

基于压缩的本地差分隐私的序列数据收集方法

金 严¹, 朱友文¹, 吴启晖²

(1. 南京航空航天大学计算机科学与技术学院, 南京 211106; 2. 南京航空航天大学电子信息工程学院, 南京 211106)

摘 要: 压缩的本地差分隐私是本地差分隐私的一种基于度量的放松形式, 它具有比本地差分隐私更好的效用性和灵活性。但是, 现有方案在序列模式捕捉和效用性方面存在不足。为了克服这些局限性, 提出了一种新颖的基于压缩的本地差分隐私的序列数据收集方法 SCM-CLDP。该方法在收集过程中充分考虑了序列数据的长度、转移等重要信息, 通过这些信息, 数据收集者能够合成接近原始数据集的隐私保护的数据集。根据扰动对象的不同, 本文提出了两种收集方法, 分别是基于值扰动的 SCM-VP 方法和基于转移扰动的 SCM-TP 方法。理论证明了 SCM-VP 和 SCM-TP 满足序列级别的压缩的本地差分隐私, 并基于两个真实数据集, 在 Markov 链模型准确性、合成数据集效用性及频繁序列模式挖掘准确性上, 与现有方案进行了对比实验。结果表明, SCM-CLDP 表现出显著的优势, 其中 SCM-VP 的性能在大多数情况下都要优于 SCM-TP。并且在最优的情况下, 相较于现有方法, SCM-CLDP 在 Markov 链模型及合成数据集分布误差方面至少降低了一个数量级。同时, SCM-CLDP 在合成数据集中各项频率排序的准确性以及频繁序列模式挖掘的准确性方面, 相较于现有方法提升了近 30%。

关键词: 压缩的本地差分隐私; 序列数据; Markov 链模型; 数据收集; 隐私保护

中图分类号: TP309 **文献标志码:** A

Sequential Data Collection Method with Condensed Local Differential Privacy

JIN Yan¹, ZHU Youwen¹, WU Qihui²

(1. College of Computer Science and Technology, Nanjing University of Aeronautics and Astronautics, Nanjing 211106, China;
2. College of Electronic and Information Engineering, Nanjing University of Aeronautics and Astronautics, Nanjing 211106, China)

Abstract: Condensed local differential privacy is a metric-based relaxation of local differential privacy with better utility and flexibility than local differential privacy. However, existing solutions are deficient in terms of sequence pattern capture and utility. To address these limitations, this paper proposes SCM-CLDP, a novel sequential data collection method based on condensed local differential privacy. SCM-CLDP fully takes into account important information such as the length and transitions of sequential data during the collection process, through which the data collector is able to synthesize privacy-preserving dataset close to the original dataset. Specifically, according to different perturbation objects, we propose two collection methods, SCM-VP based on value perturbation and SCM-TP based on transition perturbation,

基金项目: 江苏省重点研发计划(产业前瞻与关键核心技术)(BE2022068, BE2022068-1); 国家自然科学基金(62172216); 中央高校基本科研业务费项目(NP2024117); 稳定支持国防特色学科基础研究项目(ILF240061A24)。

收稿日期: 2024-06-19; **修订日期:** 2024-10-22

respectively. We theoretically prove that SCM-VP and SCM-TP satisfy sequence-level condensed local differential privacy, and comparative experiments are conducted with existing solutions based on two real datasets in terms of Markov chain model accuracy, synthetic dataset utility, and frequent sequence pattern mining accuracy. The results show that SCM-CLDP performs significantly better than the existing solutions, with SCM-VP outperforming SCM-TP in most cases. In the optimal situation, SCM-CLDP reduces the error of the Markov chain model and the distribution of the synthetic dataset by at least one order of magnitude compared to the existing method. Meanwhile, SCM-CLDP improves the accuracy of item frequency ranking of the synthetic dataset and the accuracy of frequent sequence pattern mining by nearly 30% compared to existing solutions.

Key words: condensed local differential privacy; sequential data; Markov chain model; data collection; privacy protection

引言

序列数据是具有时间相关性或逻辑相关性的一系列数据点,这些数据点按照它们出现的时间顺序或逻辑顺序排列,如网页点击数据、移动轨迹等。在实际生产生活中,序列数据的应用极其广泛,比如,从网页点击数据中分析用户行为模式、从个人轨迹数据中做旅游路线推荐^[1-2]以及从文本序列中构建语言模型^[3]等。隐私已成为当前大数据时代的重要关注点,随着个人隐私保护意识的提高,各国已逐步建立和完善了相关的个人隐私保护法律法规,如中国在2021年相继实施了《中华人民共和国数据安全法》^[4]和《中华人民共和国个人信息保护法》^[5]。序列数据中包含大量的用户隐私信息,如从网页点击数据中可推断出个人偏好、从个人轨迹数据中可以推断出家庭住址甚至健康状况^[6]。因此,在序列数据的收集中应避免直接收集用户真实的敏感序列数据,而是通过隐私保护的方法来收集。

近年来,本地差分隐私(Local differential privacy, LDP)^[7]是一种被广受关注的隐私保护模型,它是传统的依赖于可信第三方的差分隐私(Differential privacy, DP)^[8]的本地化版本,它给出了严格的可证明的隐私保护,且无需可信的第三方。LDP要求用户在本地扰动自身的隐私数据且不同用户有一定概率扰动输出相同的值,使得攻击者无法从截获到的个人数据中准确推断出隐私信息,但是数据收集者可从扰动的数据中聚合得到统计信息。目前,LDP已在工业界得到广泛应用,例如,Google公司在Chrome浏览器中嵌入LDP模块来收集浏览器默认主页设置信息^[9]、Apple公司在终端输入法中利用LDP来收集emoji表情使用信息以预测用户输入^[10]。

然而,尽管LDP在处理大规模数据集时能够利用数据量来抵消扰动带来的影响,从而显示出较高的效用性,但在小规模数据集上,其性能表现并不尽如人意。除此之外,LDP平等地保护所有数据,即真实值扰动到其他所有非真实值的概率相等,但在某些场景希望不同数据对间的隐私保护程度与数据对的距离相关,为此,有学者提出压缩的本地差分隐私(Condensed LDP, CLDP)^[11],它是一种基于度量的LDP放松形式。相比LDP,CLDP在数据规模较小的场景下表现更优。本文专注于序列数据收集问题,LDP下的序列数据收集任务已经有了广泛而深入的研究。然而,自CLDP提出以来,针对满足CLDP的序列数据收集方法研究鲜少。现存的满足CLDP的序列数据收集方法单一地逐位扰动序列的数据项,忽视了序列模式。简单来说,序列模式是数据集中经常出现的子序列,它在序列数据分析中发挥着重要作用^[12],它蕴含着数据间的内在关系和规律。对序列数据进行分析 and 建模的一种主流方法是Markov链,它是描述一系列可能事件的随机模型,其中每个事件都依赖于前面一个或多个事件,利用Markov链模型能很好地模拟数据间的转移。现已有不少利用Markov链模型来实现隐私保护的序列数据收集方法,如满足DP/LDP的频繁序列模式挖掘^[13]、轨迹合成等^[6,14-15]。

针对上述问题,本文提出了一种新颖的满足压缩的本地差分隐私的序列数据收集方法(Sequential data collection method with CLDP, SCM-CLDP)。不同于现有方法中用户直接独立地扰动序列中的各项,且数据收集者直接以收集到的扰动集作为最终得到的隐私保护数据集,SCM-CLDP中用户扰动序列并上传给数据收集者,数据收集者在扰动集的基础上先学习 Markov 链模型,再利用 Markov 链模型合成数据集。具体来说,SCM-CLDP 将序列数据的收集转为数据集的生成问题。首先,由于序列长度分布不均,为后续扰动聚合方便,根据序列长度的估计分布确定了统一的长度,这一分布在后续的序列合成中起到了控制序列结束的作用。其次,针对用户端的扰动,一种直观的做法是直接单独逐一扰动序列中的每个单值,但是希望尽可能保留原始序列模式。因此还考虑了以多个连续值为一组的转移,并将以这两种扰动对象为基础的收集方法分别命名为 SCM-VP 和 SCM-TP。为了增强序列数据的真实性,在 SCM-TP 中单独处理了序列的起初数据和终止数据。然后,数据收集者在扰动集上学习 Markov 链模型,SCM-VP 中根据序列中每一位的数据分布来进一步构建 Markov 链模型,SCM-TP 中直接估计所有转移的分布并映射到 Markov 链模型。最后,结合序列长度的估计分布和 Markov 链模型来合成隐私保护的序列数据集。理论分析表明,SCM-VP 和 SCM-TP 具有良好的隐私保护性。为验证其性能,在相同的隐私保护程度下,在多个评价指标下进行了大量的实验,旨在确保合成的隐私保护序列数据集尽可能地接近原始数据集。实验结果充分证明了 SCM-CLDP 相比于现有方法的显著优越性,且在大多数情况下 SCM-VP 优于 SCM-TP。

1 相关工作

现有的序列数据的隐私保护工作多集中在传统的差分隐私和本地差分隐私,针对问题主要是频繁序列模式挖掘和数据合成,其中,数据合成方法不局限于某一具体数据分析任务,其适用场景更加广泛。Xu 等^[16]提出了满足差分隐私频繁序列模式挖掘方法(Differentially private frequent sequence mining via sampling-based candidate pruning, PFS²),设计了基于采样的候选剪枝技术,该技术是实现 PFS²的基础。Xia 等^[17]提出了满足 LDP 的频繁序列模式挖掘方法(Sequential patterns mining with LDP, LDPSPM),该方法中,各用户随机报告一组隐私保护的序列模式。Wang 等^[18]利用前缀树实现了满足 LDP 的频繁序列挖掘方法 PrivTrie。黄硕等^[19]基于现有的满足 LDP 的频繁项集挖掘算法^[20]提出 PrivSPM(Private sequential pattern mining),该方法旨在识别 top- k 频繁序列模式并估计其频率。现有的序列数据合成方法大多针对轨迹数据。Gursoy 等^[6]提出了一个可扩展的位置轨迹合成方法 AdaTrace,该方法具有可证明的隐私保护性、对特定攻击的抵抗力和优越的效用性,在轨迹合成中利用 Markov 模型来模拟实际的轨迹移动。Wang 等^[14]提出基于自适应 Markov 模型的满足 DP 的轨迹合成方法 PrivTrace。Du 等^[15]第一个提出了 LDP 下的轨迹合成解决方案 LDPTTrace,该方法中用户本地扰动轨迹移动模式,服务端同样利用 Markov 模型来模拟轨迹运动。Guner 等^[21]分析了现有方法的局限性,提出了在 LDP 下,从序列数据中学习 Markov 链模型的方法 PRIMA,该方法中各用户在本地上单独扰动序列中的单值,然后在服务端学习 Markov 链模型。

CLDP 下现存一种序列数据保护的工作, Gursoy 等^[11]提出了 CLDP 下的序列数据保护方法 Sequence-CLDP,同时保护序列的长度和内容。序列扰动时,利用停止扰动概率 halt 和生成概率 gen 两个参数来控制扰动序列的长度。对序列中的每一项,以 halt 的概率终止,以 $1 - \text{halt}$ 的概率扰动;若序列未达到预先设定的序列最大长度,则以 gen 的概率从 $\mathcal{X}$ 中随机选择一项填充序列,以 $1 - \text{gen}$ 的概率终止。Sequence-CLDP 逐位扰动序列中的每个值,为了保护序列长度,利用 halt 和 gen 来控制序列的结束和数据项的生成,但是这增加了额外的噪声,且数据收集者不做任何后处理,这些都降低了数据的效用性。除此之外,序列模式是序列数据的重要信息,Sequence-CLDP 在扰动过程中,直接单独扰动值且在

服务端未做任何后处理,这难以反映序列数据中的重要模式。

2 基础知识与问题描述

本文所提方法基于压缩的本地差分隐私,本节给出压缩的本地差分隐私的相关概念和本文针对的问题模型。

2.1 压缩的本地差分隐私

本节介绍压缩的本地差分隐私的基本概念、常用的扰动机制以及序列级别的定义。压缩的本地差分隐私是本地差分隐私的一种放松形式,在介绍压缩的本地差分隐私之前先介绍本地差分隐私。

定义 1 ϵ -本地差分隐私^[7]:扰动算法 K 满足 ϵ -本地差分隐私,当且仅当对任意两真实数据 v_1 和 v_2 有

$$\frac{\Pr[K(v_1)=t]}{\Pr[K(v_2)=t]} \leq e^\epsilon \quad \forall t \in \text{Range}(K) \quad (1)$$

式中: ϵ 表示隐私预算; $\text{Range}(K)$ 表示算法 K 的输出域; t 表示输出域中任意值。

LDP中,用户在本地扰动隐私数据,然后将扰动结果发送给服务器方,使得攻击者无法从传输数据中推断出用户的真实数据。LDP提供了严格的可证明的隐私保证,其中隐私预算 ϵ 控制着隐私保护强度。但是,LDP提供所有数据以相同程度的隐私保护,且大都适用于大规模数据,在小规模数据下表现不如意,针对上述问题,压缩的本地差分隐私应运而生。

定义 2 α -压缩的本地差分隐私^[11]:扰动算法 K 满足 α -压缩的本地差分隐私,当且仅当对任意两真实数据 v_1 和 v_2 有

$$\frac{\Pr[K(v_1)=t]}{\Pr[K(v_2)=t]} \leq e^{\alpha \cdot \text{dist}(v_1, v_2)} \quad \forall t \in \text{Range}(K) \quad (2)$$

式中: α 为隐私预算; $\text{dist}(v_1, v_2)$ 表示 v_1, v_2 间距离。

CLDP是LDP的一种基于度量(距离函数)的放松形式,距离的引入使得CLDP能更灵活调节隐私性和效用性间的平衡。隐私预算不变,数据对间距离越小,这一数据对扰动输出相同值的概率就越接近,这一数据对间的不可区分性就越强;相反,数据对间距离越大,这一数据对扰动输出相同值的概率就越疏远,这一数据对间的可区分性就越强。在相同程度的隐私保护下,由于距离的影响,CLDP的隐私参数 α 要远小于LDP的隐私参数 ϵ ^[11]。特别地,当任意数据对间距离固定时,CLDP就退化成LDP。

指数机制是满足CLDP的扰动机制,它是DP下指数机制的变体,其基本思想是:用户根据数据域中各值和真实值间的距离对各值打分,距离越小的分数越高,以较高的概率输出分数较高的值,输出概率与分数成正比。给定数据域 $\mathcal{X}$,对于任意真实值 $x \in \mathcal{X}$,扰动输出 $y \in \mathcal{X}$ 的概率为

$$\Pr[K_{\text{ExpM}}(x)=y] = \frac{e^{-\frac{\alpha}{2} \cdot \text{dist}(x, y)}}{\sum_{z \in \mathcal{X}} e^{-\frac{\alpha}{2} \cdot \text{dist}(x, z)}} \quad (3)$$

序列数据下的隐私保护可分为事件级别的和序列级别的,序列级别的隐私保护旨在保护整个序列而非单个事件,因此采用序列级别的CLDP定义。

定义 3 序列级别的 α -压缩的本地差分隐私^[11]:扰动算法 K 满足序列级别的 α -压缩的本地差分隐私,当且仅当对任意两真实序列数据 S_1 和 S_2 有

$$\frac{\Pr[K(S_1)=t]}{\Pr[K(S_2)=t]} \leq e^{\alpha \cdot \text{dist}_{\text{seq}}(S_1, S_2)} \quad \forall t \in \text{Range}(K) \quad (4)$$

式中 $\text{dist}_{\text{seq}}(S_1, S_2)$ 表示 S_1, S_2 间距离。

从上述定义中可以看出,相较于 α -CLDP,序列级别的只是扰动对象不同,原始的 CLDP 限制了单个数据值扰动到其他数据值的概率,而序列级别的 CLDP 限制了整个序列扰动到其他序列的概率。特别地,本文将序列间的距离函数定义为 $\text{dist}_{\text{seq}}(X, Y) = \sum_{i=1}^{|X|} \text{dist}(X[i], Y[i])$, 其中 $X[i]$ 和 $Y[i]$ 分别表示序列 X 和序列 Y 的第 i 项。

2.2 问题描述

假设有 n 个用户和 1 个数据收集者,用户 u 拥有 1 敏感的序列数据 $S_u = x_i \rightarrow x_j \rightarrow x_k \rightarrow \dots$, 其中 $x_i \in X = \{x_1, x_2, \dots, x_d\}$, $d = |x|$, $S_u[i]$ 表示用户 u 的序列数据的第 i 项。数据收集者收集所有用户的序列数据,并在此数据集上做统计分析任务,如收集个人轨迹做道路流量预测、收集个人网页浏览数据做网页推荐等。然而,这些数据中往往包含有隐私信息,比如从个人轨迹中可分析出家庭住址、工作场所,甚至健康状况。用户并不希望提交自己真实的敏感序列数据给数据收集者,因此,本文关注隐私保护的序列数据收集问题。

如图 1 所示,为保护隐私,用户 u 在本地使用扰动机制对 S_u 扰动,得到 $\tilde{S}_u$,攻击者即使知道 $\tilde{S}_u$ 也无法准确推断出 S_u 。收集者对所有扰动数据 $\tilde{D}$ 做聚合,得到隐私保护的序列数据集 $\hat{D}$ 。本文

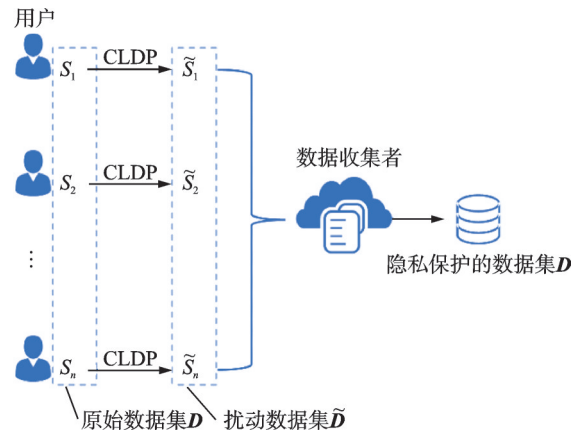


图 1 问题模型图

Fig.1 Problem model diagram

的目标是设计一个满足序列级别 CLDP 的序列数据收集方法,在保护隐私数据的同时,提供与原始序列数据集 D 高度相似的隐私保护的序列数据集 $\hat{D}$,其中,序列数据集的相似性从 Markov 模型的相似性、 X 上数据分布的相似性和频繁序列模型的相似性来综合衡量。

3 序列数据收集方法

3.1 方法概述

序列数据中各项按一定顺序排列,如轨迹序列数据中各位置点按照时间先后顺序排列,网页浏览数据中各网页按照用户浏览的先后顺序排列,文本数据中各字符、单词按照特定的顺序排列。Markov 链模型是分析研究序列数据的主要方法,它可以很好地表示各状态间的转移关系,本文提出的方法 SCM-CLDP 利用一阶 Markov 链模型来生成与原始序列数据集高度相似的隐私保护的序列数据集。

首先,收集序列长度。用户在本地图扰真实序列长度并提交给数据收集者,数据收集者估计长度分布,并选择一个合适的值作为后续序列扰动的最大长度。其次,收集序列数据,生成 Markov 链模型。用户在本地图扰序列数据,并将扰动结果发送给数据收集者。数据收集者根据所有用户的扰动结果学习 Markov 链模型。进一步地,依据不同的扰动对象设计了不同的序列数据收集方法,分别是基于值扰动的 SCM-VP 和基于转移扰动的 SCM-TP。最后,合成隐私保护的序列数据集。已知序列长度分布和该用户群体组成的序列数据集上的 Markov 链模型,数据收集者可根据此合成隐私保护的序列数据集,且每条序列数据的长度由长度分布和 Markov 链模型同时控制。

下面首先介绍序列长度的收集,然后分别介绍基于值扰动的序列数据收集方法 SCM-VP 和基于转移扰动的序列数据收集方法 SCM-TP,并在 SCM-VP 的介绍中给出序列数据集的生成算法。

3.2 序列长度收集

序列长度是序列数据集中的关键信息,利用长度分布可以更好地合成序列数据。假设序列长度所属的数据域公开已知,每个用户利用指数机制扰动其真实序列长度,其中,定义距离函数为绝对距离,数据收集者收集所有用户的扰动结果,然后利用期望最大化算法估计长度分布 L 。现已有不少工作利用期望最大化算法估计频率^[22-23],由于文章篇幅问题,本文不再详细介绍利用期望最大化算法做分布估计的具体步骤和相关特性。

由于每个用户提交的序列数据长短不一,为了简单起见,假设每个用户的序列数据长度为 l , l 是一个向用户和数据收集者公开参数。如果用户的序列数据长度大于 l ,则截断数据只扰动前 l 项,否则填充虚假项至序列长度为 l ,在后续的聚合估计过程中会忽略这些虚假项。 l 的取值影响着数据的效用性和算法的时间空间代价,随着 l 的增大,参与扰动的数据项增加,数据的效用性提高,但算法的时间空间代价也会随之增加。本文遵循了现有的 LDP 下集值数据^[15]、轨迹数据收集的设置^[24-25],将 l 设置为序列长度估计分布的 90% 分位数。

3.3 基于值扰动的序列数据收集方法

SCM-VP 中,用户的扰动对象是序列中的单项,已知序列中数据项所属数据域 $\mathcal{X}$,则序列 S 的扰动输出 $\tilde{S}$ 的概率为

$$\Pr[\tilde{S}[i]|S[i]] = \frac{e^{-\alpha \cdot \text{dist}(S[i], \tilde{S}[i])/2}}{\sum_{z \in \mathcal{X}} e^{-\alpha \cdot \text{dist}(S[i], z)/2}} \quad i = 1, 2, \dots, l \quad (4)$$

式(1)确保了序列中每个数据项都能以较高的概率扰动到离它本身较近的数据项,使得最终得到的扰动序列尽可能接近原始真实序列。数据收集者收集所有用户的扰动序列,首先采用期望最大化算法逐位估计 $\mathcal{X}$ 上各值的频率,再根据每位上的数据分布学习 Markov 链模型,最终基于 Markov 链模型和序列长度分布合成序列数据集。

令 P 表示估计频率矩阵, P_{ij} 表示序列第 i 位上值为 x_j 的估计频率,已知 P 可学习到 Markov 链模型。令 M 表示 Markov 链模型,状态集为 X , M_{ij} 表示从状态 $x_i \in X$ 到状态 $x_j \in X$ 的转移,则 M_{ij} 转移概率为序列中当前项是 x_i 且下一项是 x_j 的概率和,即

$$\Pr(M_{ij}) = \sum_{k=1}^{l-1} P_{ki} \cdot P_{(k+1)j} / \sum_{k=1}^{l-1} \sum_{t=1}^d P_{ki} \cdot P_{(k+1)t} \quad (5)$$

Markov 链模型中,初始状态为后续状态提供了初始条件,在诸如天气预测、股票价格预测和交通流量预测等应用中,初始状态的选择可能影响模型预测的准确性。初始状态和终止状态还有助于后续的数据合成。然而,式(5)仅反应了序列的内部转移,忽略了对序列开始和结束的建模,为此,单独考虑了初始转移和终止转移的收集。具体来说,增加一虚拟项 x_0 用来表示开始和结束,初始转移表示为 $\{M_{0i}: x_0 \rightarrow x_i, i = 1, 2, \dots, d\}$,类似地,终止转移表示为 $\{M_{i0}: x_i \rightarrow x_0, i = 1, 2, \dots, d\}$,那么, M_{0i} 转移概率为序列第一位是 x_i 的概率, M_{i0} 转移概率为序列最后一位是 x_i 的概率,即

$$\Pr[M_{0i}] = P_{1i}, \Pr[M_{i0}] = P_{li} \quad (6)$$

算法 1 给出了基于值扰动的 Markov 链模型学习算法,该算法中,输入为数据收集者收集到的所有用户的基于值扰动的序列扰动结果,输出为从扰动序列集中学习到的 Markov 链模型。已知 Markov 链模型和序列长度分布,可通过算法 2 合成序列数据集。首先采样序列长度得到 l' ,然后根据 M 来采样序列中的每个值,当采样的值为 x_0 或当前序列中已有 l' 个值时,则终止当前序列的合成,开始下一条序列

数据的合成。

算法1 基于值扰动的Markov链模型学习算法

输入:所有用户的扰动结果。

输出:Markov链模型 M 。

- (1) //期望最大化算法估计每一位上的频率分布
- (2) 初始化分布矩阵 P
- (3) for $t = 1$ to l do
- (4) 统计扰动序列中第 t 位上值为 x_i 的频率 $\tilde{\theta}_i$
- (5) while not converge do
- (6) E-step: $Q_i = \sum_{j=1}^d \tilde{\theta}_j \cdot \frac{\hat{\theta}_i \Pr[x_j | x_i]}{\sum_{k=1}^d \hat{\theta}_k \Pr[x_j | x_k]}, i = 1, 2, \dots, d$ // $\hat{\theta}_i$ 表示估计频率
- (7) M-step: $\hat{\theta}_i = \frac{Q_i}{\sum_{j=1}^d Q_j}$
- (8) end while
- (9) $P_{ti} = \hat{\theta}_i$
- (10) end for
- (11) //学习 Markov 链模型 M
- (12) for $i = 1$ to d do
- (13) $\Pr[M_{0i}] = P_{1i}, \Pr[M_{i0}] = P_{li}$
- (14) end for
- (15) for $i = 1$ to d do
- (16) $\Pr[M_{ij}] = \sum_{k=1}^{l-1} P_{ki} \cdot P_{k+1j}, j = 1, 2, \dots, d$
- (17) end for
- (18) 归一化频率的每一行
- (19) return M

算法2 序列数据集成算法

输入:序列长度估计分布 L , Markov 链模型 M 。

输出:序列数据集 $\hat{D}$ 。

- (1) for $u = 1$ to n do
- (2) 采样符合分布 L 的序列长度 l'
- (3) $k = 0$ //当前值为 x_k
- (4) for $i = 1$ to l' do
- (5) 以概率 $\Pr[M_{kj}]$ 采样
- (6) if $j = 0$ then //采样到虚拟项则结束
- (7) break

- (8) end if
- (9) $S'_u[i] \leftarrow x_j // S'_u[i]$ 表示用户 u 的合成序列
- (10) $k = j$
- (11) end for
- (12) $\hat{D}$ 中添加用户 u 的合成序列 S'_u
- (13) end for
- (14) return D

3.4 基于转移扰动的序列数据收集方法

在序列分析中序列的模式至关重要,因此希望尽可能保留原始的序列模式。为此,在SCM-TP中,不同于SCM-VP的值扰动,而是扰动序列中的转移模式。为了后续的描述方便,首先介绍转移集 T 。已知状态集大小为 d ,那么不考虑初始转移和终止转移共有内部转移 d^2 种转移,考虑初始转移 $\{M_{0i}: x_0 \rightarrow x_i, i = 1, 2, \dots, d\}$ 和终止转移 $\{M_{i0}: x_i \rightarrow x_0, i = 1, 2, \dots, d\}$ 共有 $d^2 + 2d$ 种转移。进一步地,将转移域定义为 $T = \{t_1, t_2, \dots, t_{|T|}\}$, 且 $|T| = d^2 + 2d$, t_{id+j} 表示转移 $x_i \rightarrow x_j$, T 上的距离函数定义为 $\text{dist}_T(t_{id+j}, t_{md+n}) = \text{dist}(x_i, x_m) + \text{dist}(x_j, x_n)$ 。由此,可进一步得到初始转移域 $T_{\text{start}} = \{t_1, t_2, \dots, t_d\}$ 、终止转移域 $T_{\text{end}} = \{t_{d^2+d+1}, t_{d^2+d+2}, \dots, t_{d^2+2d}\}$ 和内部转移域 $T_{\text{intra}} = \{t_{d+1}, t_{d+2}, \dots, t_{d^2+d}\}$ 。

用户扰动时,首先将长为 l 的序列映射到转移集 T ,其中 $|T| = l + 1$ 且 $T[i] \in T$; 然后采用指数机制分别扰动初始转移、内部转移和终止转移;最终提交扰动结果 $\tilde{T}$ 给数据收集者。下面,将分别介绍初始/终止转移和内部转移的扰动。

希望扰动不改变转移所属的子域,即初始转移的扰动域为 T_{start} , 终止转移的扰动域为 T_{end} , 内部转移的扰动域为 T_{intra} 。确定扰动域后,可得到转移的扰动概率。初始转移的扰动概率为

$$\Pr[\tilde{T}[1]|T[1]] = \frac{e^{-\alpha \cdot \text{dist}_T(T[1], \tilde{T}[1])/2}}{\sum_{t \in T_{\text{start}}} e^{-\alpha \cdot \text{dist}_T(T[1], t)/2}} \quad (7)$$

类似地,终止转移的扰动概率为

$$\Pr[\tilde{T}[l]|T[l]] = \frac{e^{-\alpha \cdot \text{dist}_T(T[l], \tilde{T}[l])/2}}{\sum_{t \in T_{\text{end}}} e^{-\alpha \cdot \text{dist}_T(T[l], t)/2}} \quad (8)$$

T 中内部转移有 $T[2], T[3], \dots, T[l-1]$, 它们的扰动概率为

$$\Pr[\tilde{T}[i]|T[i]] = \frac{e^{-\alpha \cdot \text{dist}_T(T[i], \tilde{T}[i])/2}}{\sum_{t \in T_{\text{intra}}} e^{-\alpha \cdot \text{dist}_T(T[i], t)/2}} \quad i = 2, 3, \dots, l-1 \quad (9)$$

数据收集者收集到所有用户的扰动结果,分别利用期望最大化算法估计初始转移、内部转移和终止转移的频率,然后将转移分布映射到Markov链模型,如算法3所示。已知Markov链模型和长度分布,利用算法2可合成序列数据集。

算法3 基于转移扰动的Markov链模型学习算法

输入:所有用户的扰动结果。

输出:Markov链模型 M 。

- (1) //期望最大化算法估计初始/终止转移、内部转移分布
- (2) 统计任意 $t_i \in T$ 的频率 $\tilde{\theta}_i$

- (3) //初始转移频率估计
- (4) while not converge do
- (5) E-step: $Q_i = \sum_{j=1}^d \tilde{\theta}_j \cdot \frac{\hat{\theta}_i \Pr[t_j | t_i]}{\sum_{k=1}^d \hat{\theta}_k \Pr[t_j | t_k]}, i = 1, 2, \dots, d$
- (6) M-step: $\hat{\theta}_i = \frac{Q_i}{\sum_{j=1}^d Q_j}, j = 1, 2, \dots, d$
- (7) end while
- (8) //内部转移频率估计
- (9) while not converge do
- (10) E-step: $Q_i = \sum_{j=d+1}^{d^2+d} \tilde{\theta}_j \cdot \frac{\hat{\theta}_i \Pr[t_j | t_i]}{\sum_{k=d+1}^{d^2+d} \hat{\theta}_k \Pr[t_j | t_k]}, i = d+1, d+2, \dots, d^2+d$
- (11) M-step: $\hat{\theta}_i = \frac{Q_i}{\sum_{j=d+1}^{d^2+d} Q_j}, i = d+1, d+2, \dots, d^2+d$
- (12) end while
- (13) //终止转移频率估计
- (14) while not converge do
- (15) E-step: $Q_i = \sum_{j=d^2+d+1}^{d^2+2d} \tilde{\theta}_j \cdot \frac{\hat{\theta}_i \Pr[t_j | t_i]}{\sum_{k=d^2+d+1}^{d^2+2d} \hat{\theta}_k \Pr[t_j | t_k]}, i = d^2+d+1, \dots, d^2+2d$
- (16) M-step: $\hat{\theta}_i = \frac{Q_i}{\sum_{j=d^2+d+1}^{d^2+2d} Q_j}, i = d^2+d+1, \dots, d^2+2d$
- (17) end while
- (18) //学习 Markov 链模型 M
- (19) for $i = 1$ to d do
- (20) $\Pr[M_{0i}] = \hat{\theta}_i, \Pr[M_{i0}] = \hat{\theta}_{d^2+d+i}$
- (21) end for
- (22) for $i = 1$ to d do
- (23) $\Pr[M_{ij}] = \hat{\theta}_{id+j}, j = 1, 2, \dots, d$
- (24) end for
- (25) 归一化 M 的每一行
- (26) return M

3.5 隐私性分析

SCM-VP 中用户单独扰动序列中的每一项。SCM-TP 中用户以序列中相邻两项组成的转移为单

位做扰动,考虑初始转移和终止转移,SCM-TP相当于做了两次值扰动。令SCM-VP的隐私预算为 α_{VP} ,SCM-TP的隐私预算为 α_{TP} ,下面将分别证明SCM-VP和SCM-TP满足序列级别的 α -CLDP,并给出在相同隐私保护下 α_{VP} 和 α_{TP} 间的关系。

定理1 SCM-VP算法满足序列级别的 α_{VP} -CLDP

证明:若要证明SCM-VP算法满足序列级别的 α_{VP} -CLDP,只需证明输入任意两长为 l 的序列 S_1 和 S_2 ,输出长为 l 的序列 S 的概率之比满足

$$\frac{\Pr[S|S_1]}{\Pr[S|S_2]} \leq e^{\alpha_{VP} \cdot \text{dist}_{\text{seq}}(S_1, S_2)} \quad (10)$$

证明

$$\begin{aligned} \frac{\Pr[S|S_1]}{\Pr[S|S_2]} &= \frac{\Pr[S[1]|S_1[1]]}{\Pr[S[1]|S_2[1]]} \cdot \frac{\Pr[S[2]|S_1[2]]}{\Pr[S[2]|S_2[2]]} \cdots \frac{\Pr[S[l]|S_1[l]]}{\Pr[S[l]|S_2[l]]} \leq \\ &= e^{\alpha_{VP} \cdot \text{dist}(S_1[1], S_2[1])} \cdot e^{\alpha_{VP} \cdot \text{dist}(S_1[2], S_2[2])} \cdots e^{\alpha_{VP} \cdot \text{dist}(S_1[l], S_2[l])} = e^{\alpha_{VP} \cdot \sum_{i=1}^l \text{dist}(S_1[i], S_2[i])} = e^{\alpha_{VP} \cdot \text{dist}_{\text{seq}}(S_1, S_2)} \end{aligned} \quad (11)$$

得证。SCM-VP算法满足序列级别的 α_{VP} -CLDP。

定理1给出了SCM-VP所满足的隐私性证明,由于SCM-TP每次扰动的对象是相邻两值组成的转移,因此给定隐私参数,不同于SCM-VP中实际隐私预算等于隐私参数,在SCM-TP中实际隐私预算是隐私参数的2倍。

定理2 SCM-TP算法满足序列级别的 $2\alpha_{TP}$ -CLDP

证明:若要证明SCM-TP算法满足序列级别的 α_{TP} -CLDP,只需证明输入任意两长为 l 的序列 S_1 和 S_2 ,输出长为 l 的序列 S 的概率之比满足

$$\frac{\Pr[S|S_1]}{\Pr[S|S_2]} \leq e^{2\alpha_{TP} \cdot \text{dist}_{\text{seq}}(S_1, S_2)} \quad (12)$$

SCM-TP中将长为 l 的序列转为长为 $l+1$ 的转移集上,因此有

$$\begin{aligned} \frac{\Pr[S|S_1]}{\Pr[S|S_2]} &= \frac{\Pr[T[1]|T_1[1]]}{\Pr[T[1]|T_2[1]]} \cdot \frac{\Pr[T[2]|T_1[2]]}{\Pr[T[2]|T_2[2]]} \cdots \frac{\Pr[T[l+1]|T_1[l+1]]}{\Pr[T[l+1]|T_2[l+1]]} \leq e^{\alpha_{TP} \cdot \text{dist}_T(T_1[1], T_2[1])} \times \\ &\quad e^{\alpha_{TP} \cdot \text{dist}_T(T_1[2], T_2[2])} \cdots e^{\alpha_{TP} \cdot \text{dist}_T(T_1[l+1], T_2[l+1])} = e^{\alpha_{TP} \cdot \text{dist}(S_1[1], S_2[1])} \times \\ &\quad e^{\alpha_{TP} \cdot (\text{dist}(S_1[1], S_2[1]) + \text{dist}(S_1[2], S_2[2]))} \cdots e^{\alpha_{TP} \cdot \text{dist}(S_1[l], S_2[l])} = e^{2\alpha_{TP} \cdot \sum_{i=1}^l \text{dist}(S_1[i], S_2[i])} = e^{2\alpha_{TP} \cdot \text{dist}_{\text{seq}}(S_1, S_2)} \end{aligned} \quad (13)$$

得证。SCM-TP算法满足序列级别的 $2\alpha_{TP}$ -CLDP。

通过定理1和理论2,可以进一步得到,当SCM-VP和SCM-TP均满足序列级别的 α -CLDP时,有 $\alpha_{VP} = 2\alpha_{TP} = \alpha$ 。实验中始终将 α_{TP} 设置为 $\alpha_{VP}/2$ 。

3.6 复杂度分析

在SCM-CLDP方法中,每个用户首先对其序列长度进行扰动,并将扰动结果发送给数据收集者。这里令 $|L|$ 表示序列长度区间大小。在本文中,所有扰动均采用指数机制。因此,序列长度扰动的复杂度为 $O(|L|)$ 。数据收集者在收集到所有用户的长度信息后,将利用期望最大化算法来估计长度分布。分布估计的时间复杂度主要取决于期望最大化算法的迭代次数和数据域的大小。用 r 来表示期望最大化算法的迭代次数,因此,长度分布估计的时间复杂度为 $O(|L|r)$ 。

数据收集者将长度分度的90%分位数 l 发送给各用户用于序列扰动。在序列扰动阶段,SCM-VP方法中,用户对序列中的 l 个项进行独立地扰动,每项的数据域大小为 d ,因此SCM-VP方法的扰动时

间复杂度为 $O(ld)$ 。而在 SCM-TP 方法中,用户对 $l-1$ 个内部转移以及起始和终止转移进行扰动,其中内部转移的数据域大小为 d^2 ,起始和终止转移的数据域大小为 d ,因此 SCM-TP 方法的总扰动时间复杂度为 $O(ld^2)$ 。

数据收集者在收集到所有用户的扰动序列后,将进行序列分布估计。在 SCM-VP 方法中,数据收集者逐位估计序列中每一位的数据分布,每位都需要调用一次期望最大化算法,因此总的分布估计时间复杂度为 $O(ldr)$ 。而在 SCM-TP 方法中,数据收集者直接估计转移分布,转移域大小为 $d^2 + 2d$,总的分布估计时间复杂度为 $O(d^2r)$ 。数据收集者将基于估计的序列分布生成 Markov 链模型,该过程的时间复杂度为 $O(d^2)$ 。最后利用 Markov 链模型和长度分布合成数据集,这一过程的时间复杂度为 $O(nl)$ 。

4 实验评估与验证

4.1 实验设置

本节将介绍仿真实验的设置,包括数据集设置和评价指标设置,实验环境是 Intel i5-10500 6 核 CPU、16 GB 内存、Windows11 系统,代码采用 Python 实现。

4.1.1 数据集

在两个公开数据集上进行实验,分别是 MSNBC 和 T-drive。MSNBC 数据集^[26]描述了 1999 年 9 月 28 日访问 msnbc.com 网站的用户页面访问数据。数据集中的每条序列对应一个用户在 24 h 内按时间顺序访问的页面。其中页面被分为 17 类,且被编号为 1~17。实验中,随机选取其中 100 000 条记录,序列平均长度为 4.68。T-drive 数据集^[27]包含了 2008 年 2 月 2 日至 2 月 8 日期间在北京市内的 10 357 辆出租车的 GPS 数据。实验中,只考虑经度范围在 $116.20^\circ \sim 116.60^\circ$,且纬度度范围在 $39.84^\circ \sim 40.04^\circ$ 的区域,并将该区域划分成 2×4 的网格,然后将出租车 GPS 数据映射到该网格上,网格单元编号为 1~8,最终得到的数据集中轨迹平均长度为 85.72。

4.1.2 评价指标

本文从 3 个方面来评估算法性能,分别是 Markov 链模型的效用性、合成数据集的效用性和频繁序列模式的效用性。下面将从这 3 个方面来分别介绍评价指标。

(1) Markov 链模型的效用性。Markov 链模型对合成数据集的生成至关重要,这里通过初始分布的误差和(内部)转移概率的误差来衡量 Markov 链模型的效用性。令 M 表示真实 Markov 链模型, $\hat{M}$ 表示隐私保护的 Markov 链模型,初始分布误差和转移概率误差定义如下。

(a) 初始分布误差 (Initial distribution error, IDE)。衡量隐私保护 Markov 链模型与真实模型在初始分布上的差异,计算方式为

$$\text{IDE} = \frac{1}{d} \sum_{i=1}^d \left(\Pr[M_{0i}] - \Pr[\hat{M}_{0i}] \right)^2 \quad (14)$$

(b) 转移概率误差 (Transition probability error, TPE)。衡量隐私保护 Markov 链模型与真实模型在转移概率上的差异,计算方式为

$$\text{TPE} = \frac{1}{d^2} \sum_{i=1}^d \sum_{j=1}^d \left(\Pr[M_{ij}] - \Pr[\hat{M}_{ij}] \right)^2 \quad (15)$$

(2) 合成数据集的效用性。希望合成数据集尽可能接近真实数据集,本文从数据集分布的误差和数据项排名的准确性来衡量合成数据集的效用性,其中数据项排名的准确性采用肯德尔 τ 系数来衡量。令 θ 表示真实分布且 θ_i 表示 x_i 的频率, $\hat{\theta}$ 表示合成数据集的分布且 $\hat{\theta}_i$ 表示合成数据集中 x_i 的频率,数据分布误差和肯德尔 τ 系数定义如下。

(a) 数据集分布误差 (Dataset distribution error, DDE)。衡量合成数据集与真实数据集中数据分布

的差异,计算方式为

$$\text{DDE} = \frac{1}{d} \sum_{i=1}^d (\theta_i - \hat{\theta}_i)^2 \quad (16)$$

(b) 肯德尔 τ 系数。衡量合成数据集中数据项排名的准确性。当且仅当 $\theta_i > \theta_j$ 且 $\hat{\theta}_i > \hat{\theta}_j$ 时, (x_i, x_j) 才被认为是和谐的数据对, 否则是不和谐的。令 N_c 表示和谐的数据对数, N_d 表示不和谐的数据对数, 那么有

$$\tau = \frac{N_c - N_d}{d(d-1)/2} \quad (17)$$

(3) 频繁序列模式的效用性。序列模式往往能反映数据之间的内在关系和规律, 这里考虑的序列模式是连续的有序的子序列, 并选择 top- k 频繁模式做进一步的度量。令真实数据集上 top- k 的频繁模式为 FP , 对应地, 合成数据集上的为 FP_{syn} , 利用频繁模式的 F_1 分数和频数误差来反映频繁序列模式的效用性。

(a) F_1 分数: 用于衡量二分类模型精确度的一种指标, 这里用于评估合成数据集中 top- k 频繁模式挖掘的准确性, 有

$$F_1 = 2 \times \frac{\text{Precision}(FP, FP_{\text{syn}}) \times \text{Recall}(FP, FP_{\text{syn}})}{\text{Precision}(FP, FP_{\text{syn}}) + \text{Recall}(FP, FP_{\text{syn}})} \quad (18)$$

(b) 模式频数误差 (Pattern frequency error, PFE): 衡量隐私保护的合成数据集中模式出现频数的差异。这里只考虑 top- k 的频繁模式, 统计真实的 top- k 频繁模式在真实数据集和合成数据集中频数的差异, 有

$$\text{PFE} = \frac{1}{|FP|} \sum_{fp \in FP} \frac{|c(fp) - c_{\text{syn}}(fp)|}{c(fp)} \quad (19)$$

式中: $c(fp)$ 和 $c_{\text{syn}}(fp)$ 分别表示在真实数据集和合成数据集中频繁模式 fp 的出现次数。特别地, 若 $fp \notin FP_{\text{syn}}$, 则 $c_{\text{syn}}(fp) = 0$ 。

如无特殊说明, 后文展示的所有结果为实验重复运行 20 次取平均值的结果。

4.1.3 对比方案

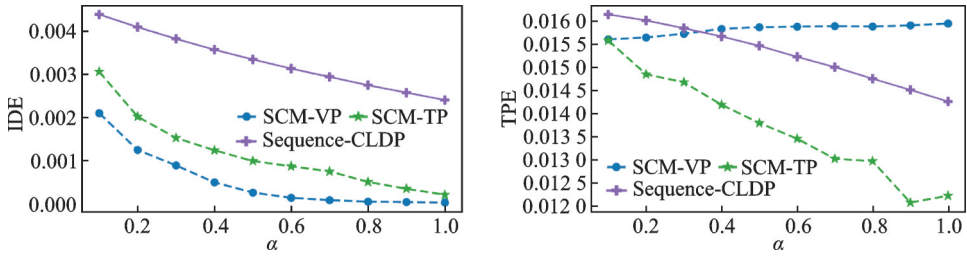
本文专注于满足压缩的本地差分隐私的序列数据收集问题, 而 Sequence-CLDP 是目前最先进的且唯一满足 CLDP 的序列数据收集方案, 因此将本文所提方法与 Sequence-CLDP^[11]进行了广泛且深入的对比, 以评估和比较本文提出的方法的性能。为了公平比较, Sequence-CLDP 同样只扰动序列的前 l 项, 且隐私预算设置与 SCM-VP 的相同。

4.2 实验结果与分析

4.2.1 Markov 链模型效用性分析

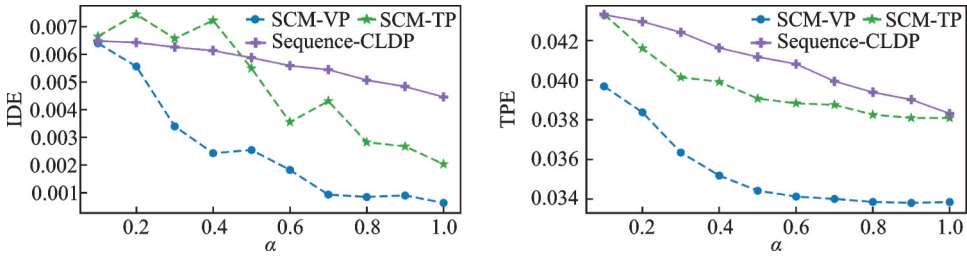
本节在不同数据集、不同隐私预算下对比了不同方案的 Markov 链模型效用性, 由于距离函数的影响, CLDP 中的隐私参数要远小于 LDP 中的隐私参数, 因此这里设置的隐私参数范围 $\alpha \in (0, 1)$ 要远小于 LDP 研究中常见的设置。实验中和现有的 Sequence-CLDP 方案进行了对比, 公平起见, 将 Sequence-CLDP 方案中的最大序列长度设置为 l , 若无特殊说明, 后续的实验均为这个设置。

MSNBC 上的实验结果如图 2 所示。从图 2 中可以清楚地看出在绝大多数情况下, 所有方案的误差都呈下降趋势。基于初始分布误差的实验结果, 可以发现 SCM-VP 的 IDE 最低, SCM-TP 次之, Sequence-CLDP 最差, SCM-TP 的 IDE 高于 SCM-VP 是因为 SCM-TP 中初始状态的实际隐私预算为 SCM-VP 的 $1/2$ 。此外, 随着 α 的增大, SCM-TP 的效果逐渐逼近 SCM-VP。在评估转移概率的误差时, 观察到当 $\alpha > 0.1$ 时, SCM-TP 的 TPE 显著低于 SCM-VP 和 Sequence-CLDP, 且随着 α 的增大, 与其他方案的差异明显增大。除此之外, 注意到 SCM-VP 的 TPE 受 α 的影响较小, 甚至随着 α 的增大而轻

图2 MSNBC数据集上初始分布误差和转移概率误差在不同 α 下的变化Fig.2 Variation of initial distribution error and transition probability error with α on the MSNBC dataset

微增大,并且当 $\alpha > 0.3$ 时,SCM-VP的效果要劣于Sequence-CLDP。

T-drive上的实验结果如图3所示。从图3中,可以发现在所有情况下,所有方案的误差都呈下降趋势,且SCM-VP始终表现最优。此外,综合考虑IDE和TPE,在绝大多数情况下,SCM-VP表现最优,其次是SCM-TP,而Sequence-CLDP表现最差。特别地,当 $\alpha \leq 0.4$ 时,Sequence-CLDP的IDE要略微低于SCM-TP,这是因为T-drive的数据集规模较小,且实际分配给SCM-TP的隐私预算是其他两方案的1/2,在隐私预算较小时无法有效弥补扰动带给SCM-TP的消极影响。与图2相比,发现MSNBC数据集上的实验结果要整体好于T-drive,且在TPE指标上,SVM-TP在MSNBC数据集上的效果明显更优,甚至远优于SVM-VP和Sequence-CLDP,猜测这是因为数据规模对于SVM-TP的影响更大,随着数据规模的增大,SVM-TP中转移分布估计的准确性随之提高,而转移分布直接且密切影响着TPE,进而TPE随之降低。

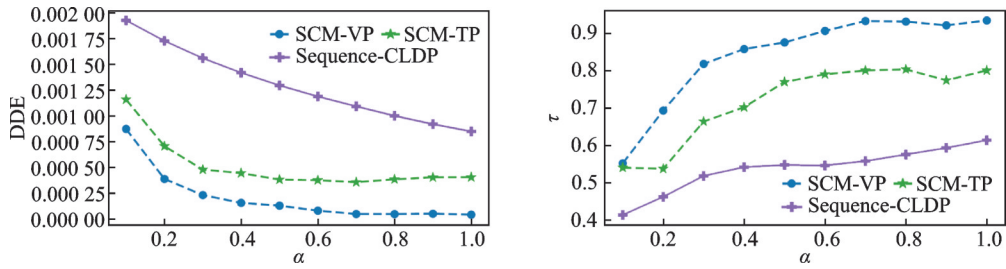
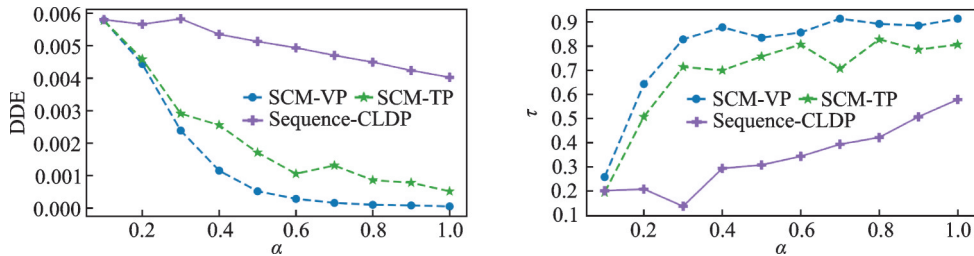
图3 T-drive数据集上初始分布误差和转移概率误差在不同 α 下的变化Fig.3 Variation of initial distribution error and transition probability error with α on the T-drive dataset

4.2.2 合成序列数据集效用性分析

下面分析利用Markov链模型合成的隐私保护数据集的效用性,图4和图5分别描述了在MSNBC数据集和T-drive数据集上的实验结果。由图4、5可以发现:在所有数据集上,随着 α 的增大,所有方案的DDE均呈下降趋势,所有方案的 τ 均呈上升趋势,并且在所有情况下SCM-VP表现最优,SCM-TP次之,Sequence-CLDP表现最差。特别地,当 $\alpha = 0.1$ 时,三者间的差距较小,在T-drive数据集上,三者的DDE和 τ 几乎重合,随着 α 的增大,SCM-CLDP方法的优势显著显现出来。结合图2,SVM-TP的转移概率误差最小,但是这一优势在合成序列数据集上并没有显现出来,这是因为初始状态定义了过程的起点,为后续状态提供了初始条件,合成序列数据是初始分布和内部转移概率共同作用的结果。

4.2.3 频繁序列模式效用性分析

这里分析了隐私保护的数据集上的频繁序列模式的效用性,实验中选取top-25的频繁模式进行评估。从图6中发现在MSNBC数据集上,SVM-TP的 F_1 分数要远高于SCM-VP和Sequence-CLDP,但是,它的频繁模式频数误差要远高于SCM-VP,这是因为虽然SCM-VP的top-25序列模式的命中率较

图4 MSNBC数据集上数据集分布误差和肯德尔 τ 系数在不同 α 下的变化Fig.4 Variation of dataset distribution error and Kendall's τ coefficient with α on the MSNBC dataset图5 T-drive数据集上数据集分布误差和肯德尔 τ 系数在不同 α 下的变化Fig.5 Variation of dataset distribution error and Kendall's τ coefficient with α on the T-drive dataset

低,但是它有着更准确的较频繁的若干序列模式频数。另外,还发现 F_1 分数上3种方案的表现和图2中TPE的一致:SCM-TP的表现始终最优,SCM-VP的 F_1 分数和TPE都随着 α 的增大而变差,当 α 较大时,SCM-VP开始劣于Sequence-CLDP。这表明 F_1 分数和内部转移概率有着密切关联。

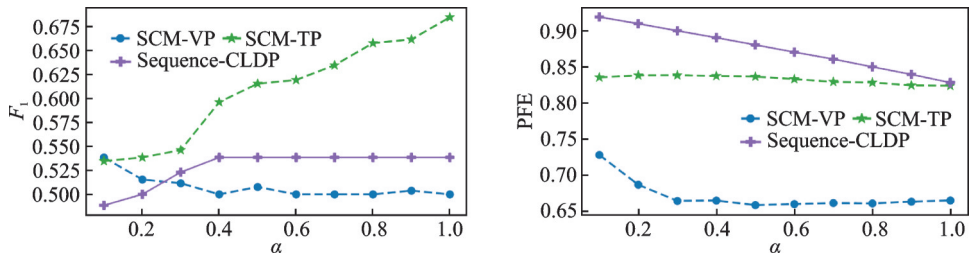
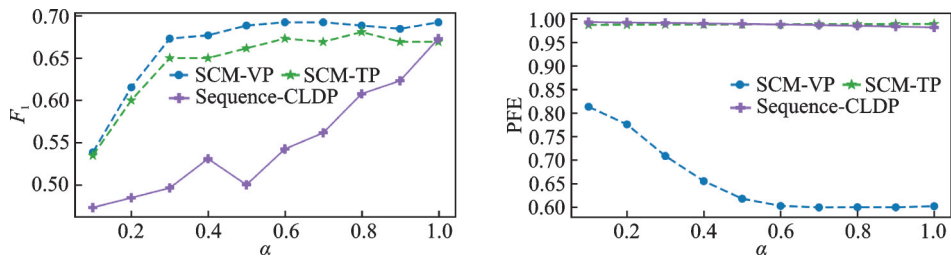
图6 MSNBC数据集上频繁模式的 F_1 分数和频数误差在不同 α 下的变化Fig.6 Variation of F_1 score and frequency error of frequent pattern with α on the MSNBC dataset

图7展示了T-drive上的实验结果,可以发现这3个方案的 F_1 分数的变化趋势与前文中的实验保持一致,即随着 α 的增大,效果逐渐变好,且SCM-VP表现最好,Sequence-CLDP表现最差。SCM-TP的

图7 T-drive数据集上频繁模式的 F_1 分数和频数误差在不同 α 下的变化Fig.7 Variation of F_1 score and frequency error of frequent pattern with α on the T-drive dataset

PFE表现与在MSNBC数据集上的一致,即虽然能挖掘出较多的准确序列模式,但是这些序列模式的频数要远远低于真实频数,在数据集规模较小时表现得尤为明显。

5 结束语

针对目前基于压缩的本地差分隐私下的序列数据收集方法忽略序列模式,且效用性仍有提升空间的问题,本文提出了SCM-CLDP方法。根据扰动对象的不同,该方法又细分为SCM-VP和SCM-TP。其中,SCM-VP是针对值扰动的序列数据收集方法,SCM-TP是针对转移扰动的序列数据收集方法,这两种方法中都需要先估计数据值或转移分布再学习Markov模型,最终由Markov模型合成数据集。首先理论证明了SCM-VP和SCM-TP满足序列级别的CLDP,然后在两个真实数据集、三方面的评价指标上和现有方案进行了广泛且深入的实验对比,实验表明SCM-CLDP在各个方面都有着显著的优越性,且在大多数情况下,SCM-VP要优于SCM-TP。本文利用一阶Markov链模型来提取序列模式,未来希望能够通过更高阶的Markov链模型以及除序列长度外的其他重要信息来提取更丰富的序列模式,以此提升合成序列集的准确性,并探究其在各种真实场景下的应用。

参考文献:

- [1] 吉根林,赵斌.时空轨迹大数据模式挖掘研究进展[J].数据采集与处理,2015,30(1):47-58.
JI Genlin, ZHAO Bin. Research progress in pattern mining for big spatio-temporal trajectories[J]. Journal of Data Acquisition and Processing, 2015, 30(1): 47-58.
- [2] 张逸凡,赵斌,孙鸿艳,等.基于时空轨迹的移动对象汇聚模式挖掘算法[J].数据采集与处理,2018,33(3):487-495.
ZHANG Yifan, ZHAO Bin, SUN Hongyan, et al. Algorithm for mining converging patterns of moving objects from spatiotemporal trajectories[J]. Journal of Data Acquisition and Processing, 2018, 33(3): 487-495.
- [3] 夏润泽,李丕绩.ChatGPT大模型技术发展与应用[J].数据采集与处理,2023,38(5):1017-1034.
XIA Runze, LI Piji. Large language model ChatGPT: Evolution and application[J]. Journal of Data Acquisition and Processing, 2023, 38(5): 1017-1034.
- [4] 全国人民代表大会常务委员会.中华人民共和国数据安全法[EB/OL].(2021-09-01).http://www.gov.cn/xinwen/2021-06/11/content_5616919.htm.
- [5] 全国人民代表大会常务委员会.中华人民共和国个人信息保护法[EB/OL].(2021-11-01).http://www.gov.cn/xinwen/2021-08/20/content_5632486.htm.
- [6] GURSOY M E, LIU L, TRUEX S, et al. Utility-aware synthesis of differentially private and attack-resilient location traces [C]//Proceedings of the 2018 ACM SIGSAC Conference on Computer and Communications Security. [S.l.]: ACM, 2018: 196-211.
- [7] DUCHI J C, JORDAN M I, WAINWRIGHT M J. Local privacy and statistical minimax rates[C]//Proceedings of 2013 IEEE 54th Annual Symposium on Foundations of Computer Science. [S.l.]: IEEE, 2013: 429-438.
- [8] DWORK C. Differential privacy[C]//Proceedings of International Colloquium on Automata, Languages, and Programming. Berlin, Heidelberg: Springer, 2006: 1-12.
- [9] ERLINGSSON Ú, PIHUR V, KOROLOVA A. RAPPOR: Randomized aggregatable privacy-preserving ordinal response [C]//Proceedings of the 2014 ACM SIGSAC Conference on Computer and Communications Security. [S.l.]: ACM, 2014: 1054-1067.
- [10] ZDZIARSKI J. Apple's differential privacy is about collecting your data—but not your data[EB/OL].(2016-06-14).<https://www.oreilly.com/radar/apples-differential-privacy-collecting-data/>.
- [11] GURSOY M E, TAMERSOY A, TRUEX S, et al. Secure and utility-aware data collection with condensed local differential privacy[J]. IEEE Transactions on Dependable and Secure Computing, 2019, 18(5): 2365-2378.
- [12] MOONEY C H, RODDICK J F. Sequential pattern mining—approaches and algorithms[J]. ACM Computing Surveys (CSUR), 2013, 45(2): 1-39.

- [13] CHEN R, ACS G, CASTELLUCCIA C. Differentially private sequential data publication via variable-length N-grams[C]// Proceedings of the 2012 ACM Conference on Computer and Communications Security. [S.l.]: ACM, 2012: 638-649.
- [14] WANG H, ZHANG Z, WANG T, et al. PrivTrace: Differentially private trajectory synthesis by adaptive Markov models [C]//Proceedings of the 32nd USENIX Security Symposium (USENIX Security 23). [S.l.]: USENIX Association, 2023: 1649-1666.
- [15] DU Y, HU Y, ZHANG Z, et al. LDPTTrace: Locally differentially private trajectory synthesis[J]. Proceedings of the VLDB Endowment, 2023, 16(8): 1897-1909.
- [16] XU S, CHENG X, SU S, et al. Differentially private frequent sequence mining[J]. IEEE Transactions on Knowledge and Data Engineering, 2016, 28(11): 2910-2926.
- [17] XIA H, HUANG W, XIONG Y, et al. Mining frequent sequential patterns with local differential privacy[J]. International Journal of Network Security, 2021, 23(5): 817-829.
- [18] WANG N, XIAO X, YANG Y, et al. PrivTrie: Effective frequent term discovery under local differential privacy[C]// Proceedings of 2018 IEEE 34th International Conference on Data Engineering (ICDE). [S.l.]: IEEE, 2018: 821-832.
- [19] 黄硕, 李艳辉, 曹建秋. 本地化差分隐私下的频繁序列模式挖掘算法 PrivSPM[J]. 计算机应用, 2023, 43(7): 2057-2064.
HUANG Shuo, LI Yanhui, CAO Jianqiu. PrivSPM: Frequent sequential pattern mining algorithm under local differential privacy[J]. Journal of Computer Applications, 2023, 43(7): 2057-2064.
- [20] WANG S, HUANG L, NIE Y, et al. PrivSet: Set-valued data analyses with locale differential privacy[C]//Proceedings of IEEE INFOCOM 2018 IEEE Conference on Computer Communications. [S.l.]: IEEE, 2018: 1088-1096.
- [21] GUNER E, GURSOY M E. Learning Markov chain models from sequential data under local differential privacy[C]// Proceedings of European Symposium on Research in Computer Security. Cham: Springer Nature Switzerland, 2023: 359-379.
- [22] LI Z, WANG T, LOPUHAÄ-ZWAKENBERG M, et al. Estimating numerical distributions under local differential privacy [C]//Proceedings of the 2020 ACM SIGMOD International Conference on Management of Data. [S.l.]: ACM, 2020: 621-635.
- [23] FANTI G, PIHUR V, ERLINGSSON Ú. Building a RAPPOR with the unknown: Privacy-preserving learning of associations and data dictionaries[J]. Proceedings on Privacy Enhancing Technologies (PoPETS), 2016(3): 41-61.
- [24] WANG T, LI N, JHA S. Locally differentially private frequent itemset mining[C]//Proceedings of 2018 IEEE Symposium on Security and Privacy (SP). [S.l.]: IEEE, 2018: 127-143.
- [25] QIN Z, YANG Y, YU T, et al. Heavy hitter estimation over set-valued data with local differential privacy[C]//Proceedings of the 2016 ACM SIGSAC Conference on Computer and Communications Security. [S.l.]: ACM, 2016: 192-203.
- [26] HECKERMAN D. MSNBC.com anonymous web data, UCI machine learning repository[EB/OL]. (1999-03-09). <https://doi.org/10.24432/C5390X>.
- [27] YUAN J, ZHENG Y, ZHANG C, et al. T-drive: driving directions based on taxi trajectories[C]//Proceedings of the 18th SIGSPATIAL International Conference on Advances in Geographic Information Systems. [S.l.]: Association for Computing Machinery, 2010: 99-108.

作者简介:



金严(2000-),女,硕士研究生,研究方向:隐私保护和数据安全,E-mail:jinyan1618@nuaa.edu.cn。



朱友文(1986-),通信作者,男,博士,教授,博士生导师,研究方向:应用密码学、人工智能安全、隐私计算,E-mail: zhuyw@nuaa.edu.cn。



吴启晖(1970-),男,博士,教授,博士生导师,研究方向:认知信息论、电磁空间频谱智能管控、天地一体化信息网络、无人机集群智能通信和安全等。

(编辑:刘彦东)