

Why this work

- Tabular data = prevalent+ for value potential Chui et al. 2018
- Generative SOTA for tabular data = DL based **unlike** SOTA for classification (tree based)
- Dissatisfaction, new directions needed for generative approaches Camino et al. 2020

Starting point

- Map key frameworks for tabular data classification onto generative world

- ↳ **Loss functions:** properness Savage, 1971 → GAN-game losses *from discriminator, tight* if discriminator *calibrated*, chi square “*universal*” to train generator
- ↳ **Models:** tree-based Breiman et al., 1984 → New *tree-based* generative models: XAI+, density in $O(\text{depth})$, likelihood w/ *missing features* in $O(\text{size})$
- ↳ **Algorithms:** boosting Kearns & Mansour, 1996 → *Boosting* algorithm for adversarial training +geometric conv. of chi square under *weak generative assumption*
New *copycat training* (unknown for DL), “boosting for free” convergence if boosted discriminator (C4.5, etc.)

Loss functions

▲ No discussion on generalisation

Models

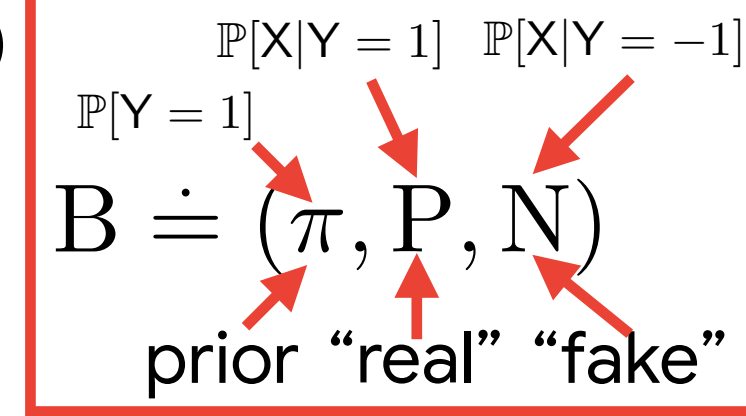
Algorithms

▲ The prior is under control, so we pick $\pi=0.5$

Discriminative

(for class probability estimation)

Binary task



Generative

(for measure comparison)

Objective: learn sampler for N that looks like P

Ideal: P

Objective: learn posterior η to estimate $P[Y=1|X]$

Ideal: $\eta^* = \pi \cdot (dP/dM)$

Bayes posterior $M \doteq \pi \cdot P + (1-\pi) \cdot N$

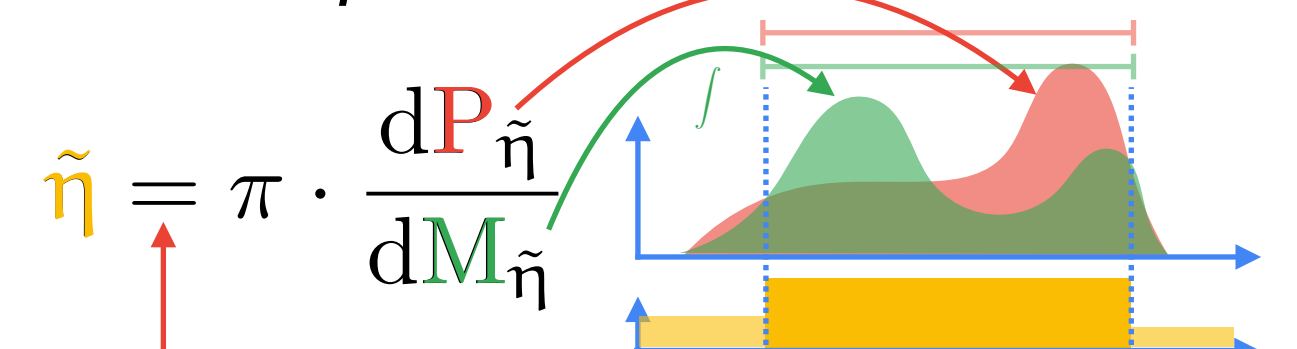
Loss: $\ell: \{-1, 1\} \times [0, 1] \rightarrow \mathbb{R}$
true class posterior estimate



$\underline{L}(p) \doteq \inf_u \mathbb{E}_{Y \sim B(p)} [\ell(Y, u)]$ Bayes risk

$p = \arg \inf \Rightarrow \ell$ **strictly proper**

Calibrated posterior:



agreement on the level sets of the posterior

- ↳ Best calibrated posterior: Bayes η^* ; worst: prior π
- ↳ Important: **any decision tree** with local posterior predictions at the leaves is calibrated

Statistical information of calibrated $\tilde{\eta}$

$$\Delta \underline{L}(\tilde{\eta}, M_{\tilde{\eta}}) = \underline{L}(\pi) - \mathbb{E}_{X \sim M_{\tilde{\eta}}} [\underline{L}(\tilde{\eta}(X))]$$

CART, C4.5, etc. (splitting criterion)

The better $\tilde{\eta}$, the **higher** $\Delta \underline{L}(\tilde{\eta}, M_{\tilde{\eta}})$

Information of Binary Task B

$$\mathbb{I}_f(P, N) \doteq \int f \left(\frac{dP}{dN} \right) dN$$

The better N , the **smaller** $\mathbb{I}_f(P, N)$

Theorem: for any *calibrated* $\tilde{\eta}$ and any *strictly proper symmetric differentiable* loss ℓ ,

$$\Delta \underline{L}(\tilde{\eta}, M_{\tilde{\eta}}) \doteq H + G \doteq \mathbb{I}_{f\pi}(P_{\tilde{\eta}}, N_{\tilde{\eta}})$$

Tight, unlike f -GAN formulation:

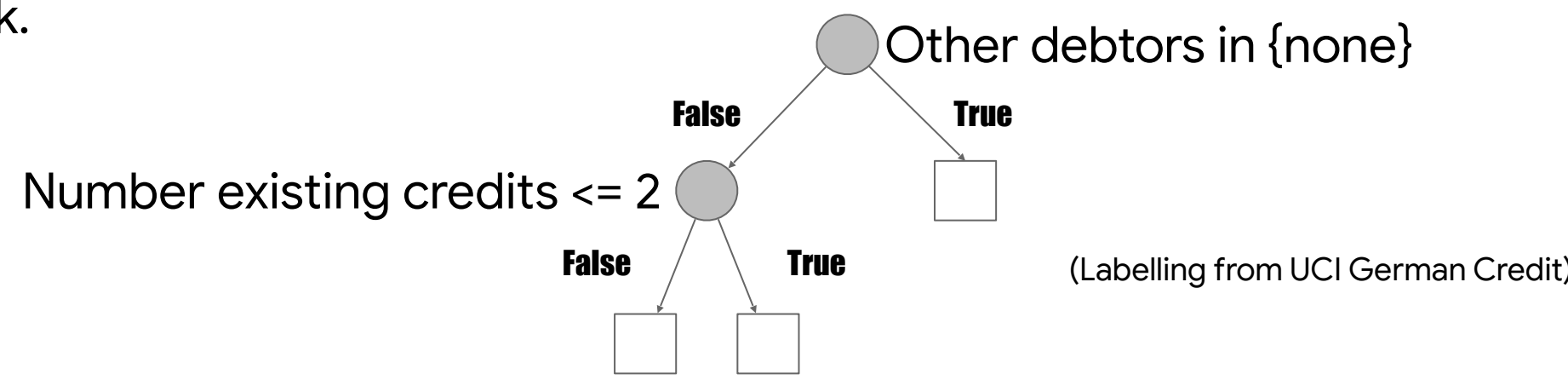
$$\mathbb{I}_f(P, N) \geq \sup_{\tilde{h}} \{ \mathbb{E}_P[\tilde{h}(X)] - \mathbb{E}_N[\tilde{h}(X)] \}$$

+Lemma: if additional property on ℓ , to minimise G , sufficient to minimise the *chi square* $\chi^2(N_{\tilde{\eta}}||P_{\tilde{\eta}})$

holds for many popular losses (log-, square-, etc.)

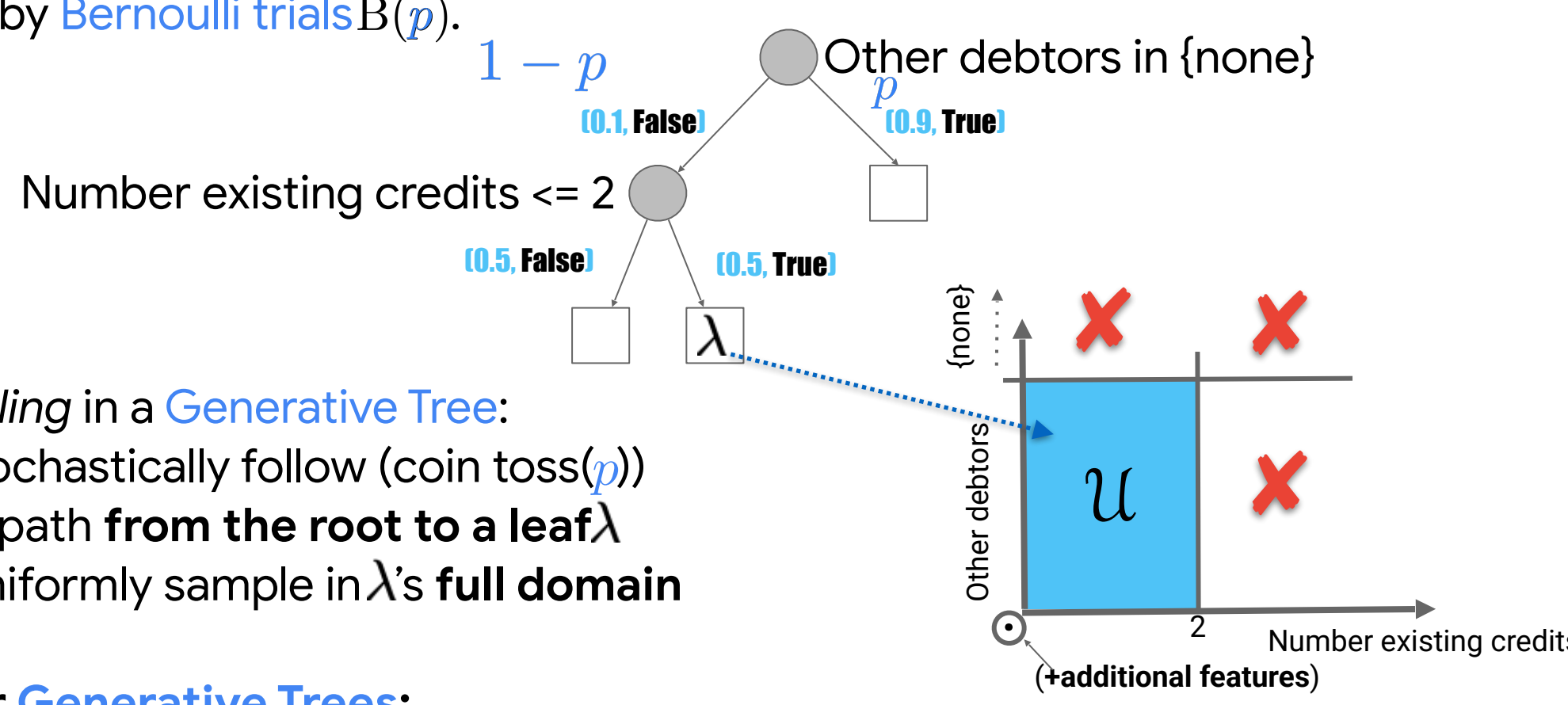
↳ gives “one loss to train N against any calibrated discriminator” — valid outside *decision trees*

Tree: a tree is a binary directed tree whose internal nodes are labeled with a test on an observation variable and outgoing arcs are labeled with truth values. Leaves are blank.



Decision Tree: A *decision tree* h is a *tree* in which leaves are labeled by values in $[0, 1]$

Generative Tree: A *generative tree* (GT) G is a *tree* in which outgoing arcs are labeled by *Bernoulli trials* $B(p)$.



↳ **Sampling** in a *Generative Tree*:

- (1) Stochastically follow (coin toss p) a path **from the root to a leaf** λ
- (2) Uniformly sample in λ 's **full domain**

Pros for Generative Trees:

- ↳ Trainable from data w/ *missing values* (see algorithms)
- ↳ XAI+: as “easy” to *interpret* as a *decision tree* (see experiments)
- ↳ *Density* computable in $O(\text{depth})$ at any point,
- ↳ *Likelihood* of missing values given partial observation in $O(\text{size})$, $P[A|B]$ equivalently easy
- ↳ Many metrics involving a GT computable *exactly* (no sampling required !): chi square, etc.

Cons for Generative Trees:

- ↳ Axis // *partition* of the support (restrictive)
- ↳ Support / domain *closed* (workarounds exist, e.g. Box-Muller transform for Gaussians)

Adversarial training

Algorithm TD-GEN(G, h)

Input: current generator G , current discriminator h ;

Output: G with a new split;

Step 1 : pick $\textcircled{S} \in \Lambda(G)$, $i \in [d]$; // leaf and variable for the current split

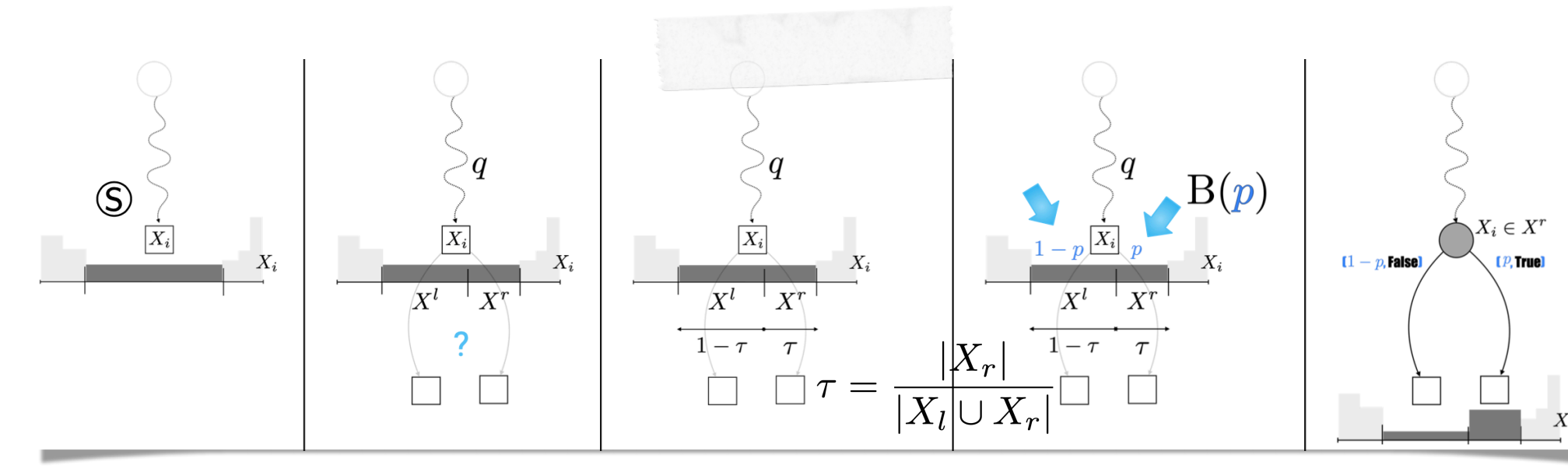
Step 2 : choose (X^l, X^r) and compute τ ; // split choice

Step 3 : compute p as

$$p \leftarrow \text{CLAMP} \left(\frac{\mu_{LL} - \mu_{LR}}{\mu_{LL} + \mu_{RR} - 2\mu_{LR}} \right);$$

parameters depend on G, h and τ (see paper)

Step 4 : replace $\textcircled{S}$ by a split as designed in Steps 1,2 w/ Bernoulli probability p as in (16);



Top-down: pick a leaf, variable, split, compute Bernoulli's p , *repeat*

Convergence if p in $(0, 1)$

↳ if $\chi^2(N_{\tilde{\eta}}||P_{\tilde{\eta}}) \geq \delta$ and $|\tau - p| \geq \varepsilon$

new old

$$\chi^2(N_{\tilde{\eta}}||P_{\tilde{\eta}}) \leq \frac{1}{1 + \delta \varepsilon^2} \cdot \chi^2(N_{\tilde{\eta}}||P_{\tilde{\eta}})$$

Convergence if p in $\{0, 1\}$ ▲ *discard support* ▲

↳ if weak generative assumption ($\delta > 0$)

($\delta=0 \Rightarrow$ complete independence between generation and discrimination)

$$\chi^2(N_{\tilde{\eta}}||P_{\tilde{\eta}}) \leq \frac{1}{1 + \delta q^2(\tau + (1 - 2\tau)p)^2} \cdot \chi^2(N_{\tilde{\eta}}||P_{\tilde{\eta}})$$

Copycat training

Generator and **Discriminator** contain a **tree**

In copycat training, the **Discriminator** h is learnt by a boosting algorithm (C4.5, CART, ...)

↳ After each iteration for h , the **Generator** G copies the update of the **tree** of h in its **tree**

and computes its Bernoulli probabilities to keep $\chi^2(N_{\tilde{\eta}}||P_{\tilde{\eta}}) = 0$

↳ After the update on G , G updates the training sample of h by resampling what has been affected by its update $\Rightarrow$ the error of h on the new sample is 50%

↳ Repeat with a new update to the **tree** of h

Pros for Copycat training

- ↳ Straightforward to implement to build the **Generator**
- ↳ Almost zero algorithmic overhead over the **Discriminator**'s boosting algorithm

Cons for Copycat training

- ↳ No privacy for the **Discriminator**

Boosting-compliant convergence “for free”

↳ boosting rates in density ratio loss for the **Generator** follow directly from

boosting rates for the **Discriminator**

↳ Let $\ell =$ *Matusita's loss* ($\underline{L}(u) = 2\sqrt{u(1-u)}$),

which means we run Kearns & Mansour's top-down DT

induction algorithm (Kearns & Mansour, STOC'96)

↳ Suppose boosting's Weak Hypothesis Assumption WHA($\delta > 0$) holds to train h

↳ For any $\varepsilon > 0$, if #iters T for h satisfies

$$T \geq \left(\frac{1}{f(\varepsilon)} \right)^{\frac{2}{\delta}} \in (0, 1)$$

then the **Generator** G satisfies

$$\mathcal{B}_{\ell}(P, G) \leq \varepsilon \cdot \mathbb{I}_{f\pi}(P, U)$$

density ratio loss

Experiments

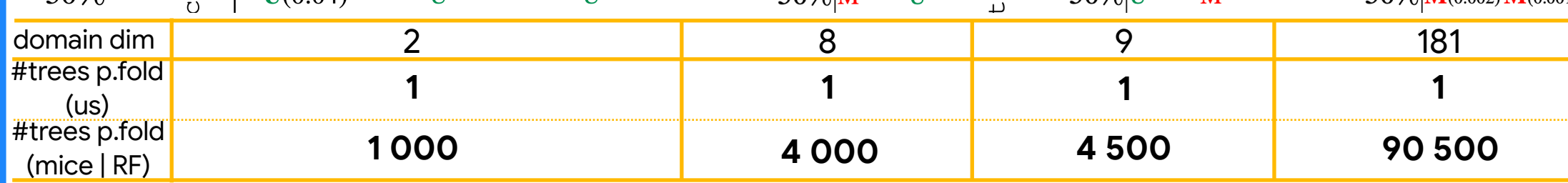
Missing data imputation

Setup:

- ↳ Remove $q\%$ values in dataset (*Missing Completely At Random*), 5-fold CV
- ↳ Impute missing values w/ algorithm, compare with domain using OT (W_2^2) metric

Some results (more in paper)

		U = we win; M = mice wins (bold = p-val < .1, indicated)							
Us vs Mice	#trees p. fold	NORM	CART	RF	Us vs Mice	#trees p. fold	NORM	CART	RF
us	1000	40	4000	45	4500	905	90500		
(q=)5%		U(0.003)	U	U(0.07)	5%	U(0.004)	U	U(0.003)	M(0.002)
10%		U(0.03)	U(0.002)	U(0.007)	10%	U(0.004)	M	10%	M(0.003)
20%		U(0.0005)	U(0.001)	U(0.001)	20%	U	U	20%	M(0.001)
50%		U(0.04)	U	U	50%	M	U	50%	M(0.002)



Observations:

- ↳ GTs can beat mice if small dim (unexpected) & always with 1 tree (vs many for mice)
- ↳ training time: GTs << mice (especially w/ RFs)
- ↳ +GTs provide generator (mice: no byproduct)

us = copycat training a GT G for max 10K splits

SOTA: mice (van Buuren, 2018)

↳ After initial guess, round-robin updates one column's missing values given the others using regression/classification *method*, iterates 5 times

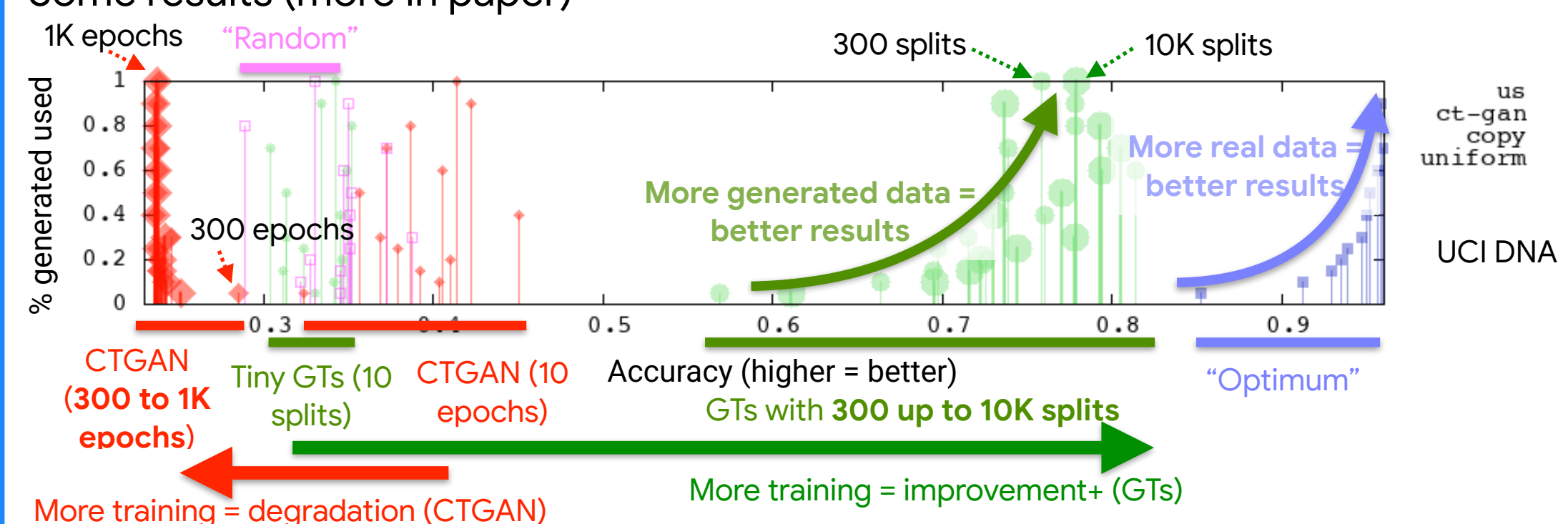
↳ We used *method* in {NORM, CART, Random Forests = RF with 100 trees each}

Generated data augmentation

Setup:

- ↳ Take a **supervised** domain (e.g. UCI iris), 5-fold CV
- ↳ Train generator, generate additional q in {5, 10, ...100}% (of training) examples, train supervised model (GBDT or RFs) from all data, get accuracy/RMSE on test
- ↳ 2 baselines: *random* = +uniform data (“worst”), *copy* = +real data (“optimum”)

Some results (more in paper)



us = copycat training a GT G for 10, 300 or 10K max splits

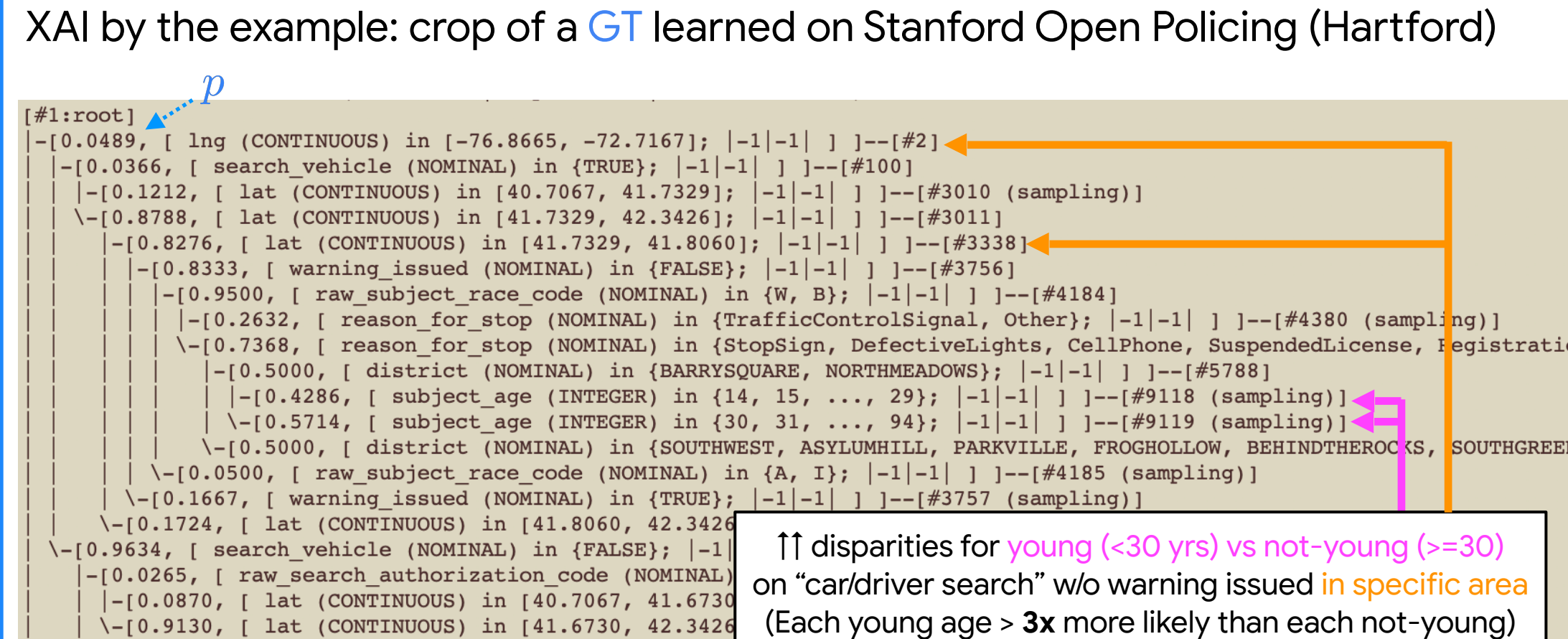
SOTA: CTGAN (Xu et al., NeurIPS'19), trained for 10, 300 or 1K epochs

... and more

Additional experiments in paper:

- ↳ Training from just generated data (vs CTGAN)
- ↳ Distinguishing real from fake data (vs CTGAN)

XAI by the example: crop of a GT learned on Stanford Open Policing (Hartford)



Conclusion and future work:

- ↳ Contribution on losses applies beyond our framework to any calibrated disc.
- ↳ Interesting avenues on alleviating the Cons for **Generative Trees**
- ↳ Experiments display edge vs CTGAN + potential use for missing data imputation