



情報系Winter Festa Episode 3

Genome Privacy CREST

Jun Sakuma, U. Tsukuba (privacy)

Koji Tsuda, U. Tokyo (stat.)

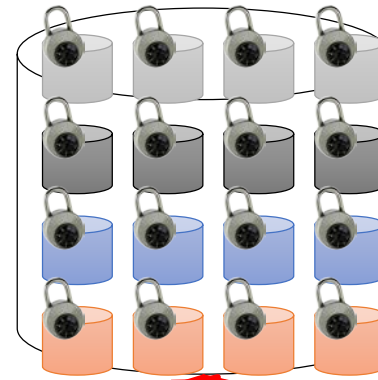
Ichiro Takeuchi, NITECH (ML)

Kunihiro Noboru, U. Tokyo (crypto)

Yoshiji Yamada, U. Mie (genetics)

目標: 機密性・プライバシー保護の障壁を取り除き、孤立したsmall dataをBIG DATAに

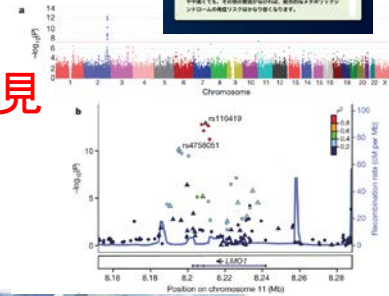
孤立したSmall data



新サービス

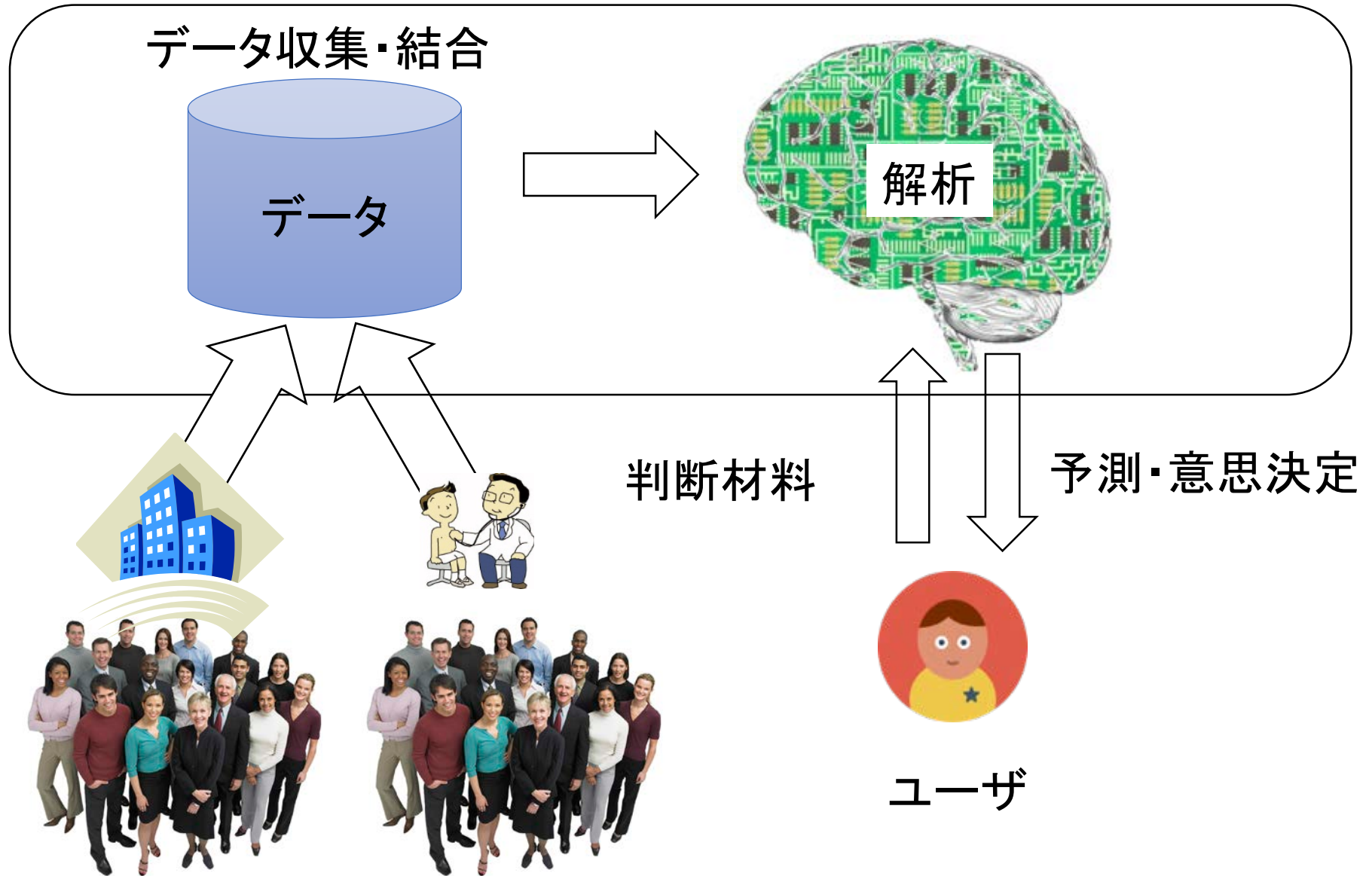


新発見



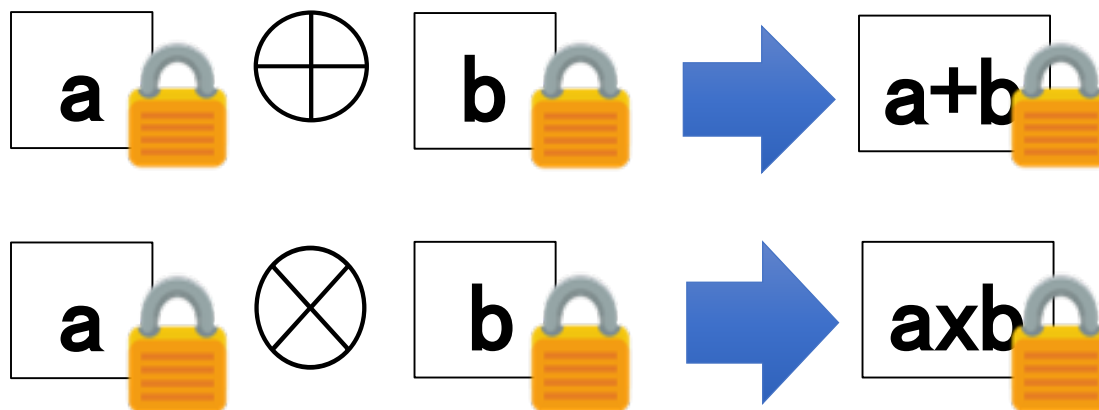
統計解析・機械学習におけるプライバシーの問題

1. 個人情報（同意なしには）結合できない



完全準同型暗号

- データを暗号化したまま加算・乗算
- 原理的には任意の演算を「データを見せずに」評価可能



Gentry Craig, “A fully homomorphic encryption scheme”,
Doctoral dissertation, Stanford University, 2009

準同型暗号を用いた医療・ゲノムデータに基づく生活習慣病予防システム

遺伝要因によるリスク

環境/臨床要因によるリスク

$$\text{risk} = \mathbf{w}^T \mathbf{x} = w_1^G x_1^G + w_2^G x_2^G + \dots + w_1^C x_1^C + w_2^C x_2^C + \dots$$

SNP1の有無
SNP1が疾病に及ぼす影響

環境/臨床要因 (e.g. 年齢、
性別、飲酒歴、喫煙歴、
血液検査結果 etc)

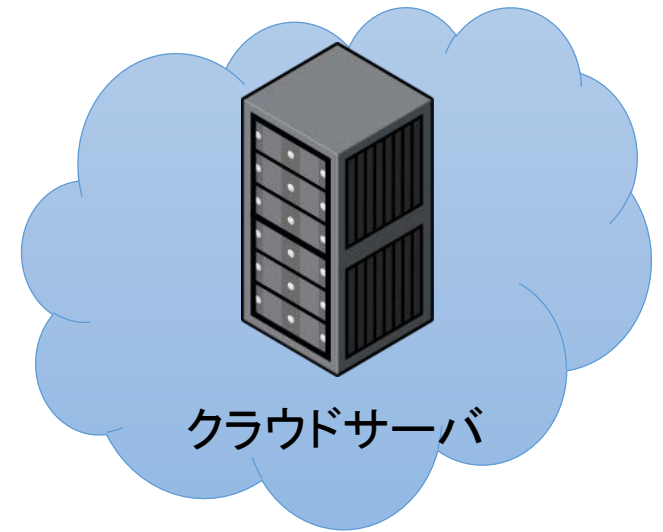
暗号文上で評価可能

その臨床要因が
疾病に及ぼす影響

$$\begin{aligned} e(\mathbf{w}^T \mathbf{x}) &= e\left(\sum_i w_i^G x_i^G + \sum_j w_j^C x_j^C\right) \\ &= \prod_i e(w_i^G x_i^G) \otimes \prod_j e(w_j^C x_j^C) = \prod_i \underbrace{e(x_i^G)^{w_i^G}}_{\text{遺伝要因の暗号文}} \otimes \prod_j \underbrace{e(x_j^C)^{w_j^C}}_{\text{臨床要因の暗号文}} \end{aligned}$$

遺伝要因の暗号文

臨床要因の暗号文

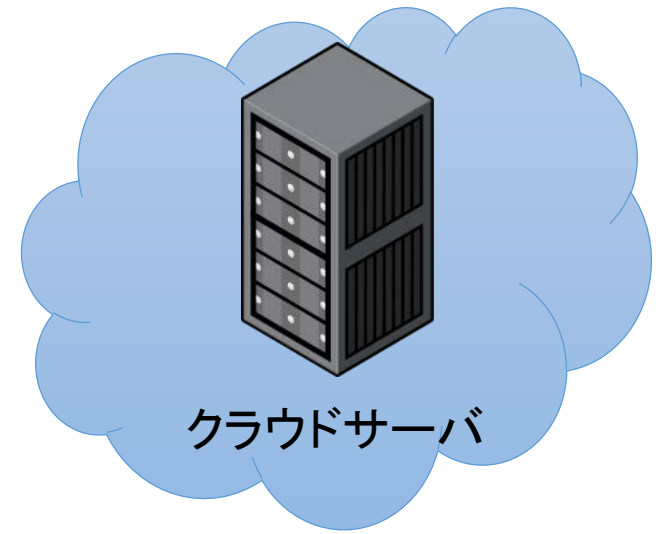


診療所



遺伝子検査機関





診療所

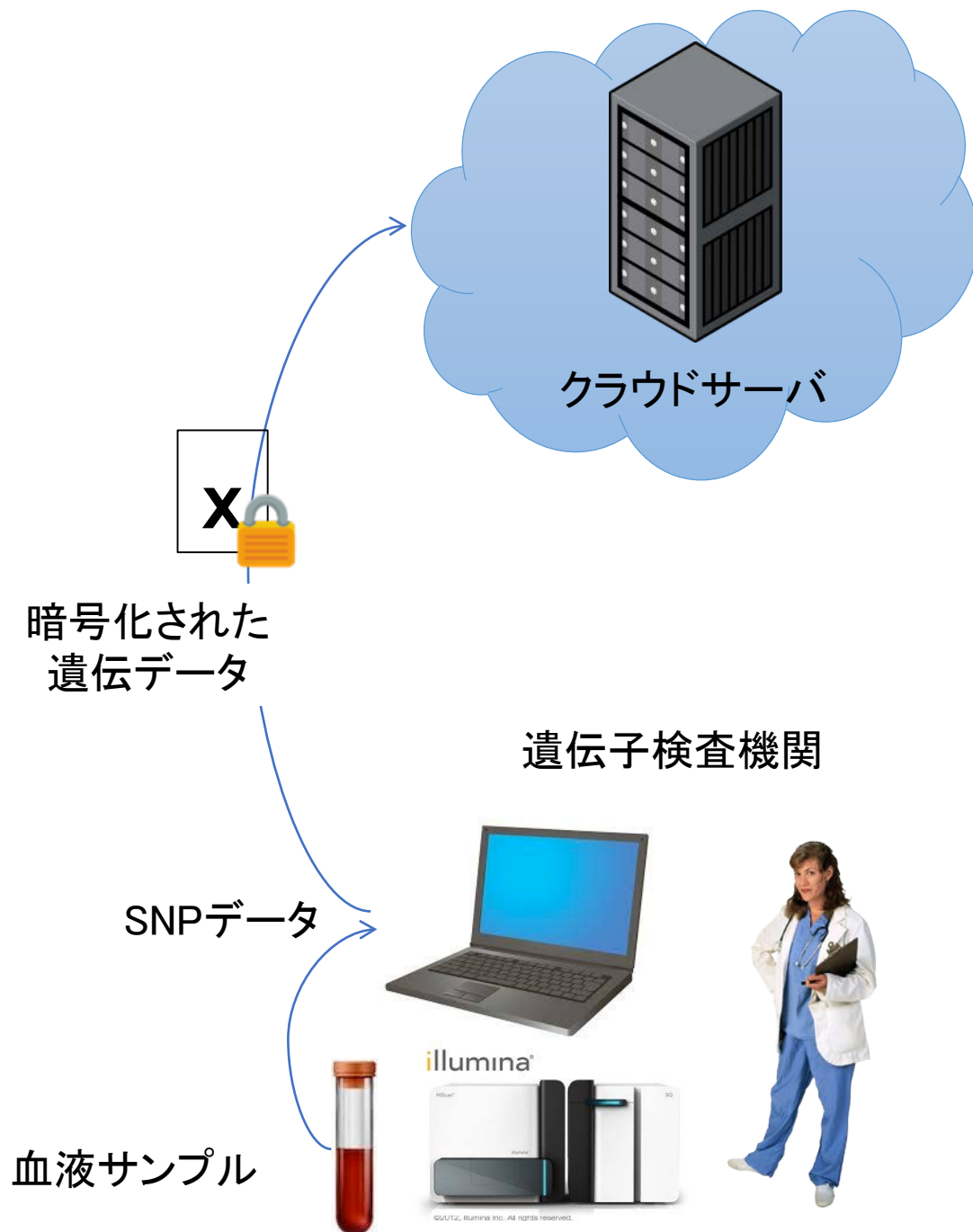


遺伝子検査機関



血液サンプル



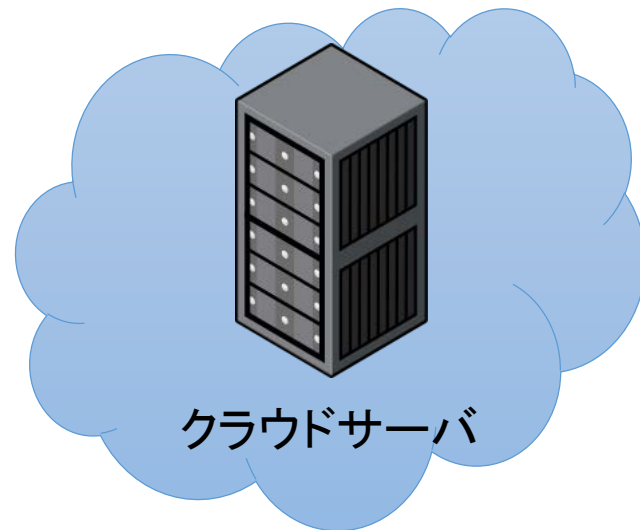
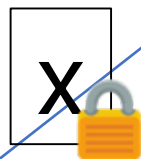




診療所



暗号化された臨床データ



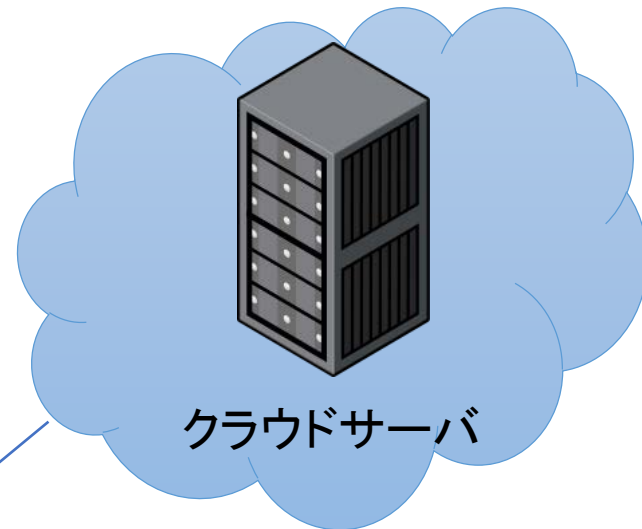
クラウドサーバ

遺伝子検査機関





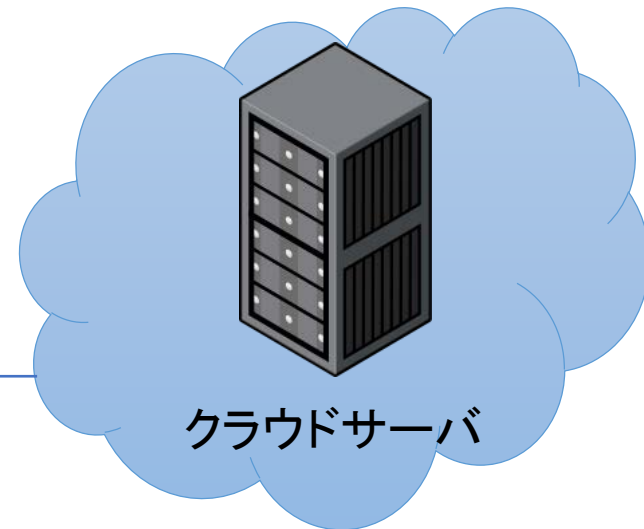
診療所



暗号化された
分析結果のレポート

遺伝子検査機関





後日総合病院へ...

診療所



現在、2病院で
実証実験中

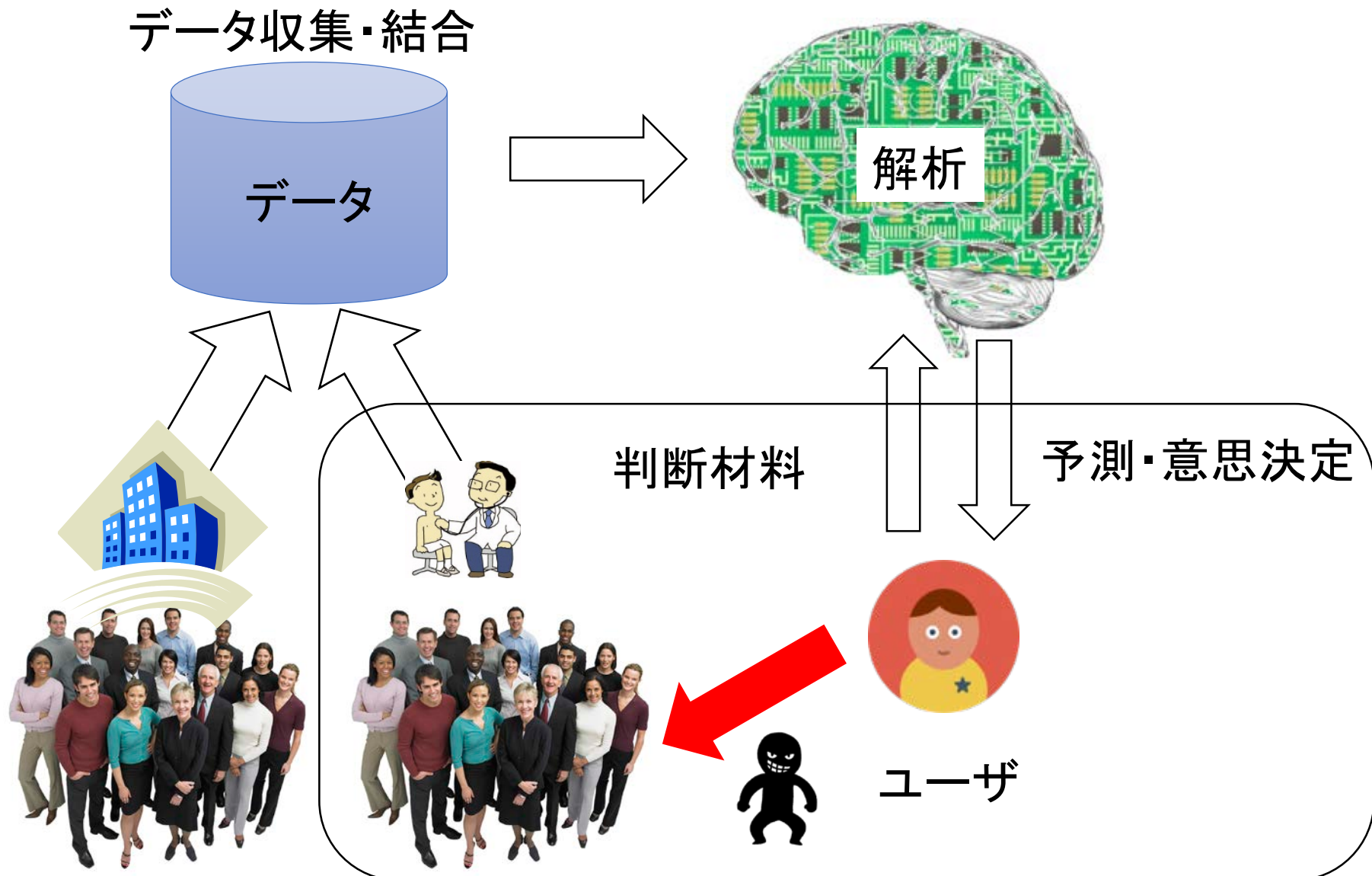
遺伝子検査機関

暗号化された
分析結果のレポート



統計解析・機械学習におけるプライバシーの問題

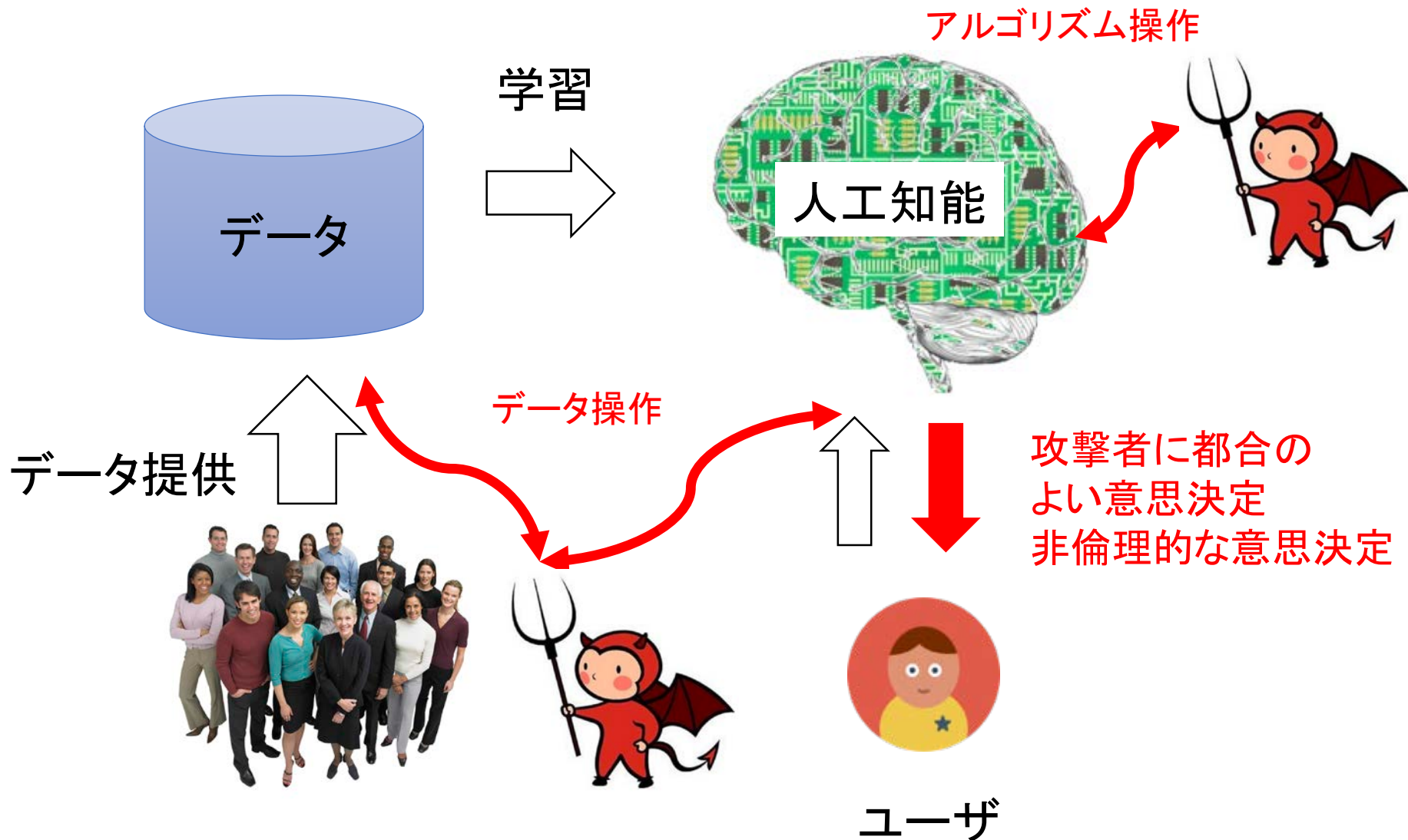
2. 公開統計情報から、個別入力が推測可能か評価できない



出力から入力が推定されないためには？

- ランダム化による保護
 - 出力をノイズで汚せば、入力はバレない
 - カイ二乗検定の差分プライバシー[Kakizaki+, ICML '17]
- 区間化による保護
 - 出力を区間化すれば、入力はバレない
 - 列挙によるリスク評価[Kusano+, ASIACCS' 17]
 - BDDによる近似リスク評価, プレビュー@竹内(東大)
 - 区間化データのERM[Hanada+, AAAI' 18], プレビュー@花田(名工大)

AIセキュリティ・プライバシー@理研AIP

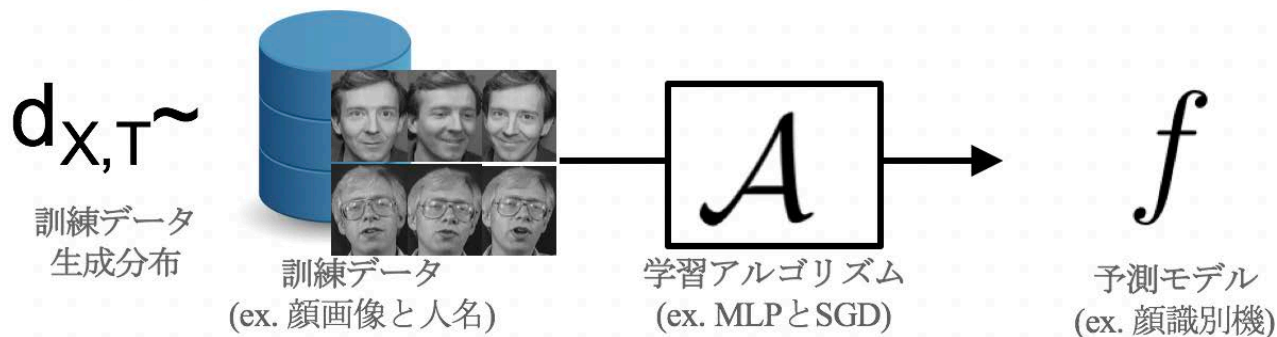


差別を起こさないバンディットの解析→レビュー@福地(筑波)

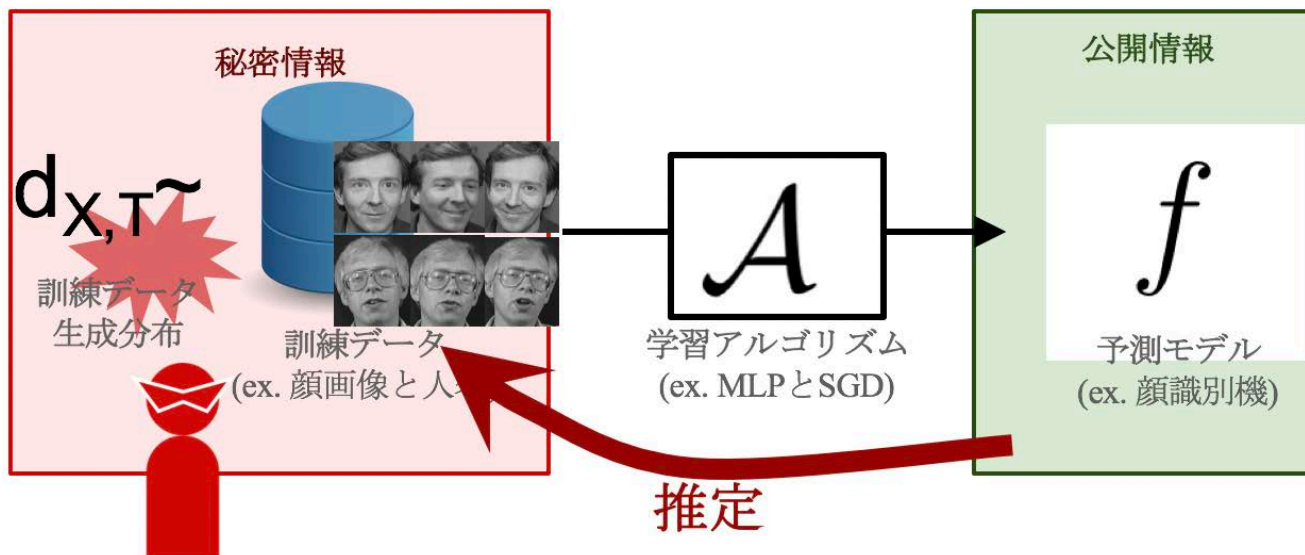
Classifier-to-Generator 攻撃

Kusano+, under review

- 通常の機械学習



- 分類器から訓練データを復元できるか？



畳み込みニューラルネットワーク の秘密計算プロトコル

Setting:

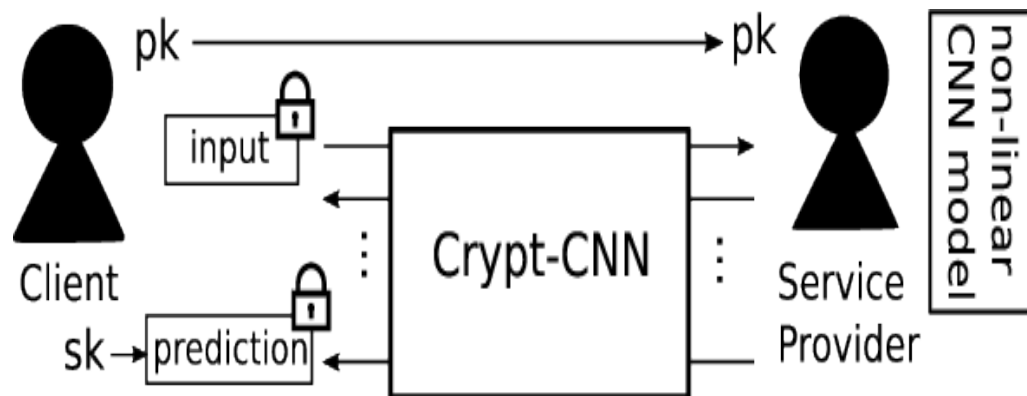
Both client's input and CNN model are private.

Objective:

To evaluate a CNN without disclosing any private information.

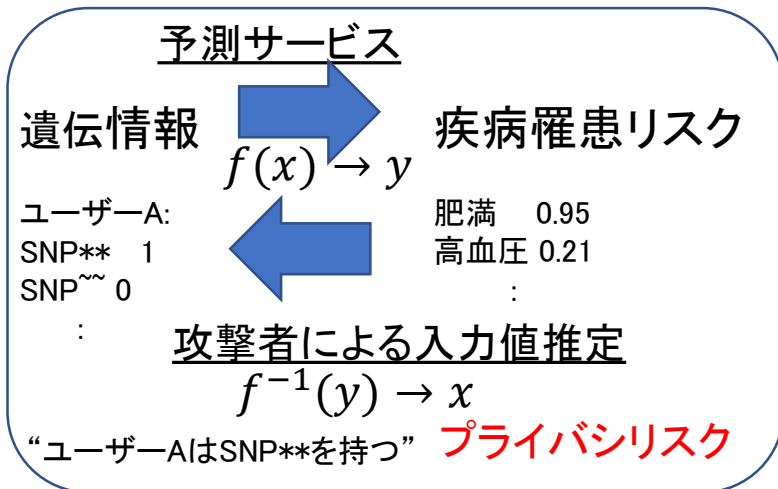
Keywords:

Homomorphic encryption, Garbled circuit, Secure computation



二分決定図を用いた、 区間値公開時のプライバシー安全性の計算

CREST 佐久間プロジェクト 東京大学 竹内 聖悟



安全な(推定されにくい)設計
(例: 数値ではなく区間値を使う)
安全性評価:

α -obscure安全[草野+2016]

目的: 高速な計算

今回の対象:

SNP情報からの疾病罹患リスク予測
区間値: $lb \leq f(x) \leq ub$, 複数の疾患

⇒ $Ax \leq b$ の解 x の列挙

$x \in \{0,1\}^d, A \in \mathbb{Z}^{m \times d}, b \in \mathbb{Z}^m$

既存手法: 二分決定図(BDD)の利用

提案: より効率的な計算手法

α -obscure安全:

$$\alpha \geq |\Pr[X_i = k | Y = y] - \Pr[X_i = k]|$$

事前分布と事後分布の差から評価

事前分布 $\Pr[X_i = k]$

事後分布 $\Pr[X_i = k | Y = y] =$

$$\frac{\sum_{x \in \{x | x \in f^{-1}[\{y\}], x_i = k\}} \Pr[X=x]}{\sum_{x \in f^{-1}[\{y\}]} \Pr[X=x]}$$

解の列挙

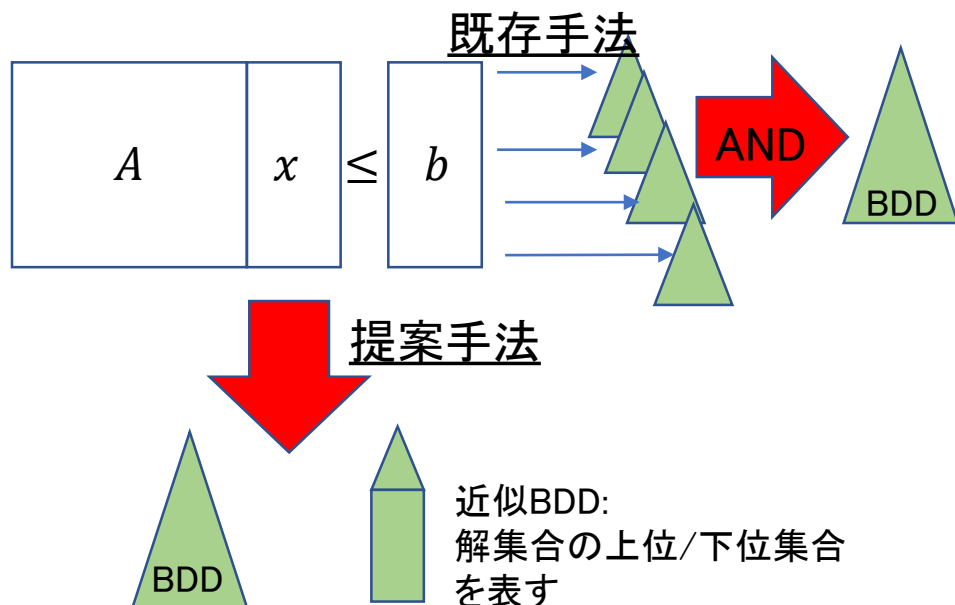
二分決定図を用いた、 区間値公開時のプライバシー安全性の計算

BDDを用いた、解の列挙:

- 既存手法 [Behle 2007]
 - 個々の不等式のBDDを構築, AND演算 $\rightarrow$ 解
 - AND演算は重い: 入力BDDサイズの乗算
- 提案手法1: 一度に全不等式を扱う
 - 関連: フロントティア法 [Kawahara+2017]
- 提案手法+: 近似BDDの構築
 - Relaxed, Restricted BDD [Bergman+2013]

	Behle	提案法	提案法(近似)
計算量			

W : BDD幅, w : 近似のパラメータ, BDD幅

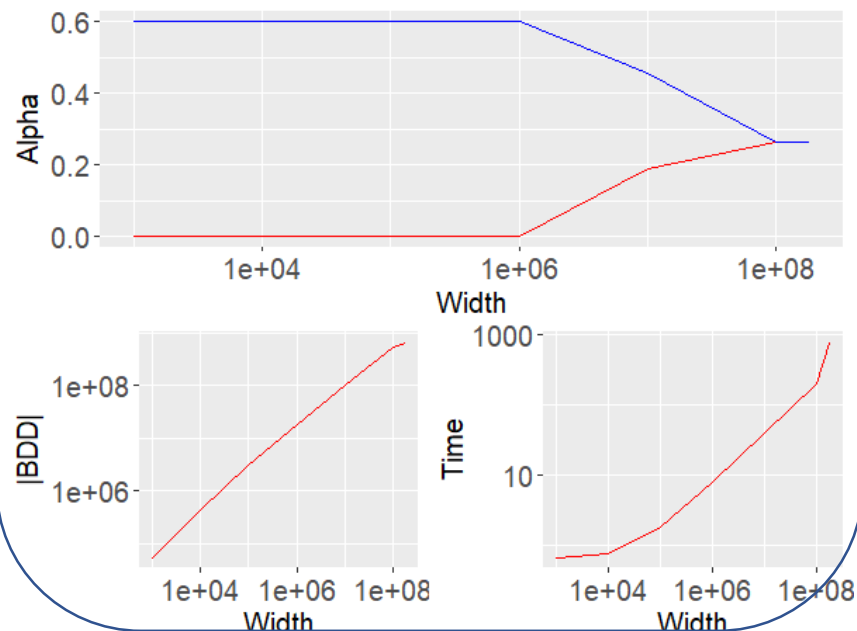


実験結果:

$m = [9, 18], d = 99$

既存手法: ≥ 24 時間 で解けず.

提案手法: $\leq 1,000$ 秒



結果的公平な文脈付きバンディット学習

福地 一斗 (筑波大), 佐久間 淳 (筑波大/JST/理研)

文脈付きバンディット

例) 広告推薦



クリック率が上がるように
広告配信を最適化したい

公平性

例) (Sweeny 13)の差別的広告に関する報告

アフリカ系の
名前

ヨーロッパ系の
名前



ネガティブな
広告



中立的な
広告

センシティブ属性(人種など)に依存しない
広告配信をおこないたい

$$\max \text{累積報酬} \quad \text{sub to} \quad \text{不公平スコア} \leq \eta$$

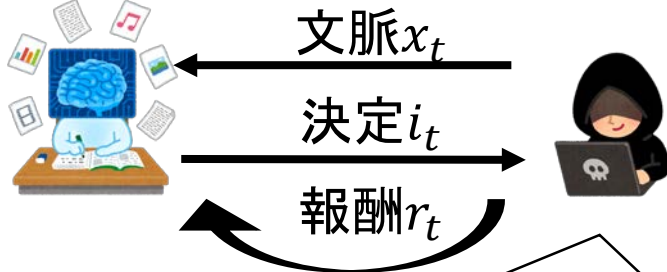
結果的公平な文脈付きバンディット学習

福地 一斗 (筑波大), 佐久間 淳 (筑波大/JST/理研)

研究目的

公平性を保証した文脈付きバンディットアルゴリズムの開発・解析

敵対的文脈

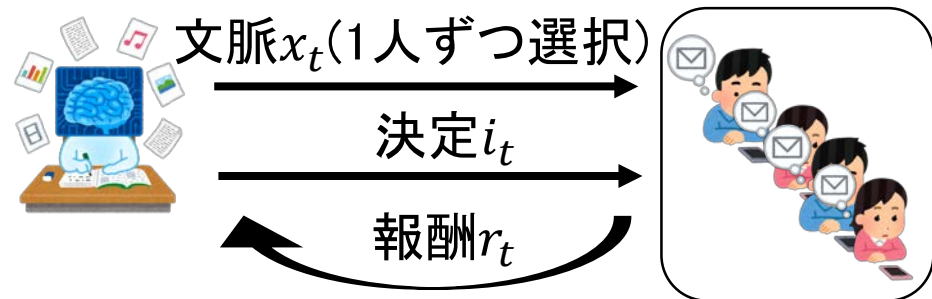


累積報酬が下がるように x_t を決定

敵対的文脈は難しすぎて
解けないことを示す

分配的文脈

登録ユーザ集合 L



- 列線系レグレットを得られる
アルゴリズムを提案・解析
- 一定条件下における最適性を証明