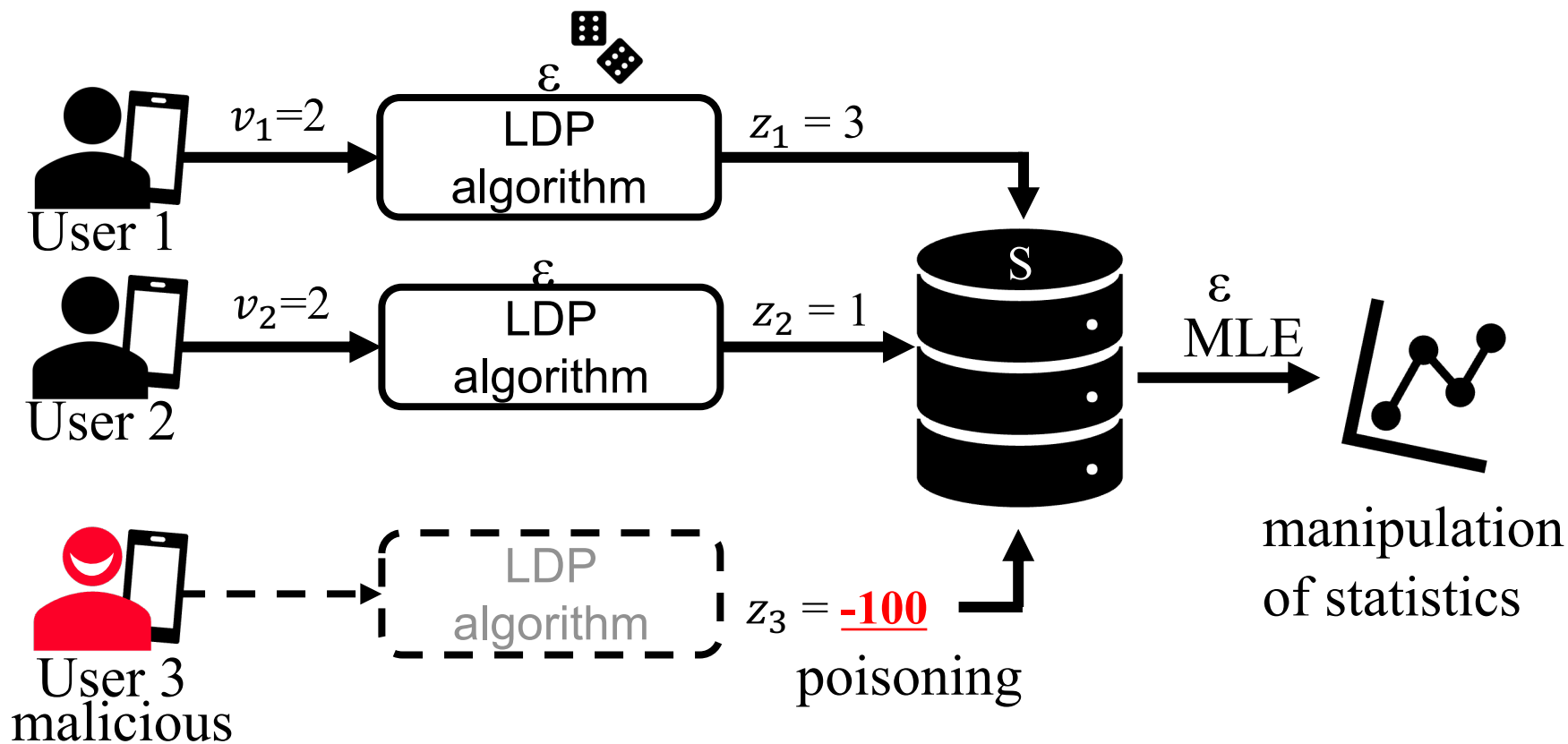


Harmless Opening for Continuous-Valued Local Differential Privacy

Hiroaki Kikuchi
Meiji University, Japan

DPM 2026, Roma

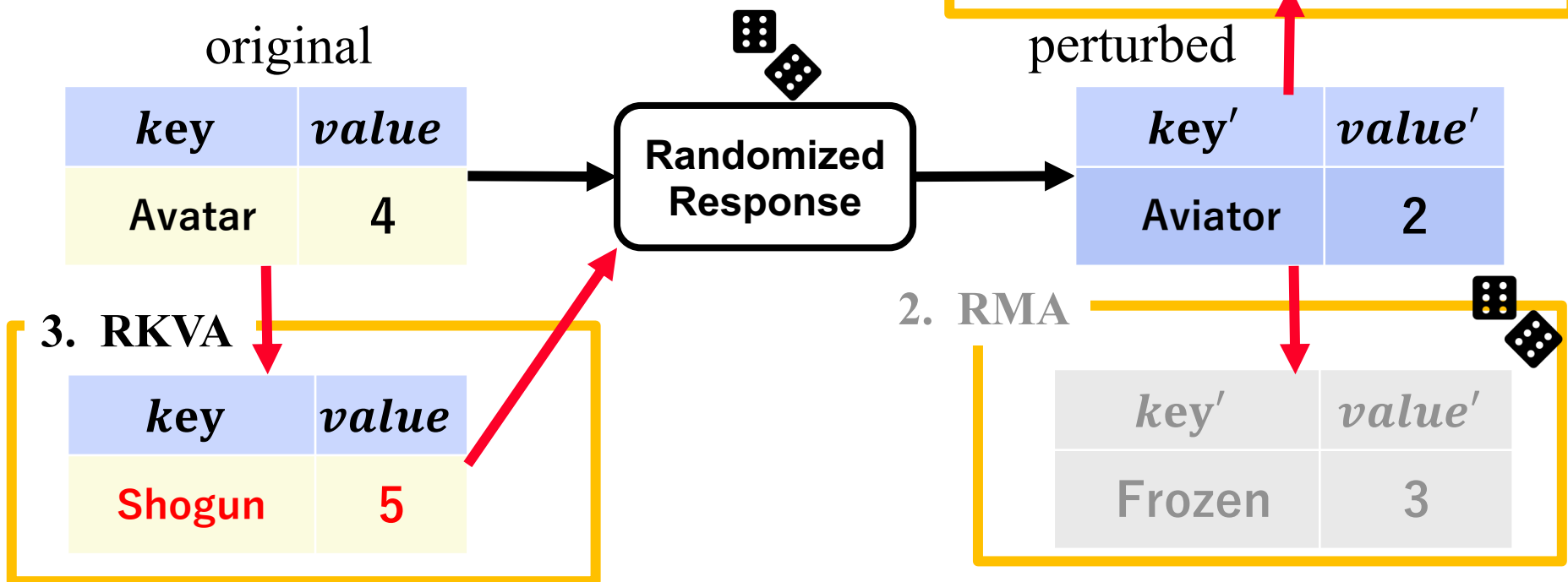
Issue of LDP: *Poisoning*



Poisoning attacks

Input Poisoning (IPA)

Output Poisoning (OPA)



Related Works

- Poisonings

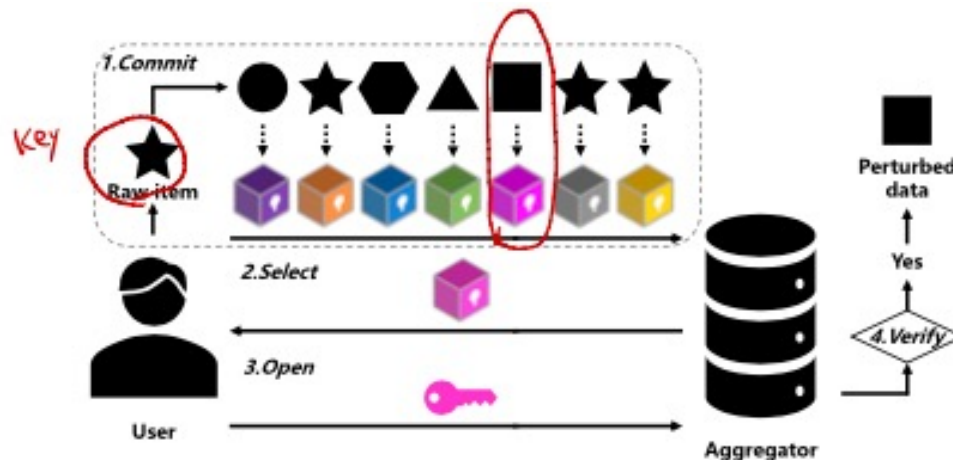
- [7] Cheu S&P21
- [8] Cao Usenix 21
- [9] Wu Usenix 22

- Robust LDP protocols

- [10] Zhao, S&P 25, RobustLDP
- Ye, S&P 25, Subset Counting Query

- Cryptographical approach

- [6] Kato, DBSec 2021
- [1] Song, IEEE Trans. IFS 2023.



“Harmless Opening”

■ VGRR (Verifiable GRR)

- Commit ℓ_1 true value and ℓ_2 fake values
- $\ell = \ell_1 + (d-1)\ell_2$
- Example $\ell=10$, $d = 3$
 - » $\Pr[x = 1] = 6/10$
 - » $\Pr[x = 2] = \Pr[x = 3] = 2/10$
 - » Aggregator can not distinguish $x = 1$ and $x = 2, 3$



VGRR (Verified GRR)

■ Pros:

- VGRR is $(\epsilon' = \ln \frac{\ell_1}{\ell_2})$ -LDP
- VGRR is **robust against OPA**
 - » robust iff $\Pr[S_{\max} = S_{\text{in}}] > 1 - \text{negl}(k)$
 - » $S_{\text{in}} = \text{IPA}$, $S_{\max} = \text{OPA}$
- VGRR is lightweight (no ZK proof)

■ Cons:

- deals with only discrete values (frequency estimation)

Our work

■ Goal

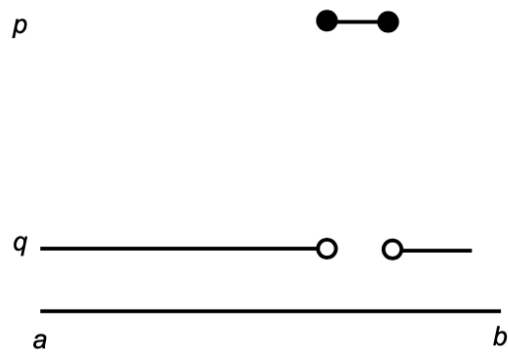
□ Robust LDP for continuous value

Task	LDP	Poisoning	Mitigation	
			probabilistic	cryptographical
Frequency estimation (categorical value)	GRR, OUE, OLH, Rappor	[7] Cheu [8] Cao [9] Wu	[8], [12]	[1] Song, [6] Kato
Mean estimation (continuous value)	[3] SR [4] PM	[10] Zhao	[10] Zhao	Our work

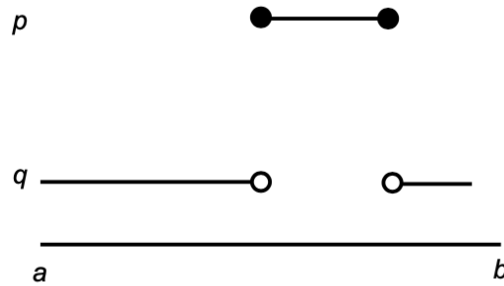
Continuous-valued LDP

■ LDP

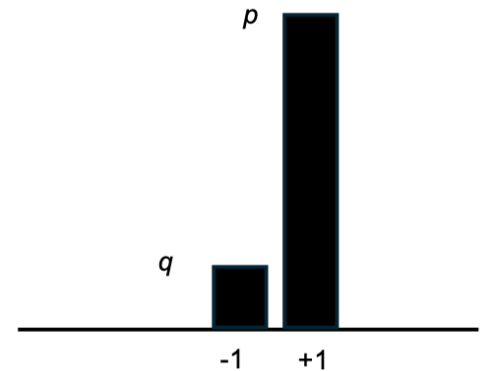
- [2] GRR (Generalized Randomized Response)
- [3] SR (Stochastic Rounding)
- [4] PM (Pricewise Mechanism)



1. GRR



3. PM



2. SR

SR (Stochastic Rounding)

- Input: x

- x in $[a, b] \rightarrow x_{\sim}$ in $[-1, 1]$

- » $x_{\sim} = -1 + (x-a)*2/(b-a)$

- Perturbation: x'

$$\Pr[\Psi_{SR(\epsilon)}(\tilde{x}) = x'] = \begin{cases} q + \frac{(p-q)(1-\tilde{x})}{2}, & \text{if } x' = -1 \\ q + \frac{(p-q)(1+\tilde{x})}{2}, & \text{if } x' = 1 \end{cases}$$

$$p = \frac{e^{\epsilon}}{1+e^{\epsilon}} \text{ and } q = 1 - p$$

- Unbiased estimator

$$\mathbb{E}\left(\frac{x'}{p-q}\right) = \tilde{x}.$$

PM (Piecewise Mechanism)

- Input

- x in $[a, b] \rightarrow \tilde{x}$ in $[-1, 1]$

- Perturbation

- x' in $[-s, +s]$

$$\Pr[\Psi_{PM(\epsilon)}(\tilde{x}) = x'] = \begin{cases} \frac{e^{\epsilon/2}(e^{\epsilon/2}-1)}{2(e^{\epsilon/2}+1)}, & \text{if } x' \in [l(\tilde{x}), r(\tilde{x})] \\ \frac{e^{\epsilon/2}-1}{2(e^{\epsilon/2}+e^{\epsilon})}, & \text{otherwise} \end{cases},$$

where $-s \leq l(\tilde{x}) < r(\tilde{x}) \leq s$, $l(\tilde{x}) = \frac{e^{\epsilon/2}\tilde{x}-1}{e^{\epsilon/2}-1}$ and $r(\tilde{x}) = \frac{e^{\epsilon/2}\tilde{x}+1}{e^{\epsilon/2}-1}$.

- Unbiased Estimator

- $E[x'] = \tilde{x}$

Threat model

- n honest users
- m malicious users
 - rate $\beta = m/(n+m)$
 - gain = target mean (**max**) – true mean
 - IPA (input = targeted value, with LDP)
 - OPA (output = targeted value, without LDP, **leaves if $x < \theta$**)

Proposed Method

■ Idea

- **binning** (making into discrete bins) and applies Harmless Opening (VGRR)

base		Harmless Opening	
LDP	encoding	Binning	
GRR	$D \rightarrow D$	$D \rightarrow d$ values $\{0, \dots, d-1\}$	VGRR-HO
SR	$[a, b] \rightarrow \{-1, +1\}$	$[a, b] \rightarrow$ d values	SR-HO
PM	$[a, b] \rightarrow [-s, +s]$	$[-s, +s] \rightarrow$ d values	PM-HO

Example 1) GRR

- 1. Commit

- $d = 5$, x in $\{1,2,3,4,5\}$, $x^\sim$ in $\{-1, -1/2, 0, 1/2, 1\}$,
 $\varepsilon = \log 2$

- $x = 4$, $x^\sim = 1/2$, $p = e^\varepsilon/(e^\varepsilon+d-1) = 2/6$, $q = 1/6$

1	2	2	3	4	5
C(1)	C(2)	C(2)	C(3)	C(4)	C(5)

- 2. A picks $j = 2$ in $\{1,..,6\}$
- 3. U opens $c(1)$ $c(2)$ $c(3)$ $c(4)$ $c(5)$
- 4. A checks C and $c(1)$ $c(2)$ $c(3)$ $c(4)$ $c(5)$

Example 2) SR-HO

- $D = \{0, 1, 2, 3, 4\}$, $d = 5$, $\varepsilon = \ln 2$, $\ell = 12$

x	$\sim x$	P^*	ℓ_{-1}	ℓ_1
0	-1	1/3	4	8
1	-1/2	5/12	5	7
2	0	1/2	6	6
3	1/2	7/12	7	5
4	1	2/3	8	4

x	1	2	3	4	5	6	7	8	9	10	11	12
0	C(-1)	C(-1)	C(-1)	C(-1)	C(1)	C(1)	C(1)	C(1)	C(1)	C(1)	C(1)	C(1)
1	C(-1)	C(-1)	C(-1)	C(-1)	C(-1)	C(1)	C(1)	C(1)	C(1)	C(1)	C(1)	C(1)

□ HO proves that $4/12 \leq p^* \leq 8/12$ for all cases

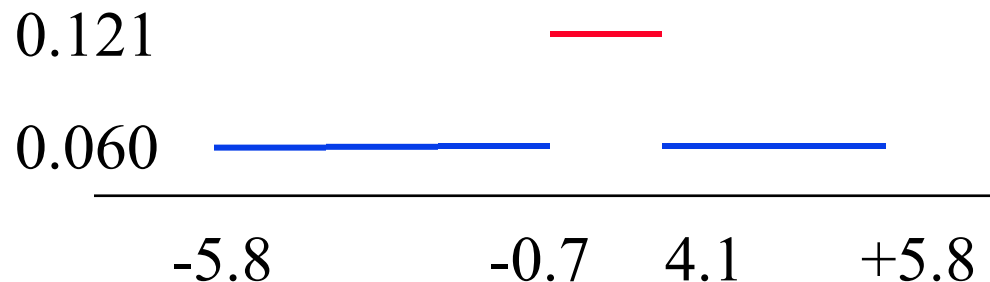
Example 3) PM-HO

■ Commit

□ $d = 5$, x in $\{0, 1, 2, 3, 4\}$, $\sim x$ in $\{-1, -1/2, 0, 1/2, 1\}$, $\varepsilon = \log 2$,

□ $x = 3$, $x^\sim = 1/2$, $s = (e^{\varepsilon/2} + 1) / (e^{\varepsilon/2} - 1) = 5.8$

□ $\ell(x) = 2x - 1 / (2 - 1) = -0.7$, $r(x) = 4.1$



Questions

- RQ1. Does Harmless Opening mitigates poisoning attacks?
- RQ2. How much accuracy reduction does come with the binning?
- RQ3. Which is the best LDP scheme in terms of accuracy and computational costs?



Security Analysis

■ Privacy

- GRR-HO, SR-HO and PM-HO satisfies ϵ -LDP.

■ Robustness

- Theorem 1 GRR-HO is robust

$$\gg \Pr[S_{\max} = S_{\text{in}}] > 1 - \text{negl}(k)$$

- Proposition SR-HO is semi-robust?

$$\gg \Pr[S_{\max} > S_{\text{in}}] > 1 - \text{negl}(k)$$

bogus

Security of SR-HO

- Lemma 1. Any report accepted by the Harmless opening satisfies $E\left[\frac{x'}{p-q}\right] \leq 1$

for **mean-maximizing** attack.

- Mean-maximizing adversary would like to commit +1 value, which is limited at most ℓ_2 (ℓ_{-1}) times in HO verification process.
- Theorem 2. SR-HO is robust for mean estimation, i.e., gain $G_{\max} = G_{\text{IPA}}$.

Evaluation

- Dataset

- UCI Adult dataset , Age [$a = 17$, $b = 90$]

- Tasks

- Mean estimation

- Metrics

- Accuracy: MSE

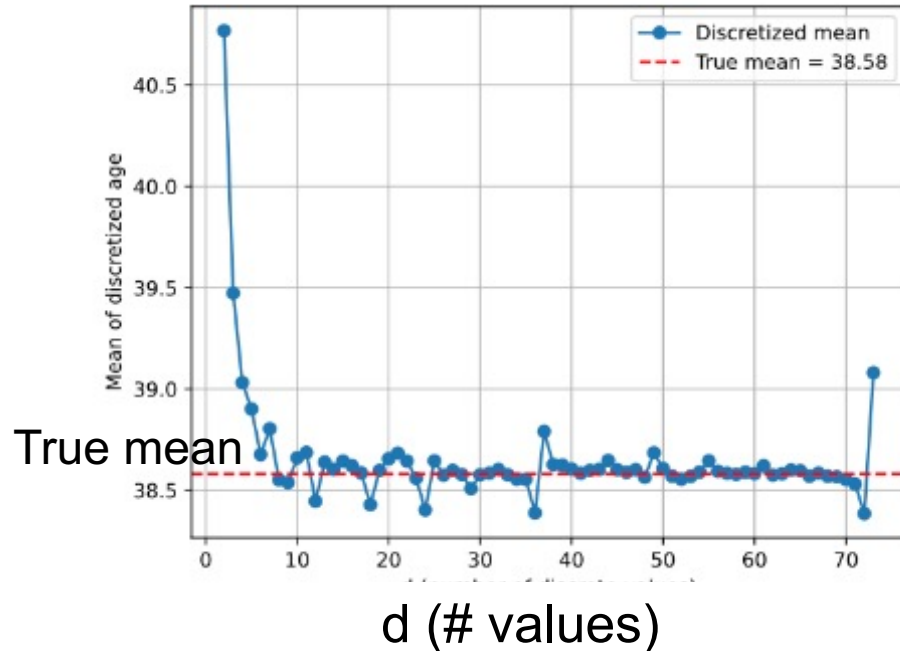
- Robustness: gain

- gain = estimate with poisoning – estimate

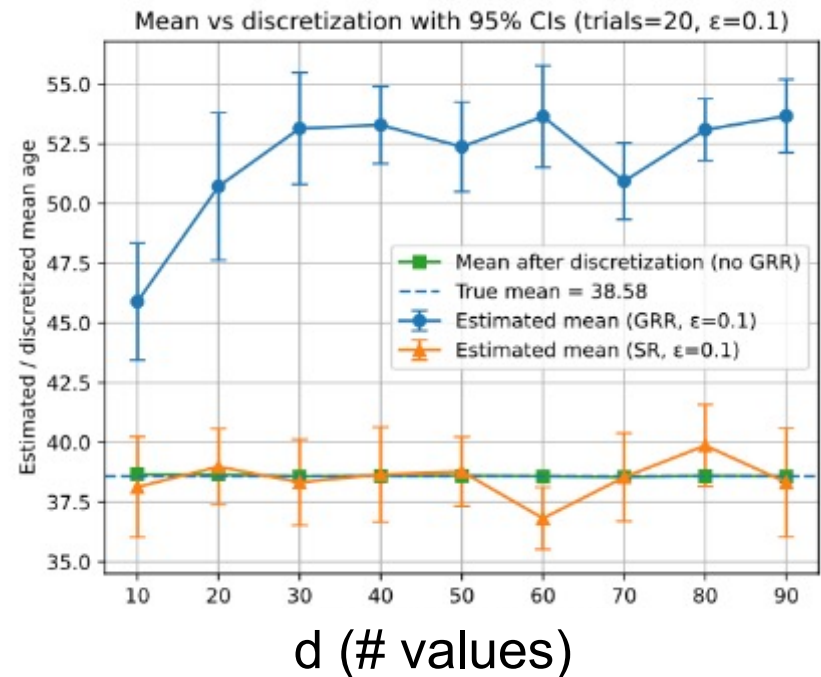
Result 1. binning

Without perturbation

Estimated mean

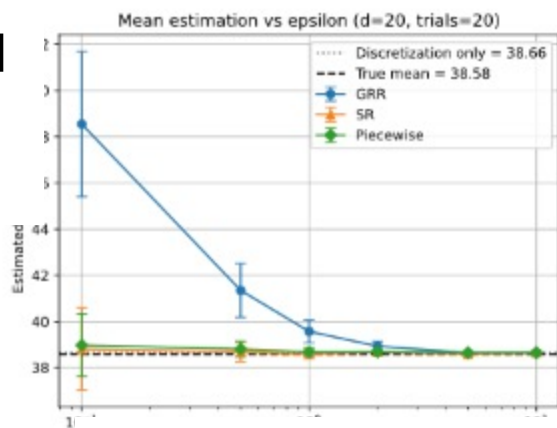


LDP (GRR, SR)

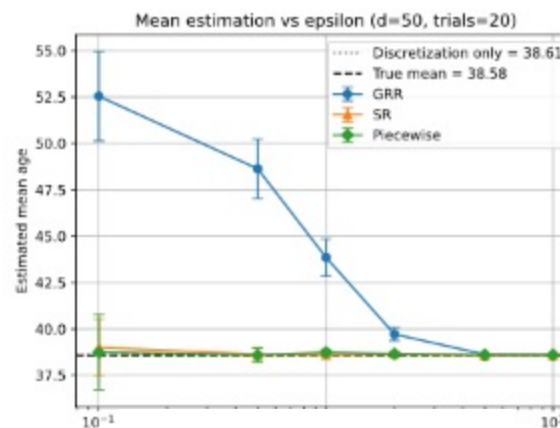


Result 2. Estimated Mean

Estimated
mean
(age)

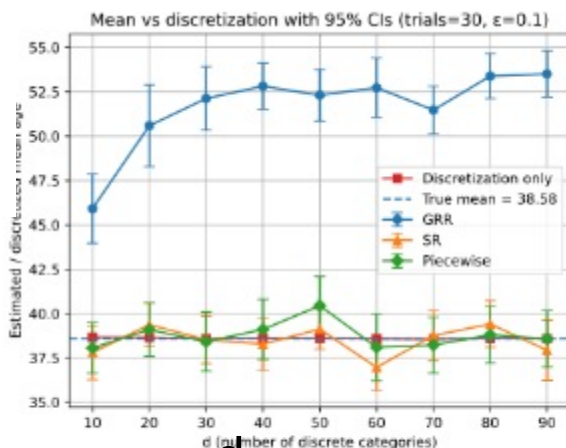


Privacy budget ϵ
(a) $d = 20$

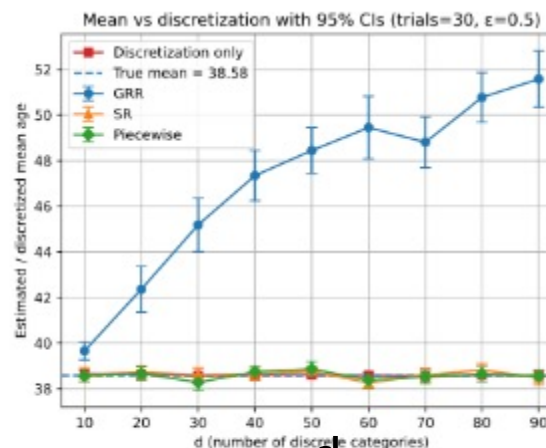


ϵ
(b) $d = 50$

Estimated
mean
(age)



d (# values)
(a) $\epsilon = 0.1$



d
(b) $\epsilon = 0.5$

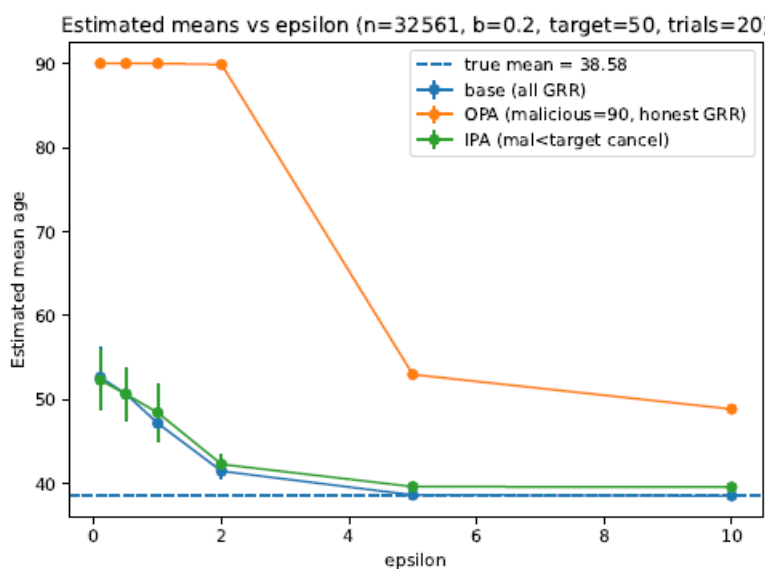
Result 3. Gain (=vulnerability)

Gain

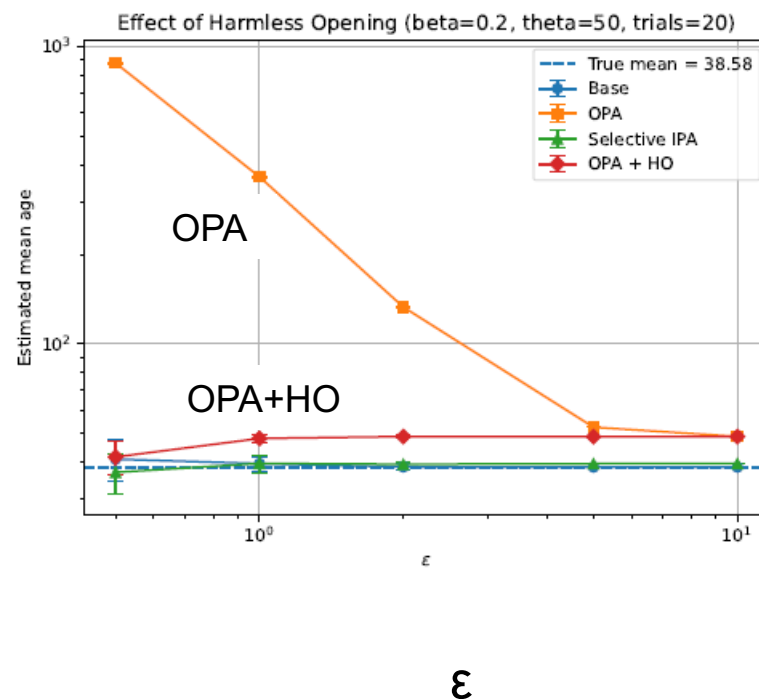
vulnerable



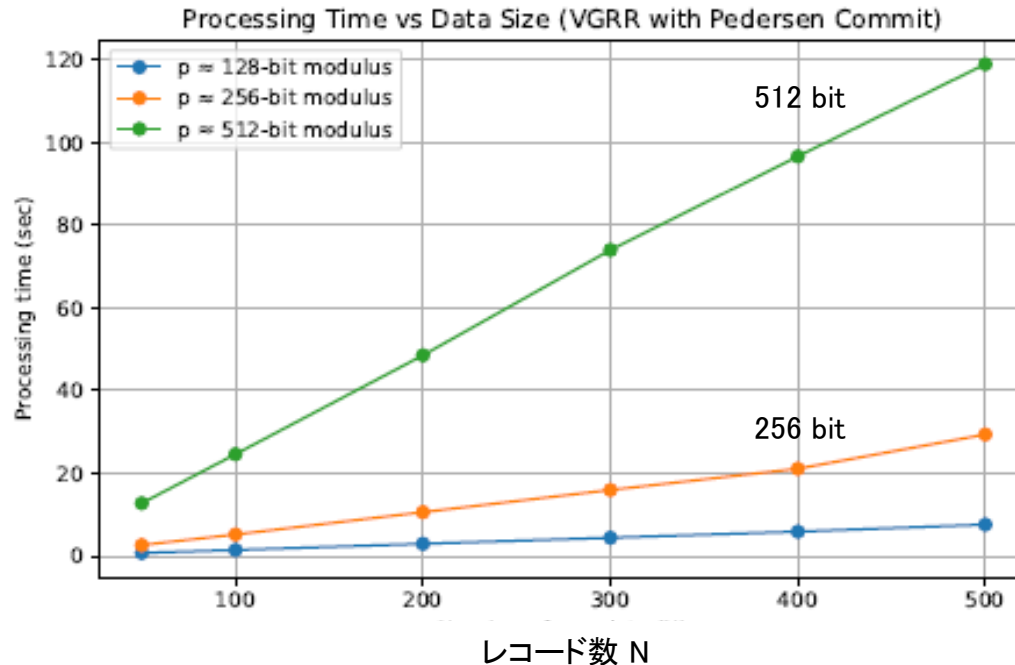
Robust



Privacy budget ϵ



Result 4 processing time



Total time used for
Pedersen
Commitment
 $d = 73$, $\ell = 200$

Per user time is
very small.

E.g. $75 \text{ s} / 300 (N)$
 $= 0.25 \text{ sec/user}$

Conclusions

- RQ1. does HO mitigate poisoning?
 - Yes. GRR and SM are robust.
- RQ2. Accuracy reduction? small
- RQ3. Best scheme? SR is robust and lightweighted
- Future work: robustness for PM and work with Robust LDP [Zhao 25]

	GRR	SR	PM
Robustness	Frequency	Mean	Mean
security	Robust	Robust	(open)
ℓ	$\ell_1 + (d-1) \ell_2$	$\ell_1 + \ell_2$	$\ell_1 s + \ell_2 s$
Cost	High	Low	High
Accuracy	Low	High	High
Gain	High	Low	Low

Pedersen Commitment

- 情報理論的安全性を保証したコミット

- Z_1, Z_2 : 有限群 (DLP仮定)

- 公開鍵: (g, h) Z_1, Z_2 の生成元

- コミットする値: ρ

- コミット: $C(\rho) = g^\rho h^\tau$

- 準同型性

$$C_K(\rho_1; \tau_1) C_K(\rho_2; \tau_2) = C_K(\rho_1 + \rho_2; \tau_1 + \tau_2).$$

準同型性を用いたランダムサンプリング

■ 初期値

□ $A: x, B: y$

□ $A \rightarrow B: C(x, r1)$

□ $B \rightarrow A: C(y, r2)$

□ $C' = C(x, r1) * C(y, r2) = C(x+y, r1+r2)$ を求めて, x, y を open

VGRR

Protocol 1 VGRR

The user $\mathcal{U}$ possesses a raw item v . The aggregator $\mathcal{A}$ obtains a value y as the perturbed data from the user $\mathcal{U}$.

Initialize: The user $\mathcal{U}$ and the aggregator $\mathcal{A}$ share the public key K and the length parameters ℓ , ℓ_1 and ℓ_2 .

1) **Commit**

- $\mathcal{U}$ constructs an ℓ -dimensional vector μ which satisfies $\theta_v^\mu = \ell_1$ and $\theta_i^\mu = \ell_2$ for $i \in [d] \setminus \{v\}$.
- A new ℓ -dimensional vector η is constructed by performing $\eta[i] \triangleq C_K(\mu[i]; \tau_i)$ for $1 \leq i \leq \ell$, where τ_i is drawn randomly from $\mathbb{Z}_{prime_2}$. The vector η is sent to $\mathcal{A}$.

2) **Select**

$\mathcal{A}$ randomly selects j from $\{1, 2, \dots, \ell\}$ and sends j to $\mathcal{U}$.

3) **Open**

$\mathcal{U}$ picks $d\ell_2$ numbers $j_1, j_2, \dots, j_{d\ell_2}$, where $j_1 = j$. For $1 \leq i \leq d\ell_2$, $\mathcal{U}$ opens the commitment $\eta[j_i]$ by sending $\mu[j_i]$ and τ_{j_i} to $\mathcal{A}$. After the $d\ell_2$ commitments $\eta[j_1], \eta[j_2], \dots, \eta[j_{d\ell_2}]$ are opened, the vector η need to satisfy $\varphi_i^\eta = \ell_2$ for $i \in [d]$.

4) **Verify**

- $\mathcal{A}$ checks if $j_1 = j$ and $\eta[j_i] = C_K(\mu[j_i]; \tau_{j_i})$ for $1 \leq i \leq d\ell_2$.
 - For $i \in [d]$, $\mathcal{A}$ checks if the vector η satisfies $\varphi_i^\eta = \ell_2$ for $i \in [d]$
 - If all the checks are affirmative, $\mathcal{A}$ takes $\mu[j]$ as the output y from $\mathcal{U}$.
-

robustの定義

- 定義2 (識別不能性)
 - A^* (真のアイテムを推測するアルゴリズム)
 - view (Aが観測する情報)
 - $\Pr(A^*(\text{view}; y) = v) - \Pr[A^*(y) = v] < \text{nelg}(k)$
である時, Mは識別不能.
- 定義3 (ロバスト)
 - $S_{\max}$ = Overall scoreの最大値.
 - $\Pr[S_{\max} = S_{\text{in}}] > 1 - \text{nelg}(k)$ である時, Mはrobust.
- 定義4 (semi-robust)
 - $\Pr[S_{\text{out}} > S_{\max} > S_{\text{in}}] > 1 - \text{nelg}(k)$

VOUE

■ VGRRのOUEバージョン

□ $\ell = 2\ell_1$, $\ell_1 > \ell_2$

□ d 個のそれぞれについて, j を選ぶ.

$v = 1 \Rightarrow$	1	1	1	1	1	0	0	0	0	0
	1	1	1	0	0	0	0	0	0	0
	1	1	1	0	0	0	0	0	0	0
$v = 2 \Rightarrow$	1	1	1	0	0	0	0	0	0	0
	1	1	1	1	1	0	0	0	0	0
	1	1	1	0	0	0	0	0	0	0
$v = 3 \Rightarrow$	1	1	1	0	0	0	0	0	0	0
	1	1	1	0	0	0	0	0	0	0
	1	1	1	1	1	0	0	0	0	0

VOUEの性質

- 定理4 (LDP)
 - $\varepsilon' = \ln \frac{2\ell_1 - \ell_2}{\ell_2}$ LDPを満たす.
- 定理5 (識別不能性)
- 定理6 (安全性, 少し弱い)
 - VOUEは semi-robust である.
 - » $S_{\text{out}} = r \ell_1 > S_{\text{max}} = r/2$

VOUE+

- 集合への所属のZKIPの導入
 - [36, Kato], [40, Benarroh 2021]
 - 未開示の元素が正しく所定の集合に属していることを証明する.
- 定理9 (頑強性)
 - VOUE+はrobustである.