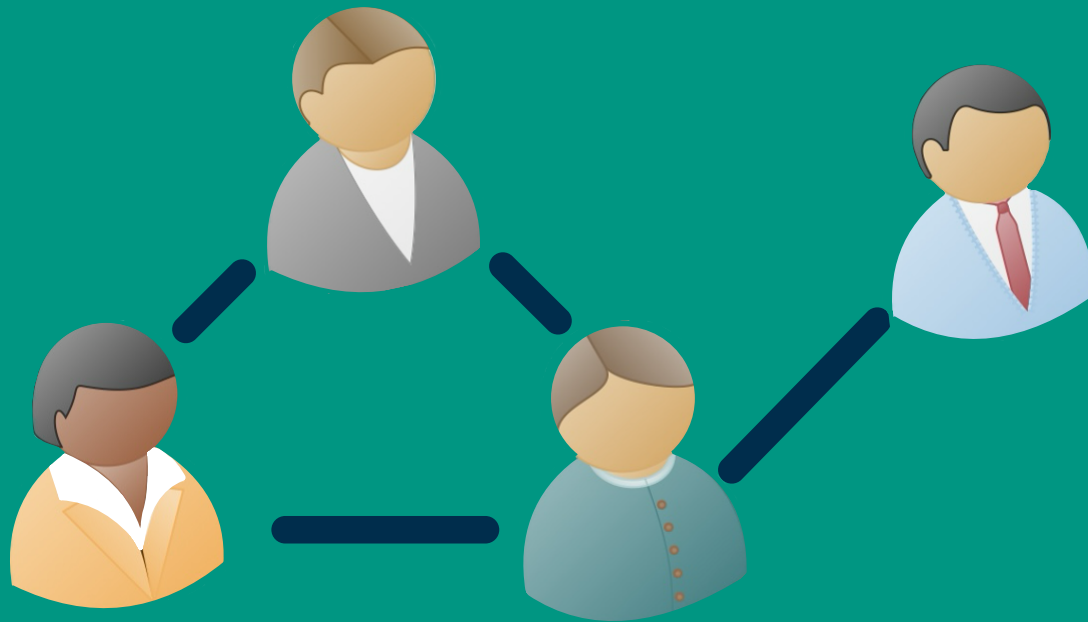




Balancing Privacy and Utility in Correlated Data

INRIA Montpellier, November 12th, 2025

Patricia Guerra-Balboa

Our General Goal

Learn population-level information without harming individual's privacy

Our General Goal

Learn population-level information without harming individual's privacy

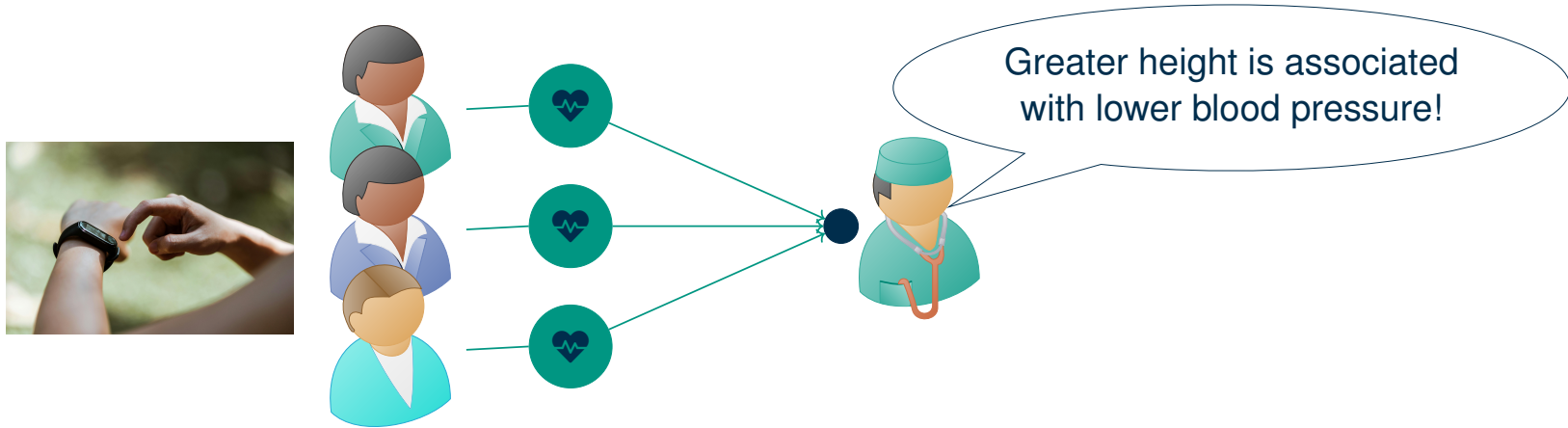


Privacy Goal: Protect Alice's location

Utility Goal: Number of cars per street

Our General Goal

Learn population-level information without harming individual's privacy



Privacy Goal: Protect Alice's activity data

Utility Goal: Correlation between height and health

The Best Tool Until Now: Differential Privacy

Idea: We want to bound participation risk.

The Best Tool Until Now: Differential Privacy

Idea: We want to bound participation risk.



- “Strongest” assumption: everybody’s record is known but the target.

The Best Tool Until Now: Differential Privacy

Idea: We want to bound participation risk.



- “Strongest” assumption: everybody’s record is known but the target.

The Best Tool Until Now: Differential Privacy

Idea: We want to bound participation risk.



$$\ln \frac{p_{\mathcal{M}}(\theta \mid x_1, \dots, x_{n-1}, \mathbf{x}_n)}{p_{\mathcal{M}}(\theta \mid x_1, \dots, x_{n-1}, \mathbf{y}_n)} \leq \epsilon$$

- “Strongest” assumption: everybody’s record is known but the target.
- The privacy leakage ϵ controls the indistinguishability level between $\mathbf{x}_n, \mathbf{y}_n$.

The Best Tool Until Now: Differential Privacy

Idea: We want to bound participation risk.

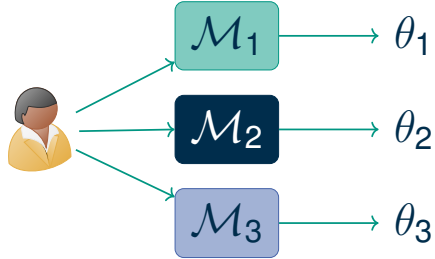


$$\ln \frac{p_{\mathcal{M}}(\theta \mid x_1, \dots, x_{n-1}, \mathbf{x}_n)}{p_{\mathcal{M}}(\theta \mid x_1, \dots, x_{n-1}, \mathbf{y}_n)} \leq \varepsilon$$

- “Strongest” assumption: everybody’s record is known but the target.
- The privacy leakage ε controls the indistinguishability level between $\mathbf{x}_n, \mathbf{y}_n$.
- But at some cost! The smaller the ε the less utility.

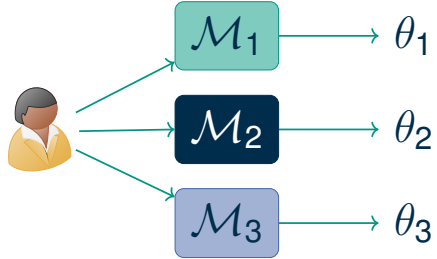
Why DP Is The Best So Far?

Composition

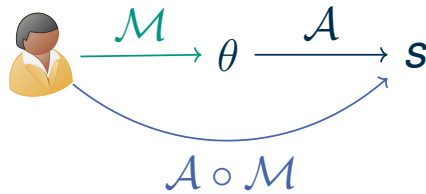


Why DP Is The Best So Far?

Composition

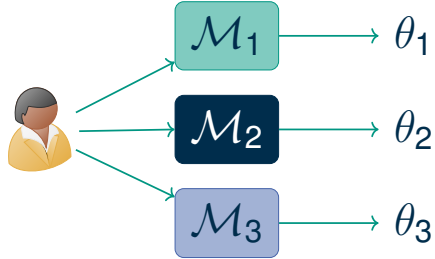


Post-processing

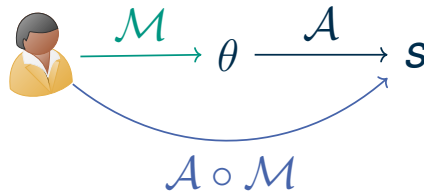


Why DP Is The Best So Far?

Composition



Post-processing

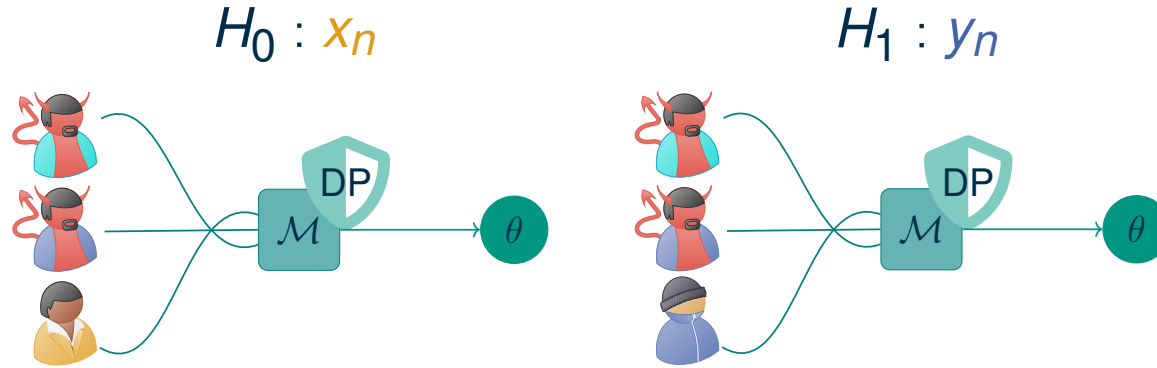


Attack Mitigation

$$\mathcal{M} \text{ } \varepsilon\text{-DP} \Rightarrow \text{Adv} \leq f(\varepsilon)$$

Membership Inference Attack Knowing D_-

The attacker receives θ and aims to distinguish between:



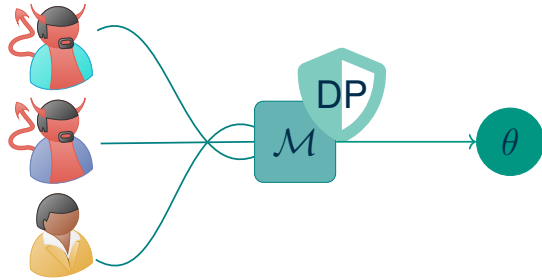
D_- is known:

$$H_0 = D_{x_n} \text{ Vs. } H_1 = D_{y_n}$$

Membership Inference Attack Knowing D_-

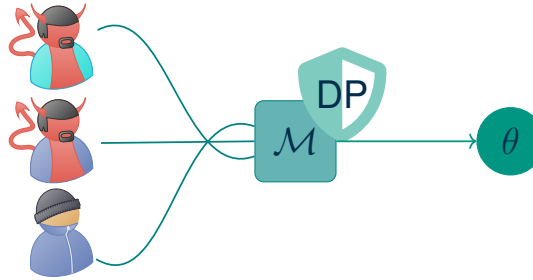
The attacker receives θ and aims to distinguish between:

$$H_0 : x_n$$



Type I error:
 $\alpha = \Pr_{A \circ \mathcal{M}}(y_n \mid D_{x_n})$

$$H_1 : y_n$$



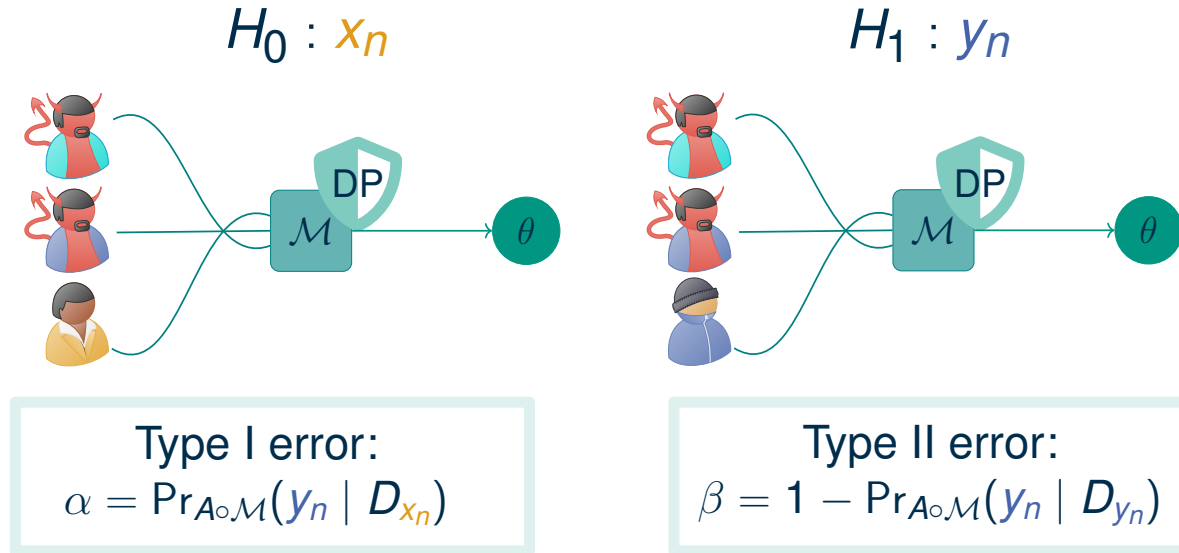
Type II error:
 $\beta = 1 - \Pr_{A \circ \mathcal{M}}(y_n \mid D_{y_n})$

D_- is known:

$$H_0 = D_{x_n} \text{ Vs. } H_1 = D_{y_n}$$

Membership Inference Attack Knowing D_-

The attacker receives θ and aims to distinguish between:



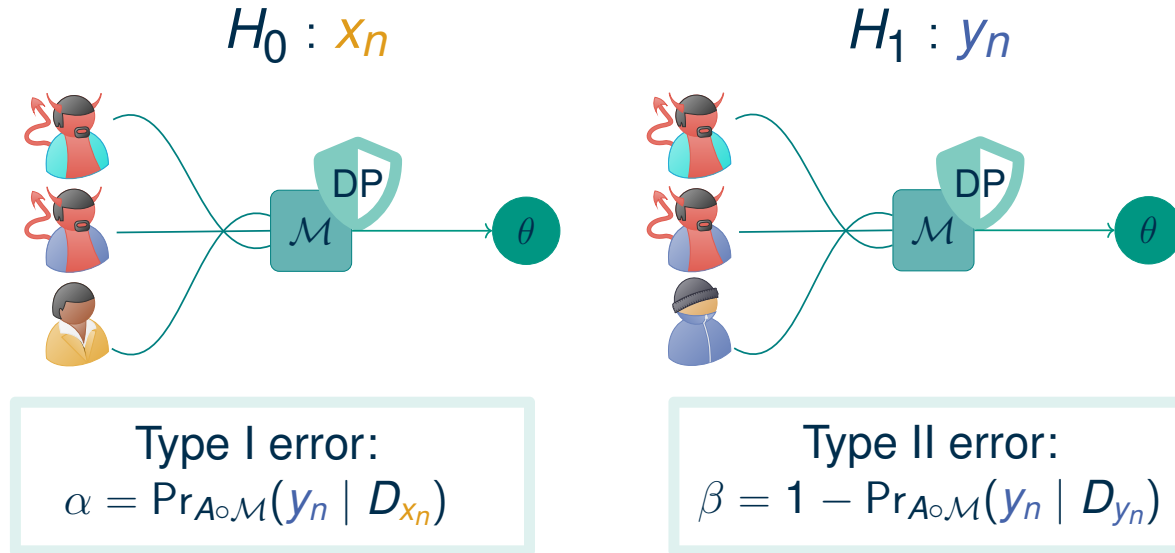
D_- is known:

$$H_0 = D_{x_n} \text{ Vs. } H_1 = D_{y_n}$$

$A \circ \mathcal{M}$ is ϵ -DP $\Rightarrow$

Membership Inference Attack Knowing D_-

The attacker receives θ and aims to distinguish between:



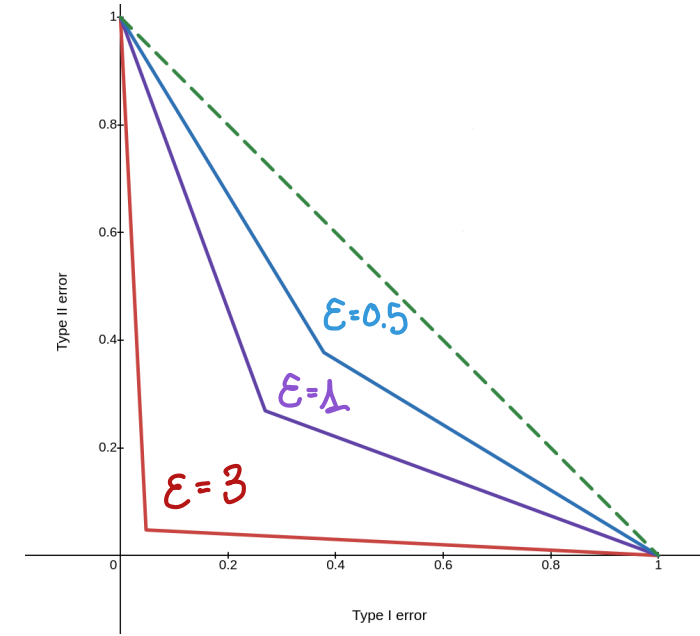
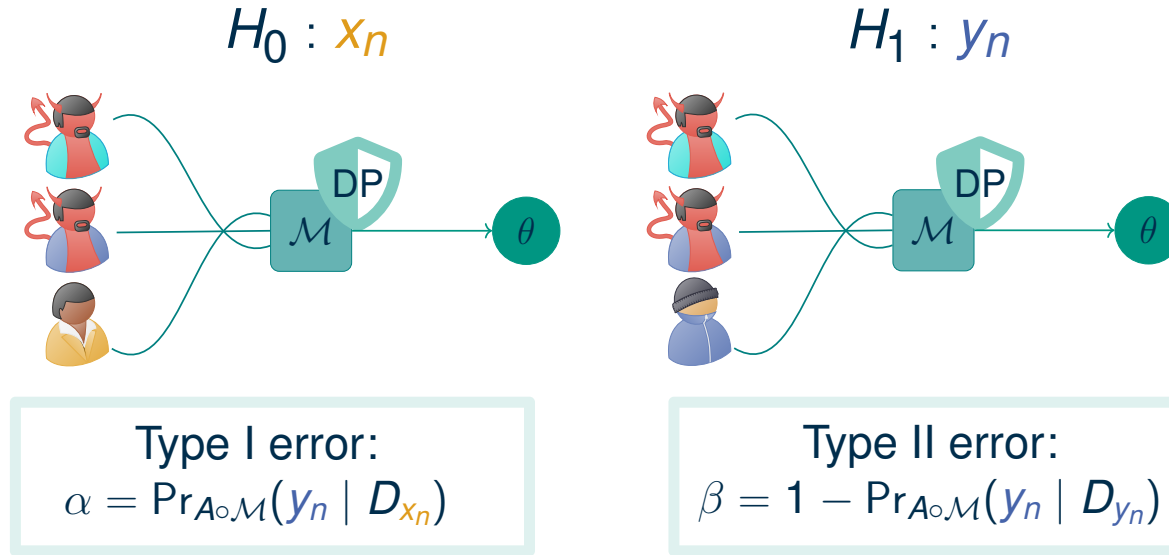
D_- is known:

$$H_0 = D_{x_n} \text{ Vs. } H_1 = D_{y_n}$$

$$A \circ \mathcal{M} \text{ is } \varepsilon\text{-DP} \Rightarrow \begin{aligned} 1 - \beta &\leq e^\varepsilon \alpha \\ \alpha &\leq e^\varepsilon (1 - \beta) \end{aligned}$$

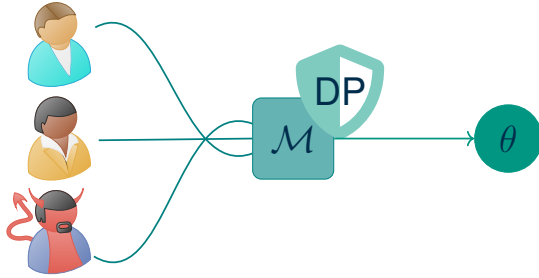
Membership Inference Attack Knowing D_-

The attacker receives θ and aims to distinguish between:

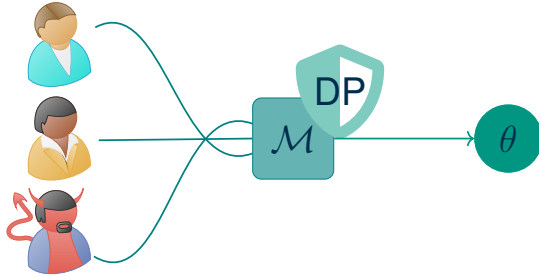


$$A \circ \mathcal{M} \text{ is } \epsilon\text{-DP} \Rightarrow \begin{aligned} 1 - \beta &\leq e^\epsilon \alpha \\ \alpha &\leq e^\epsilon (1 - \beta) \end{aligned} \Rightarrow \beta \geq \max\{1 - e^\epsilon \alpha, e^\epsilon (1 - \alpha)\}$$

What about other attacker models?



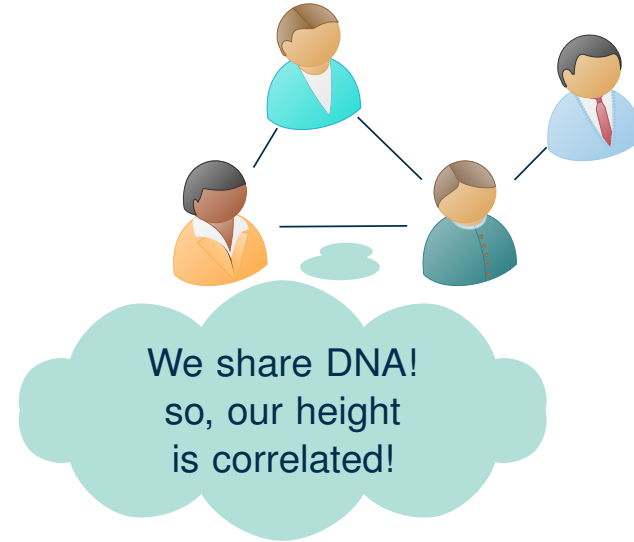
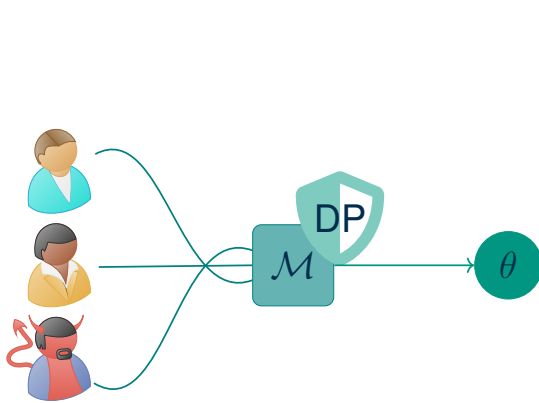
What about other attacker models?



Statistical Independence

The strongest attacker is the worst-case one, and we have at least the same protection.

What about other attacker models?



Statistical Independence

**Dependencies between
Records**

The strongest attacker is the worst-case one, and we have at least the same protection.

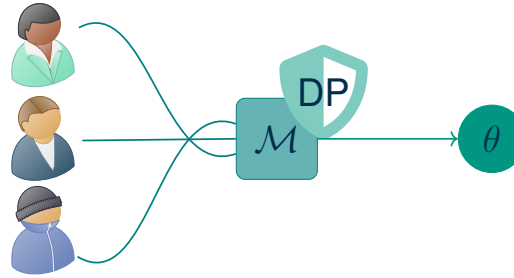
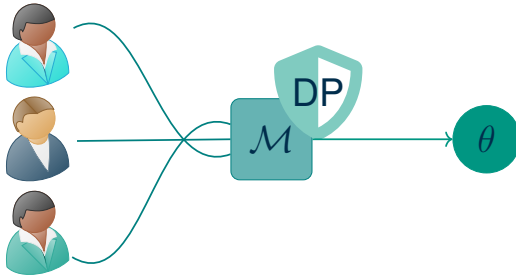
DP interpretation does not hold anymore.

Membership Inference Attack With Dependencies

The attacker receives θ and aims to distinguish between:

$$H_0 : x_n$$

$$H_1 : y_n$$



D_- is **unknown**:

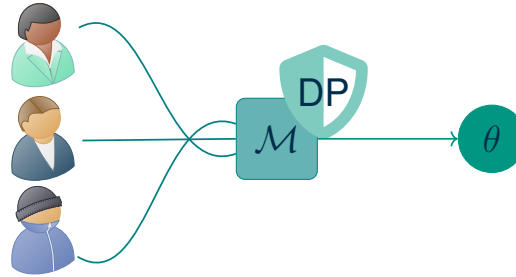
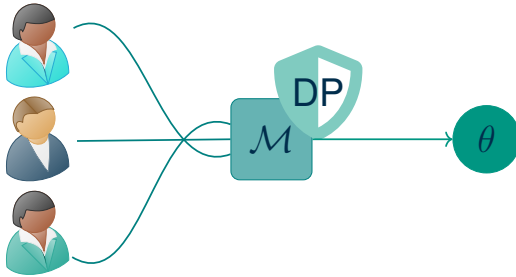
$$H_0 = \{D : x_n \in D\} \text{ Vs. } H_1 = \{D : y_n \in D\}$$

Membership Inference Attack With Dependencies

The attacker receives θ and aims to distinguish between:

$$H_0 : x_n$$

$$H_1 : y_n$$



D_- is **unknown**:

$$H_0 = \{D : x_n \in D\} \text{ Vs. } H_1 = \{D : y_n \in D\}$$

Type I error:

$$\alpha = \Pr_{A \circ \mathcal{M}}(y_n \mid x_n)$$

Type II error:

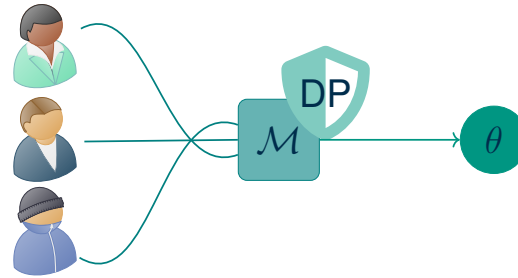
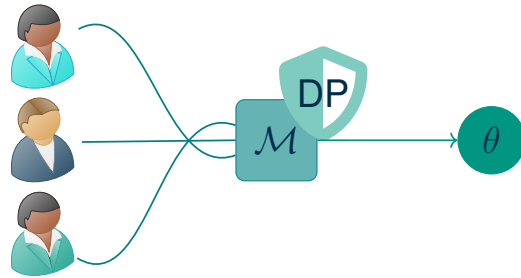
$$\beta = 1 - \Pr_{A \circ \mathcal{M}}(y_n \mid y_n)$$

Membership Inference Attack With Dependencies

The attacker receives θ and aims to distinguish between:

$$H_0 : x_n$$

$$H_1 : y_n$$



D_- is **unknown**:

$$H_0 = \{D : x_n \in D\} \text{ Vs. } H_1 = \{D : y_n \in D\}$$

Type I error:

$$\begin{aligned} \alpha &= \Pr_{A \circ \mathcal{M}}(y_n | x_n) \\ &= \sum_{D_-} \Pr_{A \circ \mathcal{M}}(y_n | D_{x_n}) \pi(D_- | x_n) \end{aligned}$$

Type II error:

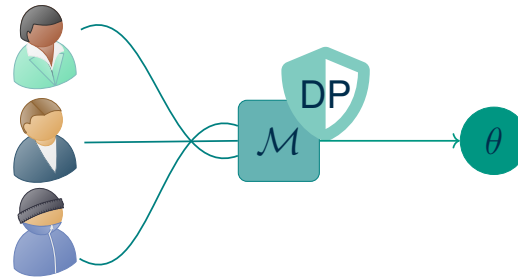
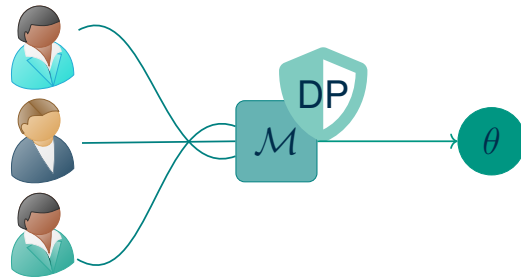
$$\begin{aligned} \beta &= 1 - \Pr_{A \circ \mathcal{M}}(y_n | y_n) \\ &= 1 - \sum_{D_-} \Pr_{A \circ \mathcal{M}}(y_n | D_{y_n}) \pi(D_- | y_n) \end{aligned}$$

Membership Inference Attack With Dependencies

The attacker receives θ and aims to distinguish between:

$$H_0 : x_n$$

$$H_1 : y_n$$



D_- is **unknown**:

$$H_0 = \{D : x_n \in D\} \text{ Vs. } H_1 = \{D : y_n \in D\}$$

Type I error:

$$\begin{aligned} \alpha &= \Pr_{A \circ \mathcal{M}}(y_n | x_n) \\ &= \sum_{D_-} \Pr_{A \circ \mathcal{M}}(y_n | D_{x_n}) \pi(D_- | x_n) \end{aligned}$$

Type II error:

$$\begin{aligned} \beta &= 1 - \Pr_{A \circ \mathcal{M}}(y_n | y_n) \\ &= 1 - \sum_{D_-} \Pr_{A \circ \mathcal{M}}(y_n | D_{y_n}) \pi(D_- | y_n) \end{aligned}$$

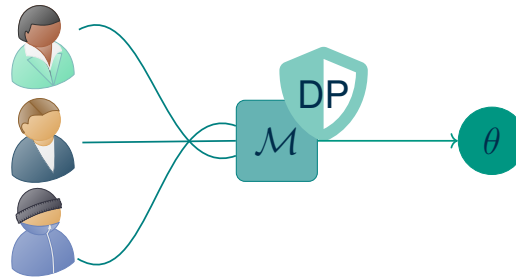
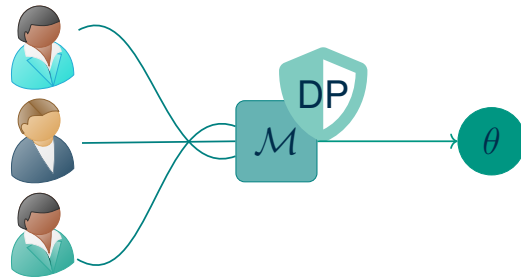
$$A \circ \mathcal{M} \text{ is } \varepsilon\text{-DP} \quad \Rightarrow \quad 1 - \beta \leq \sum_{D_-} e^\varepsilon \Pr_{A \circ \mathcal{M}}(y_n | D_{x_n}) \pi(D_- | y_n)$$

Membership Inference Attack With Dependencies

The attacker receives θ and aims to distinguish between:

$$H_0 : x_n$$

$$H_1 : y_n$$



D_- is **unknown**:

$$H_0 = \{D : x_n \in D\} \text{ Vs. } H_1 = \{D : y_n \in D\}$$

Type I error:

$$\begin{aligned} \alpha &= \Pr_{A \circ \mathcal{M}}(y_n | x_n) \\ &= \sum_{D_-} \Pr_{A \circ \mathcal{M}}(y_n | D_{x_n}) \pi(D_- | x_n) \end{aligned}$$

Type II error:

$$\begin{aligned} \beta &= 1 - \Pr_{A \circ \mathcal{M}}(y_n | y_n) \\ &= 1 - \sum_{D_-} \Pr_{A \circ \mathcal{M}}(y_n | D_{y_n}) \pi(D_- | y_n) \end{aligned}$$

$$A \circ \mathcal{M} \text{ is } \varepsilon\text{-DP} \Rightarrow 1 - \beta \leq \sum_{D_-} e^\varepsilon \Pr_{A \circ \mathcal{M}}(y_n | D_{x_n}) \pi(D_- | y_n) = e^\varepsilon \sum_{D_-} \Pr_{A \circ \mathcal{M}}(y_n | D_{x_n}) \pi(D_- | y_n) \neq e^\varepsilon \alpha$$

Standard DP Underestimates Participation Risk

Differential Privacy fails to measure privacy leakage under correlation

  Empirically confirmed

  Theoretically exposed

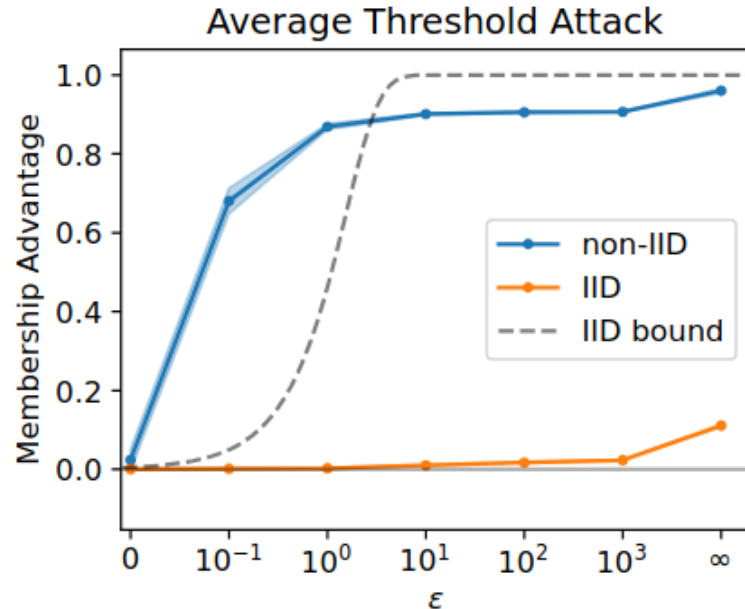


Figure: Humphries et al. 2023 MIA breaks DP guarantees.

Standard DP Underestimates Participation Risk

Differential Privacy fails to measure privacy leakage under correlation

  Empirically confirmed

  Theoretically exposed

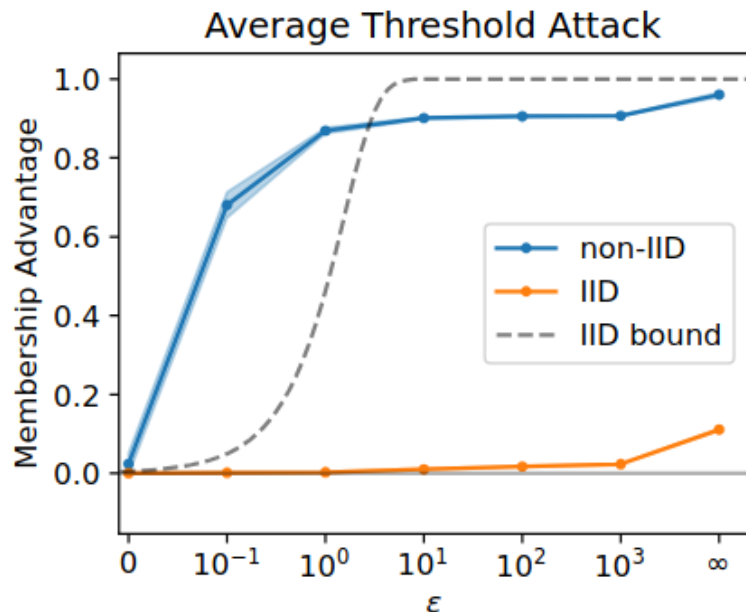
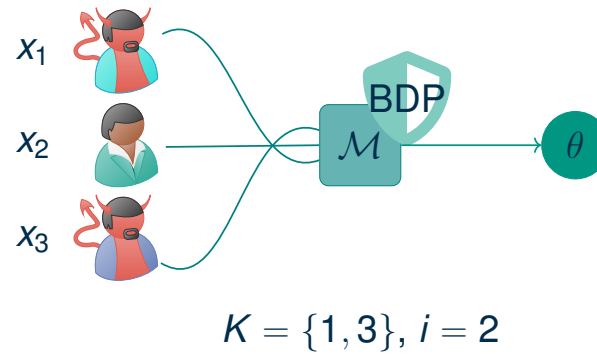


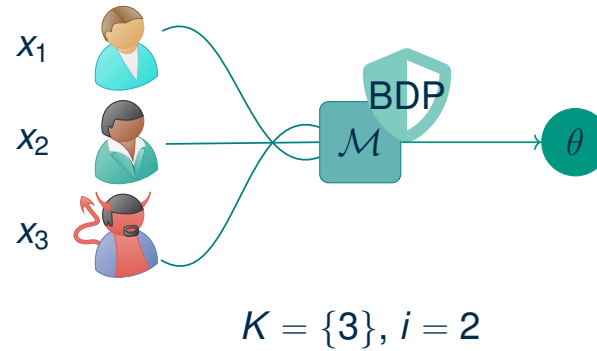
Figure: Humphries et al. 2023 MIA breaks DP guarantees.

New enhanced notion: Bayesian Differential Privacy

Proposed Solution: Bayesian Differential Privacy



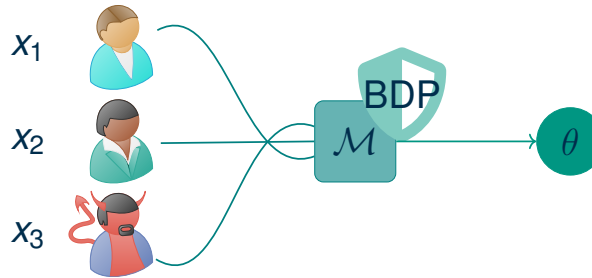
Proposed Solution: Bayesian Differential Privacy



Proposed Solution: Bayesian Differential Privacy

Bayesian DP leakage (Yang et al. 2017)

$$\text{BDPL}_{(K,i)} = \sup_{x_i, x'_i, \mathbf{x}_K, S} \ln \frac{\Pr_{\mathcal{M}}[Y \in S \mid \mathbf{X}_K = \mathbf{x}_K, X_i = x_i]}{\Pr_{\mathcal{M}}[Y \in S \mid \mathbf{X}_K = \mathbf{x}_K, X_i = x'_i]}, \text{ then } \varepsilon = \sup_{K,i} \text{BDPL}_{(K,i)}.$$



Proposed Solution: Bayesian Differential Privacy

Bayesian DP leakage (Yang et al. 2017)

$$\text{BDPL}_{(K,i)} = \sup_{x_i, x'_i, \mathbf{x}_K, S} \ln \frac{\Pr_{\mathcal{M}}[Y \in S \mid \mathbf{X}_K = \mathbf{x}_K, X_i = x_i]}{\Pr_{\mathcal{M}}[Y \in S \mid \mathbf{X}_K = \mathbf{x}_K, X_i = x'_i]}, \text{ then } \varepsilon = \sup_{K,i} \text{BDPL}_{(K,i)}.$$

We check our MIA again, with $H_0 : x_n \in D$, $H_1 : y_n \in D$ ($K = \emptyset$)

Proposed Solution: Bayesian Differential Privacy

Bayesian DP leakage (Yang et al. 2017)

$$\text{BDPL}_{(K,i)} = \sup_{x_i, x'_i, \mathbf{x}_K, S} \ln \frac{\Pr_{\mathcal{M}}[Y \in S \mid \mathbf{X}_K = \mathbf{x}_K, X_i = x_i]}{\Pr_{\mathcal{M}}[Y \in S \mid \mathbf{X}_K = \mathbf{x}_K, X_i = x'_i]}, \text{ then } \varepsilon = \sup_{K,i} \text{BDPL}_{(K,i)}.$$

We check our MIA again, with $H_0 : x_n \in D$, $H_1 : y_n \in D$ ($K = \emptyset$)

Type I error:

$$\alpha = \Pr_{A \circ \mathcal{M}}(y_n \mid x_n)$$

Type II error:

$$\beta = 1 - \Pr_{A \circ \mathcal{M}}(y_n \mid y_n)$$

Proposed Solution: Bayesian Differential Privacy

Bayesian DP leakage (Yang et al. 2017)

$$\text{BDPL}_{(K,i)} = \sup_{x_i, x'_i, \mathbf{x}_K, S} \ln \frac{\Pr_{\mathcal{M}}[Y \in S \mid \mathbf{X}_K = \mathbf{x}_K, X_i = x_i]}{\Pr_{\mathcal{M}}[Y \in S \mid \mathbf{X}_K = \mathbf{x}_K, X_i = x'_i]}, \text{ then } \varepsilon = \sup_{K,i} \text{BDPL}_{(K,i)}.$$

We check our MIA again, with $H_0 : \mathbf{x}_n \in D$, $H_1 : \mathbf{y}_n \in D$ ($K = \emptyset$)

Type I error:

$$\alpha = \Pr_{A \circ \mathcal{M}}(\mathbf{y}_n \mid \mathbf{x}_n)$$

Type II error:

$$\beta = 1 - \Pr_{A \circ \mathcal{M}}(\mathbf{y}_n \mid \mathbf{y}_n)$$

$A \circ \mathcal{M}$ is
 ε -BDP



$$1 - \beta \leq e^\varepsilon \alpha$$

$$\alpha \leq e^\varepsilon (1 - \beta)$$

Proposed Solution: Bayesian Differential Privacy

Bayesian DP leakage (Yang et al. 2017)

$$\text{BDPL}_{(K,i)} = \sup_{x_i, x'_i, \mathbf{x}_K, S} \ln \frac{\Pr_{\mathcal{M}}[Y \in S \mid \mathbf{X}_K = \mathbf{x}_K, X_i = x_i]}{\Pr_{\mathcal{M}}[Y \in S \mid \mathbf{X}_K = \mathbf{x}_K, X_i = x'_i]}, \text{ then } \varepsilon = \sup_{K,i} \text{BDPL}_{(K,i)}.$$

We check our MIA again, with $H_0 : x_n \in D$, $H_1 : y_n \in D$ ($K = \emptyset$)

Type I error:

$$\alpha = \Pr_{A \circ \mathcal{M}}(y_n \mid x_n)$$

Type II error:

$$\beta = 1 - \Pr_{A \circ \mathcal{M}}(y_n \mid y_n)$$

$A \circ \mathcal{M}$ is
 ε -BDP



$$1 - \beta \leq e^\varepsilon \alpha$$

$$\alpha \leq e^\varepsilon (1 - \beta)$$



$$\beta \geq \max\{1 - e^\varepsilon \alpha, e^\varepsilon (1 - \alpha)\}$$



Proposed Solution: Bayesian Differential Privacy

Bayesian DP leakage (Yang et al. 2017)

$$\text{BDPL}_{(K,i)} = \sup_{x_i, x'_i, \mathbf{x}_K, S} \ln \frac{\Pr_{\mathcal{M}}[Y \in S \mid \mathbf{X}_K = \mathbf{x}_K, X_i = x_i]}{\Pr_{\mathcal{M}}[Y \in S \mid \mathbf{X}_K = \mathbf{x}_K, X_i = x'_i]}, \text{ then } \varepsilon = \sup_{K,i} \text{BDPL}_{(K,i)}.$$

_____Privacy_____

- ✓ Effective measure and resistance to correlation-based attacks.

Proposed Solution: Bayesian Differential Privacy

Bayesian DP leakage (Yang et al. 2017)

$$\text{BDPL}_{(K,i)} = \sup_{x_i, x'_i, \mathbf{x}_K, S} \ln \frac{\Pr_{\mathcal{M}}[Y \in S \mid \mathbf{X}_K = \mathbf{x}_K, X_i = x_i]}{\Pr_{\mathcal{M}}[Y \in S \mid \mathbf{X}_K = \mathbf{x}_K, X_i = x'_i]}, \text{ then } \varepsilon = \sup_{K,i} \text{BDPL}_{(K,i)}.$$

Privacy

- ✓ Effective measure and resistance to correlation-based attacks.
- ✓ Good properties: post-processing & composition.
 - While other correlation-aware notions (General Pufferfish framework) don't!

Proposed Solution: Bayesian Differential Privacy

Bayesian DP leakage (Yang et al. 2017)

$$\text{BDPL}_{(K,i)} = \sup_{x_i, x'_i, \mathbf{x}_K, S} \ln \frac{\Pr_{\mathcal{M}}[Y \in S \mid \mathbf{X}_K = \mathbf{x}_K, X_i = x_i]}{\Pr_{\mathcal{M}}[Y \in S \mid \mathbf{X}_K = \mathbf{x}_K, X_i = x'_i]}, \text{ then } \varepsilon = \sup_{K,i} \text{BDPL}_{(K,i)}.$$

Privacy

- ✓ Effective measure and resistance to correlation-based attacks.
- ✓ Good properties: post-processing & composition.
 - While other correlation-aware notions (General Pufferfish framework) don't!

Utility

- ✗ Poor utility (methods based on group privacy).

Proposed Solution: Bayesian Differential Privacy

Bayesian DP leakage (Yang et al. 2017)

$$\text{BDPL}_{(K,i)} = \sup_{x_i, x'_i, \mathbf{x}_K, S} \ln \frac{\Pr_{\mathcal{M}}[Y \in S \mid \mathbf{X}_K = \mathbf{x}_K, X_i = x_i]}{\Pr_{\mathcal{M}}[Y \in S \mid \mathbf{X}_K = \mathbf{x}_K, X_i = x'_i]}, \text{ then } \varepsilon = \sup_{K,i} \text{BDPL}_{(K,i)}.$$

Privacy

- ✓ Effective measure and resistance to correlation-based attacks.
- ✓ Good properties: post-processing & composition.
 - While other correlation-aware notions (General Pufferfish framework) don't!

Utility

- ✗ Poor utility (methods based on group privacy).
- ✗ Computationally intractable methods (computing the Wasserstein distance).

Proposed Solution: Bayesian Differential Privacy

Bayesian DP leakage (Yang et al. 2017)

$$\text{BDPL}_{(K,i)} = \sup_{x_i, x'_i, \mathbf{x}_K, S} \ln \frac{\Pr_{\mathcal{M}}[Y \in S \mid \mathbf{X}_K = \mathbf{x}_K, X_i = x_i]}{\Pr_{\mathcal{M}}[Y \in S \mid \mathbf{X}_K = \mathbf{x}_K, X_i = x'_i]}, \text{ then } \varepsilon = \sup_{K,i} \text{BDPL}_{(K,i)}.$$

Privacy

- ✓ Effective measure and resistance to correlation-based attacks.
- ✓ Good properties: post-processing & composition.
 - While other correlation-aware notions (General Pufferfish framework) don't!

Utility

- ✗ Poor utility (methods based on group privacy).
- ✗ Computationally intractable methods (computing the Wasserstein distance).
- ✗ Limited applicability (lazy, binary, stationary Markov chains).

Our Research Question

Can we reduce utility loss while still retaining the privacy guarantees of BDP?

Our methodology: Understanding how DP leakage relates to BDP leakage:

ϵ -DP $\Rightarrow$??-BDP.

Against Arbitrary Correlations It Is Impossible

Kifer and Machanavajjhala 2014:

Pufferfish (including BDP)

&

$\Rightarrow$

Free-lunch Privacy

$\Rightarrow$

No utility.

arbitrary correlation

Against Arbitrary Correlations It Is Impossible

Kifer and Machanavajjhala 2014:

Pufferfish (including BDP)

&

$\Rightarrow$

Free-lunch Privacy

$\Rightarrow$

No utility.

arbitrary correlation

We express this in term of (α, β) -accuracy for any numerical target query f :

(α, β) -accuracy

$$\Pr(|f(D) - \mathcal{M}(D)| \geq \alpha) \leq \beta$$

$1 - \beta = \text{confidence}$

$\alpha = \text{error interval}$



Against Arbitrary Correlations It Is Impossible

Kifer and Machanavajjhala 2014:

Pufferfish (including BDP)

&

$\Rightarrow$

Free-lunch Privacy

$\Rightarrow$

No utility.

arbitrary correlation

We express this in term of (α, β) -accuracy for any numerical target query f :

(α, β) -accuracy

$$\Pr(|f(D) - \mathcal{M}(D)| \geq \alpha) \leq \beta$$

$1 - \beta = \text{confidence}$

$\alpha = \text{error interval}$

Our result (informal):

$$\beta < \frac{1}{e^\epsilon + 1} \Rightarrow \alpha > \frac{1}{2} \text{Range}(f).$$



Out of 100
how many
are infected?



Against Arbitrary Correlations It Is Impossible

Kifer and Machanavajjhala 2014:

Pufferfish (including BDP)

&

$\Rightarrow$

Free-lunch Privacy

$\Rightarrow$

No utility.

arbitrary correlation

We express this in term of (α, β) -accuracy for any numerical target query f :

(α, β) -accuracy

$$\Pr(|f(D) - \mathcal{M}(D)| \geq \alpha) \leq \beta$$

$1 - \beta = \text{confidence}$

$\alpha = \text{error interval}$

Our result (informal):

$$\beta < \frac{1}{e^\epsilon + 1} \Rightarrow \alpha > \frac{1}{2} \text{Range}(f).$$



Out of 100
how many
are infected?



Against Arbitrary Correlations It Is Impossible

Kifer and Machanavajjhala 2014:

Pufferfish (including BDP)

&

$\Rightarrow$

Free-lunch Privacy

$\Rightarrow$

No utility.

arbitrary correlation

We express this in term of (α, β) -accuracy for any numerical target query f :

(α, β) -accuracy

$$\Pr(|f(D) - \mathcal{M}(D)| \geq \alpha) \leq \beta$$

$1 - \beta = \text{confidence}$

$\alpha = \text{error interval}$

Our result (informal):

$$\beta < \frac{1}{e^\epsilon + 1} \Rightarrow \alpha > \frac{1}{2} \text{Range}(f).$$



Out of 100
how many
are infected?



Is it 0 or
100? 🤔



Few Correlated Records, Same Disaster

Our result (informal)

Privacy decreases linearly proportional to number of correlated records:

$$\epsilon\text{-DP} \Rightarrow m\epsilon\text{-BDP}$$

How does it impact utility?

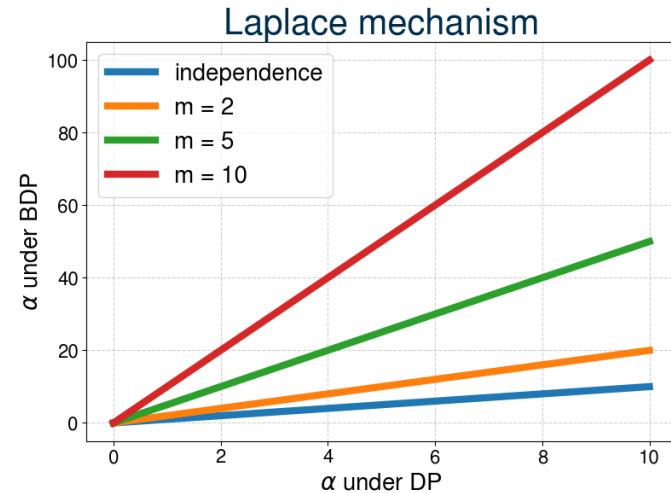


Figure: For the same confidence level, the upper bound on the query error α increases sharply.

Few Correlated Records, Same Disaster

Our result (informal)

Privacy decreases linearly proportional to number of correlated records:

$$\epsilon\text{-DP} \Rightarrow m\epsilon\text{-BDP}$$

This result is tight! Even if $\rho \rightarrow 0$.

How does it impact utility?

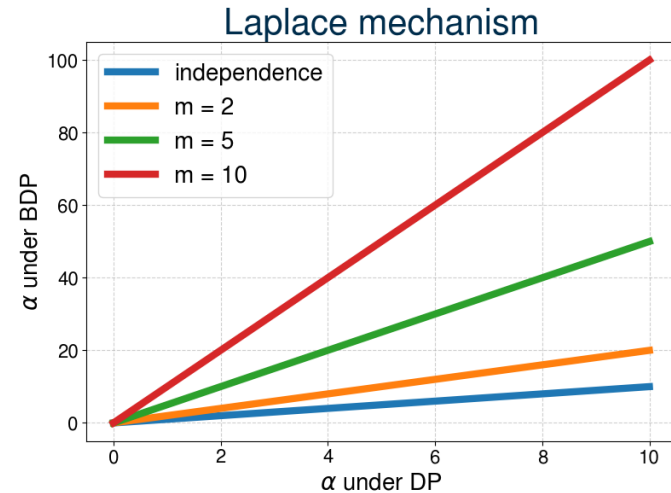


Figure: For the same confidence level, the upper bound on the query error α increases sharply.

Few Correlated Records, Same Disaster

Our result (informal)

Privacy decreases linearly proportional to number of correlated records:

$$\epsilon\text{-DP} \Rightarrow m\epsilon\text{-BDP}$$

This result is tight! Even if $\rho \rightarrow 0$.

Conclusion:

We need to target specific correlation models π to obtain utility

How does it impact utility?

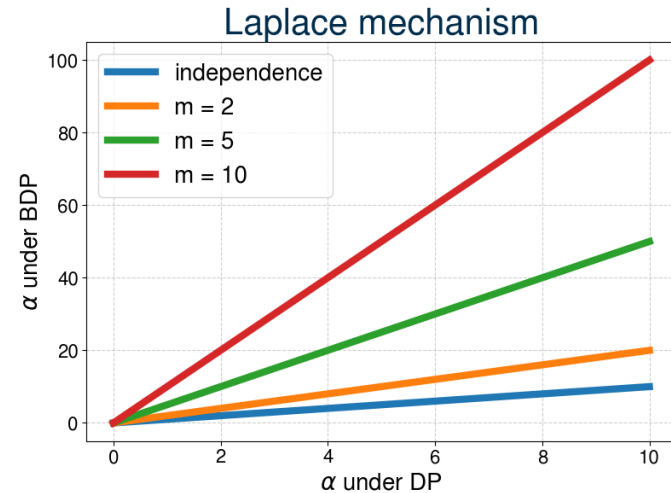


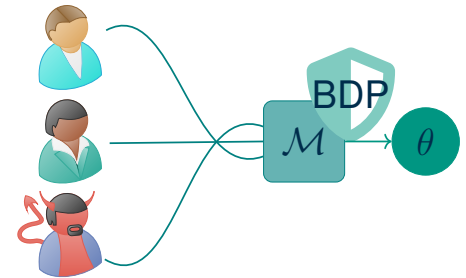
Figure: For the same confidence level, the upper bound on the query error α increases sharply.

New Strategy

Adjust the noise of DP mechanisms to obtain useful BDP mechanisms targeting specific priors π .

Assumptions:

- Global setting: All data is collected by a trusted data curator that applies the mechanism.
- The attacker does not have more knowledge about π than the data curator.

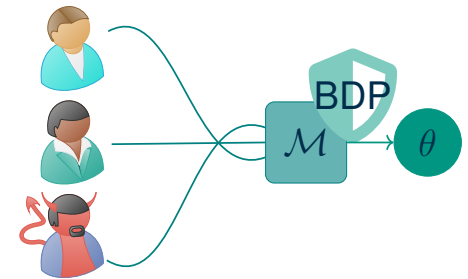


New Strategy

Adjust the noise of DP mechanisms to obtain useful BDP mechanisms targeting specific priors π .

Assumptions:

- Global setting: All data is collected by a trusted data curator that applies the mechanism.
- The attacker does not have more knowledge about π than the data curator.



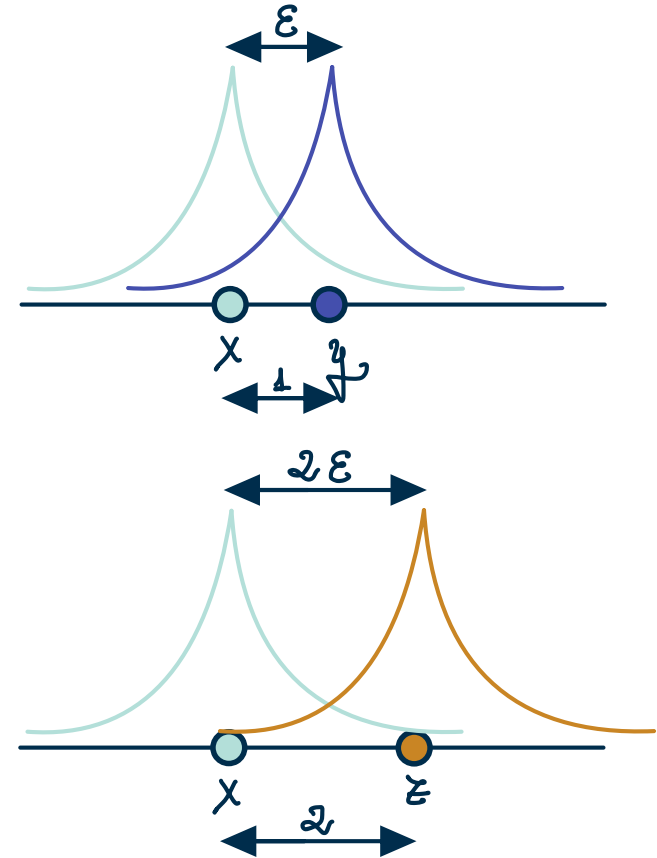
Multivariate Gaussian

Markov Chains

Multivariate Gaussian Correlation

Main Result (Informal)

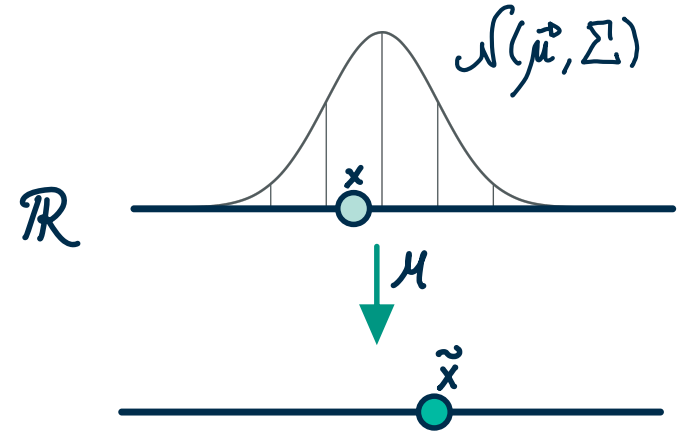
- Let $\mathcal{M}$ be an $\varepsilon \ell_1$ -private mechanism,



Multivariate Gaussian Correlation

Main Result (Informal)

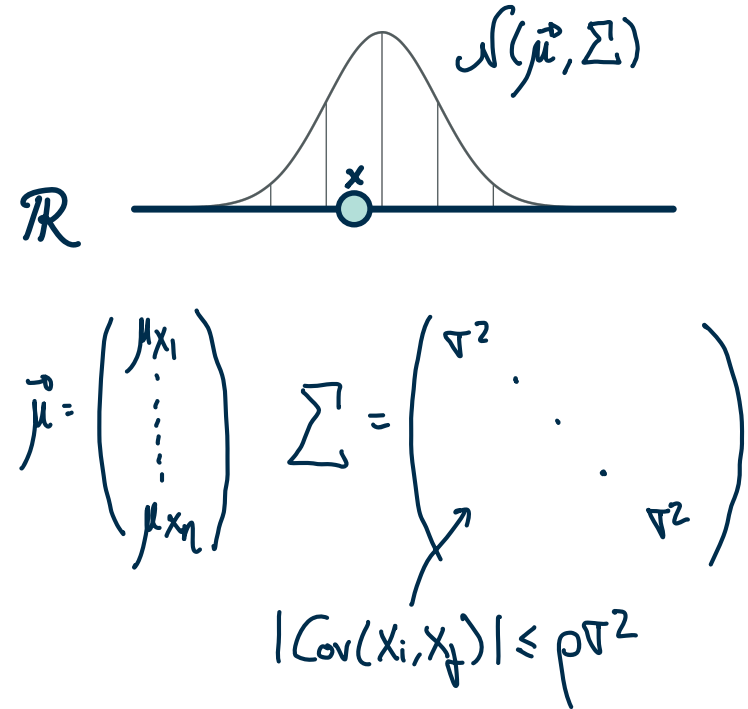
- Let $\mathcal{M}$ be an $\varepsilon\ell_1$ -private mechanism,
- input data drawn from a multivariate Gaussian distribution



Multivariate Gaussian Correlation

Main Result (Informal)

- Let $\mathcal{M}$ be an $\varepsilon\ell_1$ -private mechanism,
- input data drawn from a multivariate Gaussian distribution
- $\rho(m-2) < 1$ is the maximum correlation coefficient.

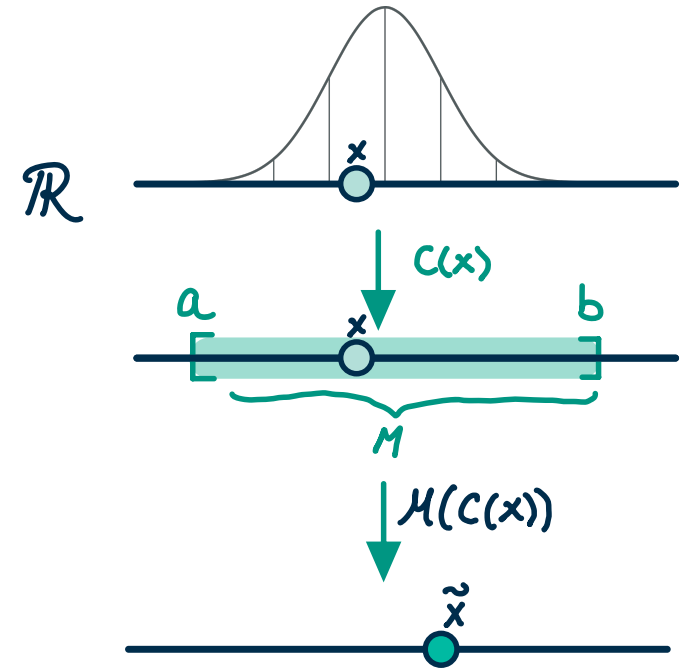


Multivariate Gaussian Correlation

Main Result (Informal)

- Let $\mathcal{M}$ be an $\varepsilon\ell_1$ -private mechanism,
- input data drawn from a multivariate Gaussian distribution
- $\rho(m-2) < 1$ is the maximum correlation coefficient.

Then, using clipping as preprocessing step, $c_i(D)_i = \max(a, \min(b, D_i))$, we obtain $\mathcal{M}_i$ satisfying



Multivariate Gaussian Correlation

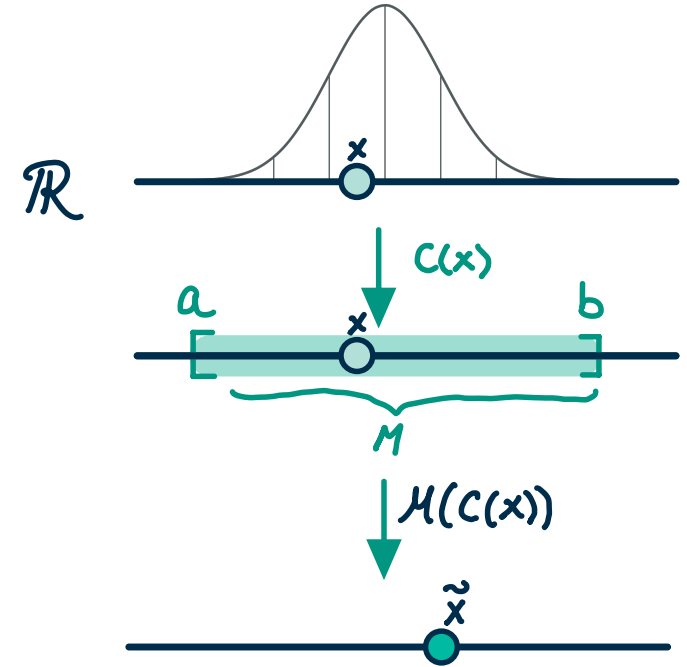
Main Result (Informal)

- Let $\mathcal{M}$ be an $\varepsilon\ell_1$ -private mechanism,
- input data drawn from a multivariate Gaussian distribution
- $\rho(m-2) < 1$ is the maximum correlation coefficient.

Then, using clipping as preprocessing step, $c_l(D)_i = \max(a, \min(b, D_i))$, we obtain $\mathcal{M}_l$ satisfying

$$\text{BDPL}(\mathcal{M}_l) \leq \left(\frac{m^2}{4(\frac{1}{\rho} - m + 2)} + 1 \right) M\varepsilon.$$

where M is the diameter of the interval $I = [a, b]$



Multivariate Gaussian Correlation (Impact on Real Databases)

- **Use-case:** Sum queries with Laplace mechanism. $\theta = f(D) + Z$ with $Z \sim \text{Lap}(b)$.

Multivariate Gaussian Correlation (Impact on Real Databases)

- **Use-case:** Sum queries with Laplace mechanism. $\theta = f(D) + Z$ with $Z \sim \text{Lap}(b)$.
- **Strategy:** We calibrate b to obtain BDP using our theorem.

Multivariate Gaussian Correlation (Impact on Real Databases)

- **Use-case:** Sum queries with Laplace mechanism. $\theta = f(D) + Z$ with $Z \sim \text{Lap}(b)$.
- **Strategy:** We calibrate b to obtain BDP using our theorem.
- **Utility metric:** We set $\beta = 0.05$ (i.e., 95% confidence) and measure (α, β) -accuracy, both theoretically (—) and empirically ($\times$).

Multivariate Gaussian Correlation (Impact on Real Databases)

- **Use-case:** Sum queries with Laplace mechanism. $\theta = f(D) + Z$ with $Z \sim \text{Lap}(b)$.
- **Strategy:** We calibrate b to obtain BDP using our theorem.
- **Utility metric:** We set $\beta = 0.05$ (i.e., 95% confidence) and measure (α, β) -accuracy, both theoretically (—) and empirically ($\times$).

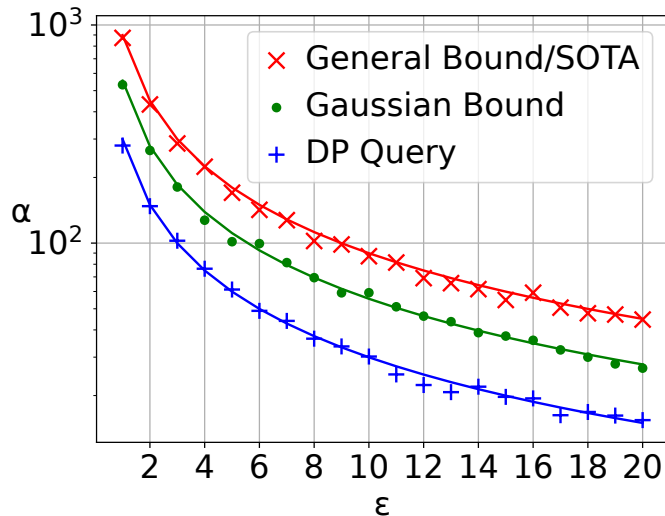


Figure: Galton, $n = 897$ $m = 3$

Key takeaway:

Substantial utility gains compared to the general bound!

- More experiments with different real and synthetic datasets in our paper show similar results.

Markov Chain Correlation Model

Main result (Informal)

- Let $\mathcal{M}$ be an ε -DP mechanism,
- input data sampled from Markov chain with transition matrix $P \in \mathbb{R}^{s \times s}$ and initial distribution $w \in \mathbb{R}^s$ with the following properties:

(H1) For all $x, y \in \mathcal{S}$ we have $P_{x,y} > 0$ and, (H2) $wP = w$.

Then, $\mathcal{M}$ is an $(\varepsilon + 4 \ln \gamma)$ -BDP mechanism where $\gamma = \frac{\max_{x,y \in \mathcal{S}} P_{xy}}{\min_{x,y \in \mathcal{S}} P_{xy}}$.

Markov Chain Correlation Model

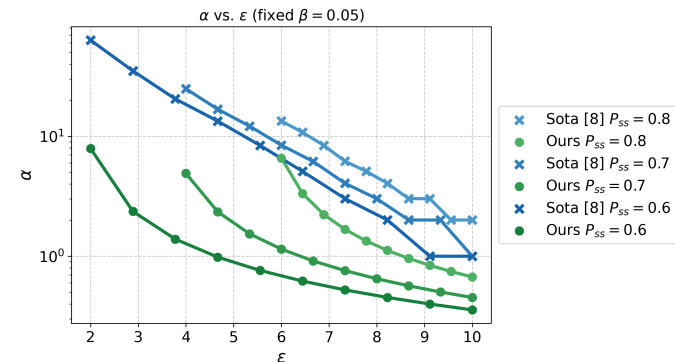
Main result (Informal)

- Let $\mathcal{M}$ be an ε -DP mechanism,
- input data sampled from Markov chain with transition matrix $P \in \mathbb{R}^{s \times s}$ and initial distribution $w \in \mathbb{R}^s$ with the following properties:

(H1) For all $x, y \in \mathcal{S}$ we have $P_{x,y} > 0$ and, (H2) $wP = w$.

Then, $\mathcal{M}$ is an $(\varepsilon + 4 \ln \gamma)$ -BDP mechanism where $\gamma = \frac{\max_{x,y \in \mathcal{S}} P_{xy}}{\min_{x,y \in \mathcal{S}} P_{xy}}$.

Previous mechanism	Ours
$P_{xy} > 0$	$P_{xy} > 0$
stationary	stationary
lazy	
binary	
symmetric	
$\varepsilon' > 0$	$\varepsilon' > 4 \ln(\gamma)$



Markov Chain Correlation Model (Impact on Real Databases)

- **Use-case:** Counting queries with Laplace mechanism. $\theta = f(D) + Z$ with $Z \sim \text{Lap}(b)$.

Markov Chain Correlation Model (Impact on Real Databases)

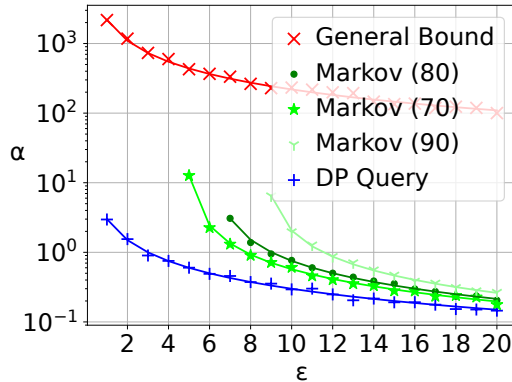
- **Use-case:** Counting queries with Laplace mechanism. $\theta = f(D) + Z$ with $Z \sim \text{Lap}(b)$.
- **Strategy:** We calibrate b to obtain BDP using our theorem.

Markov Chain Correlation Model (Impact on Real Databases)

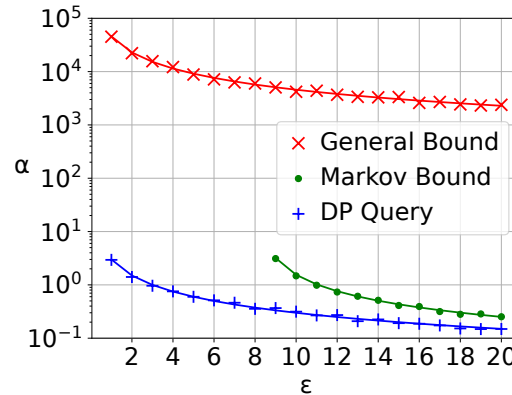
- **Use-case:** Counting queries with Laplace mechanism. $\theta = f(D) + Z$ with $Z \sim \text{Lap}(b)$.
- **Strategy:** We calibrate b to obtain BDP using our theorem.
- **Utility metric:** We set $\beta = 0.05$ (i.e., 95% confidence) and measure $-(\alpha, \beta)$ -accuracy, $\times$ upper bound of a $(1 - \beta)$ confidence interval for the absolute query error.

Markov Chain Correlation Model (Impact on Real Databases)

- **Use-case:** Counting queries with Laplace mechanism. $\theta = f(D) + Z$ with $Z \sim \text{Lap}(b)$.
- **Strategy:** We calibrate b to obtain BDP using our theorem.
- **Utility metric:** We set $\beta = 0.05$ (i.e., 95% confidence) and measure $-(\alpha, \beta)$ -accuracy, $\times$ upper bound of a $(1 - \beta)$ confidence interval for the absolute query error.



(g) Electricity, $n = 731$.



(h) Activity, $n = 17\,568$.

Key takeaway:

- **Substantial utility gains** compared to the general bound!
- Markov bound independent of n
 $\Rightarrow$ **huge improvement for large datasets.**

Conclusion

- ✓ We provide a **feasible method** to generate a **BDP mechanism** by **recalibrating** existing DP methods, tailored to **Gaussian** and **Markov** models.

Conclusion

- ✓ We provide a **feasible method** to generate a **BDP mechanism** by **recalibrating** existing DP methods, tailored to **Gaussian** and **Markov** models.
- ✓ We offer **significantly better utility than prior results**.

Conclusion

- ✓ We provide a **feasible method** to generate a **BDP mechanism** by **recalibrating** existing DP methods, tailored to **Gaussian** and **Markov** models.
- ✓ We offer **significantly better utility than prior results**.

Key takeaway:

BDP becomes usable when correlations are structured.



Conclusion

- ✓ We provide a **feasible method** to generate a **BDP mechanism** by **recalibrating** existing DP methods, tailored to **Gaussian** and **Markov** models.
- ✓ We offer **significantly better utility than prior results**.

Key takeaway:

BDP becomes usable when correlations are structured.

Future Work:

- Other distributions ?
- Can we build methods from scratch instead or recycling ?
- What if we calibrate directly to the attack advantage ?



Conclusion

- ✓ We provide a **feasible method** to generate a **BDP mechanism** by **recalibrating** existing DP methods, tailored to **Gaussian** and **Markov** models.
- ✓ We offer **significantly better utility than prior results**.

Key takeaway:

BDP becomes usable when correlations are structured.

Future Work:

- Other distributions ?
- Can we build methods from scratch instead or recycling ?
- What if we calibrate directly to the attack advantage ?



Paper



Code

Backup Slides

Experiment Details

Database	n	m	Parameters	Sensitivity
Galton	897	3	$\rho = 0.275$	$\Delta q = 254cm$
FamilyIQ	868	2	$\rho = 0.4483$	$\Delta q = 120$
SyntheticIQ	20000	2	$\rho = 0.45$	$\Delta q = 120$
Activity	17568	n	$\gamma = 7.54$	$\Delta q = 1$
Activity Single Day	288	n	$\gamma = 7.54$	$\Delta q = 1$
Electricity	731	n	70 kWh, $\gamma = 3.29$ 80 kWh, $\gamma = 4.49$ 90 kWh, $\gamma = 8.43$	$\Delta q = 1$

Table: Data description. m is the max number of correlated records and n the total amount.

Multivariate Gaussian More Results

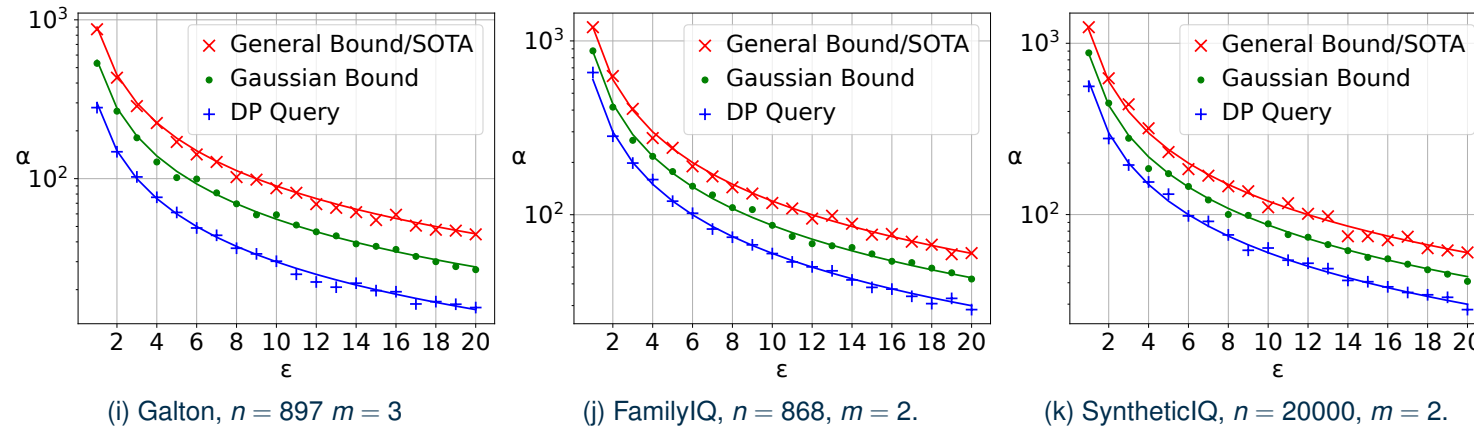


Figure: Gaussian data results. Lines show theoretical error at $\beta = 5\%$ and markers indicate empirical 95% upper bounds.