分析结果 } Value_{tk} = \sum_{j=1}^m w_{ij} \times Value_{tjk}$$

24. END FOR

25. //综合决策指标泄露风险

$$26. PL_t = \frac{k_0}{\max(b'_k \times Value_{tk}) + \alpha}$$

27. END FOR

RETURN $PL = \{PL_1, PL_2, \dots, PL_N\}$

5 均衡差分隐私保护直方图发布方法

传统的直方图桶计数隐私保护可以保障直方图的绝对隐私不被泄漏, 然而对于直方图之间关联隐私关注程度不够, 容易导致隐私泄露. 因此, 本文结合上节隐私风险量化结果, 通过博弈思想获得隐私预算分配, 实现直接隐私与间接隐私的保护.

为了实现均衡差分隐私保护, 同时保障数据可用性, 本文结合指数机制与噪声机制的特点构建隐私保护直方图发布方法. 首先, 基于直接隐私泄露与本文所提间接隐私泄露量化结果建立隐私保护预算的 Nash 博弈模型实现对直接隐私泄露与间接隐私泄露的保护预算进行均衡分配(为了均衡保护); 基于上述直接隐私保护预算分配结果, 建立推断攻击与隐私保护策略的斯坦克尔伯格博弈模型, 通过线性规划求解得到最优的直方图保护策略; 最后, 基于上述间接隐私预算分配结果, 利用间接隐私保护预算, 结合指数机制的动态桶划分机制实现间接隐私保护直方图发布方法(间接隐私保护).

5.1 隐私预算分配策略

结合差分隐私的思想, 为了实现保护直接与间接隐私全面保护, 如何合理的分配隐私预算法是目前需要解决的关键问题. 由于隐私预算的总和是一个固定值, 如果提高对于直接隐私的保护力度势必会造成间接隐私的损失, 反之亦然. 在此过程中, 直接与间接隐私保护对于隐私预算的需求存在一种竞争关系, 因为是有有限次的竞争^[31-33], 同时在预算总和的限制下, 必定可以求得均衡最优解. 我们将这种关系看作一种博弈, 同时可以将其 Nash 均衡解作为结果, 实现约束下的均衡保护.

一般地, 隐私泄露风险大的直方图对于隐私保护强度的需求相比于其他直方图要高一些, 风险小的必定要相对低一些. 由于隐私预算与保护强度成正比, 因此, 对于隐私泄露风险与隐私预算也呈正比的关系. 基于上述分析, 我们将直方图隐私泄露量化结果, 包括基于信息熵的直接隐私泄露风险量化结

果 $IPL^{[6]}$ 与本文所提间接隐私风险量化结果 PL 看作 Nash 博弈中决策主体, 其相应权重 v 看作主体决策策略, 构建博弈模型. 并且博弈模型的 Nash 均衡解即为隐私分配结果. 可以看出, 通过博弈发挥不同决策的优势, 又可以通过增加 Nash 均衡约束提高对于竞争策略的控制, 即达到均衡直接与间接隐私的目的, 最终可以得到均衡最优的策略, 即最优的权重. 将权重作为隐私分配因子, 进行分配, 实现直接和间接的均衡保护, 即

$$\begin{aligned} \min \quad & \sum_{j=1}^2 (\Omega_j - \Omega_j^*)^2 v_j^2 \\ \text{s. t.} \quad & \sum_{j=1}^2 v_j = 1 \\ & 0 \leq v_j \leq 1 \end{aligned} \quad (25)$$

其中, Ω_j 指代 H_i 中不同类型隐私量化结果 ($i=1, 2$), 分别代表直接隐私泄露风险量化结果与间接隐私泄露风险量化结果, Ω_j^* 代表决策主体中的理想决策点, 即最优值. 此时的最优值为博弈均衡的理想点^[26-27].

为了求解模型中的最优权重, 结合拉格朗日方法建立如下等式

$$L = \sum_{j=1}^2 (\Omega_j - \Omega_j^*)^2 v_j^2 + \left(\sum_{j=1}^2 v_j - 1 \right) \quad (26)$$

对上式求偏导, 并使 $\frac{\partial L}{\partial v} = 0$, 求得最优权重.

为了避免竞争过度, 导致博弈最优解求解困难, 以 β 作为约束计算出各决策主体的最优均衡权重. 完全竞争模型中计算出的权重是各决策单元的完全优势最优权重, 获取针对不同决策最优 Nash 均衡, 同时通过递减 Nash 约束参数. 其中每次迭代的约束参数可以表示为

$$a_i = \sum_{j=1}^2 (\Omega_j - \Omega_j^*)^2 v_j^2 \quad (27)$$

基于上述分析, 最优 Nash 约束参数为

$$\beta = \min(a_i) \quad (28)$$

最后, 构建 Nash 均衡约束因子不等式, 建立约束模型

$$\begin{aligned} \min \quad & a = \sum_{j=1}^2 (\Omega_j - \Omega_j^*)^2 v_j^2 \\ \text{s. t.} \quad & \sum_{j=1}^2 (\Omega_j - \Omega_j^*)^2 v_j^2 \leq \beta \\ & \sum_{j=1}^n v_j = 1 \\ & 0 < v_j \leq 1 \end{aligned} \quad (29)$$

通过上述过程不断重复, 在任一决策单元优势最大

化的基础上, 其它各单元在约束参数不断迭代递减的同时增大自身的收益, 最终实现了所有单元的竞争最优的 Nash 均衡, 进而得到最优的权重, 即均衡预算分配因子 v_j , 进而求得隐私预算分配 ϵ_1, ϵ_2 .

5.2 直接隐私保护策略

基于上述分配的隐私预算 ϵ_1 , 进行直方图直接隐私保护. 针对直方图序列中任一直方图 H_i , 考虑直接隐私攻击与保护的数据发布模式可以看作 3 个主要过程: (1) 收集原始数据, 获得其分布概率 $\pi(H_i)$; (2) 保护直方图数据, 在直方图数据上执行扰动策略 $P\{\tilde{H}_i | H_i\}$; (3) 通过推断攻击获取原始数据 $Q(\hat{H}_i | \tilde{H}_i)$, 表示如下:

$$\pi(H_i) \rightarrow H_i \rightarrow P\{\tilde{H}_i | H_i\} \rightarrow \tilde{H}_i \rightarrow Q(\hat{H}_i | \tilde{H}_i) \rightarrow \hat{H}_i,$$

其中, $\tilde{H}_i$ 代表扰动后的直方图, $\hat{H}_i$ 代表攻击者重建直方图.

最大限度的直接隐私保护机制是通过加噪机制模糊数据的真实计数实现的, 其在发布模型中主要体现在观察值和敏感数据之间的相关性最小化. 然而, 越强大的保护机制对于数据的直接可用性损失的越大, 以至于对于用户发布数据的最初目的产生负面影响. 隐私性与可用性的平衡一直是隐私保护领域的关键挑战. 为了解决这个问题, 我们将满足隐私约束并且可用性最大的保护策略作为优先考虑.

根据上述分析, 直接隐私保护通过在隐私数据添加不同的值实现扰动, 这里通过借鉴旁路攻击的思想^[36], 将保护机制看作用户与保护者之间的旁路信息, 利用旁路通道信息量化方法进行评保护程度, 表示为 $P\{\tilde{H}_i | H_i\}$.

当在执行保护策略后, 必然会对可用性造成一定损坏. 直接可用性损失可以表示为隐私保护措施后(这里指代添加噪声)对于数据的信息表达的衰减, 使用直接可用性距离函数 $c(H_i, \tilde{H}_i)$ 表达, 即

$$c(H_i, \tilde{H}_i) = |\tilde{H}_i - H_i|, \quad \tilde{H}_i \in \tilde{H}, H_i \in H \quad (30)$$

据此, 则执行保护策略 $P\{\tilde{H}_i | H_i\}$ 后, 直方图数据的可用性损失期望可以表示为

$$\sum_{H_i \in H} \pi(H_i) \sum_{\tilde{H}_i \in \tilde{H}} P(\tilde{H}_i | H_i) \cdot c(H_i, \tilde{H}_i) \quad (31)$$

最大化保护效果可以看作最小化可用性损失, 则最优保护策略的求解可以表示为

$$P^*(\tilde{H} | H) =$$

$$\arg \min \sum_{H_i \in H} \pi(H_i) \sum_{\tilde{H}_i \in \tilde{H}} P(H_i | \tilde{H}_i) \cdot c(H_i, \tilde{H}_i) \quad (32)$$

同时, 为了实现差分隐私的保护, 根据差分隐私的定义变形子式, 将输出算法改变为隐私保护后的

查询函数构建隐私约束实现。根据差分隐私的经典定义,对于给定根据保护预算 ϵ 的前提下,采用隐私保护策略 $P(\tilde{H}_i|H)$ 后,对于任意 H_i 属于 H 满足如下条件,即满足差分隐私:

$$P(\tilde{H}_i|H_i) \leq e^\epsilon \cdot P(\tilde{H}_i|H'_i) \quad (33)$$

然而这种经典定义存在一个显著问题,过度依赖不同数据集的差异性,导致隐私保护效果不佳。例如对于 0,0.01 与 2000,2000.01 采用相同程度的隐私保护方法是明显不正确的。应综合考虑数值本身的差异程度进行加噪才能实现更有效的隐私保护。因此,学者提出了一种扩展差分隐私定义^[29-30]。

$$P(\tilde{H}_i|H) \leq e^\epsilon \cdot d_1(H, H') \cdot P(\tilde{H}_i|H') \quad (34)$$

其中, $d_1(H, H')$ 代表原始数据 H 与相似数据 H' 之间的距离。文献[29]中指出,此隐私保护模型的噪声敏感度仅依赖于两者之间的数值距离,即两个秘密之间的距离 d_1 越小,它们的混淆概率分布函数(执行差分隐私保护)就越难以区分,同时更能满足不同数据的隐私需求。因此,这种隐私模型又称为 d_1 敏感度的差分隐私保护模型。同样地,扩展隐私模型实现差分隐私的基本方法也是通过添加拉普拉斯噪声实现,噪声大小依赖于全局敏感度。为了降低敏感度,采用距离 d_1 敏感度进行衡量,其中 d_1 距离度量方法为标准欧式距离^[28-29],即

$$d_1(H_i, H'_i) = \frac{d_E(H_i, H'_i)}{\sum_{H_i \in H, H'_i \in H'} d_E(H_i, H'_i)},$$

其中, $d_E(H_i, H'_i)$ 代表欧式距离的度量。通过上述公式,可以看出 $0 \leq d_1 \leq 1$, 则上述隐私保护模型也符合 ϵ 的差分隐私。根据全局敏感度的定义,在此度量标准下同样满足数据高可用的需求,即改变任一直方图桶计数,其查询输出改变最大不超过 1。例如对于计数范围 0~5000 的直方图进行查询,采用传统欧式距离的查询敏感度为 5000,添加的噪声容易存在过大的问题;采用本文所提标准欧式距离,其敏感度为 1,添加噪声符合差分隐私的要求并具有更好的隐私保护效果。

结合上述最大化保护效果的目标函数与差分隐私约束函数,我们获得了基于可用性与隐私性均衡的保护策略。然而,当前保护策略仅可以针对固定背景攻击进行保护,难以保障对于任意攻击下的隐私保护效果都是最好的。同时,攻击者通过攻击希望获取的数据与隐私数据误差尽量小,保护者希望在可用性保障的前提下两者之间的误差最大。在上述过程中,隐私攻击通过对发布数据重建隐私数据实现,

同样建模为旁路攻击与信息窃取模型进行量化,表示 $Q(\hat{H}_i|\tilde{H}_i)$ 。隐私攻击效果可以通过窃取敏感信息的完整程度进行表达,本文利用重建数据 $\tilde{H}$ 与原始数据 H 的误差距离 d_2 进行衡量。因此,隐私攻击攻击期望可以表示为

$$\sum_{H_i \in H} \pi(H_i) \sum_{\tilde{H}_i \in \tilde{H}} P(H_i|\tilde{H}_i) \cdot Q(\hat{H}_i|\tilde{H}_i) \cdot d_2(\hat{H}_i, H'_i) \quad (35)$$

最优的隐私保护策略是在针对当前最优攻击下保护效果最佳的策略。在此过程中,针对每一次对手的最优攻击,用户希望获得最优的保护策略。对于每一种保护策略,都有一个推理攻击来提高保护效果,同时为用户带来一定回报。因此,获取所有的可能攻击便可以获得最优的保护策略,然而枚举所有对用户攻击来寻找最优策略是不可行的。本文通过将这种寻优问题看作一个保护者与攻击者的博弈过程,即保护者通过选择最优保护策略 P 来领导博弈,攻击者通过攻击策略 Q 来跟随。上述过程的最优解为对于最优攻击 Q 的最优保护策略 P 。

上述博弈过程可以看作一个 Stackberg 博弈过程,在基于差分隐私保护约束的前提下,保护者首先执行保护策略作为领导者,得到保护后的直方图发布^[22,24];然后,将扰动后的数据作为攻击的输入得到攻击结果,执行推断攻击作为追溯者;最后,根据模型约束,循环执行领导者策略与追随者策略,达到均衡,得到一个双方最优的模型,并据此获得最优的保护策略。在此过程中,保护者作为领导者首先采取保护策略。根据上述知识可知最优保护策略为隐私约束下最大化可用性的策略。因此,执行保护策略的收益包括隐私保护收益与可用性损失支出。为了获取最优保护策略,我们通过线性规划方法建立收益模型,即以最小化保护策略的数据损失期望为目标函数与满足差分隐私约束来实现,进而建立线性模型:

$$\begin{aligned} \min & \sum_{H_i \in H} \pi(H_i) \sum_{\tilde{H}_i \in \tilde{H}} P(\tilde{H}_i|H_i) \cdot c(H_i, \tilde{H}_i) \\ \text{s. t. } & P(\tilde{H}_i|H_i) \leq e^\epsilon \cdot d_1(H_i, H'_i) \cdot P(\tilde{H}_i|H'_i) \\ & \sum_{\tilde{H}_i \in \tilde{H}} P(\tilde{H}_i|H_i) = 1 \\ & P(\tilde{H}_i|H_i) \geq 0 \end{aligned} \quad (36)$$

当执行保护策略数据发布后,攻击者基于发布后的数据 $\tilde{H}_i$,采用攻击策略对于隐私数据进行重建。其收益包括获得隐私信息的程度收益与攻击花费支出,其中,隐私收益通过攻击期望进行衡量,攻击花费使用攻击概率约束实现

$$\begin{aligned}
& \min \sum_{\tilde{H}_t \in \tilde{H}} \pi(\tilde{H}_t) \sum_{H_t \in \tilde{H}, \hat{H}_t \in \hat{H}} P(H_t | \tilde{H}_t) \cdot Q(\hat{H}_t | \tilde{H}_t) \cdot d_2(\tilde{H}_t, \hat{H}_t) \\
& \text{s. t.} \quad \sum_{\tilde{H}_t \in \tilde{H}, \hat{H}_t \in \hat{H}} Q(\hat{H}_t | \tilde{H}_t) = 1 \\
& \quad Q(\hat{H}_t | \tilde{H}_t) \geq 0
\end{aligned} \quad (37)$$

通过上述模型的迭代,求解最优的 Q 必须以 P 最优为前提,同时求解 P 存在同样的问题. 上述两种线性规划模型的求解可以通过求解最优解上限获得^[13], 据此获得保护策略对直方图直接隐私进行保护.

5.3 间接隐私保护策略

为了防止间接隐私泄露,我们基于隐私预算分配 ϵ_1 , 在直接隐私保护的基础上,进行间接隐私的保护. 第3节定义指出,直方图发布后面临多种关联隐私窃取攻击. 为了应对这种情况,通过隐藏直方图之间的关联关系实现关联隐私保护,具体基于隐私泄露风险实现关联直方图的聚类后发布. 基于可用性损失最小化的指数机制,可以实现直方图聚类,保障差分隐私^[11]. 然而,对于关联隐私的保护上述方法存在以下不足:(1) 忽略了直方图的关联可用性容易导致直方图关联关系损失过大的问题;(2) 难以同时保障直接可用性与关联可用性损失最小. 为了解决上述问题,我们提出一种选择权重的直方图间接关联隐私保护方法. 主要工作包括:(1) 基于关联可用性损失最小化建立选择权重,结合指数机制实现直方图候选数据集的构建;(2) 通过与当前直方图聚类误差判断是否进行聚类划分;(3) 最后使用聚类簇均值代表原始数据进行发布.

假设初始聚类簇 C_i , 首先,基于计算 $\tilde{H}_i^*$ 与当前聚类簇 C_i 进行聚类的选择权重. 假设聚类簇 C_i 中桶计数 $\bar{C}_i$ 与 $\tilde{H}_i^*$, $\tilde{H}_i^* \in \tilde{H}$ 关联距离为 $d_3(\tilde{H}_i^*, C_i)$, 基于定义2可知

$$d_3(\tilde{H}_i^*, C_i) = \Pr(\tilde{H}_i^* | \bar{C}_i) \quad (38)$$

其中, $\Pr(\tilde{H}_i^* | \bar{C}_i)$ 代表直方图转移概率.

对于直方图序列 $\tilde{H}$, 关联隐私保护需求最大化即加入聚类的关联性越大, 则攻击者利用其推断出隐私的概率越小. 则关联隐私可用性最大化关联隐私保护的线性规划模型为

$$\begin{aligned}
& \max \sum_{H_{i+1} \in \tilde{H}, H_i \in \tilde{H}} \pi(\tilde{H}_i^*) O(\tilde{H}_i^* | \tilde{H}) d_3(\tilde{H}_i^*, C_i) \\
& \text{s. t.} \quad \sum_{\tilde{H}_i^* \in \tilde{H}} O(\tilde{H}_i^* | \tilde{H}) = 1 \\
& \quad O(\tilde{H}_i^* | \tilde{H}) \geq 0
\end{aligned} \quad (39)$$

其中, $\pi(\tilde{H}_i^*)$ 代表直接扰动后的直方图数据的分布概率. $O(\tilde{H}_i^* | \tilde{H})$ 代表选择权重, 反映直方图 $\tilde{H}_i^*$ 被选择的概率. 关联可用性越大的直方图作为聚类候选数据集其关联隐私泄露风险愈小, 则其选择权重越高.

同时, 选择隐私需求大的数据必定会造成单直方图之间的关联可用性损失. 为了均衡可用性损失与关联隐私泄露风险, 设置关联可用性损失阈值作为可忍受的单直方图可用性损失下界, 加入到上述模型中. 其中, 通过关联熵的变化衡量关联可用性损失. 可用性损失熵表示为

$$\text{Loss} = \left| \sum_{\tilde{H}_i, \tilde{H}_{i+1} \in \tilde{H}} \frac{1}{p(\tilde{H}_i | \tilde{H}_{i+1})} \log p(\tilde{H}_i | \tilde{H}_{i+1}) - \sum_{\tilde{H}'_i, \tilde{H}'_{i+1} \in \tilde{H}'} \frac{1}{p(\tilde{H}'_i | \tilde{H}'_{i+1})} \log p(\tilde{H}'_i | \tilde{H}'_{i+1}) \right|,$$

其中, $\tilde{H}$ 代表 $\tilde{H}_i$ 未聚类的发布序列, $p(\tilde{H}_i | \tilde{H}_{i+1})$ 代表 $\tilde{H}$ 中 $\tilde{H}_i$ 与 $\tilde{H}_{i+1}$ 之间的关联关系, $\tilde{H}'$ 代表聚类 $\tilde{H}_i$ 的直方图序列, 即直方图缺失, $p(\tilde{H}'_i | \tilde{H}'_{i+1})$ 代表 $\tilde{H}'$ 中 $\tilde{H}'_i$ 与 $\tilde{H}'_{i+1}$ 的关联关系.

基于上述分析, 构建可用性最大化关联隐私保护模型如下:

$$\begin{aligned}
& \max \sum_{H_{i+1} \in \tilde{H}, H_i \in \tilde{H}} \pi(\tilde{H}_i^*) O(\tilde{H}_i^* | \tilde{H}) d_3(\tilde{H}_i^*, C_i) \\
& \text{s. t.} \quad \sum_{\tilde{H}_i^* \in \tilde{H}} O(\tilde{H}_i^* | \tilde{H}) = 1 \\
& \quad O(\tilde{H}_i^* | \tilde{H}) \geq 0 \\
& \quad \text{loss} \geq \eta
\end{aligned} \quad (40)$$

其中, η 为关联阈值.

通过上述计算, 获取每一个与聚类簇有关联关系直方图选择权重 O . 结合指数机制, 则直方图 $\tilde{H}_i^*$ 被抽取进行聚类的概率为

$$\pi(\tilde{H}_i^*) = \frac{\exp\left(\frac{\epsilon_2 u(\tilde{H}, \tilde{H}_i^*)}{\Delta u}\right)}{\sum_{\tilde{H}_i^* \in \tilde{H}} \exp\left(\frac{\epsilon_2 u(\tilde{H}, \tilde{H}_i^*)}{\Delta u}\right)},$$

其中, 基于式(40)获取的最优权重 O 构建的打分函数计算如下

$$u(\tilde{H}, \tilde{H}_i^*) = |O(\tilde{H}_i^* | \tilde{H})|.$$

可以看出, 由于 $0 \leq O(\tilde{H}_i^* | \tilde{H}) \leq 1$, 改变一个直方图计数使得打分函数最大改变为 1, 即全局敏感度 $\Delta u = 1$.

为了获取发布序列 H' , 我们根据选择概率递减的顺序进行聚类数据的选择. 聚类函数应尽可能避

免可用性损失过大. 为了减小桶聚类后直接可用性损失, 每次选择与当前聚类簇实现最小误差的直方图进行聚类, 具体根据直方图聚类簇的误差下界进行判断.

定理 1. 在原始直方图中 H 聚类形成的直方图聚类簇 $C = \{C_1, C_2, \dots, C_N\}$, 使用均值 $\bar{C}_t$ 代替原始值会引入两种误差, 包括均值代表原始值引起的重构误差以及第一步中添加拉普拉斯噪声引起误差. 对于直方图聚类簇 $C_t, 1 \leq t \leq N$ 中所携带误差的下界为

$$\text{error}(C_t) \leq \sum_{H_i \in C_t} (\tilde{H}_i) - \bar{C}_t \quad (41)$$

证明. 利用绝对误差度量划分误差, 标准绝对误差度量噪声误差, 则划分误差

$$\text{error}(C_t) = E \left(\sum_{H_i \in C_t} (H_i - \bar{H}_t) \right),$$

$$\bar{H}_t = \frac{\bar{C}_t + \text{lap}(1/\epsilon_1)}{|C_t|},$$

则噪声误差

$$\text{error}(C_t) = E \left(\sum_{H_i \in C_t} \left(H_i + \frac{\text{lap}(1/\epsilon_1)}{|C_t|} - \frac{\bar{C}_t}{|C_t|} \right) \right).$$

由 $|a+b| \leq |a| + |b|$ 定理可知 $\text{error}(C_t)$ 满足:

$$\begin{aligned} \text{error}(C_t) &\leq E \left(\sum_{H_i \in C_t} \left(H_i + \frac{\text{lap}(1/\epsilon_1)}{|C_t|} - \frac{\bar{C}_t}{|C_t|} \right) \right) \\ &= \sum_{H_i \in C_t} (H_i + 1/\epsilon_1) - \bar{C}_t \\ &= \sum_{\tilde{H}_i \in C_t} (\tilde{H}_i) - \bar{C}_t. \end{aligned} \quad \text{证毕.}$$

因此, 对于即将加入分组的 H_t^* , 其加入后应使得误差减小, 即满足式(41). 若不满足, 则利用选择直方图作为新的初始聚类簇重新进行聚类划分, 直到所有直方图数据都完成聚类停止. 结合定义 5 可知, 当前算法满足差分隐私要求.

基于上述分析, 下面给出均衡差分隐私发布算法 BDPP (Balanced Differential Privacy Publishing Algorithm) 方法伪代码实现.

算法 2. BDPP 算法.

输入: 直方图序列 H

隐私预算 ϵ

直方图直接隐私序列 IPL

直方图间接隐私序列 PL

迭代次数 $Time$

输出: 直方图发布序列 H'

1. FOR EACH H_t IN H DO
2. //隐私预算分配策略

3. $\Omega^* = \max(PL_t, IPL_t)$

4. FOR EACH Ω_t IN Ω DO

5. (1) 构建约束因子的拉格朗日模型

$$6. L = \sum_{j=1}^2 (\Omega_j - \Omega^*)^2 v_j^2 + \left(\sum_{j=1}^2 v_j - 1 \right)$$

7. (2) 通过求偏导 $\frac{\partial L}{\partial v} = 0$ 获取最优约束 v

8. (3) 使用式(29)建立隐私预算博弈模型, 并获得预算分配权重

9. $v_1, v_2 = \text{BudgetGame}(v, PL_t, IPL_t)$ //预算分配函数

10. END FOR

11. $\epsilon_1 = v_1 \times \epsilon, \epsilon_2 = v_2 \times \epsilon$

12. //直接隐私保护策略, 获取直接隐私保护直方图 $\tilde{H}$

13. FOR EACH $time$ IN $Time$ DO

14. 获取防守概率 P , 攻击概率 Q , 约束条件 Z, Z'

15. $Z, Z', P, Q = \text{GetParameter}(H)$

16. //使用式(36)建立领导者博弈模型, 获取最优保护策略

17. $P(\tilde{H}_t | H_t) = \text{LeaderStrategy}(\epsilon_1, P, Z, time)$
//领导者策略函数

18. //使用式(37)建立跟随者博弈模型, 获取最优攻击策略

$Q(H_t | \tilde{H}_t) = \text{FollowerStrategy}(\epsilon_1, Q, Z', time)$
//追随者策略函数

19. END FOR

20. //间接隐私保护策略, 获取间接隐私保护直方图聚类簇 C

21. FOR EACH $\tilde{H}_t$ IN $\tilde{H}$ DO

22. (1) 计算关联可用性损失约束阈值

$$\text{Loss} = \left| \sum_{\tilde{H}_t, \tilde{H}_{t+1} \in \tilde{H}} \frac{1}{p(\tilde{H}_t | \tilde{H}_{t+1})} \log p(\tilde{H}_t | \tilde{H}_{t+1}) - \sum_{\tilde{H}'_t, \tilde{H}'_{t+1} \in \tilde{H}} \frac{1}{p(\tilde{H}'_t | \tilde{H}'_{t+1})} \log p(\tilde{H}'_t | \tilde{H}'_{t+1}) \right|$$

23. (2) 使用式(40)构建线性规划模型, 获取选择权重

24. $\sigma^*(\tilde{H}_t | \tilde{H}) = \text{CanStrategy}(\text{Loss})$ //计算权重函数

25. (3) 使用式(41)确定直方图 H_t 的打分函数 $u(\tilde{H}, \tilde{H}_t^*)$

26. $u(\tilde{H}, \tilde{H}_t^*) = |O(\tilde{H}_t^* | \tilde{H})|$

27. (4) 以概率 $\pi(\tilde{H}_t^*) = \frac{\exp\left(\frac{\epsilon_2 u(\tilde{H}, \tilde{H}_t^*)}{\Delta u}\right)}{\sum_{H_t^* \in H} \exp\left(\frac{\epsilon_2 u(\tilde{H}, \tilde{H}_t^*)}{\Delta u}\right)}$ 抽取

直方图 $\tilde{H}_t^*$

28. (5) 确定 $\tilde{H}_t^*$ 聚类误差下界 $\text{error}(\tilde{H}_t^*) \leq \sum_{\tilde{H}_i^* \in C_t} (\tilde{H}_i) - \bar{C}_t$

29. (6) 基于误差下界通过聚类下界进行贪心聚类

30. $C = \text{GreedCluster}(\text{error}, \tilde{H})$ //聚类函数

31. END FOR

Return $H' = C$

可以看出,算法中最耗时的运算当属线性规划(第 17 步)求最优解,其最坏复杂度为 $O(mn^2)$, n 为线性规划变换中变量个数, m 为线性规划中约束数量. 则 N 个直方图所花费的时间复杂度为 $O(mn^2N)$, 则利用复杂度组合差分隐私机制的复杂度可得 BDPP 算法整体时间复杂度为 $O(mn^2N)$.

5.4 隐私性与可用性分析

定理 2. BDPP 算法满足 ϵ -差分隐私

证明. 首先,第 2~11 步,根据隐私预算分配策略,得到隐私分配结果 $\epsilon = \epsilon_1 + \epsilon_2$, 然后执行直方图保护算法. 第 12~18 步根据博弈结果添加独立噪声,虽然使用扩展隐私定义,其最终敏感度限制为 1,根据扩展隐私的定义,可知候选数据集中的数据满足 ϵ_1 差分隐私. 第 19~31 步基于选择权重,采用指数机制选择直方图数据,同时根据定理 2 可知直方图发布数据满足 ϵ_2 差分隐私. 最后,基于定理 3 组合策略,本文采取的隐私保护策略,算法 2 满足 ϵ -差分隐私.

综上所述,BDPP 算法满足 ϵ -差分隐私保护.

同时,算法实现主要通过加噪与指数机制来实现. 加噪机制主要通过线性规划集合拉普拉斯噪声实现,随着数据规模的变换,由于线性规划算法的执行开销大,发布算法的执行开销会增长的非常明显,但仍可以执行,根据定理 1 发布误差保持在 $\sum_{H_i \in C_i} (H_i + 1/\epsilon_1) - \bar{C}_i$, 相对于 AHP 算法携带的误差是比较小的,所以算法是可行的.

6 实验设置与结果分析

6.1 实验使用数据集

实验所用硬件环境为 Intel(R)Core(TM) i7 CPU, 3.19 GHz, 4 GB 内存,操作系统平台为 Windows 10 操作系统,使用 Python 编程语言实现所有方法. 实验所用数据集为新西兰 2006 年人口街区的普查数据 WakeTakere 数据集,我们选取其中 7725 个街区数据作为直方图的桶;以及合成 2004~2010 年 Google Trends 与 American Online 搜索关键字为“奥巴马”的搜索日志的 Search_log 数据集,其中囊括了 32768 条数据,将其划分为与 WakeTakere 相近的桶进行验证:

首先,由于数据集缺乏认知能力的标记,结合认知理论^[34]与直方图发布相关理论^[35]设定认知标记.

具体采集 15 个用户反复 10 次的平均认知能力作为认知结果. 同时,背景知识设置为 0.15~0.90, 6 个等级,认知范围设定为 5~50 个直方图范围,在此基础上进行反复测试 10 次的平均认知结果作为用户对于该直方图的认知能力. 基于上述设置可以得到直方图数据的认知能力分布如表 1.

表 1 不同数据集认知能力分布情况

数据集	CI					
	0.15	0.30	0.45	0.60	0.75	0.90
WakeTakere	1231	1011	1002	1235	1685	1201
Search_log	1607	1275	1168	1405	1110	1343

然后,使用旁路攻击相关概率模型^[36]与隐私保护的数据发布^[37]的方法,计算原始分布概率、原始直方图、发布直方图、重建直方图相应分布概率. 并且本文所提线性规划方法都通过公开线性规划平台进行运算.

最后,实现其它主要参数设置如下:数据规模为 10%~70%,稀疏度为 0.05~0.35,迭代次数上限为 150 次,查询范围 100~700,其余参数对应于其后的解释.

6.2 评价指标

本文采用均方误差与 KL 距离、间接隐私泄露风险综合隐私泄露风险 4 个衡量指标评价实验结果.

(1) KL 距离 (KL-Divergence). 用来描述原始直方图 H 与发布后的直方图 H' 概率分布的距离,即直方图分布差异度的期望,计算公式如下:

$$KLD(H|H') = \sum_{H'_i \in H', H_j \in H} \ln\left(\frac{H'_i}{H_j}\right).$$

(2) 均方误差 (Mean Variance Error). 假设有原始直方图 H 及一个计数查询向量 Q , 则均方误差代表原始直方图查询输出 $Q(H)$ 与执行隐私保护算法后的直方图查询 $Q(H')$ 输出差异程度,即

$$MSE(H, H') = (Q(H') - Q(H))^2 / |Q|,$$

其中, $|Q|$ 代表查询数量.

(3) 间接隐私泄露风险 (Privacy Leakage) 代表直方图之间的关联隐私泄露,通过本文所提算法 1 计算,即

$$PL(H_t) = QPLMDM(H_t).$$

(4) 综合隐私泄露风险 (Synthesize Privacy Leakage) 代表直方图的综合隐私泄露情况,对直接 IPL 与间接隐私泄露 PL 进行加权平均计算

$$SPL(H_t) = \gamma_1 \cdot PL(H_t) + \gamma_2 \cdot IPL(H_t),$$

其中 γ_1, γ_2 代表权重, 设为 0.5.

6.3 结果分析

6.3.1 隐私风险量化效果分析

本次实验目的是检验不同参数对于间接隐私风险量化方法的影响, 主要通过不同数据规模下与不同稀疏度情况下的风险量化结果验证所提方法的有效性. 经过反复实验, 相关参数设置如下: 当稀疏度作为变量时, 数据规模为 10%、20%、30%、40%、50%、60%、70%; 当数据规模作为变量时, 分布稀疏度设置为 0.05、0.10、0.15、0.10、0.25、0.30、0.35,

同时隐私预算为 0.1.

从图 3 中可以看出随着发布数据规模的上升, 在不同数据集中, 间接隐私泄露风险也逐渐变大; 这是由于基于一方面直方图数量的上升间接关联随之升高; 另一方面, 随着直方图数据的增高, 与攻击背景的相关性也随之变大. 同时, 从图 4 中也可以看出随着数据稀疏度的增长, 间接隐私也随之增高, 即数据相似数量增多, 则其关联关系也愈发明显. 上述分析证明间接隐私量化方法对于不同数据规模、稀疏程度具有鲁棒性.

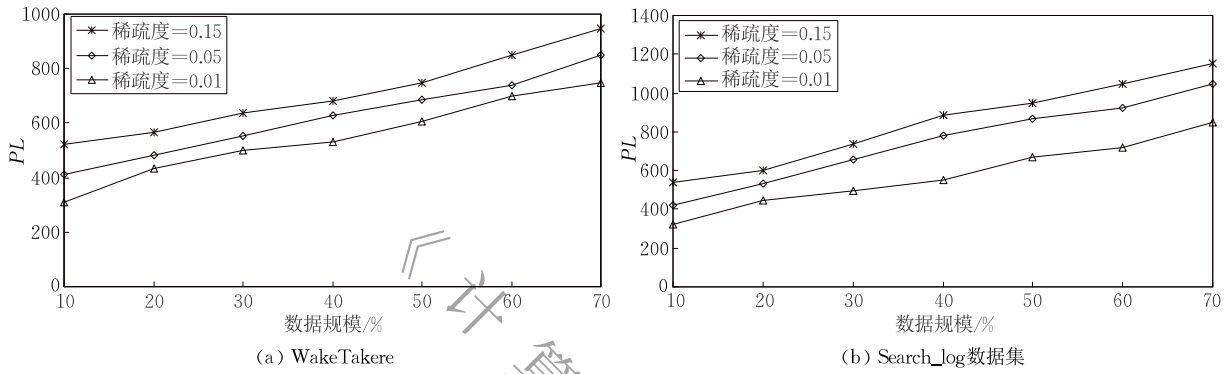


图 3 不同数据规模条件下间接隐私量化分析

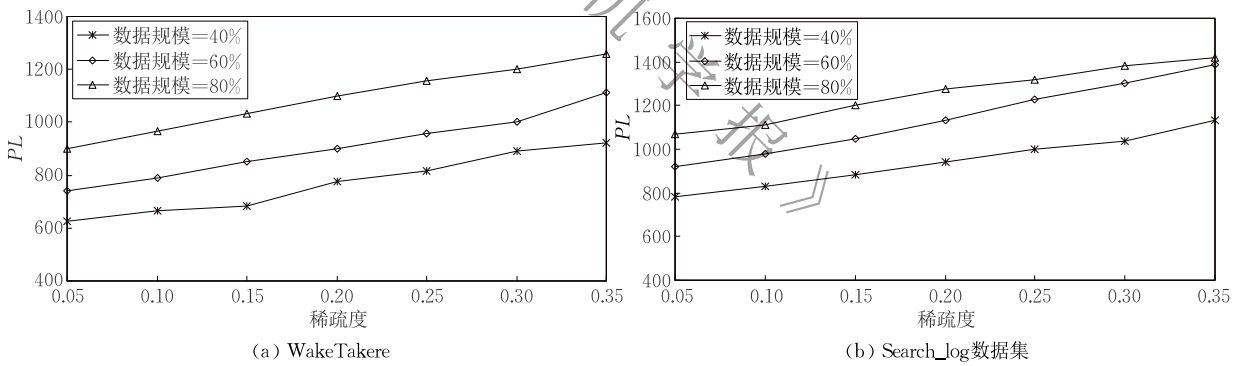


图 4 不同稀疏度条件下间接隐私量化分析

6.3.2 隐私保护算法效果分析

(1) 不同参数对于隐私保护效果分析

本次实验目的是为了分析不同参数对于本文所提保护算法的影响, 主要对于不同查询范围、不同隐私保护需求下的隐私保护效果进行分析. 其中, 隐私预算参数代表隐私保护级别. 借鉴文献[37]中的实验设计思路与反复测试结果, 我们设定查询范围为 100, 200, 300, 400, 500, 600, 700, 隐私预算为 0.01, 0.1, 1 下的参数进行验证.

从图 5~图 7 中可以看出随着隐私预算(从 0.01 到 1)的增大, 数据隐私泄露风险在逐渐降低, 说明 BDPP 方法可以按需进行保护, 并且具有良好的效果;

同时, 图 5~图 7 中显示, SPL 、 MSE 都是先急剧增大, 随后保持平稳. SPL 、 MSE 曲线增大是由于范围增大导致数据量的增大, 造成了 MSE 、 SPL 的累积; 同时, 范围的增大一定程度上增大了攻击者的攻击能力, 增大了隐私泄露风险; 后 SPL 、 MSE 曲线平稳是由于 BDPP 根据攻击能力而获取最优策略的能力很好的适应了攻击的增大, 加强了保护, 使得发布直方图关联隐私泄露减缓. 另一方面, 伴随泄露风险的增大, BDPP 基于泄露风险结果按需均衡直接与间接隐私的保护力度、隐私需求与可用需求, 从而降低了发布误差, 使得 MSE 变得慢慢平稳. 可以看出, BDPP 算法对于查询范围的变化有着较好的适应性.

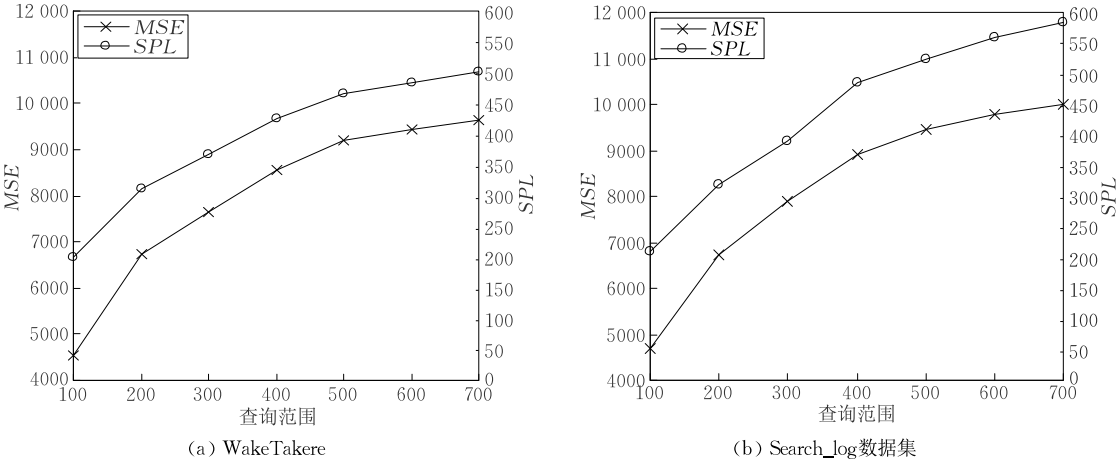


图 5 当 $\epsilon=0.1$ 时,不同数据集隐私保护算法鲁棒性分析

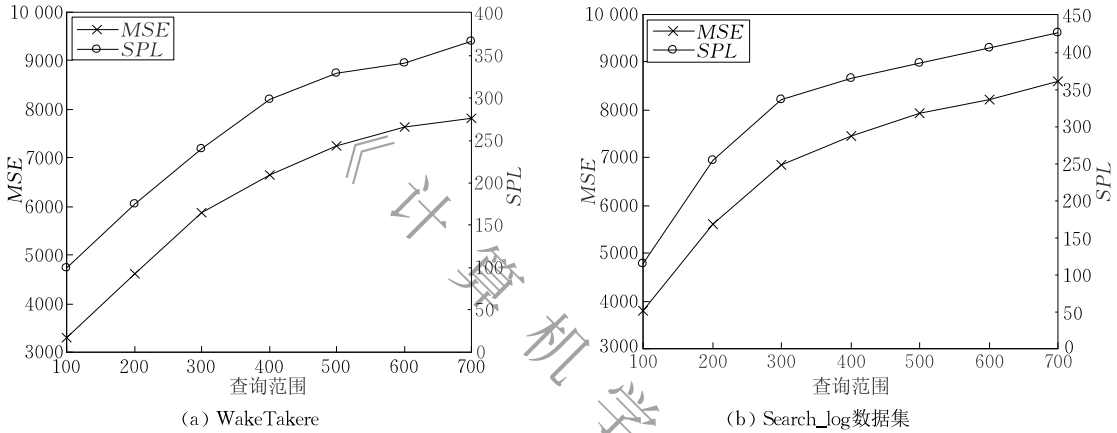


图 6 当 $\epsilon=0.01$ 时,不同数据集隐私保护算法鲁棒性分析

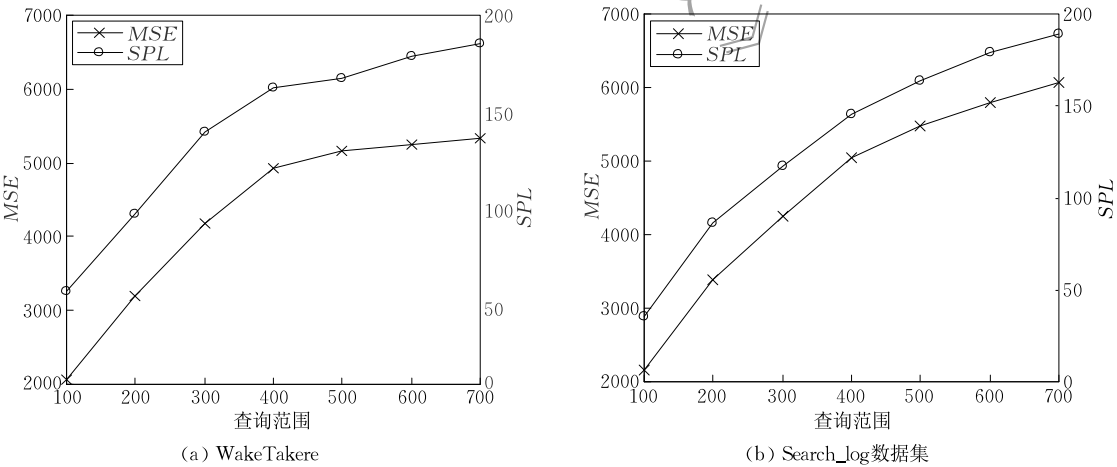


图 7 当 $\epsilon=1$ 时,不同数据集隐私保护算法鲁棒性分析

(2) 不同保护算法保护效果分析

本次实验目的是为了对比本文所提算法与其它主流算法的保护效果. 其中,选取了 AHP^[11]、GS^[38]、DiffHR^[25]直方图常用的隐私保护方法进行对比验证. AHP 算法使用噪声机制对直方图数据加噪,据此构建平衡重构误差与 Lap 误差的聚类函数形

成直方图聚类簇,形成直方图发布序列. DiffHR 算法结合指数机制的蒙特卡洛逼近的直方图数据排序与贪心划分的直方图簇,加噪后构建直方图发布序列. GS 算法通过求解全局最优划分函数,并据此构建等宽直方图聚类簇,加噪后实现直方图发布. 根据反复测试设置查询范围、隐私预算两个因素作为

变量,以 MSE 、 KL -距离、 PL 、 SPL 作为评价指标进行对比验证直方图发布效果. 其中,泄露风险代表了对于间接隐私与关联隐私泄露的保护效果,主要为了佐证其余两个隐私保护可用性指标的分析; MSE 代表保护后的单个直方图的误差累积情况,更多的关注个体误差; KL -距离代表两者分布误差累积情况,更多的关注整体误差. 通过反复实验,设定查询范围为 100、200、300、400、500、600、700,隐私预算

为 0.01、0.1、1 进行对比验证.

① 基于 PL 量化结果的保护效果对比

从图 8 可以看出随着隐私保护预算(隐私级别)的增大,不同方法发布的直方图关联隐私泄露情况都有所缓解;同时,随着保护预算的增大,BDPP 算法相对于 AHP、DiffHR 算法都具有明显的保护效果,而 AHP、DiffHR 由于忽略了关联隐私保护的重要性,虽有效果,却不明显.

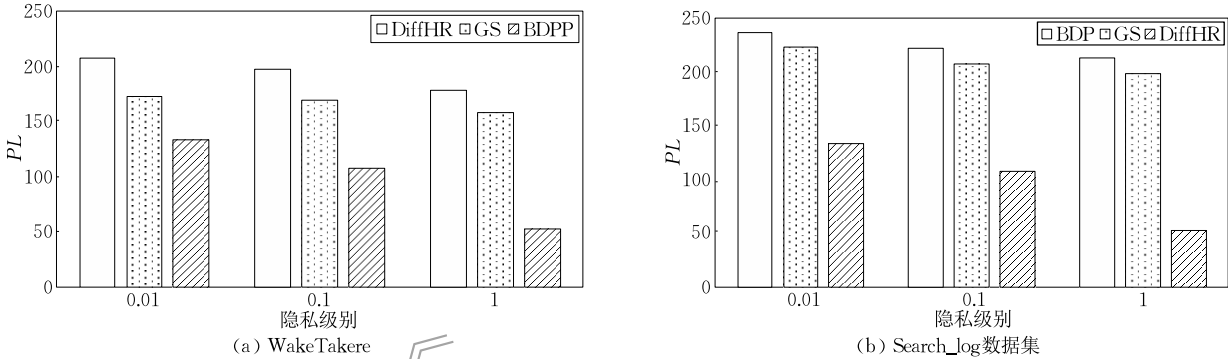


图 8 不同数据集关联隐私保护情况分析

② 基于 MSE 评价的保护效果对比

从图 9~图 11 显示出随着隐私保护预算的增大,直方图数据的 MSE 也随之降低,可用性也增高. 同时,在图 9 中可以看出本文提出的 BDPP 算法相对于 AHP、DiffHR 算法取得了不错的效果. 同时可以看出,在不同隐私预算情况下本文的隐私方法

在不同数据集上的效果明显优于 DiffHR 算法,并且也比 AHP 方法略好. 主要原因是 BDPP 同时考虑到了直接与间接隐私的泄露情况,根据攻击情况自适应调节隐私预算分配,提高隐私预算分配精准度,提高了隐私保护效果. 尤其可以看出,在 $\epsilon=1$ 的情况下,本文所提 BDPP 算法的效果完全优于同

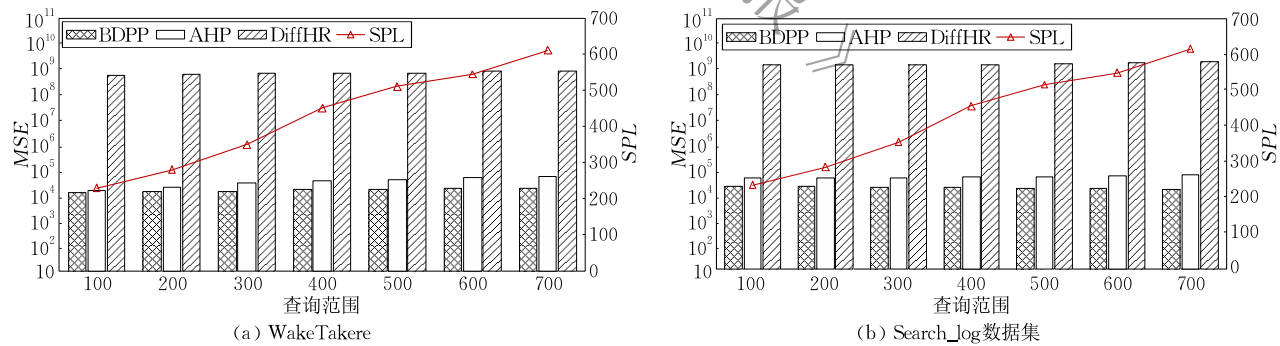


图 9 当 $\epsilon=0.01$ 时,不同数据集隐私保护算法鲁棒性分析

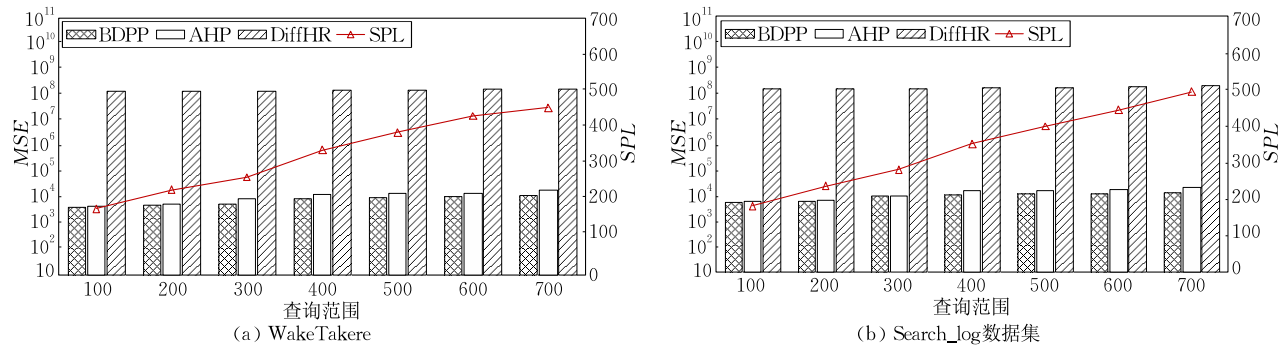


图 10 当 $\epsilon=0.1$ 时,不同数据集隐私保护算法鲁棒性分析

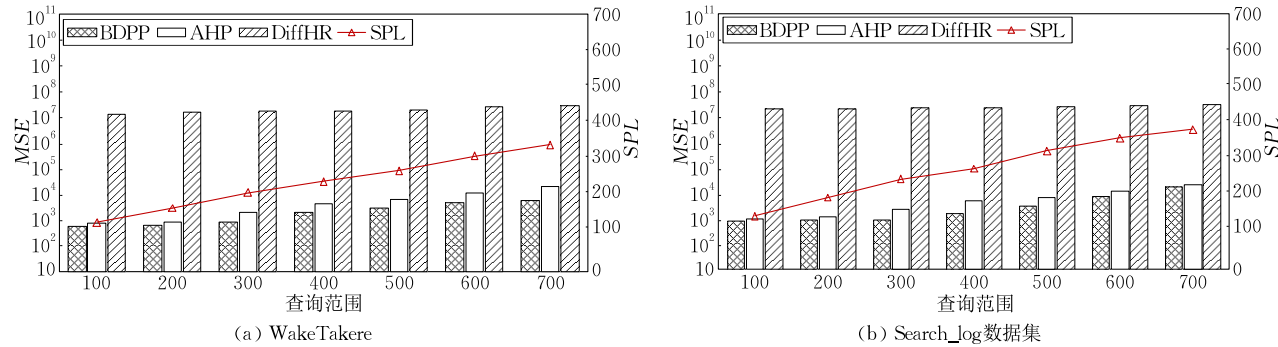


图 11 当 $\epsilon=1$ 时,不同数据集隐私保护算法鲁棒性分析

类算法,可以保证数据的隐私与可用性的均衡.同时可以看出随着关联隐私泄露风险的增大,BDPP 方法的 MSE 结果也具有相对最好的效果.

③ 基于 KL-距离的保护效果对比

从表 2 中可以看出,我们对比在不同隐私预算情况下,实验所得 KL-距离的结果.在不同隐私保护预算的情况下,本文所提算法仅次于 GS 方法的效果,特别是在隐私预算 $\epsilon=1$ 的情况下,本文所提算法的保护效果约是 AHP 的 2 倍.同时,GS 算法采用高度复杂性取得了较好的结果,而 BDPP 用远低于 GS 算法的复杂度取得了相近的效果.这是由于 BDPP 根据情况自动调整保护策略,并且采取了更加高效的差分隐私保护方法,在保护个体可用性的同时使得总体分布距离也实现了较低的误差.

表 2 不同差分隐私保护算法距离熵评价结果对比分析

数据集	方法			
	GS	BDPP	DiffHR	AHP
$\epsilon=0.01$				
WakeTakere	0.535	2.011	2.589	3.241
Search_log	0.941	3.188	4.102	4.141
$\epsilon=0.1$				
WakeTakere	0.487	1.988	2.089	3.002
Search_log	0.618	2.768	3.502	3.818
$\epsilon=1$				
WakeTakere	0.348	1.281	1.654	2.241
Search_log	0.241	0.641	0.983	1.054

7 总结与展望

对于直方图数据发布的隐私保护研究成为数据发布领域的研究热点.然而,因其忽略了直方图之间的关联隐私,也会导致不能完全防止隐私泄露的问题.为此,本文基于上述问题开展研究.针对缺乏关联隐私保护的问题,基于马尔科夫模型提出一种关联隐私定义,在此基础上提出面向关联隐私泄露隐私风险的多指标量化方法,并结合评估结果与博弈

理论实现了均衡全面的关联隐私保护,最后通过实验验证了我们的方法具有较强的鲁棒性.

在未来研究中,我们将针对流式直方图数据发布的关联隐私保护问题进行研究.同时,本文所提算法的解决效率并不是最优,因此怎样提高算法效率使其可以推广到高维直方图数据的隐私保护进行研究也是未来的主要方向.

参 考 文 献

[1] Baldimtsi F, Camenisch J, Dubovitskaya M, et al. Accumulators with applications to anonymity-preserving revocation// Proceedings of the 2017 IEEE European Symposium on Security and Privacy (EuroS&P). Paris, France, 2017: 301-315

[2] Li Chun-Ta, Wu Tsu-Yang, Chen Chin-Ling, et al. An efficient user authentication and user anonymity scheme with provably security for IoT-based medical care system. Sensors, 2017, 17(7): 1482-1489

[3] Caballero-Gil C, Molina-Gil J, Hernández-Serrano J, et al. Providing k -anonymity and revocation in ubiquitous VANETs. Ad Hoc Networks, 2016, 36(P2): 482-494

[4] Liu F, Li T. A clustering-anonymity privacy-preserving method for wearable IoT devices. Security and Communication Networks, 2018, 2018: 1-8

[5] Dwork C. Differential privacy//International Colloquium on Automata, Languages, and Programming. Berlin, Heidelberg: Springer, 2006: 1-12

[6] Hay M, Rastogi V, Miklau G, et al. Boosting the accuracy of differentially private histograms through consistency. Proceedings of the VLDB Endowment, 2010, 3(1-2): 1021-1032

[7] Xiao X, Wang G, Gehrke J. Differential privacy via wavelet transforms. IEEE Transactions on Knowledge and Data Engineering, 2010, 23(8): 1200-1214

[8] Dwork C, McSherry F, Nissim K, et al. Calibrating noise to sensitivity in private data analysis//Conference on Theory of Cryptography. Berlin Heidelberg: Springer-Verlag, 2006: 265-284

- [9] Albarghouthi A, Hsu J. Synthesizing coupling proofs of differential privacy. *Proceedings of the ACM on Programming Languages*, 2017, 2(POPL): 1-30
- [10] Eibl G, Engel D. Differential privacy for real smart metering data. *Computer Science-Research and Development*, 2017, 32(1-2): 173-182
- [11] Zhang X, Chen R, Xu J, et al. Towards accurate histogram publication under differential privacy//*Proceedings of the 2014 SIAM International Conference on Data Mining*. Philadelphia, USA, 2014: 587-595
- [12] Xiao Y, Xiong L. Protecting locations with differential privacy under temporal correlations//*Proceedings of the 22nd ACM SIGSAC Conference on Computer and Communications Security*. Denver, USA, 2015: 1298-1309
- [13] Ou L, Qin Z, Liu Y, et al. Multi-user location correlation protection with differential privacy//*Proceedings of the 2016 IEEE 22nd International Conference on Parallel and Distributed Systems (ICPADS)*. Wuhan, China, 2017: 422-429
- [14] Huo Zheng, Meng Xiao-Feng. A trajectory data publication method under differential privacy. *Chinese Journal of Computers*, 2018, 41(2): 400-412(in Chinese)
(霍峥, 孟小峰. 一种满足差分隐私的轨迹数据发布方法. *计算机学报*, 2018, 41(2): 400-412)
- [15] Chen J, Ma H, Zhao D, et al. Correlated differential privacy protection for mobile crowdsensing. *IEEE Transactions on Big Data*, 2017, PP(99): 1-1
- [16] Zhang Xiao-Jian, Shao Chao, Meng Xiao-Feng. Accurate histogram release under differential privacy. *Journal of Computer Research and Development*, 2016, 53(5): 1106-1117 (in Chinese)
(张啸剑, 邵超, 孟小峰. 差分隐私下一种精确直方图发布方法. *计算机研究与发展*, 2016, 53(5): 1106-1117)
- [17] McSherry F, Talwar K. Mechanism design via differential privacy//*Proceedings of the 48th Annual IEEE Symposium on Foundations of Computer Science*. Washington: IEEE Computer Society, 2007: 94-103
- [18] Cao Y, Yoshikawa M, Xiao Y, et al. Quantifying differential privacy in continuous data release under temporal correlations. *IEEE Transactions on Knowledge & Data Engineering*, 2017, PP(99): 1-1
- [19] Liu H, Li X, Li H, et al. Spatiotemporal correlation-aware dummy-based privacy protection scheme for location-based services//*Proceedings of the IEEE Conference on Computer Communications (INFOCOM 2017)*. Atlanta, USA, 2017: 1-9
- [20] Gao Quan-Li, Gao Ling, Yang Jian-Feng, et al. A preference elicitation method based on users' cognitive behavior for context-aware recommender system. *Chinese Journal of Computers* 2015, 38(9):1767-1776(in Chinese)
(高全力, 高岭, 杨建锋等. 上下文感知推荐系统中基于用户认知行为的偏好获取方法. *计算机学报*, 2015, 38(9): 1767-1776)
- [21] Fu Zhi-Yao, Gao Ling, Sun Qian, et al. Evaluation of vulnerability severity based on rough sets and attributes reduction. *Journal of Computer Research and Development*, 2016, 53(5): 1009-1017(in Chinese)
(付志耀, 高岭, 孙蓁等. 基于粗糙集的漏洞属性约简及严重性评估. *计算机研究与发展*, 2016, 53(5): 1009-1017)
- [22] Shokri R. Privacy games: Optimal protection mechanism design for Bayesian and differential privacy. *CoRR*, vol. abs/1402.3426, 2014
- [23] Conitzer V, Sandholm T. Computing the optimal strategy to commit to//*Proceedings of the 7th ACM Conference on Electronic Commerce*. Ann Arbor, USA, 2006: 82-90
- [24] Ma Q, Zhang S, Zhu T, et al. PLP: Protecting location privacy against correlation analyze attack in crowdsensing. *IEEE Transactions on Mobile Computing*, 2015, PP(99): 1-1
- [25] Wu Yun-Cheng, Chen Hong, Zhao Su-Yun, et al. Differentially privacy trajectory protection based on spatial and temporal correlation. *Chinese Journal of Computers*, 2018, 41(2): 309-322(in Chinese)
(吴云乘, 陈红, 赵素云等. 一种基于时空相关性的差分隐私轨迹保护机制. *计算机学报*, 2018, 41(2): 309-322)
- [26] Jiang Wei, Fang Bing-Xing, Tian Zhi-Hong, et al. Research on defense strategies selection based on attack-defense stochastic game model. *Journal of Computer Research and Development*, 2010, 47(10): 1714-1723(in Chinese)
(姜伟, 方滨兴, 田志宏等. 基于攻防随机博弈模型的防御策略选取研究. *计算机研究与发展*, 2010, 47(10): 1714-1723)
- [27] Jiang Wei, Fang Bingxing, Tian Zhi-Hong, et al. Evaluating network security and optimal active defense based on attack-defense game model. *Chinese Journal of Computers*, 2009, 32(4): 817-827(in Chinese)
(姜伟, 方滨兴, 田志宏等. 基于攻防博弈模型的网络安全测评和最优主动防御. *计算机学报*, 2009, 32(4): 817-827)
- [28] Reed J, Pierce B C. Distance makes the types grow stronger: A calculus for differential privacy//*Proceedings of the 15th ACM SIGPLAN International Conference on Functional Programming*. Baltimore, USA, 2010: 157-168
- [29] Chatzikokolakis K, Andrés M E, Bordenabe N E, et al. Broadening the scope of differential privacy using metrics//*Proceedings of the International Symposium on Privacy Enhancing Technologies Symposium*. Berlin, Heidelberg: Springer, 2013: 82-102
- [30] Shokri R. Privacy games: Optimal protection mechanism design for Bayesian and differential privacy. *CoRR*, vol. abs/1402.3426, 2014
- [31] Barthe G, Köpf B, Olmedo F, et al. Probabilistic relational reasoning for differential privacy. *ACM Transactions on Programming Languages and Systems*, 2013, 35(3): 1-49
- [32] Kirschner P A, Park B, Malone S, et al. Toward a cognitive theory of multimedia assessment (CTMMA). *Learning, Design, and Technology: An International Compendium of Theory, Research, Practice, and Policy*, 2017: 1-23

[33] Kang Jian, Wu Yingjie, Huang Siyong, et al. Algorithm for differential privacy histogram publication with non-uniform private budget. *Journal of Frontiers of Computer Science and Technology*, 2016, 10(6): 0786-13

[34] Alvim M S, Chatzikokolakis K, Palamidessi C, et al. Measuring information leakage using generalized gain functions//*Proceedings of the 2012 IEEE 25th Computer Security Foundations Symposium*. Cambridge, USA, 2012: 265-279

[35] Köpf B, Basin D A. An information-theoretic model for adaptive side-channel attacks//*Proceedings of the 2007 ACM Conference on Computer and Communications Security (CCS 2007)*. Alexandria, USA, 2007: 286-296

[36] Zhang X J, Meng X F. Streaming histogram publication method with differential privacy. *Journal of Software*, 2016, 27(2): 381-393

[37] Zhang X, Chen R, Xu J, et al. Towards accurate histogram publication under differential privacy//*Proceedings of the 2014 SIAM International Conference on Data Mining*. Society for Industrial and Applied Mathematics, 2014: 587-595

[38] Kellaris G, Papadopoulos S. Practical differential privacy via grouping and smoothing. *Proceedings of the VLDB Endowment*, 2013, 6(5): 301-312



GAO Ling, Ph. D. , professor. His research interests include privacy protection, edge computing and computer

Background

As a commonly used data publishing method, histogram data publishing has received extensive attention due to its intuitive and convenient features. The k -anonymous privacy protection method has the following problems as a traditional method: (1) the privacy protection effect is greatly affected by the data distribution, and the protection effect of the sparse data is not ideal; (2) the privacy protection method needs to obtain the background knowledge of the attacker, and the acquisition difficulty is relatively large.

In order to solve the above problems, the differential privacy was born. There are many research scholars have done research on histogram publishing methods based on differential privacy protection. Many research in this area is concentrated on the method of combining the exponential mechanism and the noise mechanism to achieve tree division and recombination, and achieves certain effects. (1) tree division; the basic idea is based on the layered structure of the tree structure to achieve the noise is added to different data levels. (2) recombination; the basic idea is clustering of related histogram data by combining histogram data based on the reconstruct error minimization. But the existing differential

network.

WANG Hai, Ph. D. , professor. His research interests include privacy protection and edge computing.

GUO Hong-Bo, Ph. D. His research interests include privacy protection and medical image processing.

ZHENG Jie, Ph. D. His research interests include privacy protection and edge computing, heterogeneous networks.

privacy protection methods suffer from the privacy protection for histogram associations, efficient and accurate privacy disclosure assessment methods, and balanced protection for association and statistical privacy problems.

In this paper, we design a balanced differential privacy protection method for histogram data publishing by introducing the associated privacy leakage assessment and quantization mechanism. First, we define indirect association privacy and related concepts. Second, we propose a comprehensive decision-making method that combines multiple privacy loss indicators to quantify the risk of privacy disclosure. Finally, based on the game idea and combined with linear programming, we realized the effective distribution of privacy budget balanced distribution and direct and indirect privacy. The validity and robustness of the proposed method are verified by experiments. And the protection performance of the proposed method is better than the method proposed in the paper.

This work is supported by the National Natural Science Foundation of China under Grant Nos. 61672426, 61572401 and the National Key Research and Development Program under Grant Nos. 2019YFC1521400.