
Adversarial Random Forests for Density Estimation and Generative Modeling

David S. Watson
King's College London

Kristin Blesch
Leibniz Institute for Prevention
Research and Epidemiology – BIPS,
University of Bremen

Jan Kapar
Leibniz Institute for Prevention
Research and Epidemiology – BIPS,
University of Bremen

Marvin N. Wright
Leibniz Institute for Prevention
Research and Epidemiology – BIPS,
University of Bremen,
University of Copenhagen

Abstract

We propose methods for density estimation and data synthesis using a novel form of unsupervised random forests. Inspired by generative adversarial networks, we implement a recursive procedure in which trees gradually learn structural properties of the data through alternating rounds of generation and discrimination. The method is provably consistent under minimal assumptions. Unlike classic tree-based alternatives, our approach provides smooth (un)conditional densities and allows for fully synthetic data generation. We achieve comparable or superior performance to state-of-the-art probabilistic circuits and deep learning models on various tabular data benchmarks while executing about two orders of magnitude faster on average. An accompanying R package, `arf`, is available on CRAN.

1 INTRODUCTION

Density estimation is a fundamental unsupervised learning task, an essential subroutine in various methods for data imputation (Efron, 1994; Rubin, 1996), clustering (Bramer, 2007; Rokach and Maimon, 2005), anomaly detection (Chandola et al., 2009; Pang et al., 2021), and classification (Lugosi and Nobel, 1996; Vincent and Bengio,

2002). One important application for density estimators is generative modeling, where we aim to create synthetic samples that mimic the characteristics of real data. These simulations can be used to test the robustness of classifiers (Song et al., 2018; Buzhinsky et al., 2021), augment training sets (Ravuri and Vinyals, 2019; Lopez et al., 2018), or study complex systems without compromising the privacy of data subjects (Augenstein et al., 2020; Yelmen et al., 2021).

The current state of the art in generative modeling relies on deep neural networks, which have proven remarkably adept at synthesizing images, audio, and even video data. Architectures built on variational autoencoders (VAEs) (Kingma and Welling, 2013) and generative adversarial networks (GANs) (Goodfellow et al., 2014) have dominated the field for the last decade. Recent advances in normalizing flows (Papamakarios et al., 2021) and diffusion models (Ramesh et al., 2022) have sparked considerable interest. While these algorithms are highly effective with structured data, they can struggle in tabular settings with continuous and categorical covariates. Even when successful, deep learning models are notoriously data-hungry and require extensive tuning.

Another major drawback of these deep learning methods is that they do not generally permit tractable inference for tasks such as marginalization and conditioning, which are essential for coherent probabilistic reasoning. A family of hierarchical mixture models known as probabilistic circuits (PCs) (Vergari et al., 2020; Choi et al., 2020) are better suited to such problems. Despite their attractive theoretical properties, existing PCs can also be slow to train and are often far less expressive than unconstrained neural networks.

We introduce an adversarial random forest algorithm that vastly simplifies the task of density estimation and data synthesis. Our method naturally accommodates mixed data

in tabular settings, and performs well on small and large datasets using the computational resources of a standard laptop. It compares favorably with deep learning models while executing some 100 times faster on average. It can be compiled into a PC for efficient and exact probabilistic inference.

Following a brief discussion of related work (Sect. 2), we review relevant notation and background on random forests (Sect. 3). We motivate our method with theoretical results that guarantee convergence under reasonable assumptions (Sect. 4), and illustrate its performance on a range of benchmark tasks (Sect. 5). We conclude with a discussion (Sect. 6) and directions for future work (Sect. 7).

2 RELATED WORK

A random forest (RF) is a bootstrap-aggregated (bagged) ensemble of independently randomized trees (Breiman, 2001), typically built using the greedy classification and regression tree (CART) algorithm (Breiman et al., 1984). RFs are extremely popular and effective, widely used in areas like bioinformatics (Chen and Ishwaran, 2012), remote sensing (Belgiu and Drăguț, 2016), and ecology (Cutler et al., 2007), as well as more generic prediction tasks (Fernández-Delgado et al., 2014). Advantages include their efficiency (RFs are embarrassingly parallelizable), ease of use (they require minimal tuning), and ability to adapt to sparse signals (uninformative features are rarely selected for splits).

It is well-known that tree-based models can approximate joint distributions. Several authors advocate using leaf nodes of CART trees as piecewise constant density estimators (Ram and Gray, 2011; Wu et al., 2014; Wen and Hang, 2022). While this method provably converges on the true density in the limit of infinite data, finite sample performance is inevitably rough and discontinuous. Smooth results can be obtained by fitting a distribution within each leaf, e.g. via kernel density estimation (KDE) or maximum likelihood estimation (MLE) (Smyth et al., 1995; Gray and Moore, 2003; Loh, 2009; Ram and Gray, 2011; Criminisi et al., 2012), a version of which we develop further below. Existing methods have mostly been limited to supervised trees rather than unsupervised forests, and are often inefficient in high dimensions.

Another strategy, better suited to high-dimensional settings, uses Chow-Liu trees (Chow and Liu, 1968) to learn a second-order approximation to the underlying joint distribution (Bach and Jordan, 2003; Liu et al., 2011; Rahman et al., 2014). Whereas these methods estimate a series of bivariate densities over the full support of the data, we attempt to solve a larger number of simpler tasks, modeling univariate densities in relatively small subregions.

Despite the popularity of tree-based density estimators, they are rarely if ever used for fully synthetic data generation.

Instead, they are commonly used for *conditional* density estimation and data imputation (Stekhoven and Bühlmann, 2011; Tang and Ishwaran, 2017; Correia et al., 2020; Lundberg et al., 2020; Hothorn and Zeileis, 2021; Čevič et al., 2022). We highlight that methods optimized for this task are often ill-suited to generative modeling, since their reliance on supervised signals limits their ability to capture dependencies between features with little predictive value for the selected outcome variable(s).

Another family of methods for density estimation and data synthesis is based on probabilistic graphical models (PGMs) (Lauritzen, 1996; Koller and Friedman, 2009), e.g. Bayesian networks (Pearl and Russell, 2003; Darwiche, 2009). Learning graph structure is difficult in practice, which is why most methods impose restrictive parametric assumptions for tractability (Heckerman et al., 1995; Drton and Maathuis, 2017). PCs replace the representational semantics of PGMs with an operational semantics, encoding answers to probabilistic queries in the structural alignment of sum and product nodes. This class of computational graphs subsumes sum-product networks (Poon and Domingos, 2011), cutset networks (Rahman et al., 2014), and probabilistic sentential decision diagrams (Kisa et al., 2014), among others. Correia et al. (2020) show that RFs instantiate smooth, decomposable, deterministic PCs, thereby enabling efficient marginalization and maximization.

Deep learning approaches to generative modeling became popular with the introduction of VAEs (Kingma and Welling, 2013) and GANs (Goodfellow et al., 2014), which jointly optimize parameters for network pairs—encoder-decoder and generator-discriminator, respectively—via stochastic gradient descent. Various extensions of these approaches have been developed (Higgins et al., 2017; Arjovsky et al., 2017), including some designed for mixed data in the tabular setting (Choi et al., 2017; Jordon et al., 2019; Xu et al., 2019). While the evidence lower bound of a VAE approximates the data likelihood, there is no straightforward way to compute this quantity with GANs. More recent work in neural density estimation includes autoregressive networks (van den Oord et al., 2016; Ramesh et al., 2021; Roy et al., 2021), normalizing flows (Kobyzev et al., 2021; Papamakarios et al., 2021; Lee et al., 2022), and diffusion models (Kingma et al., 2021; Song et al., 2021; Ramesh et al., 2022). These methods are generally optimized for structured data such as images or audio, where they often attain state-of-the-art results.

3 BACKGROUND

Consider the binary classification setting with training data $\mathcal{D} = \{(\mathbf{x}_i, y_i)\}_{i=1}^n$, where $\mathbf{x}_i \in \mathcal{X} \subset \mathbb{R}^d$ and $y_i \in \mathcal{Y} = \{0, 1\}$. Samples are independent and identically distributed according to some fixed but unknown distribution P with density p . The classic RF algorithm takes B bootstrap sam-

ples of size n from $\mathcal{D}$ and fits a binary decision tree for each, in which observations are recursively partitioned according to some optimization target (e.g., Gini index) evaluated on a random subset of features at each node. The size of this subset is controlled by the `mtree` parameter, conventionally set at $\lfloor \sqrt{d} \rfloor$ for classification. Resulting splits are literals of the form $X_j \bowtie x$ for some $X_j, j \in [d] = \{1, \dots, d\}$, and value $x \in \mathcal{X}_j$, where $\bowtie \in \{=, <\}$ (the former for categorical, the latter for continuous variables). Data pass to left or right child nodes depending on whether they satisfy the literal. Splits continue until some stopping criterion is met (e.g., purity). Terminal nodes, a.k.a. *leaves*, describe hyperrectangles in feature space with boundaries given by the learned splits. These disjoint cells collectively cover all of $\mathcal{X}$. Each leaf is associated with a label $\hat{y} \in [0, 1]$, representing either the frequency of positive outcomes (soft labels) or the majority class (hard labels) within that cell. Because trees are grown on independent bootstraps, an average of n/e samples are excluded from each tree. This so-called “out-of-bag” (OOB) data can be used to estimate empirical risk without need for cross-validation.

Each new datapoint $\mathbf{x}$ falls into exactly one leaf in each tree. Predictions are computed by aggregating over the trees, e.g. by tallying votes across all B basis functions of the ensemble. Let θ_b^ℓ denote the conjunction of literals that characterize membership in leaf $\ell \in [L_b]$, where L_b is the number of leaves in tree $b \in [B]$, with corresponding hyperrectangular subspace $\mathcal{X}_b^\ell \subset \mathcal{X}$. Each leaf has some nonnegative volume and diameter, denoted $\text{vol}(\mathcal{X}_b^\ell)$ and $\text{diam}(\mathcal{X}_b^\ell)$, where the latter represents the longest line segment contained in $\mathcal{X}_b^\ell$. Let n_b be the number of training samples for tree b (not necessarily equal to n) and n_b^ℓ the number of samples that fall into leaf ℓ of b . The ratio n_b^ℓ/n_b represents an empirical estimate of the leaf’s coverage $p(\theta_b^\ell)$, i.e. the probability that a random $\mathbf{x}$ falls within $\mathcal{X}_b^\ell$. A tree is fully parametrized by $\theta_b = \bigcup_{\ell=1}^{L_b} \theta_b^\ell$, and the complete forest by $\Theta = \bigcup_{b=1}^B \theta_b$.

Many variations of the classic algorithm exist, including a number of simplified versions designed to be more amenable to statistical analysis. See (Biau and Scornet, 2016) for an overview. Common sources of variation include how observations are randomized across trees (e.g., by subsampling or bootstrapping) and how splits are selected (e.g., uniformly or according to some adaptive procedure).

Our method builds on the *unsupervised* random forest (URF) algorithm (Shi and Horvath, 2006).¹ This procedure creates a synthetic dataset $\tilde{\mathbf{X}}$ of n observations by independently sampling from the marginals of $\mathbf{X}$, i.e. $\tilde{\mathbf{x}} \sim \prod_{j=1}^d P(X_j)$. A RF classifier is trained to distinguish between $\mathbf{X}$ and $\tilde{\mathbf{X}}$, with labels indicating whether samples are original ($Y = 1$)

¹Not to be confused with RF variants that employ non-adaptive splits, which are sometimes also referred to as unsupervised, since they ignore the response variable. See, e.g., Genuer (2012).

or synthetic ($Y = 0$). The method has expected accuracy $1/2$ in the worst case, corresponding to a dataset in which all features are mutually independent. However, if dependencies are present, then a consistent learning procedure will converge on expected accuracy $1/2 + \delta$ for some $\delta > 0$ as n grows (Kim et al., 2021).

4 ADVERSARIAL RANDOM FORESTS

We introduce a recursive variant of URFs, which we call *adversarial random forest* (ARF). The goal of this algorithm is to render data jointly independent within each leaf. We achieve this by first fitting an ordinary URF $f^{(0)}$ with synthetic data $\tilde{\mathbf{X}}^{(0)}$. We compute the coverage of each leaf w.r.t. original data, then generate a new synthetic dataset by sampling from marginals within random leaves selected with probability proportional to this coverage. Call the resulting $n \times d$ matrix $\tilde{\mathbf{X}}^{(1)}$. A new classifier $f^{(1)}$ is trained to distinguish $\mathbf{X}$ from $\tilde{\mathbf{X}}^{(1)}$. If OOB accuracy for this model is sufficiently low, then the ARF has converged, and we move forward with splits from $f^{(0)}$. Otherwise, we iterate the procedure, drawing a new synthetic dataset from the splits learned by $f^{(1)}$ and evaluating performance via a new classifier. The loop repeats until convergence (see Alg. 1).

ARFs bear some obvious resemblance to GANs. The “generator” is a simple sampling scheme that draws from the marginals in adaptively selected subregions; the “discriminator” is a RF classifier. The result can be understood as a zero-sum game in which adversaries take turns increasing and decreasing label uncertainty at each round. However, beyond this conceptual link between our method and GANs lie some important differences. Both generator and discriminator share the same parameters in our algorithm. Indeed, our generator does not, strictly speaking, *learn* anything; it merely exploits what the discriminator has learned. This means that ARFs cannot be used for adversarial attacks of the sort made famous by GANs, which involve separately parametrized networks for each model. Moreover, the synthetic data generated by ARFs is relatively naïve, consisting of bootstrap samples drawn from subsets of the original observations. That is because our goal is not (yet) to generate new data, but merely to learn an independence-inducing partition. Empirically, we find that this is often achieved in just a single round even with the tolerance δ set to 0.

Formally, we seek a set of splits Θ such that, for all trees b , leaves ℓ , and samples $\mathbf{x}$, we have $p(\mathbf{x}|\theta_b^\ell) = \prod_{j=1}^d p(x_j|\theta_b^\ell)$. Call this the *local independence criterion*. Our first result states that ARFs converge on this identity in the limit. We assume:

- (A1) The feature domain is limited to $\mathcal{X} = [0, 1]^d$, with joint density p bounded away from 0 and ∞ .
- (A2) At each round, the target function $P(Y = 1|\mathbf{x})$ is Lipschitz-continuous. The Lipschitz constant may

Algorithm 1 ADVERSARIAL RANDOM FOREST

Input: Training data $\mathbf{X}$, tolerance δ
Output: Random forest classifier $f^{(0)}$

```

Sample  $\tilde{\mathbf{X}}^{(0)} \sim \prod_{j=1}^d P(X_j)$ 
 $\mathbf{X}^+ \leftarrow \text{row.append}(\mathbf{X}, \tilde{\mathbf{X}}^{(0)})$ 
 $Y \leftarrow \text{row.append}(\mathbf{1}_n, \mathbf{0}_n)$ 
 $f^{(0)} \leftarrow \text{RANDOMFOREST}(\mathbf{X}^+, Y)$ 
if  $\text{Acc}(f^{(0)}) > 1/2 + \delta$  then
    converged  $\leftarrow$  FALSE
    while not converged do
        for all  $b \in [B^{(0)}], \ell \in [L_b^{(0)}]$  do
             $q(\theta_b^\ell) \leftarrow \frac{2}{n_b} \sum_{i: \mathbf{x}_i \in \mathcal{X}_b^\ell} y_i$ 
        end for
        for  $i \in [n]$  do
            Sample tree  $b \in [B^{(0)}]$  uniformly
            Sample leaf  $\ell \in [L_b^{(0)}]$  w.p.  $q(\theta_b^\ell)$ 
            Sample  $\tilde{\mathbf{x}}_i^{(1)} \sim \prod_{j=1}^d P(X_j | \theta_b^\ell)$ 
        end for
         $\mathbf{X}^+ \leftarrow \text{row.append}(\mathbf{X}, \tilde{\mathbf{X}}^{(1)})$ 
         $f^{(1)} \leftarrow \text{RANDOMFOREST}(\mathbf{X}^+, Y)$ 
        if  $\text{Acc}(f^{(1)}) \leq 1/2 + \delta$  then
            converged  $\leftarrow$  TRUE
        else
             $f^{(0)} \leftarrow f^{(1)}$ 
        end if
    end while
end if
    
```

vary with from one round to the next, but it does not increase faster than $1/\max_{\ell, b} (\text{diam}(\mathcal{X}_b^\ell))$.

- (A3) Trees satisfy the following conditions: (i) training data for each tree is split into two subsets: one to learn split parameters, the other to assign leaf labels; (ii) trees are grown on subsamples rather than bootstraps, with subsample size n_b satisfying $n_b \rightarrow \infty, n_b/n \rightarrow 0$ as $n \rightarrow \infty$; (iii) at each internal node, the probability that a tree splits on any given X_j is bounded from below by some $\pi > 0$; (iv) every split puts at least a fraction $\gamma \in (0, 0.5]$ of the available observations into each child node; (v) for each tree b , the total number of leaves L_b satisfies $L_b \rightarrow \infty, L_b/n \rightarrow 0$ as $n \rightarrow \infty$; and (vi) predictions are made with soft labels both within leaves and across trees, i.e. by averaging rather than voting.

(A1) is simply for notational convenience, and can be replaced w.l.o.g. by bounding the feature domain with arbitrary constants. Lipschitz continuity is a common learning theoretic assumption widely used in the analysis of RFs. (A2)’s extra condition regarding the Lipschitz constant controls the variation in smoothness over adversarial training rounds. (A3) imposes standard regularity conditions for RFs (Meinshausen, 2006; Biau, 2012; Denil et al., 2014; Scornet, 2016; Wager and Athey, 2018). With these assumptions in place, we have the following result (see Appx. A for all proofs).

Algorithm 2 FORDE

Input: ARF classifier f , training data $\mathbf{X} \in \mathbb{R}^{n \times d}$
Output: Estimated density q

```

for all  $b \in [B], \ell \in [L_b]$  do
     $q(\theta_b^\ell) \leftarrow \frac{2}{n_b} \sum_{i: \mathbf{x}_i \in \mathcal{X}_b^\ell} y_i$ 
    for  $j \in [d]$  do
         $\psi_{b,j}^\ell \leftarrow$  estimated parameter(s) for  $p(x_j | \theta_b^\ell)$ 
         $q(\cdot; \psi_{b,j}^\ell) \leftarrow$  corresponding pdf/pmf
    end for
end for
    
```

Algorithm 3 FORGE

Input: FORDE model q , target sample size m
Output: Synthetic dataset $\tilde{\mathbf{X}} \in \mathbb{R}^{m \times d}$

```

for  $i \in [m]$  do
    Sample tree  $b \in [B]$  uniformly
    Sample leaf  $\ell \in [L_b]$  w.p.  $q(\theta_b^\ell)$ 
    for  $j \in [d]$  do
        Sample data  $\tilde{x}_{ij} \sim q(\cdot; \psi_{b,j}^\ell)$ 
    end for
end for
    
```

Theorem 1 (Convergence). *Under (A1)-(A3), ARFs converge in probability on the local independence criterion. Let Θ_n be the parameters of an ARF trained on a sample of size n . Then for all $\mathbf{x} \in \mathcal{X}$, $\theta_b^\ell \in \Theta_n$, and $\epsilon > 0$:*

$$\lim_{n \rightarrow \infty} \mathbb{P} \left[\left| p(\mathbf{x} | \theta_b^\ell) - \prod_{j=1}^d p(x_j | \theta_b^\ell) \right| \geq \epsilon \right] = 0.$$

4.1 Density Estimation and Data Synthesis

ARFs are the basis for two further algorithms, FORests for Density Estimation (FORDE) and FORests for GEnerative modeling (FORGE). We present pseudocode for both (see Algs. 2 and 3). The key point to recognize is that, under the local independence criterion, joint densities can be learned by running d separate univariate estimators within each leaf. This is exponentially easier than multivariate density estimation, which suffers from the notorious *curse of dimensionality*. Summarizing the challenges with estimating joint densities, one recent textbook on KDE concludes that “nonparametric methods for kernel density problems should not be used for high-dimensional data and it seems that a feasible dimensionality should not exceed five or six...” (Gramacki, 2018, p. 60). By contrast, our method scales much better with data dimensionality, exploiting the flexibility of ARFs to learn an independence-inducing partition that renders density estimation relatively straightforward.

Of course, this does not *escape* the curse of dimensionality so much as relocate it. The cost for this move is potentially deep trees and/or many ARF training rounds, especially when dependencies between covariates are strong or com-

plex. However, deep forests are generally more efficient than deep neural networks in terms of data and computation, and our experiments suggest that ARF convergence is usually fast even for $\delta = 0$ (see Sect. 5).

With our ARF in hand, the algorithm proceeds as follows. For each tree b , we record the split criteria θ_b^ℓ and empirical coverage $q(\theta_b^\ell)$ of each leaf ℓ . Call these the *leaf parameters*. Then we estimate *distribution parameters* $\psi_{b,j}^\ell$ independently for each (original) X_j within $\mathcal{X}_b^\ell$, e.g. the kernel bandwidth for KDE or class probabilities for MLE with categorical data. In the continuous case, $\psi_{b,j}^\ell$ must either encode leaf bounds (e.g., via a truncated normal distribution with extrema given by θ_b^ℓ) or include a normalization constant to ensure integration to unity. The generative model then follows a simple two-step procedure. First, sample a tree uniformly from $[B]$ and a leaf from that tree with probability $q(\theta_b^\ell)$, just as we do to construct synthetic data within the recursive loop of the ARF algorithm. Next, sample data for each feature X_j according to the density/mass function parametrized by $\psi_{b,j}^\ell$. We repeat this procedure until the target number of synthetic samples has been generated.

We are deliberately agnostic about how distribution parameters $\psi_{b,j}^\ell$ should be learned, as this will tend to vary across features. In our theoretical analysis, we restrict focus to continuous variables and consider a flexible family of KDE methods. In our experiments, we use MLE for continuous data, effectively implementing a truncated Gaussian mixture model, and Bayesian inference for categorical variables, to avoid extreme probabilities when values are unobserved but not beyond the support of a given leaf. Under local independence, distribution learning is completely modular, so different methods can coexist without issue. We revisit this topic in Sect. 6.

Our estimated density takes the following form:

$$q(\mathbf{x}) = \frac{1}{B} \sum_{\ell, b: \mathbf{x} \in \mathcal{X}_b^\ell} q(\theta_b^\ell) \prod_{j=1}^d q(x_j; \psi_{b,j}^\ell). \quad (1)$$

Compare this with the true density:

$$p(\mathbf{x}) = \frac{1}{B} \sum_{\ell, b: \mathbf{x} \in \mathcal{X}_b^\ell} p(\theta_b^\ell) p(\mathbf{x} | \theta_b^\ell). \quad (2)$$

In both cases, the density evaluated at a given point is just a coverage-weighted average of its density in all leaves whose split criteria it satisfies.

Because we are concerned with L_2 -consistency, our loss function is the mean integrated squared error (MISE)², defined as:

$$\text{MISE}(p, q) := \mathbb{E} \left[\int_{\mathcal{X}} (p(\mathbf{x}) - q(\mathbf{x}))^2 d\mathbf{x} \right].$$

²Alternative loss functions may also be suitable, e.g. the Kullback-Leibler divergence or the Wasserstein distance.

We require one extra assumption, imposing standard conditions for KDE consistency (Silverman, 1986):

- (A4) The true density function p is smooth. Specifically, its second derivative p'' is finite, continuous, square integrable, and ultimately monotone.

Our method admits three potential sources of error, quantified by the following residuals:

$$\epsilon_1 := \epsilon_1(\ell, b) := p(\theta_b^\ell) - q(\theta_b^\ell) \quad (3)$$

$$\epsilon_2 := \epsilon_2(\ell, b, \mathbf{x}) := \prod_{j=1}^d p(x_j | \theta_b^\ell) - \prod_{j=1}^d q(x_j; \psi_{b,j}^\ell) \quad (4)$$

$$\epsilon_3 := \epsilon_3(\ell, b, \mathbf{x}) := p(\mathbf{x} | \theta_b^\ell) - \prod_{j=1}^d p(x_j | \theta_b^\ell) \quad (5)$$

We refer to these as errors of *coverage*, *density*, and *convergence*, respectively. Observe that ϵ_1 is a random variable that depends on ℓ and b , while ϵ_2, ϵ_3 are random variables depending on ℓ, b and $\mathbf{x}$. We suppress the dependencies for ease of notation.

Lemma 1. The error of our estimator satisfies the following bound:

$$\text{MISE}(p, q) \leq 2B^{-2} \mathbb{E} \left[\int_{\mathcal{X}} \alpha^2 + \beta^2 d\mathbf{x} \right],$$

where

$$\alpha := \sum_{\ell, b: \mathbf{x} \in \mathcal{X}_b^\ell} p(\theta_b^\ell) \epsilon_3 \quad \text{and}$$

$$\beta := \sum_{\ell, b: \mathbf{x} \in \mathcal{X}_b^\ell} \left(p(\theta_b^\ell) \epsilon_2 + \epsilon_1 \prod_{j=1}^d p(x_j | \theta_b^\ell) - \epsilon_1 \epsilon_2 \right).$$

This lemma establishes that total error is bounded by a quadratic function of $\epsilon_1, \epsilon_2, \epsilon_3$. We know by Thm. 1 that errors of convergence vanish in the limit. Our next result states that the same holds for errors of coverage and density.

Theorem 2 (Consistency). *Under assumptions (A1)-(A4), FORDE is L_2 -consistent. Let q_n denote the joint density estimated on a training sample of size n . Then we have:*

$$\lim_{n \rightarrow \infty} \text{MISE}(p, q_n) = 0.$$

Our consistency proof is fundamentally unlike those of piecewise constant density estimators with CART trees (Ram and Gray, 2011; Wu et al., 2014; Correia et al., 2020), which essentially treat base learners as adaptive histograms and rely on tree-wise convergence when leaf volume goes to zero (Devroye et al., 1996; Lugosi and Nobel, 1996). Alternative methods that perform KDE or MLE within each leaf do not come with theoretical guarantees (Smyth et al., 1995;

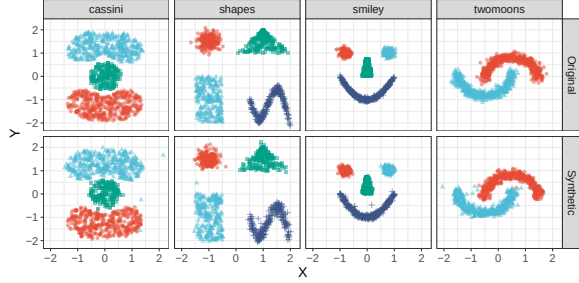


Figure 1: Visual examples. Original (top) and synthetic (bottom) data are presented for four three-dimensional problems with two continuous covariates and one categorical feature.

Gray and Moore, 2003; Loh, 2009; Ram and Gray, 2011). Recently, consistency has been shown for some RF-based conditional density estimators (Hothorn and Zeileis, 2021; Čevič et al., 2022). However, these results do not extend to the unconditional case, since features with little predictive value for the outcome variable(s) are unlikely to be selected for splits. The resulting models will therefore fail to detect dependencies between features deemed uninformative for the given prediction task. In the rare case that authors use some form of unsupervised splits, they make no effort to factorize the distribution and are therefore subject to the curse of dimensionality (Criminisi et al., 2012; Feng and Zhou, 2018). By contrast, our method exploits ARFs to find regions of local independence, and univariate density estimation to compute marginals within each leaf. Though our consistency result comes at the cost of some extra assumptions, we argue that this is a fair price to pay for improved performance in finite samples.

5 EXPERIMENTS

In this section, we present results from a wide range of experiments conducted on simulated and real-world datasets. We use 100 trees for density estimation tasks and 20 for data synthesis. Increasing this parameter tends to improve performance for FORDE, but appears to have less of an impact on FORGE. Trees are grown until purity or a minimum node size of two (with just a single sample, variance is undefined). In all cases, we set the slack parameter $\delta = 0$ and use the default `mtry` = $\lfloor \sqrt{d} \rfloor$. For more details on hyperparameters and datasets see Appx. B. Code for reproducing all results is available online at https://github.com/bips-hb/arf_paper.

5.1 Simulation

FORGE recreates visual patterns. We begin with a simple proof of concept experiment, illustrating our method on a handful of low-dimensional datasets that allow for easy visual assessment. The *cassini*, *shapes*, *smiley*, and *twomoons* problems are all three-dimensional examples

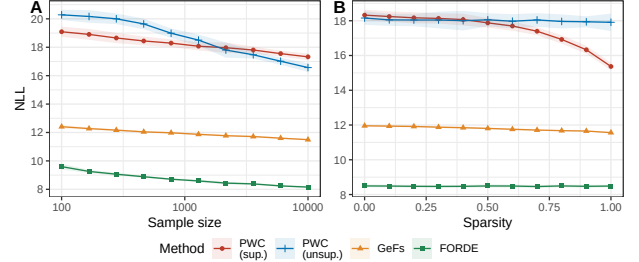


Figure 2: Negative log-likelihood (NLL) measured in nats on a test set for varying sample size (A) and sparsity (B). Lower is better. Shading represents standard errors.

that combine two continuous covariates with a categorical class label. We simulate $n = 2000$ samples from each data generating process (see Fig. 1, top row) and estimate densities using FORDE. We proceed to FORGE a synthetic dataset of $n = 1000$ samples (Fig. 1, bottom row) and compare results. We find that the model consistently approximates its target distribution with high fidelity. Classes are clearly distinguished in all cases, and the visual form of the original data is immediately recognizable. A few stray samples are evident on close inspection. Such anomalies can be mitigated with a larger training set.

FORDE outperforms alternative CART-based methods.

We simulate data from a multivariate Gaussian distribution $\mathbf{X} \sim \mathcal{N}(0, \Sigma)$, with Toeplitz covariance matrix $\Sigma_{ij} = 0.9^{|i-j|}$ and fixed $d = 10$. To compare against supervised methods, we also simulate a binary target $Y \sim \text{Bern}([1 + \exp(-\mathbf{X}\beta)]^{-1})$, where the coefficient vector β contains a varying proportion of 0's (non-informative features) and 1's (informative features). Performance is evaluated by the negative log-likelihood (NLL) on a test set of $n_{\text{tst}} = 1000$. We compare our method to piecewise constant (PWC) estimators with supervised and unsupervised split criteria,³ as well as generative forests (GeFs), a RF-based smooth density estimation procedure (Correia et al., 2020).

Fig. 2 shows the average NLL over 20 replicates for varying sample sizes (A) and levels of sparsity (B). For the former, we fix the proportion of informative features at 0.5; for the latter, we fix $n_{\text{trn}} = 2000$. We find that PWC methods fare poorly, with much greater NLL in all settings. This is likely due to the unrealistic uniformity assumption, according to which the corner of a hyperrectangle is no less probable than the center. GeFs, which also use a Gaussian mixture model to estimate densities, perform better in this experiment. However, FORDE dominates throughout.

³Supervised PWC is simply an ensemble version of the classic method (Gray and Moore, 2003; Ram and Gray, 2011; Wu et al., 2014). To the best of our knowledge, no one has previously proposed *unsupervised* PWC density estimation with CART trees. This can be understood as a variant of our approach in which all marginals are uniform within each leaf.

Table 1: Average NLL on the Twenty Datasets benchmark for five PC models and FORDE. Winning results in bold.

Dataset	EiNet	RAT-SPN	PGC	Strudel	CMCLT	FORDE
nlts	6.02	6.01	6.05	6.07	5.99	6.01
msnbc	6.12	6.04	6.06	6.04	6.05	6.10
kdd	2.18	2.13	2.14	2.14	2.12	2.13
plants	13.68	13.44	13.52	13.22	12.26	12.26
audio	39.88	39.96	40.21	42.40	39.02	39.74
jester	52.56	52.97	53.54	54.24	51.94	52.8
netflix	56.64	56.85	57.42	57.93	55.31	56.67
accidents	35.59	35.49	30.46	29.05	28.69	33.85
retail	10.92	10.91	10.84	10.83	10.82	10.93
pumsb	31.95	31.53	29.56	24.39	23.71	28.4
dna	96.09	97.23	80.82	87.15	84.91	91.85
kosarek	11.03	10.89	10.72	10.70	10.56	10.84
msweb	10.03	10.12	9.98	9.74	9.62	9.72
book	34.74	34.68	34.11	34.49	33.75	34.85
movie	51.71	53.63	53.15	53.72	49.23	50.86
webkb	157.28	157.53	155.23	154.83	147.77	153.45
reuters	87.37	87.37	87.65	86.35	81.17	84.15
20ng	153.94	152.06	154.03	153.87	148.17	155.51
bbc	248.33	252.14	254.81	256.53	242.83	240.31
ad	26.27	48.47	21.65	16.52	14.76	21.80
Avg. rank	4.5	4.2	4	3.6	1.2	3.3

Panel (B) clearly illustrates that unsupervised methods are unaffected by changes in signal sparsity, since their splits are independent of the outcome variable Y . By contrast, sparsity appears to benefit the supervised methods. This can be explained by the fact that splits are random when features are uninformative, a strategy that is known to work well in noisy settings (Geurts et al., 2006; Genuer, 2012).

5.2 Real Data

FORDE is competitive with alternative PCs. Building on Correia et al. (2020)’s observation that RFs can be compiled into probabilistic circuits, we compare the performance of FORDE to that of five leading PCs on the Twenty Datasets benchmark (Van Haaren and Davis, 2012), a heterogeneous collection of tasks ranging from retail to biology that is widely used to evaluate tractable probabilistic models. Each dataset is randomly split into training (70%), validation (10%), and test sets (20%). Competitors include Einsum networks (EiNet) (Peharz et al., 2020a), random sum-product networks (RAT-SPN) (Peharz et al., 2020b), probabilistic generating circuits (PGC) (Zhang et al., 2021), Strudel (Dang et al., 2022), and continuous mixtures of Chow-Liu trees (CMCLT) (Correia et al., 2023). We report the average NLL on the test set for each model in Table 1. Though the recently proposed CMCLT algorithm generally dominates in this experiment, FORDE attains top performance on two datasets and is never far behind the state of the art. Its average rank of 3.3 places it second overall.

FORGE generates realistic tabular data. To evaluate the performance of FORGE on real-world datasets, we recreate a benchmarking pipeline originally proposed by Xu et al. (2019). They introduce the conditional tabular GAN (CTGAN) and tabular VAE (TVAE), two deep learning algorithms for generative modeling with mixed continuous and categorical features. We include three additional state-of-

the-art tabular GAN architectures for comparison: invertible tabular GAN (IT-GAN) (Lee et al., 2021), regularized compound conditional GAN (RCC-GAN) (Esmailpour et al., 2022), and a differentially private conditional tabular GAN (CTAB-GAN+) (Zhao et al., 2022).

A complete summary of the experimental setup is presented in Appx. B. Briefly, we take five benchmark datasets for classification and partition the samples into training and test sets, which we denote by $\mathbf{Z}_{\text{trn}} = (\mathbf{X}_{\text{trn}}, Y_{\text{trn}})$ and $\mathbf{Z}_{\text{tst}} = (\mathbf{X}_{\text{tst}}, Y_{\text{tst}})$, respectively. $\mathbf{Z}_{\text{trn}}$ is used as input to a series of generative models, each of which creates a synthetic training set $\tilde{\mathbf{Z}}_{\text{trn}}$ of the same sample size as the original. Several classifiers are then trained on $\tilde{\mathbf{Z}}_{\text{trn}}$ and evaluated on $\mathbf{Z}_{\text{tst}}$, with performance metrics averaged across learners. Results are benchmarked against the same set of algorithms, now trained on the original data $\mathbf{Z}_{\text{trn}}$. We refer to this model as the *oracle*, since it should perform no worse in expectation than any classifier trained on synthetic data. However, if the generative model approximates its target with high fidelity, then differences between the oracle and its competitors should be negligible.⁴ Similar approaches are widely used in the evaluation of GANs (Yang et al., 2017; Shmelkov et al., 2018; Santurkar et al., 2018); for a critical discussion, see Ravuri and Vinyals (2019).

Results are reported in Table 2, where we average over five trials of data synthesis and subsequent supervised learning. We include information on each dataset, including the cardinality of the response variable, the training/test sample size, and dimensionality of the feature space. Performance is evaluated via accuracy and F1-score (or F1 macro-score for multiclass problems), as well as wall time. FORGE fares well in this experiment, attaining the top accuracy and F1-score in three out of five tasks. On a fourth, the highly imbalanced `credit` dataset, the only models that do better in terms of accuracy receive F1-scores of 0, suggesting that they entirely ignore the minority class. Only FORGE and RCC-GAN strike a reasonable balance between sensitivity and specificity on this task. Perhaps most impressive, FORGE executes over 60 times faster than its nearest competitor on average, and over 100 times faster than the second fastest method. (We omit results for algorithms that fail to converge in 24 hours of training time.) Differences in compute time would be even more dramatic if these deep learning algorithms were configured with a CPU backend (we used GPUs here), or if FORGE were run using more extensive parallelization (we distribute the job across 10 cores). This comparison also obscures the extra time required to tune hyperparameters for these complex models, whereas our method is an off-the-shelf solution that works

⁴Note that the so-called “oracle” is not necessarily optimal w.r.t. the true data generating process—other models may have lower risk—but it should be optimal w.r.t. a given function class-dataset pair. If logistic regression attains 60% test accuracy training on $\mathbf{Z}_{\text{trn}}$, then it should do about the same training on $\tilde{\mathbf{Z}}_{\text{trn}}$, regardless of how much better a well-tuned MLP may perform.

Table 2: Performance on the Xu et al. (2019) benchmark for five deep learning models and FORGE. We report average results across five replicates $\pm$ standard errors. Winning results in bold.

Dataset	Model	Accuracy $\pm$ SE	F1 $\pm$ SE	Time (sec)
adult classes = 2 $n_{\text{trn}} = 23\text{k}$ $n_{\text{test}} = 10\text{k}$ $d = 14$	Oracle	0.828 ± 0.006	0.884 ± 0.004	
	FORGE	0.819 ± 0.006	0.877 ± 0.005	2.9
	CTGAN	0.786 ± 0.020	0.853 ± 0.019	263.3
	CTAB-GAN+	0.808 ± 0.008	0.869 ± 0.006	561.6
	IT-GAN	0.794 ± 0.005	0.853 ± 0.005	3435.6
	RCC-GAN	0.770 ± 0.015	0.841 ± 0.015	8823.0
	TVAE	0.804 ± 0.007	0.865 ± 0.006	115.1
census classes = 2 $n_{\text{trn}} = 200\text{k}$ $n_{\text{test}} = 100\text{k}$ $d = 40$	Oracle	0.922 ± 0.002	0.957 ± 0.001	
	FORGE	0.903 ± 0.019	0.946 ± 0.012	53.2
	CTGAN	0.916 ± 0.015	0.954 ± 0.009	4287.8
	CTAB-GAN+	0.912 ± 0.026	0.952 ± 0.016	10182.1
	IT-GAN	NA	NA	>24hr
	RCC-GAN	0.900 ± 0.016	0.944 ± 0.011	8908.6
	TVAE	0.928 ± 0.007	0.961 ± 0.004	1814.9
covertype classes = 7 $n_{\text{trn}} = 481\text{k}$ $n_{\text{test}} = 100\text{k}$ $d = 54$	Oracle	0.895 ± 0.000	0.838 ± 0.000	
	FORGE	0.707 ± 0.006	0.549 ± 0.006	103.5
	CTGAN	0.633 ± 0.009	0.400 ± 0.009	13387.2
	CTAB-GAN+	NA	NA	>24hr
	IT-GAN	NA	NA	>24hr
	RCC-GAN	NA	NA	>24hr
	TVAE	0.698 ± 0.013	0.459 ± 0.013	4882.0
credit classes = 2 $n_{\text{trn}} = 264\text{k}$ $n_{\text{test}} = 20\text{k}$ $d = 30$	Oracle	0.997 ± 0.001	0.607 ± 0.029	
	FORGE	0.995 ± 0.001	0.527 ± 0.036	32.2
	CTGAN	0.881 ± 0.099	0.047 ± 0.031	4898.0
	CTAB-GAN+	0.998 ± 0.000	0.000 ± 0.000	7497.3
	IT-GAN	NA	NA	>24hr
	RCC-GAN	0.993 ± 0.003	0.569 ± 0.056	10608.4
	TVAE	0.998 ± 0.000	0.000 ± 0.000	3847.6
intrusion classes = 5 $n_{\text{trn}} = 394\text{k}$ $n_{\text{test}} = 100\text{k}$ $d = 40$	Oracle	0.998 ± 0.001	0.833 ± 0.001	
	FORGE	0.993 ± 0.001	0.656 ± 0.001	68.2
	CTGAN	0.944 ± 0.088	0.645 ± 0.088	8749.3
	CTAB-GAN+	NA	NA	>24hr
	IT-GAN	NA	NA	>24hr
	RCC-GAN	NA	NA	>24hr
	TVAE	0.990 ± 0.002	0.598 ± 0.002	4306.0

well with default settings.

5.3 Runtime

To further demonstrate the computational efficiency of our pipeline relative to deep learning methods, we conduct a runtime experiment using the smallest dataset above, `adult`. By repeatedly sampling stratified subsets—varying both sample size n and dimensionality d —and measuring the time needed to train a generative model and synthesize data from it, we illustrate how complexity scales with n and d . For this experiment, we ran the three fastest deep learning competitors—CTGAN, TVAE, and CTAB-GAN+—with both CPU and GPU backends. We use default parameters for all algorithms, which include automated parallelization over all available cores (24 in this experiment).

Fig. 3 shows the results. FORGE clearly dominates in training time (see panels A and C), executing orders of magnitude faster than the competition (note the log scale). For those with limited access to GPUs, deep learning methods may be completely infeasible for large datasets. Even when GPUs are available, FORGE still scales far better, completing the full pipeline about 35 times faster than TVAE, 85 times faster than CTGAN, and nearly 200 times faster than CTAB-GAN+ in this example. Other methods appear to gen-

erate samples more quickly than FORGE (see panels B and D), but this computation is trivial compared to training. Interestingly, our method is a faster sampler when measured in processing time (see Fig. 4, Appx. B.4), suggesting that it could outperform competitors here too with more efficient parallelization. Note that FORGE attains the highest accuracy and F1-score of all methods for the `adult` dataset, so this speedup need not come at the cost of performance.

6 DISCUSSION

ARFs enable fast, accurate density estimation and data synthesis. However, the method is not without its limitations. First, it is not tailored to structured data such as images or text, for which deep learning models have proven especially effective. See Appx. B.5 for a comparison to state-of-the-art models on the MNIST dataset, where convolutional GANs clearly outperform FORGE, as expected. Where our method excels, by contrast, is in speed and flexibility.

We caution that our convergence guarantees have no implications for finite sample performance. Though ARFs only required a few rounds of training in most of our experiments, it is entirely possible that discriminator accuracy increase from one round to the next, or that Alg. 1 fail to terminate altogether for some datasets. (In practice, this behavior is mitigated by increasing δ or setting some maximum number of iterations.) For instance, on MNIST, we generally find accuracy plateauing around 65% after five rounds with little improvement thereafter. Of course, the same caveats apply to any asymptotic guarantee. Finite sample results are rare in the RF literature, although there has been some recent work in this area (Gao et al., 2022).

Another potential difficulty for our approach is selecting an optimal density estimation subroutine. KDE relies on a smoothness assumption (A4), while MLE requires a (local) parametric model. Bayesian inference imposes a prior distribution, which may bias results. All three methods will struggle when their assumptions are violated. Resampling alternatives such as permutations or bootstrapping do not produce any data that was not observed in the training set and may therefore raise privacy concerns. No approach is generally guaranteed to strike the optimal balance between efficiency, accuracy, and privacy, and so the choice of which combination of methods to employ is irreducibly context-dependent.

We emphasize that our method performs well in a range of settings without any model tuning. However, we acknowledge that optimal performance likely depends on RF parameters (Scornet, 2017; Probst et al., 2019). In particular, there is an inherent trade-off between the goals of minimizing errors of density (ϵ_2) and errors of convergence (ϵ_3) in finite samples. Grow trees too deep, and leaves will not contain enough data to accurately estimate marginal densities; grow trees too shallow, and ARFs may not satisfy

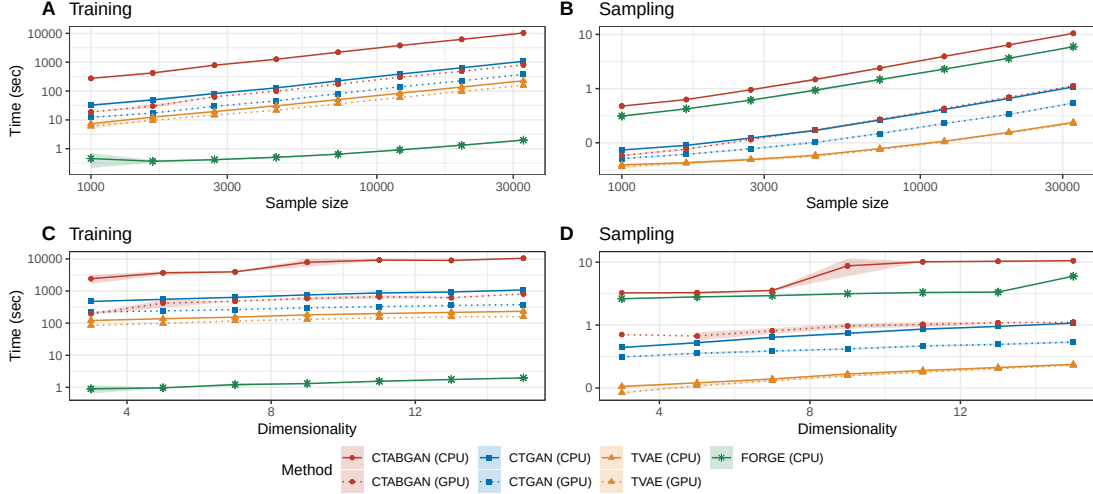


Figure 3: Complexity curves, evaluated using stratified subsamples of the `adult` dataset. (A): Training time as a function of sample size. (B): Sampling time as a function of sample size. (C): Training time as a function of dimensionality. (D): Sampling time as a function of dimensionality.

the local independence criterion. Meanwhile, the `mtry` parameter has been shown to control sparsity in low signal-to-noise regimes (Mentch and Zhou, 2020). Smaller values may therefore be appropriate when ϵ_3 is large, in order to regularize the forest. Adding more trees tends to improve density estimates, though this incurs extra computational cost in both time and memory (Probst and Boulesteix, 2017). Despite these considerations, we reiterate that ARFs do remarkably well with default parameters.

The ethical implications of generative models are potentially fraught. Deepfakes have attracted particular attention in this regard (de Ruiter, 2021; Öhman, 2020; Diakopoulos and Johnson, 2021), as they can deceive their audience into believing that people said or did things they never in fact said or did. These dangers are most acute with convolutional neural networks or other architectures optimized for visual and audio data. Despite and in full awareness of these concerns, we point out that generative models also present a valuable ethical opportunity, since they may preserve the privacy of data subjects by creating datasets that preserve statistical signals without exposing the personal information of individuals. However, the privacy-utility trade-off can be unpredictable with synthetic data (Stadler et al., 2022). As with all powerful technologies, caution is advised and regulatory frameworks are welcome.

7 CONCLUSION

We have introduced a novel procedure for learning joint densities and generating synthetic data using a recursive, adversarial variant of unsupervised random forests. The method is provably consistent under reasonable assumptions, and performs well in experiments on simulated and real-world examples. Our FORDE algorithm is more accurate

than other CART-based density estimators and compares favorably to leading PC algorithms. Our FORGE algorithm is competitive with deep learning models for data generation on tabular data benchmarks, and routinely executes some 100 times faster. An R package, `arf`, is available on CRAN. A Python implementation is forthcoming.

Future work will explore further applications for these methods, such as anomaly detection, clustering, and classification, as well as potential connections with differential privacy (Dwork, 2008). Though we have focused in this work on unconditional density estimation tasks, it is straightforward to compute arbitrary conditional probabilities with ARFs by reducing the event space to just those leaves that satisfy some logical constraint(s). More complex functionals may be estimated with just a few additional steps—e.g. (conditional) quantiles, CDFs, and copulas—thereby linking these methods with recent work on functional regression with random forests (Hothorn and Zeileis, 2021; Fu et al., 2021; Čevič et al., 2022). Alternative tree-based solutions based on gradient boosting also warrant further exploration, especially given promising recent developments in this area (Friedman, 2020; Gao and Hastie, 2022).

Acknowledgments

MNW and KB received funding from the German Research Foundation (DFG), Emmy Noether Grant 437611051. MNW and JK received funding from the U Bremen Research Alliance/AI Center for Health Care, financially supported by the Federal State of Bremen. We are grateful to Cassio de Campos, Gennaro Gala, Robert Peharz, and Alvaro H.C. Correia for their feedback on an earlier draft of this manuscript.

References

- Arjovsky, M., Chintala, S., and Bottou, L. (2017). Wasserstein generative adversarial networks. In *Proceedings of the 34th International Conference on Machine Learning*, page 214–223.
- Athey, S., Tibshirani, J., and Wager, S. (2019). Generalized random forests. *Ann. Statist.*, 47(2):1148–1178.
- Augenstein, S., McMahan, H. B., Ramage, D., Ramaswamy, S., Kairouz, P., Chen, M., Mathews, R., and y Arcas, B. A. (2020). Generative models for effective ML on private, decentralized datasets. In *International Conference on Learning Representations*.
- Bach, F. R. and Jordan, M. I. (2003). Beyond independent components: Trees and clusters. *J. Mach. Learn. Res.*, 4:1205–1233.
- Belgiu, M. and Drăguț, L. (2016). Random forest in remote sensing: A review of applications and future directions. *ISPRS J. Photogramm. Remote Sens.*, 114:24–31.
- Biau, G. (2012). Analysis of a random forests model. *J. Mach. Learn. Res.*, 13:1063–1095.
- Biau, G. and Devroye, L. (2010). On the layered nearest neighbour estimate, the bagged nearest neighbour estimate and the random forest method in regression and classification. *J. Multivar. Anal.*, 101(10):2499–2518.
- Biau, G., Devroye, L., and Lugosi, G. (2008). Consistency of random forests and other averaging classifiers. *J. Mach. Learn. Res.*, 9(66):2015–2033.
- Biau, G. and Scornet, E. (2016). A random forest guided tour. *TEST*, 25(2):197–227.
- Blackard, J. (1998). Covertypes. UCI Machine Learning Repository.
- Bramer, M. (2007). *Clustering*, pages 221–238. Springer, London, UK.
- Breiman, L. (2001). Random forests. *Mach. Learn.*, 45(1):1–33.
- Breiman, L. (2004). Consistency for a simple model of random forests. Technical Report 670, Statistics Department, UC Berkeley, Berkeley, CA.
- Breiman, L., Friedman, J., Stone, C. J., and Olshen, R. A. (1984). *Classification and Regression Trees*. Taylor & Francis, Boca Raton, FL.
- Buzhinsky, I., Nerinovskiy, A., and Tripakis, S. (2021). Metrics and methods for robustness evaluation of neural networks with generative models. *Mach. Learn.*
- Ćevic, D., Michel, L., Näf, J., Bühlmann, P., and Meinhäusen, N. (2022). Distributional random forests: Heterogeneity adjustment and multivariate distributional regression. *J. Mach. Learn. Res.*, 23(333).
- Chandola, V., Banerjee, A., and Kumar, V. (2009). Anomaly detection: A survey. *ACM Comput. Surv.*, 41(3).
- Chen, X. and Ishwaran, H. (2012). Random forests for genomic data analysis. *Genomics*, 99(6):323–329.
- Choi, E., Biswal, S., Malin, B., Duke, J., Stewart, W. F., and Sun, J. (2017). Generating multi-label discrete patient records using generative adversarial networks. In *Proceedings of the 2nd Machine Learning for Healthcare Conference*, pages 286–305.
- Choi, Y., Vergari, A., and Van den Broeck, G. (2020). Probabilistic circuits: A unifying framework for tractable probabilistic models. Technical Report, University of California, Los Angeles.
- Chow, C. and Liu, C. (1968). Approximating discrete probability distributions with dependence trees. *IEEE Trans. Inf. Theory*, 14(3):462–467.
- Correia, A., Gala, G., Quaeghebeur, E., de Campos, C., and Peharz, R. (2023). Continuous mixtures of tractable probabilistic models. Proceedings of the 37th AAAI Conference.
- Correia, A., Peharz, R., and de Campos, C. P. (2020). Joins in random forests. In *Advances in Neural Information Processing Systems*, volume 33, pages 11404–11415.
- Criminisi, A., Shotton, J., and Konukoglu, E. (2012). *Decision Forests: A Unified Framework for Classification, Regression, Density Estimation, Manifold Learning and Semi-Supervised Learning*, volume 7, pages 81–227. NOW Publishers, Norwell, MA.
- Cutler, D. R., Edwards Jr., T. C., Beard, K. H., Cutler, A., Hess, K. T., Gibson, J., and Lawler, J. J. (2007). Random forests for classification in ecology. *Ecology*, 88(11):2783–2792.
- Dang, M., Vergari, A., and Van den Broeck, G. (2022). Strudel: A fast and accurate learner of structured-decomposable probabilistic circuits. *Int. J. Approx. Reason.*, 140:92–115.
- Darwiche, A. (2009). *Modeling and Reasoning with Bayesian Networks*. Cambridge University Press, New York.
- de Ruiter, A. (2021). The distinct wrong of deepfakes. *Philos. Technol.*, 34(4):1311–1332.
- Denil, M., Matheson, D., and De Freitas, N. (2014). Narrowing the gap: Random forests in theory and in practice. In *Proceedings of the 31st International Conference on Machine Learning*, pages 665–673.
- Devroye, L., Györfi, L., and Lugosi, G. (1996). *A Probabilistic Theory of Pattern Recognition*. Springer-Verlag, New York.
- Diakopoulos, N. and Johnson, D. (2021). Anticipating and addressing the ethical implications of deepfakes in the context of elections. *New Media Soc.*, 23(7):2072–2098.
- Drton, M. and Maathuis, M. H. (2017). Structure learning in graphical modeling. *Annu. Rev. Stat. Appl.*, 4(1):365–393.

- Dua, D. and Graff, C. (2019). UCI machine learning repository.
- Dwork, C. (2008). Differential privacy: A survey of results. In *Theory and Applications of Models of Computation*, volume 4978, pages 1–19, Berlin, Heidelberg. Springer.
- Efron, B. (1994). Missing data, imputation, and the bootstrap. *J. Am. Stat. Assoc.*, 89(426):463–475.
- Esmailpour, M., Chaalia, N., Abusitta, A., Devailly, F.-X., Maazoun, W., and Cardinal, P. (2022). RCC-GAN: Regularized compound conditional gan for large-scale tabular data synthesis. *arXiv preprint*, 2205.11693.
- Feng, J. and Zhou, Z.-H. (2018). Autoencoder by forest. In *Proceedings of the 32nd AAAI Conference on Artificial Intelligence*.
- Fernández-Delgado, M., Cernadas, E., Barro, S., and Amorim, D. (2014). Do we need hundreds of classifiers to solve real world classification problems? *J. Mach. Learn. Res.*, 15(90):3133–3181.
- Friedman, J. H. (2020). Contrast trees and distribution boosting. *Proc. Natl. Acad. Sci.*, 117(35):21175–21184.
- Fu, G., Dai, X., and Liang, Y. (2021). Functional random forests for curve response. *Sci. Rep.*, 11(1):24159.
- Gao, W., Xu, F., and Zhou, Z.-H. (2022). Towards convergence rate analysis of random forests for classification. *Artif. Intell.*, 313(C).
- Gao, Z. and Hastie, T. (2022). LinCDE: Conditional density estimation via Lindsey’s method. *J. Mach. Learn. Res.*, 23(52):1–55.
- Genuer, R. (2012). Variance reduction in purely random forests. *J. Nonparametr. Stat.*, 24(3):543–562.
- Geurts, P., Ernst, D., and Wehenkel, L. (2006). Extremely randomized trees. *Mach. Learn.*, 63(1):3–42.
- Goodfellow, I. J., Pouget-Abadie, J., Mirza, M., Xu, B., Warde-Farley, D., Ozair, S., Courville, A., and Bengio, Y. (2014). Generative adversarial nets. In *Advances in Neural Information Processing Systems*, volume 27, page 2672–2680.
- Gramacki, A. (2018). *Nonparametric Kernel Density Estimation and Its Computational Aspects*. Springer, Cham.
- Gray, A. G. and Moore, A. W. (2003). Nonparametric density estimation: Toward computational tractability. In *Proceedings of the 2003 SIAM International Conference on Data Mining (SDM)*, pages 203–211.
- Györfi, L., Kohler, M., Krzyżak, A., and Walk, H. (2002). *A Distribution-Free Theory of Nonparametric Regression*. Springer-Verlag, New York.
- Heckerman, D., Geiger, D., and Chickering, D. M. (1995). Learning Bayesian networks: The combination of knowledge and statistical data. *Mach. Learn.*, 20(3):197–243.
- Higgins, I., Matthey, L., Pal, A., Burgess, C., Glorot, X., Botvinick, M., Mohamed, S., and Lerchner, A. (2017). β -VAE: Learning basic visual concepts with a constrained variational framework. In *International Conference on Learning Representations*.
- Hothorn, T. and Zeileis, A. (2021). Predictive distribution modeling using transformation forests. *J. Comput. Graph. Stat.*, 30(4):1181–1196.
- Jordon, J., Yoon, J., and van der Schaar, M. (2019). PATE-GAN: generating synthetic data with differential privacy guarantees. In *International Conference on Learning Representations*.
- Kim, I., Ramdas, A., Singh, A., and Wasserman, L. (2021). Classification accuracy as a proxy for two-sample testing. *Ann. Stat.*, 49(1):411 – 434.
- Kingma, D., Salimans, T., Poole, B., and Ho, J. (2021). Variational diffusion models. In *Advances in Neural Information Processing Systems*, volume 34, pages 21696–21707.
- Kingma, D. and Welling, M. (2013). Auto-encoding variational Bayes. In *International Conference on Learning Representations*.
- Kisa, D., Van den Broeck, G., Choi, A., and Darwiche, A. (2014). Probabilistic sentential decision diagrams. In *14th International Conference on the Principles of Knowledge Representation and Reasoning*.
- Kobyzev, I., Prince, S. J., and Brubaker, M. A. (2021). Normalizing flows: An introduction and review of current methods. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 43(11):3964–3979.
- Koller, D. and Friedman, N. (2009). *Probabilistic Graphical Models*. The MIT Press, Cambridge, MA.
- Lauritzen, S. L. (1996). *Graphical Models*. Oxford Statistical Science Series. Clarendon Press, Oxford.
- LeCun, Y., Bottou, L., Bengio, Y., and Haffner, P. (1998). Gradient-based learning applied to document recognition. *Proceedings of the IEEE*, 86(11):2278–2324.
- Lee, J., Hyeong, J., Jeon, J., Park, N., and Cho, J. (2021). Invertible tabular GANs: Killing two birds with one stone for tabular data synthesis. In *Advances in Neural Information Processing Systems*, volume 34, pages 4263–4273.
- Lee, J., Kim, M., Jeong, Y., and Ro, Y. (2022). Differentially private normalizing flows for synthetic tabular data generation. *Proceedings of the AAAI Conference on Artificial Intelligence*, 36(7):7345–7353.
- Liu, H., Xu, M., Gu, H., Gupta, A., Lafferty, J., and Wasserman, L. (2011). Forest density estimation. *J. Mach. Learn. Res.*, 12(25):907–951.
- Loh, W.-Y. (2009). Improving the precision of classification trees. *Ann. Appl. Stat.*, 3(4):1710 – 1737.
- Lopez, R., Regier, J., Cole, M. B., Jordan, M. I., and Yosef, N. (2018). Deep generative modeling for single-cell transcriptomics. *Nat. Methods*, 15(12):1053–1058.

- Lugosi, G. and Nobel, A. (1996). Consistency of data-driven histogram methods for density estimation and classification. *Ann. Stat.*, 24(2):687 – 706.
- Lundberg, S. M., Erion, G., Chen, H., DeGrave, A., Prutkin, J. M., Nair, B., Katz, R., Himmelfarb, J., Bansal, N., and Lee, S.-I. (2020). From local explanations to global understanding with explainable AI for trees. *Nat. Mach. Intell.*, 2(1):56–67.
- Malley, J., Kruppa, J., Dasgupta, A., Malley, K., and Ziegler, A. (2012). Probability machines: Consistent probability estimation using nonparametric learning machines. *Methods Inf. Med.*, 51(1):74–81.
- Mandelbrot, B. B. (1982). *The Fractal Geometry of Nature*. W.H. Freeman & Co., New York.
- Meinshausen, N. (2006). Quantile regression forests. *J. Mach. Learn. Res.*, 7:983–999.
- Mentch, L. and Hooker, G. (2016). Quantifying uncertainty in random forests via confidence intervals and hypothesis tests. *J. Mach. Learn. Res.*, 17(26).
- Mentch, L. and Zhou, S. (2020). Randomization as regularization: A degrees of freedom explanation for random forest success. *J. Mach. Learn. Res.*, 21(171).
- Mirza, M. and Osindero, S. (2014). Conditional generative adversarial nets. *arXiv preprint*, 1411.1784.
- Öhman, C. (2020). Introducing the pervert’s dilemma: a contribution to the critique of deepfake pornography. *Ethics Inf. Technol.*, 22(2):133–140.
- Pang, G., Shen, C., Cao, L., and Hengel, A. V. D. (2021). Deep learning for anomaly detection: A review. *ACM Comput. Surv.*, 54(2).
- Papamakarios, G., Nalisnick, E., Rezende, D. J., Mohamed, S., and Lakshminarayanan, B. (2021). Normalizing flows for probabilistic modeling and inference. *J. Mach. Learn. Res.*, 22(57):1–64.
- Pearl, J. and Russell, S. (2003). Bayesian networks. In Arbib, M. A., editor, *Handbook of Brain Theory and Neural Networks*, pages 157–160. The MIT Press, Cambridge, MA.
- Peharz, R., Lang, S., Vergari, A., Stelzner, K., Molina, A., Trapp, M., Van Den Broeck, G., Kersting, K., and Ghahramani, Z. (2020a). Einsum networks: Fast and scalable learning of tractable probabilistic circuits. In *Proceedings of the 37th International Conference on Machine Learning*, pages 7563–7574.
- Peharz, R., Vergari, A., Stelzner, K., Molina, A., Shao, X., Trapp, M., Kersting, K., and Ghahramani, Z. (2020b). Random sum-product networks: A simple and effective approach to probabilistic deep learning. In *Proceedings of the 35th Conference on Uncertainty in Artificial Intelligence*, volume 115, pages 334–344.
- Peng, W., Coleman, T., and Mentch, L. (2022). Rates of convergence for random forests via generalized U-statistics. *Electron. J. Stat.*, 16(1):232 – 292.
- Poon, H. and Domingos, P. (2011). Sum-product networks: A new deep architecture. In *Proceedings of the 27th Conference on Uncertainty in Artificial Intelligence*, page 337–346.
- Probst, P. and Boulesteix, A.-L. (2017). To tune or not to tune the number of trees in random forest. *J. Mach. Learn. Res.*, 18(1):6673–6690.
- Probst, P., Wright, M. N., and Boulesteix, A.-L. (2019). Hyperparameters and tuning strategies for random forest. *WIREs Data Mining and Knowledge Discovery*, 9(3):e1301.
- Rahman, T., Kothalkar, P., and Gogate, V. (2014). Cutset networks: A simple, tractable, and scalable approach for improving the accuracy of Chow-Liu trees. In *Machine Learning and Knowledge Discovery in Databases*, pages 630–645, Berlin, Heidelberg. Springer Berlin Heidelberg.
- Ram, P. and Gray, A. G. (2011). Density estimation trees. In *Proceedings of the 17th ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*, page 627–635.
- Ramesh, A., Dhariwal, P., Nichol, A., Chu, C., and Chen, M. (2022). Hierarchical text-conditional image generation with CLIP latents. *arXiv preprint*, 2204.06125.
- Ramesh, A., Pavlov, M., Goh, G., Gray, S., Voss, C., Radford, A., Chen, M., and Sutskever, I. (2021). Zero-shot text-to-image generation. *arXiv preprint*, 2102.12092.
- Ravuri, S. and Vinyals, O. (2019). Classification accuracy score for conditional generative models. In *Advances in Neural Information Processing Systems*, volume 32.
- Rokach, L. and Maimon, O. (2005). *Clustering Methods*, pages 321–352. Springer US, Boston, MA.
- Roy, A., Saffar, M., Vaswani, A., and Grangier, D. (2021). Efficient content-based sparse attention with routing transformers. *Transactions of the Association for Computational Linguistics*, 9:53–68.
- Rubin, D. B. (1996). Multiple imputation after 18+ years. *J. Am. Stat. Assoc.*, 91(434):473–489.
- Santurkar, S., Schmidt, L., and Madry, A. (2018). A classification-based study of covariate shift in GAN distributions. In *Proceedings of the 35th International Conference on Machine Learning*, volume 80, pages 4480–4489.
- Scornet, E. (2016). On the asymptotics of random forests. *J. Multivar. Anal.*, 146:72–83.
- Scornet, E. (2017). Tuning parameters in random forests. *ESAIM: Procs*, 60:144–162.
- Scornet, E., Biau, G., and Vert, J. P. (2015). Consistency of random forests. *Ann. Statist.*, 43(4):1716–1741.

- Shi, T. and Horvath, S. (2006). Unsupervised learning with random forest predictors. *J. Comput. Graph. Stat.*, 15(1):118–138.
- Shmelkov, K., Schmid, C., and Alahari, K. (2018). How good is my GAN? In *Computer Vision – ECCV 2018*, pages 218–234, Cham. Springer International Publishing.
- Silverman, B. (1986). *Density Estimation for Statistics and Data Analysis*. Chapman & Hall, London.
- Smyth, P., Gray, A. G., and Fayyad, U. M. (1995). Retrofitting decision tree classifiers using kernel density estimation. In *Proceedings of the 12th International Conference on International Conference on Machine Learning*, page 506–514.
- Song, Y., Shu, R., Kushman, N., and Ermon, S. (2018). Constructing unrestricted adversarial examples with generative models. In *Advances in Neural Information Processing Systems*, volume 31.
- Song, Y., Sohl-Dickstein, J., Kingma, D., Kumar, A., Erman, S., and Poole, B. (2021). Score-based generative modeling through stochastic differential equations. In *International Conference on Learning Representations*.
- Stadler, T., Oprisanu, B., and Troncoso, C. (2022). Synthetic data – anonymisation groundhog day. In *31st USENIX Security Symposium*, pages 1451–1468.
- Stekhoven, D. J. and Bühlmann, P. (2011). MissForest—non-parametric missing value imputation for mixed-type data. *Bioinformatics*, 28(1):112–118.
- Stone, C. J. (1977). Consistent nonparametric regression. *Ann. Statist.*, 5(4):595 – 620.
- Tang, C., Garreau, D., and von Luxburg, U. (2018). When do random forests fail? In *Advances in Neural Information Processing Systems*, volume 31.
- Tang, F. and Ishwaran, H. (2017). Random forest missing data algorithms. *Stat. Anal. Data Min.*, 10(6):363–377.
- van den Oord, A., Kalchbrenner, N., and Kavukcuoglu, K. (2016). Pixel recurrent neural networks. In *Proceedings of The 33rd International Conference on Machine Learning*, volume 48, pages 1747–1756.
- Van Haaren, J. and Davis, J. (2012). Markov network structure learning: A randomized feature generation approach. In *Proceedings of the AAAI Conference on Artificial Intelligence*, volume 26, pages 1148–1154.
- Vapnik, V. and Chervonenkis, A. (1971). On the uniform convergence of relative frequencies of events to their probabilities. *Theory of Probability and its Applications*, 16:264 – 280.
- Vergari, A., Choi, Y., Peharz, R., and Van den Broeck, G. (2020). Probabilistic circuits: Representations, inference, learning and applications. In *Tutorial at the 34th AAAI Conference on Artificial Intelligence*.
- Vincent, P. and Bengio, Y. (2002). Manifold Parzen windows. In *Advances in Neural Information Processing Systems*, volume 15.
- Wager, S. and Athey, S. (2018). Estimation and inference of heterogeneous treatment effects using random forests. *J. Am. Stat. Assoc.*, 113(523):1228–1242.
- Wager, S. and Walther, G. (2015). Adaptive concentration of regression trees, with application to random forests. *arXiv preprint*, 1503.06388.
- Wand, M. and Jones, M. (1994). *Kernel Smoothing*. Chapman & Hall, Boca Raton, FL.
- Weierstrass, K. (1895). Über continuirliche Functionen eines reellen Arguments, die für keinen Werth des letzteren einen bestimmten Differentialquotient besitzen. In *Mathematische Werke von Karl Weierstrass*, pages 71–74. Mayer & Mueller, Berlin.
- Wen, H. and Hang, H. (2022). Random forest density estimation. In *Proceedings of the 39th International Conference on Machine Learning*, pages 23701–23722.
- Worldline and the ML Group of ULB (2013). Credit card fraud detection data. license: Open database.
- Wu, K., Zhang, K., Fan, W., Edwards, A., and Yu, P. S. (2014). RS-Forest: A rapid density estimator for streaming anomaly detection. In *2014 IEEE International Conference on Data Mining*, pages 600–609.
- Xu, L., Skoularidou, M., Cuesta-Infante, A., and Veeramachaneni, K. (2019). Modeling tabular data using conditional GAN. In *Advances in Neural Information Processing Systems*, volume 32.
- Yang, J., Kannan, A., Batra, D., and Parikh, D. (2017). LR-GAN: Layered recursive generative adversarial networks for image generation. In *International Conference on Learning Representations*.
- Yelmen, B., Decelle, A., Ongaro, L., Marnetto, D., Tallec, C., Montinaro, F., Furtlehner, C., Pagani, L., and Jay, F. (2021). Creating artificial human genomes using generative neural networks. *PLOS Genetics*, 17(2):1–22.
- Zhang, H., Juba, B., and Van Den Broeck, G. (2021). Probabilistic generating circuits. In *Proceedings of the 38th International Conference on Machine Learning*, pages 12447–12457.
- Zhao, Z., Kunar, A., Birke, R., and Chen, L. Y. (2022). CTAB-GAN+: Enhancing tabular data synthesis. *arXiv preprint*, 2204.00401.

A PROOFS

A.1 Proof of Thm. 1

To secure the result, we must show that (a) the discriminator reliably converges on the Bayes risk at each iteration t ; and (b) the generator’s sampling strategy drives original and synthetic data closer together, ultimately taking the Bayes risk to $1/2$ as $n, t \rightarrow \infty$. (For the purposes of this proof, we set the tolerance parameter δ to 0.)

Take (a) first. This amounts to a consistency requirement for RFs. The consistency of RF classifiers has been demonstrated under various assumptions about splitting rules and stopping criteria (Breiman, 2004; Biau et al., 2008; Biau and Devroye, 2010; Gao et al., 2022), but these results generally require trees to be grown to purity or even completion (i.e., $n_b^\ell = 1$ for all ℓ, b). However, this would turn the generator’s sampling strategy into a simple copy-paste operation and make intra-leaf density estimation impossible. We therefore follow Malley et al. (2012) in observing that regression procedures constitute probability machines, since $P(Y = 1|\mathbf{x}) = \mathbb{E}[Y|\mathbf{x}]$ for $Y \in \{0, 1\}$.

For simplicity, we focus on the single tree case, as the consistency of the ensemble follows from the consistency of the base method (Biau et al., 2008). We define $\eta^{(t)}(\mathbf{x}) := P(Y = 1|\mathbf{x}, t)$ as the target function for fixed t . Let $f_n^{(t)}(\mathbf{x})$ be a tree trained according to (A1)-(A3) on a sample of size n at iteration t . Since L_2 -consistency entails classifier consistency using the soft labeling approach of (A3).(vi), our goal in this section is to show that, for all $t \in \mathbb{N}$, we have:

$$\lim_{n \rightarrow \infty} \mathbb{E} \left[\left(f_n^{(t)}(\mathbf{x}) - \eta^{(t)}(\mathbf{x}) \right)^2 \right] = 0.$$

Consistency for RF regression has been established for several variants of the algorithm, occasionally under some constraints on the data generating process (Genuer, 2012; Scornet et al., 2015; Wager and Walther, 2015; Biau and Scornet, 2016). Recent work in this area has tended to focus on asymptotic normality (Mentch and Hooker, 2016; Wager and Athey, 2018; Athey et al., 2019; Peng et al., 2022), which requires additional assumptions. These often include an upper bound on leaf sample size, which would complicate our analysis in Thm. 2. To avoid unnecessary difficulties, we borrow selectively from Meinshausen (2006), Biau (2012), Denil et al. (2014), Scornet (2016), and Wager and Athey (2018), striking a delicate balance between theoretical parsimony and fidelity to the classic RF algorithm.⁵

For full details, we refer readers to the original texts. The main point to recognize is that, under assumptions (A1)-(A3), RFs satisfy the conditions of Stone’s theorem (Stone, 1977), which guarantees the universal consistency of a large family of local averaging methods. Devroye et al. (1996, Thm. 6.1) and Györfi et al. (2002, Thm. 4.2) show that partitioning estimators such as decision trees qualify provided that (1) $\text{diam}(\mathcal{X}_\ell) \rightarrow_p 0$ and (2) $n_\ell \rightarrow_p \infty$ for all ℓ as $n \rightarrow \infty$, effectively creating leaves of infinite density. The former is derived by Meinshausen (2006, Lemma 2) under (A3).(iii) and (A3).(iv); the latter follows trivially from (A3).(v). Thus RF discriminators weakly converge on the Bayes risk in the large sample limit, completing part (a) of the proof.

Desideratum (b) effectively says that original and synthetic data become indistinguishable as n and t increase. Recall that at $t = 0$, we generate synthetic data $\tilde{\mathbf{X}}^{(0)} \sim \prod_{j=1}^d P(X_j)$, which becomes input to the discriminator $f_n^{(0)}$. Let $\boldsymbol{\theta}^{(0)}$ denote the resulting splits once the discriminator has converged. In subsequent rounds, synthetic data are sampled according to $\tilde{\mathbf{X}}^{(t+1)} \sim \prod_{j=1}^d P(X_j|\theta_\ell^{(t)})P(\theta_\ell^{(t)})$. (The consistency of coverage estimates is treated separately in Appx. A.3.) We proceed to train a new discriminator and repeat the process.

Let P^* be the target distribution and $P^{(t)}$ the synthetic distribution at round t . For all $t \geq 1$, the input data to the discriminator $f_n^{(t)}$ is the dataset $\mathcal{D}_n^{(t)} \sim 0.5P^* + 0.5P^{(t-1)}$. Our goal in this section is to show that, as $n, t \rightarrow \infty$:

$$\sup_{\mathbf{x} \in \mathcal{D}_n^{(t)}} |\eta^{(t)}(\mathbf{x}) - 1/2| \rightarrow_p 0.$$

An apparent challenge to our recursive strategy for generating synthetic data is posed by self-similar distributions, in which dependencies replicate at ever finer resolutions, as in some fractal equations (Mandelbrot, 1982). For instance, let g be the Weierstrass function (Weierstrass, 1895), and say that $X_2 = g(X_1)$. Then the generative model will tend to produce

⁵Several authors have conjectured that RF consistency may not require honesty (A3).(i) or subsampling (A3).(ii) after all. Empirical performance certainly seems unencumbered by these requirements. However, both come with major theoretical advantages—the former by making predictions conditionally independent of the training data while preserving some form of adaptive splits, the latter by avoiding thorny issues arising from duplicated samples when bootstrapping. See Biau (2012, Rmk. 8), Wager and Athey (2018, Appx. B), and Tang et al. (2018) for a discussion.

off-manifold data at each iteration t , no matter how small $\text{vol}(\mathcal{X}_\ell)$ becomes. However, this only shows that convergence can fail for finite t . Since the discriminator is consistent, it will accurately identify synthetic points in round $t + 1$, pruning the space still further.

Let $[L^{(t)}]$ be the leaves of the discriminator $f_n^{(t)}$, and define the maximum leaf diameter $m_t := \max_{\ell \in [L^{(t)}]} \text{diam}(\mathcal{X}_\ell)$. We say that two samples are *neighbors* in $f_n^{(t)}$ if the model places them in the same leaf. We show that, as $n, t \rightarrow \infty$, conditional probabilities for neighboring samples converge—including, crucially, original and synthetic counterparts. Our Lipschitz condition (A2) states that for all $\mathbf{x}, \mathbf{x}'$, we have:

$$|\eta^{(t)}(\mathbf{x}) - \eta^{(t)}(\mathbf{x}')| \leq c_t \|\mathbf{x} - \mathbf{x}'\|_2,$$

where c_t denotes the Lipschitz constant at round t . Suppose that $\mathbf{x}$ and $\mathbf{x}'$ are neighbors. Then we can replace the second factor on the rhs with m_t , since the L_2 distance between neighbors cannot exceed the maximum leaf diameter at round t . Meinshausen (2006)’s aforementioned Lemma 2 ensures that this value goes to zero in probability as rounds increase. This could in principle be offset by a sufficient increase in c_t over training rounds, but the second condition of (A2) prevents this, imposing the constraint that $c_t = o(m_t^{-1})$. Thus, for observations in the same leaf, $c_t m_t \rightarrow_p 0$ as $t \rightarrow \infty$. Because original and synthetic samples are equinumerous in all leaves following the generative step, each original sample has a synthetic counterpart to which it is arbitrarily close in L_2 space as t grows large. Since no feature values are sufficient to distinguish between the two classes in any region, all conditional probabilities go to 1/2, and Bayes risk therefore also converges to 1/2 in probability. This concludes the proof.

A.2 Proof of Lemma 1

Define the first-order approximation to p satisfying local independence:

$$\hat{p}(\mathbf{x}) := \frac{1}{B} \sum_{\ell, b: \mathbf{x} \in \mathcal{X}_b^\ell} p(\theta_b^\ell) \prod_{j=1}^d p(x_j | \theta_b^\ell).$$

We also define the root integrated squared error (RISE), i.e. the Euclidean distance between probability densities:

$$\text{RISE}(p, q) := \left(\int_{\mathcal{X}} \left(p(\mathbf{x}) - q(\mathbf{x}) \right)^2 d\mathbf{x} \right)^{1/2}.$$

By the triangle inequality, we have:

$$\text{RISE}(p, q) \leq \text{RISE}(p, \hat{p}) + \text{RISE}(\hat{p}, q).$$

Squaring both sides, we get:

$$\text{ISE}(p, q) \leq \text{ISE}(p, \hat{p}) + \text{ISE}(\hat{p}, q) + 2 \text{RISE}(p, \hat{p}) \text{RISE}(\hat{p}, q).$$

Adding a nonnegative value to the rhs, we can reduce the expression:

$$\begin{aligned} \text{ISE}(p, q) &\leq \text{ISE}(p, \hat{p}) + \text{ISE}(\hat{p}, q) + 2 \text{RISE}(p, \hat{p}) \text{RISE}(\hat{p}, q) + (\text{RISE}(p, \hat{p}) - \text{RISE}(\hat{p}, q))^2 \\ &= 2(\text{ISE}(p, \hat{p}) + \text{ISE}(\hat{p}, q)). \end{aligned}$$

Now observe that we can rewrite both ISE formulae in terms of our predefined residuals (Eqs. 3-5):

$$\begin{aligned} \text{ISE}(p, \hat{p}) &= \int_{\mathcal{X}} \left(\frac{1}{B} \sum_{\ell, b: \mathbf{x} \in \mathcal{X}_b^\ell} \left(p(\mathbf{x} | \theta_b^\ell) p(\theta_b^\ell) - \prod_{j=1}^d p(x_j | \theta_b^\ell) p(\theta_b^\ell) \right) \right)^2 d\mathbf{x} \\ &= \frac{1}{B^2} \int_{\mathcal{X}} \left(\sum_{\ell, b: \mathbf{x} \in \mathcal{X}_b^\ell} p(\theta_b^\ell) \epsilon_3 \right)^2 d\mathbf{x}. \\ \text{ISE}(\hat{p}, q) &= \int_{\mathcal{X}} \left(\frac{1}{B} \sum_{\ell, b: \mathbf{x} \in \mathcal{X}_b^\ell} \left(\prod_{j=1}^d p(x_j | \theta_b^\ell) p(\theta_b^\ell) - \prod_{j=1}^d q(x_j; \theta_{b,j}^\ell) q(\theta_b^\ell) \right) \right)^2 d\mathbf{x} \\ &= \frac{1}{B^2} \int_{\mathcal{X}} \left(\sum_{\ell, b: \mathbf{x} \in \mathcal{X}_b^\ell} \left(p(\theta_b^\ell) \epsilon_2 + \epsilon_1 \prod_{j=1}^d p(x_j | \theta_b^\ell) - \epsilon_1 \epsilon_2 \right) \right)^2 d\mathbf{x}. \end{aligned}$$

We replace the interior squared terms for ease of presentation:

$$\begin{aligned}\alpha &:= \sum_{\ell, b: \mathbf{x} \in \mathcal{X}_b^\ell} p(\theta_b^\ell) \epsilon_3 \\ \beta &:= \sum_{\ell, b: \mathbf{x} \in \mathcal{X}_b^\ell} \left(p(\theta_b^\ell) \epsilon_2 + \epsilon_1 \prod_{j=1}^d p(x_j | \theta_b^\ell) - \epsilon_1 \epsilon_2 \right).\end{aligned}$$

Finally, we take expectations on both sides:

$$\text{MISE}(p, q) \leq 2B^{-2} \mathbb{E} \left[\int_{\mathcal{X}} \alpha^2 + \beta^2 d\mathbf{x} \right],$$

where we have exploited the linearity of expectation to pull the factor outside of the bracketed term, and the monotonicity of expectation to preserve the inequality.

A.3 Proof of Theorem 2

Lemma 1 states that error is bounded by a quadratic function of $\epsilon_1, \epsilon_2, \epsilon_3$. Thus for L_2 -consistency, it suffices to show that $\mathbb{E}[\epsilon_j^2] \rightarrow 0$, for $j \in \{1, 2, 3\}$. Since this is already established by Thm. 1 for $j = 3$, we focus here on errors of coverage and density. Start with ϵ_1 . A general version of the Glivenko-Cantelli theorem (Vapnik and Chervonenkis, 1971) guarantees uniform convergence of empirical proportions to population proportions. Let $\mathcal{L}$ denote the set of all possible hyperrectangular subspaces induced by axis-aligned splits on $\mathcal{X}$. Then the following holds with probability 1:

$$\lim_{n \rightarrow \infty} \sup_{\ell \in \mathcal{L}} |p(\theta^\ell) - q_n(\theta^\ell)| = 0.$$

Next, take ϵ_2 . (A4) guarantees that p satisfies the consistency conditions for univariate KDE (Silverman, 1986; Wand and Jones, 1994; Gramacki, 2018), while condition (v) of (A3) ensures that within-leaf sample size increases even as leaf volume goes to zero (Meinshausen, 2006, Lemma 2). Our kernel is a nonnegative function $K : \mathbb{R}^d \rightarrow \mathbb{R}$ that integrates to 1, parametrized by the bandwidth h :

$$p_h(\mathbf{x}) = \frac{1}{nh} \sum_{i=1}^n K\left(\frac{\mathbf{x} - \mathbf{x}_i}{h}\right).$$

Using standard arguments, we take a Taylor series expansion of the MISE and minimize the *asymptotic* MISE (AMISE):

$$\text{AMISE}(p, p_h) = \frac{1}{nh} R(K) + \frac{1}{4} h^4 \mu_2(K)^2 R(p''),$$

where

$$\begin{aligned}R(K) &= \int K(x)^2 dx, \\ \mu_2(K) &= \int x^2 K(x) dx, \text{ and} \\ R(p'') &= \int p''(x)^2 dx.\end{aligned}$$

For example values of these variables under specific kernels, see (Wand and Jones, 1994, Appx. B). Under (A4), it can be shown that

$$\text{MISE}(p, p_h) = \text{AMISE}(p, p_h) + o((nh)^{-1} + h^4).$$

Thus if $(nh)^{-1} \rightarrow 0$ and $h \rightarrow 0$ as $n \rightarrow \infty$, the asymptotic approximation is exact and $\mathbb{E}[\epsilon_2^2] \rightarrow 0$.

These results, combined with the proof of Thm. 1 (see Appx. A.1), establish that errors of coverage, density, and convergence all vanish in the limit. Thus $\mathbb{E}[\epsilon_j^2] \rightarrow 0$ for $j \in \{1, 2, 3\}$, and the proof is complete.

B EXPERIMENTS

Our experiments do not include any personal data, as defined in Article 4(1) of the European Union’s General Data Protection Regulation. All data are either simulated or from publicly available resources. We performed all experiments on a dedicated 64-bit Linux platform running Ubuntu 20.04 with an AMD Ryzen Threadripper 3960X (24 cores, 48 threads) CPU, 256 gigabyte RAM and two NVIDIA Titan RTX GPUs. We used R version 4.1.2 and Python version 3.7.12. Further details on the environment setup are provided in the supplemental code.

B.1 Simulations

The `cassini`, `shapes`, and `smiley` simulations are all available in the `mlbench` R package; the `twomoons` problem is available in the `fdm2id` R package. Default parameters were used throughout, with fixed sample size $n = 2000$.

B.2 Twenty Datasets

The Twenty Datasets benchmark was originally proposed by Van Haaren and Davis (2012). A conventional training/validation/test split is widely used in the PC literature. Because our method does not include any hyperparameter search, we combine training and validation sets into a single training set. We downloaded the data from <https://github.com/joshuacnf/Probabilistic-Generating-Circuits/tree/main/data> and include the directory in our project GitHub repository for completeness. All datasets are Boolean, with sample size and dimensionality given in Table 3.

Table 3: Summary of datasets included in the Twenty Datasets benchmark.

Dataset	Train	Validation	Test	Dimensions
<code>nlts</code>	16181	2157	3236	16
<code>msnbc</code>	291326	38843	58265	17
<code>kdd</code>	180092	19907	34955	64
<code>plants</code>	17412	2321	3482	69
<code>audio</code>	15000	2000	3000	100
<code>jester</code>	9000	1000	4116	100
<code>netflix</code>	15000	2000	3000	100
<code>accidents</code>	12758	1700	2551	111
<code>retail</code>	22041	2938	4408	135
<code>pumsb</code>	12262	1635	2452	163
<code>dna</code>	1600	400	1186	180
<code>kosarek</code>	33375	4450	6675	190
<code>msweb</code>	29441	3270	5000	294
<code>book</code>	8700	1159	1739	500
<code>movie</code>	4524	1002	591	500
<code>webkb</code>	2803	558	838	839
<code>reuters</code>	6532	1028	1540	889
<code>20ng</code>	11293	3764	3764	910
<code>bbc</code>	1670	225	330	1058
<code>ad</code>	2461	327	491	1556

Results for competitors are reported in the cited papers:

- Einsum networks (Peharz et al., 2020a)
- Random sum-product networks (Peharz et al., 2020b)
- Probabilistic generating circuits (Zhang et al., 2021)
- Strudel (Dang et al., 2022)
- Continuous mixtures of Chow-Liu trees (Correia et al., 2023).

B.3 Tabular GANs

For benchmarking generative models on real-world data, we use the benchmarking pipeline proposed by Xu et al. (2019). In detail, the workflow is as follows:

1. Load classification datasets used in Xu et al. (2019), namely `adult`, `census`, `credit`, `covertypes`, `intrusion`, `mnist12`, and `mnist28`. Note that the type of prediction task does not affect the process of synthetic data generation, so we omit the single regression example (`news`) for greater consistency.

2. Split the data into training and test sets (see Table 4 for details).
3. Train the generative models FORGE (number of trees = 10, minimum node size = 5), CTGAN⁶ (batch size = 500, epochs = 300), TVAE⁷ (batch size = 500, epochs = 300), CTAB-GAN+⁸ (batch size = 500, epochs = 150), IT-GAN⁹ (batch size = 2000, epochs = 300) and RCC-GAN¹⁰ (batch size = 500, epochs = 300).
4. Generate a synthetic dataset of the same size as the training set using each of the generative models trained in step (3), measuring the wall time needed to execute this task.
5. Train a set of supervised learning algorithms (see Table 4 for details): (a) on the real training data set (i.e., the *Oracle*); and (b) on the synthetic training datasets generated by FORGE, CTGAN, TVAE, CTAB-GAN+ and RCC-GAN.
6. Evaluate the performance of the learning algorithms from step (5) on the test set.
7. For each dataset, average performance metrics (accuracy, F1-scores) across learners. We report F1-scores for the positive class, e.g. '>50k' for adult, '+50000' for census and '1' for credit.

Table 4: Benchmark Setup. Supervised learning algorithms for prediction: (A) Adaboost, estimators = 50, (B) Decision Tree, tree depth for binary/multiclass target = 15/30, (C) Logistic Regression, (D) MLP, hidden layers for binary/multiclass target = 50/100

Dataset	Train/Test	Learner	Link to dataset
adult (Dua and Graff, 2019)	23k/10k	A,B,C,D	http://archive.ics.uci.edu/ml/datasets/adult
census (Dua and Graff, 2019)	200k/100k	A,B,D	https://archive.ics.uci.edu/ml/datasets/census+income
covertypes (Blackard, 1998)	481k/100k	A,D	https://archive.ics.uci.edu/ml/datasets/covertypes
credit (Worldline and the ML Group of ULB, 2013)	264k/20k	A,B,D	https://www.kaggle.com/mlg-ulb/creditcardfraud
intrusion (Dua and Graff, 2019)	394k/100k	A,D	http://archive.ics.uci.edu/ml/datasets/kdd+cup+1999+data
mnist12 (LeCun et al., 1998)	60k/10k	A,D	http://yann.lecun.com/exdb/mnist/index.html
mnist28 (LeCun et al., 1998)	60k/10k	A,D	http://yann.lecun.com/exdb/mnist/index.html

B.4 Run Time

In order to evaluate the run time efficiency of FORGE, we chose to focus on the smallest dataset of the benchmark study in Sect. 5.2, namely *adult*. We (A) drew stratified subsamples and (B) drew covariate subsets. For step (B), the target variable is always included. We select an equal number of continuous/categorical covariates when possible and use all $n = 32,561$ instances. Results in terms of processing time are visualized in Fig. 4.

B.5 Image Data

We include results on the *mnist12* and *mnist28* datasets here, both included in the original Xu et al. (2019) pipeline. Benchmarking against CTGAN and TVAE (other methods proved too slow to test), we find that FORGE outperforms both competitors in accuracy, F1-score, and speed (see Table 5).

However, since MNIST is not a tabular data problem, perhaps a more relevant comparison would be against convolutional networks specifically designed for image data. We train a conditional GAN with convolutional layers (Mirza and Osindero, 2014) and find that the resulting cGAN clearly outperforms FORGE (see Fig. 5). This result is expected, given that our method is not optimized for image data. It also illustrates a limitation of our approach, which excels in speed and flexibility but is no match for deep learning methods on structured datasets.

⁶https://sdv.dev/SDV/api_reference/tabular/ctgan.html. MIT License.

⁷https://sdv.dev/SDV/api_reference/tabular/tvae.html. MIT License.

⁸<https://github.com/Team-TUD/CTAB-GAN-Plus>

⁹<https://github.com/leejaehoon2016/ITGAN>. Samsung SDS Public License V1.0.

¹⁰<https://github.com/EsmaeilpourMohammad/RccGAN>

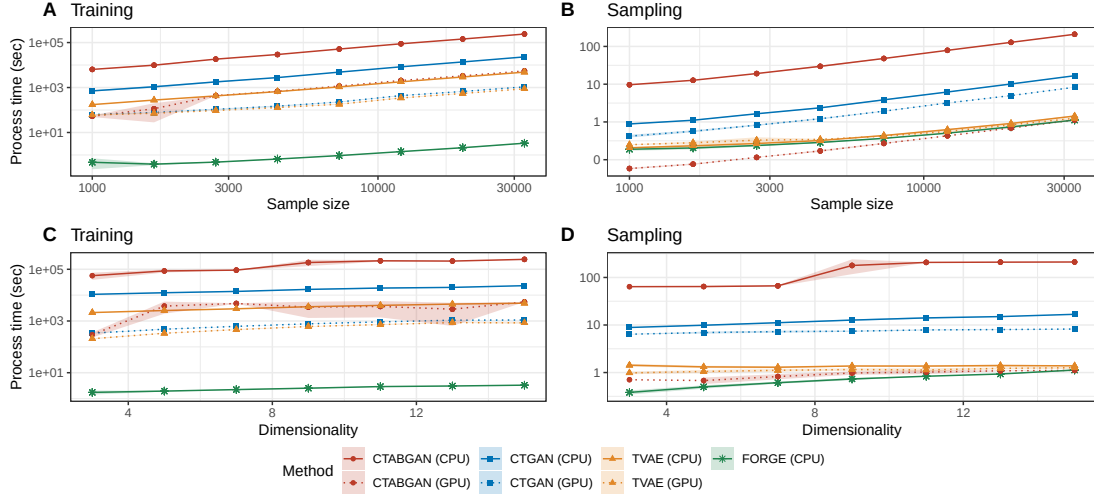


Figure 4: Complexity curves. (A): Processing time as a function of sample size, using stratified subsamples of the adult dataset. (B): Processing time as a function of dimensionality, using random features from the adult dataset.

Table 5: Performance on *mnist* datasets from the Xu et al. (2019) benchmark for CTGAN and TVAE vs. FORGE. We report average results across five replicates $\pm$ the associated standard error. Winning results in bold.

Dataset	Model	Accuracy $\pm$ SE	F1 $\pm$ SE	Time (sec)
<i>mnist</i> 12 classes = 10 $n = 70,000$ $d = 144$	Oracle	0.892 ± 0.003	0.891 ± 0.003	
	FORGE	0.799 ± 0.007	0.795 ± 0.007	32.3
	CTGAN	0.172 ± 0.032	0.138 ± 0.032	2737.4
	TVAE	0.763 ± 0.002	0.761 ± 0.002	1143.8
<i>mnist</i> 28 classes = 10 $n = 70,000$ $d = 784$	Oracle	0.918 ± 0.002	0.917 ± 0.002	
	FORGE	0.729 ± 0.008	0.723 ± 0.008	169.5
	CTGAN	0.197 ± 0.051	0.167 ± 0.051	14415.4
	TVAE	0.698 ± 0.016	0.697 ± 0.016	5056.0

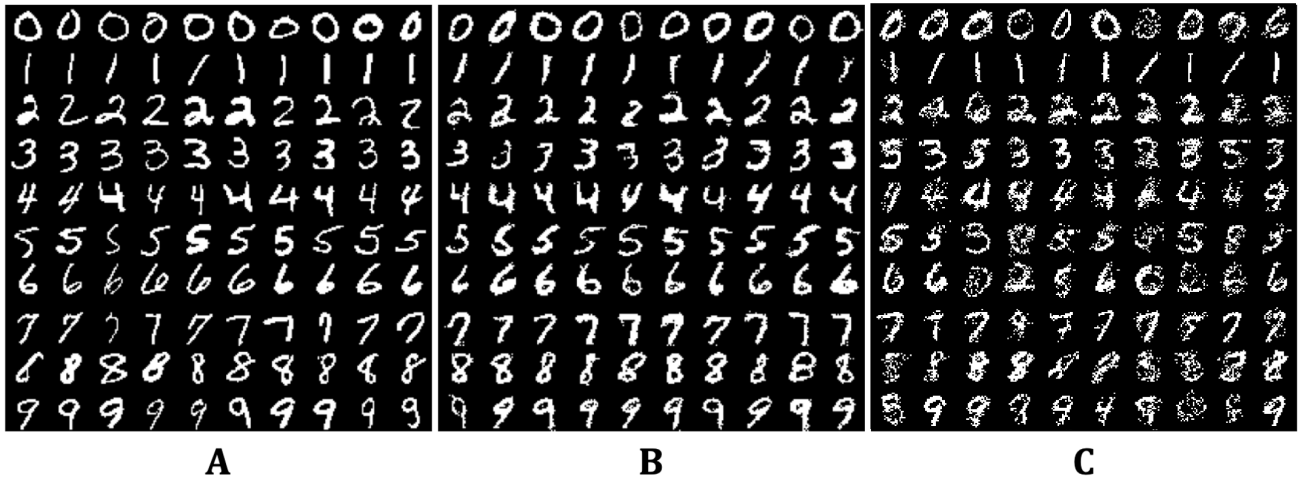


Figure 5: Results from *mnist*28 experiment. (A): Original samples. (B): Samples generated by cGAN. (C): Samples generated by FORGE.