
The Tree Ensemble Layer: Differentiability meets Conditional Computation

Hussein Hazimeh¹ Natalia Ponomareva² Petros Mol² Zhenyu Tan³ Rahul Mazumder¹

Abstract

Neural networks and tree ensembles are state-of-the-art learners, each with its unique statistical and computational advantages. We aim to combine these advantages by introducing a new layer for neural networks, composed of an ensemble of differentiable decision trees (a.k.a. soft trees). While differentiable trees demonstrate promising results in the literature, they are typically slow in training and inference as they do not support conditional computation. We mitigate this issue by introducing a new sparse activation function for sample routing, and implement true conditional computation by developing specialized forward and backward propagation algorithms that exploit sparsity. Our efficient algorithms pave the way for jointly training over deep and wide tree ensembles using first-order methods (e.g., SGD). Experiments on 23 classification datasets indicate over 10x speed-ups compared to the differentiable trees used in the literature and over 20x reduction in the number of parameters compared to gradient boosted trees, while maintaining competitive performance. Moreover, experiments on CIFAR, MNIST, and Fashion MNIST indicate that replacing dense layers in CNNs with our tree layer reduces the test loss by 7-53% and the number of parameters by 8x. We provide an open-source TensorFlow implementation with a Keras API.

1. Introduction

Decision tree ensembles have proven very successful in various machine learning applications. Indeed, they are often referred to as the best “off-the-shelf” learners (Hastie et al., 2009), as they exhibit several appealing properties such as ease of tuning, robustness to outliers, and interpretability

(Hastie et al., 2009; Chen & Guestrin, 2016). Another natural property in trees is *conditional computation*, which refers to their ability to route each sample through a small number of nodes (specifically, a single root-to-leaf path). Conditional computation can be broadly defined as the ability of a model to activate only a small part of its architecture in an input-dependent fashion (Bengio et al., 2015). This can lead to both computational benefits and enhanced statistical properties. On the computation front, routing samples through a small part of the tree leads to substantial training and inference speed-ups compared to methods that do not route samples. Statistically, conditional computation offers the flexibility to reduce the number of parameters used by each sample, which can act as a regularizer (Breiman et al., 1983; Hastie et al., 2009; Bengio et al., 2015).

However, the performance of trees relies on feature engineering, since they lack a good mechanism for representation learning (Bengio et al., 2013). This is an area in which neural networks (NNs) excel, especially in speech and image recognition applications (Bengio et al., 2013; He et al., 2015; Yu & Deng, 2016). However, NNs do not naturally support conditional computation and are harder to tune.

In this work, we combine the advantages of neural networks and tree ensembles by designing a hybrid model. Specifically, we propose the *Tree Ensemble Layer (TEL)* for neural networks. This layer is an additive model of differentiable decision trees, can be inserted anywhere in a neural network, and is trained along with the rest of the network using gradient-based optimization methods (e.g., SGD). While differentiable trees in the literature show promising results, especially in the context of neural networks, e.g., Kontschieder et al. (2015); Frosst & Hinton (2017), they do not offer true conditional computation. We equip TEL with a novel mechanism to perform conditional computation, during both training and inference. We make this possible by introducing a new sparse activation function for sample routing, along with specialized forward and backward propagation algorithms that exploit sparsity. Experiments on 23 real datasets indicate that TEL achieves over 10x speed-ups compared to the current differentiable trees, without sacrificing predictive performance.

Our algorithms pave the way for jointly optimizing over both wide and deep tree ensembles. Here joint optimization

¹Massachusetts Institute of Technology ²Google Research

³Google Brain. Correspondence to: Hussein Hazimeh <hazimeh@mit.edu>.

refers to updating all the trees simultaneously (e.g., using first-order methods like SGD). This has been a major computational challenge prior to our work. For example, jointly optimizing over classical (non-differentiable) decision trees is a hard combinatorial problem (Hastie et al., 2009). Even with differentiable trees, the training complexity grows exponentially with the tree depth, making joint optimization difficult (Kontschieder et al., 2015). A common approach is to train tree ensembles using greedy “stage-wise” procedures, where only one tree is updated at a time and never updated again—this is a main principle in gradient boosted decision trees (GBDT) (Friedman, 2001)¹. We hypothesize that joint optimization yields more compact and expressive ensembles than GBDT. Our experiments confirm this, indicating that TEL can achieve over 20x reduction in model size. This can have important implications for interpretability, latency and storage requirements during inference.

Contributions: Our contributions can be summarized as follows: **(i)** We design a new differentiable activation function for trees which allows for routing samples through small parts of the tree (similar to classical trees). **(ii)** We realize conditional computation by developing specialized forward and backward propagation algorithms that exploit sparsity to achieve an optimal time complexity. Notably, the complexity of our backward pass can be independent of the tree depth and is generally better than that of the forward pass—this is not possible in backpropagation for neural networks. **(iii)** We perform experiments on a collection of 26 real datasets, which confirm TEL as a competitive alternative to current differentiable trees, GBDT, and dense layers in CNNs. **(iv)** We provide an open-source TensorFlow implementation of TEL along with a Keras interface².

Related Work: Table 1 summarizes the most relevant related work. Differentiable decision trees (a.k.a. soft trees) are an instance of the Hierarchical Mixture of Experts introduced by Jordan & Jacobs (1994). The internal nodes of these trees act as routers, sending samples to the left and right with different proportions. This framework does not support conditional computation as each sample is processed in all the tree nodes. Our work avoids this issue by allowing each sample to be routed through small parts of the tree, without losing differentiability. A number of recent works have used soft trees in the context of deep learning. For example, Kontschieder et al. (2015) equipped soft trees with neural representations and used alternating minimization to learn the feature representations and the

Table 1: Related work on conditional computation

Paper	CT	CI	DO	Model/Optim
Kontschieder et al. (2015)	N	N	Y	Soft tree/Alter
Ioannou et al. (2016)	N	H	Y	Tree-NN/SGD
Frosst & Hinton (2017)	N	H	Y	Soft tree/SGD
Zoran et al. (2017)	N	H	N	Soft tree/Alter
Shazeer et al. (2017)	H	Y	N	Tree-NN/SGD
Tanno et al. (2018)	N	H	Y	Soft tree/SGD
Biau et al. (2019)	H	N	Y	Tree-NN/SGD
Hehn et al. (2019)	N	H	Y	Soft tree/SGD
Our method	Y	Y	Y	Soft tree/SGD

H is heuristic (e.g., training model is different from inference), *CT* is conditional training. *CI* is conditional inference. *DO* indicates whether the objective function is differentiable. *Soft tree* refers to a differentiable tree, whereas *Tree-NN* refers to NNs with a tree-like structure. *Optim* stands for optimization (SGD or alternating minimization).

leaf outputs. Hehn et al. (2019) extended Kontschieder et al. (2015)’s approach to allow for conditional inference and growing trees level-by-level. Frosst & Hinton (2017) trained a (single) soft tree using SGD and leveraged a deep neural network to expand the dataset used in training the tree. Zoran et al. (2017) also leveraged a tree structure with a routing mechanism similar to soft trees, in order to equip the k-nearest neighbors algorithm with neural representations. All of these works have observed that computation in a soft tree can be expensive. Thus, in practice, heuristics are used to speed up inference, e.g., Frosst & Hinton (2017) uses the root-to-leaf path with the highest probability during inference, leading to discrepancy between the models used in training and inference. Instead of making a tree differentiable, Jernite et al. (2017) hypothesized about properties the best tree should have, and introduced a pseudo-objective that encourages balanced and pure splits. They optimized using SGD along with intermediate processing steps.

Another line of work introduces tree-like structure to NNs via some routing mechanism. For example, Ioannou et al. (2016) employed tree-shaped CNNs with branches as weight matrices with sparse block diagonal structure. Shazeer et al. (2017) created the Sparsely-Gated Mixture-of-Experts layer where samples are routed to subnetworks selected by a trainable gating network. Biau et al. (2019) represented a decision tree using a 3-layer neural network and combined CART and SGD for training. Tanno et al. (2018) looked into adaptively growing an NN with routing nodes for performing tree-like conditional computations. However, in these works, the inference model is either different from training or the router is not differentiable (but still trained using SGD)—see Table 1 for details.

¹There are follow-up works on GBDT which update the leaves of all trees simultaneously, e.g., see Johnson & Zhang (2013). However, our approach allows for updating both the internal node and leaf weights simultaneously.

²https://github.com/google-research/google-research/tree/master/tf_trees

2. The Tree Ensemble Layer

TEL is an additive model of differentiable decision trees. In this section, we introduce TEL formally and then discuss the routing mechanism used in our trees. For simplicity, we assume that TEL is used as a standalone layer. Training trees with other layers will be discussed in Section 3.

We assume a supervised learning setting, with input space $\mathcal{X} \subseteq \mathbb{R}^p$ and output space $\mathcal{Y} \subseteq \mathbb{R}^k$. For example, in the case of regression (with a single output) $k = 1$, while in classification k depends on the number of classes. Let m be the number of trees in the ensemble, and let $T^{(j)} : \mathcal{X} \rightarrow \mathbb{R}^k$ be the j th tree in the ensemble. For an input sample $x \in \mathbb{R}^p$, the output of the layer is a sum over all the tree outputs:

$$\mathcal{T}(x) = T^{(1)}(x) + T^{(2)}(x) + \dots + T^{(m)}(x). \quad (1)$$

The output of the layer, $\mathcal{T}(x)$, is a vector in $\mathbb{R}^k$ containing raw predictions. In the case of classification, mapping from raw predictions to $\mathcal{Y}$ can be done by applying a softmax and returning the class with the highest probability. Next, we introduce the key building block of the approach: the differentiable decision tree.

The Differentiable Decision Tree: Classical decision trees perform *hard routing*, i.e., a sample is routed to exactly one direction at every internal node. Hard routing introduces discontinuities in the loss function, making trees unamenable to continuous optimization. Therefore, trees are usually built in a greedy fashion. In this section, we present an enhancement of the soft trees proposed by Jordan & Jacobs (1994) and utilized in Kotschieder et al. (2015); Frosst & Hinton (2017); Hehn et al. (2019). Soft trees are a variant of decision trees that perform *soft routing*, where every internal node can route the sample to the left and right simultaneously, with different proportions. This routing mechanism makes soft trees differentiable, so learning can be done using gradient-based methods. Soft trees cannot route a sample exclusively to the left or to the right, making conditional computation impossible. Subsequently, we introduce a new activation function for soft trees, which allows conditional computation while preserving differentiability.

We consider a single tree in the additive model (1), and denote the tree by T (we drop the superscript to simplify the notation). Recall that T takes an input sample and returns an output vector (logit), i.e., $T : \mathcal{X} \subseteq \mathbb{R}^p \rightarrow \mathbb{R}^k$. Moreover, we assume that T is a perfect binary tree with depth d . We use the sets $\mathcal{I}$ and $\mathcal{L}$ to denote the internal (split) nodes and the leaves of the tree, respectively. For any node $i \in \mathcal{I} \cup \mathcal{L}$, we define $\mathcal{A}(i)$ as its set of ancestors and use the notation $\{x \rightarrow i\}$ for the event that a sample $x \in \mathbb{R}^p$ reaches i . A summary of the notation used in this paper can be found in Table A.1 in the appendix.

Soft Routing: Internal tree nodes perform soft routing,

where a sample is routed left and right with different proportions. We will introduce soft routing using a probabilistic model. While we use probability to model the routing process, we will see that the final prediction of the tree is an expectation over the leaves, making T a deterministic function. Unlike classical decision trees which use axis-aligned splits, soft trees are based on hyperplane (a.k.a. oblique) splits (Murthy et al., 1994), where a linear combination of the features is used in making routing decisions. Particularly, each internal node $i \in \mathcal{I}$ is associated with a trainable weight vector $w_i \in \mathbb{R}^p$ that defines the node’s hyperplane split. Let $\mathcal{S} : \mathbb{R} \rightarrow [0, 1]$ be an activation function. Given a sample $x \in \mathbb{R}^p$, the probability that internal node i routes x to the left is defined by $\mathcal{S}(\langle w_i, x \rangle)$.

Now we discuss how to model the probability that x reaches a certain leaf l . Let $[l \swarrow i]$ (resp. $[i \searrow l]$) denote the event that leaf l belongs to the left (resp. right) subtree of node $i \in \mathcal{I}$. Assuming that the routing decision made at each internal node in the tree is independent of the other nodes, the probability that x reaches l is given by:

$$P(\{x \rightarrow l\}) = \prod_{i \in \mathcal{A}(l)} r_{i,l}(x), \quad (2)$$

where $r_{i,l}(x)$ is the probability of node i routing x towards the subtree containing leaf l , i.e., $r_{i,l}(x) := \mathcal{S}(\langle x, w_i \rangle)^{\mathbb{1}[l \swarrow i]} (1 - \mathcal{S}(\langle x, w_i \rangle))^{\mathbb{1}[i \searrow l]}$. Next, we define how the root-to-leaf probabilities in (2) can be used to make the final prediction of the tree.

Prediction: As with classical decision trees, we assume that each leaf stores a weight vector $o_l \in \mathbb{R}^k$ (learned during training). Note that, during a forward pass, o_l is a constant vector, meaning that it is not a function of the input sample(s). For a sample $x \in \mathbb{R}^p$, we define the prediction of the tree as the expected value of the leaf outputs, i.e.,

$$T(x) = \sum_{l \in \mathcal{L}} P(\{x \rightarrow l\}) o_l. \quad (3)$$

Activation Functions: In soft routing, the internal nodes use an activation function $\mathcal{S}$ in order to compute the routing probabilities. The logistic (a.k.a. sigmoid) function is the common choice for $\mathcal{S}$ in the literature on soft trees (see Jordan & Jacobs (1994); Kotschieder et al. (2015); Frosst & Hinton (2017); Tanno et al. (2018); Hehn et al. (2019)). While the logistic function can output arbitrarily small values, it cannot output an exact zero. This implies that any sample x will reach every node in the tree with a positive probability (as evident from (2)). Thus, computing the output of the tree in (3) will require computation over every node in the tree, an operation which is exponential in tree depth.

We propose a novel *smooth-step activation function*, which can output exact zeros and ones, thus allowing for true conditional computation. Our smooth-step function is S-

shaped and continuously differentiable, similar to the logistic function. Let γ be a non-negative scalar parameter. The smooth-step function is a cubic polynomial in the interval $[-\gamma/2, \gamma/2]$, 0 to the left of the interval, and 1 to the right. More formally, we assume that the function takes the parametric form $\mathcal{S}(t) = at^3 + bt^2 + ct + d$ for $t \in [-\gamma/2, \gamma/2]$, where a, b, c, d are scalar parameters. We then solve for the parameters under the following continuity and differentiability constraints: (i) $\mathcal{S}(-\gamma/2) = 0$, (ii) $\mathcal{S}(\gamma/2) = 1$, (iii) $\mathcal{S}'(t)|_{t=-\gamma/2} = \mathcal{S}'(t)|_{t=\gamma/2} = 0$. This leads to:

$$\mathcal{S}(t) = \begin{cases} 0 & \text{if } t \leq -\gamma/2 \\ -\frac{2}{\gamma^3}t^3 + \frac{3}{2\gamma}t + \frac{1}{2} & \text{if } -\gamma/2 \leq t \leq \gamma/2 \\ 1 & \text{if } t \geq \gamma/2 \end{cases} \quad (4)$$

By construction, the smooth-step function in (4) is continuously differentiable for any $t \in \mathbb{R}$ (including $-\gamma/2$ and $\gamma/2$). In Figure 1, we plot the smooth-step (with $\gamma = 1$) and logistic activation functions; the logistic function here takes the form $(1 + e^{-6t})^{-1}$, i.e., it is a rescaled variant of the standard logistic function, so that the two functions are on similar scales. The two functions can be very close in the middle of the fractional region. The main difference is that the smooth-step function outputs exact zero and one, whereas the logistic function converges to these asymptotically.

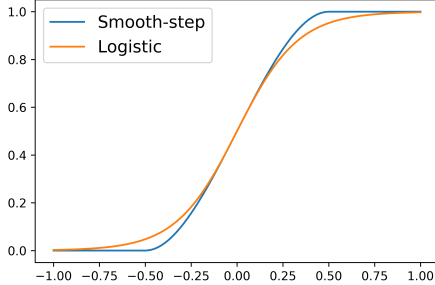


Figure 1: Smooth-step vs. Logistic $(1 + e^{-6t})^{-1}$.

Outside $[-\gamma/2, \gamma/2]$, the smooth-step function performs hard routing, similar to classical decision trees. The choice of γ controls the fraction of samples that are hard routed. A very small γ can lead to many zero gradients in the internal nodes, whereas a very large γ might limit the extent of conditional computation. In our experiments, we use batch normalization (Ioffe & Szegedy, 2015) before the tree layer so that the inputs to the smooth-step function remain centered and bounded. This turns out to be very effective in preventing the internal nodes from having zero gradients, at least in the first few training epochs. Moreover, we view γ as a hyperparameter, which we tune over the range $[10^{-4}, 1]$. This range works well for balancing the training performance and conditional computation across the 26 datasets we used (see Section 4).

For a given sample x , we say that a node i is reachable if $P(x \rightarrow i) > 0$. The number of reachable leaves directly controls the extent of conditional computation. In Figure 2, we plot the average number of reachable leaves (per sample) as a function of the training epochs, for a single tree of depth 10 (i.e., with 1024 leaves) and different γ 's. This is for the diabetes dataset (Olson et al., 2017), using Adam (Kingma & Ba, 2014) for optimization (see the appendix for details). The figure shows that for small enough γ (e.g., $\gamma \leq 1$), the number of reachable leaves rapidly converges to 1 during training (note that the y-axis is on a log scale). We observed this behavior on all the datasets in our experiments.

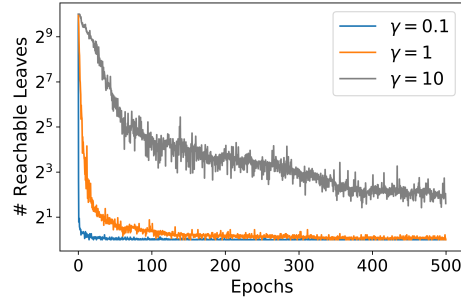


Figure 2: Number of reachable leaves (per sample) during training a tree of depth 10.

We note that variants of the smooth-step function are popular in computer graphics (Ebert et al., 2003; Rost et al., 2009). However, to our knowledge, the smooth-step function has not been used in soft trees or neural networks. It is also worth mentioning that the cubic polynomial used for interpolation in (4) can be substituted with higher-order polynomials (e.g., polynomial of degree 5, where the first and second derivatives vanish at $\pm\gamma/2$). The algorithms we propose in Section 3 directly apply to the case of higher-order polynomials.

In the next section, we show how the sparsity in the smooth-step function and in its gradient can be exploited to develop efficient forward and backward propagation algorithms.

3. Conditional Computation

We propose using first-order optimization methods (e.g., SGD and its variants) to optimize TEL. A main computational bottleneck in this case is the gradient computation, whose time and memory complexities can grow exponentially in the tree depth. This has hindered training large tree ensembles in the literature. In this section, we develop efficient forward and backward propagation algorithms for TEL by exploiting the sparsity in both the smooth-step function and its gradient. We show that our algorithms have optimal time complexity and discuss cases where they run significantly faster than standard backpropagation.

Setup: We assume a general setting where TEL is a hidden layer. Without loss of generality, we consider only one sample and one tree. Let $x \in \mathbb{R}^p$ be the input to TEL and denote the tree output by $T(x) \in \mathbb{R}^k$, where $T(x)$ is defined in (3). We use the same notation as in Section 2, and we collect the leaf vectors $o_l, l \in \mathcal{L}$ into the matrix $O \in \mathbb{R}^{|\mathcal{L}| \times k}$ and the internal node weights $w_i, i \in \mathcal{I}$ into the matrix $W \in \mathbb{R}^{|\mathcal{I}| \times p}$. Moreover, for a differentiable function $h(z)$ which maps $\mathbb{R}^s \rightarrow \mathbb{R}^u$, we denote its Jacobian by $\frac{\partial h}{\partial z} \in \mathbb{R}^{u \times s}$. Let L be the loss function to be optimized (e.g., cross-entropy). Our goal is to efficiently compute the following three gradients: $\frac{\partial L}{\partial O}$, $\frac{\partial L}{\partial W}$, and $\frac{\partial L}{\partial T}$. The first two gradients are needed by the optimizer to update O and W . The third gradient is used to continue the backpropagation in the layers preceding TEL. We assume that a backpropagation algorithm has already computed the gradients associated with the layers after TEL and has computed $\frac{\partial L}{\partial T}$.

Number of Reachable Nodes: To exploit conditional computation effectively, each sample should reach a relatively small number of leaves. This can be enforced by choosing the parameter γ of the smooth-step function to be sufficiently small. When analyzing the complexity of the forward and backward passes below, we will assume that the sample x reaches U leaves and N internal nodes.

3.1. Conditional Forward Pass

Prior to computing the gradients, a forward pass over the tree is required. This entails computing expression (3), which is a sum of probabilities over all the root-to-leaf paths in T . Our algorithm exploits the following observation: if a certain edge on the path to leaf l has a zero probability, then $P(x \rightarrow l) = 0$ so there is no need to continue evaluation along that path. Thus, we traverse the tree starting from the root, and every time a node outputs a 0 probability on one side, we ignore all of its descendants lying on that side. The summation in (3) is then performed only over the leaves reached by the traversal. We present the conditional forward pass in Algorithm 1, where for any internal node i , we denote the left and right children by $left(i)$ and $right(i)$.

Time Complexity: The algorithm visits each reachable node in the tree once. Every reachable internal node requires $\mathcal{O}(p)$ operations to compute $\mathcal{S}(\langle w_i, x \rangle)$, whereas each reachable leaf requires $\mathcal{O}(k)$ operations to update the output variable. Thus, the overall complexity is $\mathcal{O}(Np + Uk)$ (recall that N and U are the number of reachable internal nodes and leaves, respectively). This is in contrast to a dense forward pass³, whose complexity is $\mathcal{O}(2^d p + 2^d k)$ (recall that d is the depth). As long as γ is chosen so that U is sub-exponential⁴ in d , the conditional forward pass has

³By dense forward pass, we mean evaluating the tree without conditional computation (as in a standard forward pass).

⁴A function $f(t)$ is sub-exp. in t if $\lim_{t \rightarrow \infty} \log(f(t))/t = 0$.

Algorithm 1 Conditional Forward Pass

```

1: Input: Sample  $x \in \mathbb{R}^p$  and tree parameters  $W$  and  $O$ .
2: Output:  $T(x)$ 
3: {For any node  $i$ ,  $i.prob$  denotes  $P(x \rightarrow i)$ .}
4: { $to\_traverse$  is a stack for traversing nodes.}
5:  $output \leftarrow 0, to\_traverse \leftarrow \{root\}, root.prob \leftarrow 1$ 
6: while  $to\_traverse$  is not empty do
7:   Remove a node  $i$  from  $to\_traverse$ 
8:   if  $i$  is an internal node then
9:      $left(i).prob = i.prob * \mathcal{S}(\langle w_i, x \rangle)$ 
10:     $right(i).prob = i.prob * (1 - \mathcal{S}(\langle w_i, x \rangle))$ 
11:    if  $\mathcal{S}(\langle w_i, x \rangle) > 0$ , add  $left(i)$  to  $to\_traverse$ 
12:    if  $\mathcal{S}(\langle w_i, x \rangle) < 1$ , add  $right(i)$  to  $to\_traverse$ 
13:   else
14:      $output \leftarrow output + i.prob * o_i$ 
15:   end if
16: end while
    
```

a better complexity than the dense pass (this holds since $N = \mathcal{O}(Ud)$, implying that N is also sub-exponential in d).

Memory Complexity: The memory complexity for inference and training is $\mathcal{O}(d)$ and $\mathcal{O}(d + U)$, respectively. See the appendix for a detailed analysis. This is in contrast to a dense forward pass, whose complexity in training is $\mathcal{O}(2^d)$.

3.2. Conditional Backward Pass

Here we develop a backward pass algorithm to efficiently compute the three gradients: $\frac{\partial L}{\partial O}$, $\frac{\partial L}{\partial W}$, and $\frac{\partial L}{\partial T}$, assuming that $\frac{\partial L}{\partial T}$ is available from a backpropagation algorithm. In what follows, we will see that as long as U is sufficiently small, the gradients $\frac{\partial L}{\partial O}$ and $\frac{\partial L}{\partial W}$ will be sparse, and $\frac{\partial L}{\partial x}$ can be computed by considering only a small number of nodes in the tree. Let $\mathcal{R}$ be the set of leaves reached by Algorithm 1. The following set turns out to be critical in understanding the sparsity structure in the problem: $\mathcal{F} := \{i \in \mathcal{I} \mid i \in \mathcal{A}(l), l \in \mathcal{R}, 0 < \mathcal{S}(\langle x, w_i \rangle) < 1\}$. In words, $\mathcal{F}$ is the set of ancestors of the reachable leaves, whose activation is fractional.

In Theorem 1, we show how the three gradients can be computed by only considering the internal nodes in $\mathcal{F}$ and leaves in $\mathcal{R}$. Moreover, the theorem presents sufficient conditions for which the gradients are zero; in particular, $\frac{\partial L}{\partial w_i} = 0$ for every internal node $i \in \mathcal{F}^c$ and $\frac{\partial L}{\partial o_l} = 0$ for every leaf $l \in \mathcal{R}^c$ (where A^c is the complement of a set A).

Theorem 1. Define $\mu_1(x, i) = \frac{\partial \mathcal{S}(\langle x, w_i \rangle)}{\partial \langle x, w_i \rangle} / \mathcal{S}(\langle x, w_i \rangle)$, $\mu_2(x, i) = \frac{\partial \mathcal{S}(\langle x, w_i \rangle)}{\partial \langle x, w_i \rangle} / (1 - \mathcal{S}(\langle x, w_i \rangle))$, and $g(l) = P(\{x \rightarrow l\}) \langle \frac{\partial L}{\partial T}, o_l \rangle$. The gradients needed for backpropa-

gation can be expressed as follows:

$$\begin{aligned}\frac{\partial L}{\partial x} &= \sum_{i \in \mathcal{F}} w_i^T \left[\mu_1(x, i) \sum_{l \in \mathcal{R} \mid [l \prec i]} g(l) - \mu_2(x, i) \sum_{l \in \mathcal{R} \mid [i \searrow l]} g(l) \right] \\ \frac{\partial L}{\partial w_i} &= \begin{cases} 0 & i \in \mathcal{F}^c \\ x^T \left[\mu_1(x, i) \sum_{l \in \mathcal{R} \mid [l \prec i]} g(l) - \mu_2(x, i) \sum_{l \in \mathcal{R} \mid [i \searrow l]} g(l) \right] & \text{o.w.} \end{cases} \\ \frac{\partial L}{\partial o_l} &= \frac{\partial L}{\partial T} P(\{x \rightarrow l\}), \forall l \in \mathcal{L}\end{aligned}$$

In Theorem 1, the quantities $\mu_1(x, i)$ and $\mu_2(x, i)$ can be obtained in $\mathcal{O}(1)$ since in Algorithm 1 we store $\langle x, w_i \rangle$ for every $i \in \mathcal{F}$. Moreover, $P(\{x \rightarrow l\})$ is stored in Algorithm 1 for every reachable leaf. However, a direct evaluation of these gradients leads to a suboptimal time complexity because the terms $\sum_{l \in \mathcal{R} \mid [l \prec i]} g(l)$ and $\sum_{l \in \mathcal{R} \mid [i \searrow l]} g(l)$ will be computed from scratch for every node $i \in \mathcal{F}$. Our conditional backward pass traverses a *fractional tree*, composed of only the nodes in $\mathcal{F}$ and $\mathcal{R}$, while deploying smart book-keeping to compute these sums during the traversal and avoid recomputation. We define the fractional tree below.

Definition 1. Let $T_{\text{reachable}}$ be the tree traversed by the conditional forward pass (Algorithm 1). We define the fractional tree $T_{\text{fractional}}$ as the result of the following two operations: (i) remove every internal node $i \in \mathcal{F}^c$ from $T_{\text{reachable}}$ and (ii) connect every node with no parent to its closest ancestor.

In Section C.1 of the appendix, we provide an example of how the fractional tree is constructed. $T_{\text{fractional}}$ is a binary tree with U leaves and $|\mathcal{F}|$ internal nodes, each with exactly 2 children. It can be readily seen that $|\mathcal{F}| = U - 1$; this relation is useful for analyzing the complexity of the conditional backward pass. Note that $T_{\text{fractional}}$ can be constructed on-the-fly while performing the conditional forward pass (without affecting its complexity). In Algorithm 2, we present the conditional backward pass, which traverses the fractional tree once and returns $\frac{\partial L}{\partial x}$ and any (potentially) non-zero entries in $\frac{\partial L}{\partial O}$ and $\frac{\partial L}{\partial W}$.

Time Complexity: The worst-case complexity of the algorithm is $\mathcal{O}(Up + Uk)$, whereas the best-case complexity is $\mathcal{O}(k)$ (corresponds to $U = 1$), and in the worst case, the number of non-zero entries in the three gradients is $\mathcal{O}(Up + Uk)$ —see the appendix for analysis. Thus, the complexity is optimal, in the sense that it matches the number of non-zero gradient entries, in the worst case. The worst-case complexity is generally lower than the $\mathcal{O}(Np + Uk)$ complexity of the conditional forward pass. This is because we always have $U = \mathcal{O}(N)$, and there can be many cases where N grows faster than U . For example, consider a tree with only two reachable leaves ($U = 2$) and where the root is the (only) fractional node, then N grows linearly with the depth d . As long as U is sub-exponential in d , Algorithm 2’s

Algorithm 2 Conditional Backward Pass

```

1: Input: Sample  $x \in \mathbb{R}^p$ , tree parameters, and  $\frac{\partial L}{\partial T}$ .
2: Output:  $\frac{\partial L}{\partial x}$  and (potential) non-zeros in  $\frac{\partial L}{\partial W}$  and  $\frac{\partial L}{\partial O}$ .
3:  $\frac{\partial L}{\partial x} = 0$ 
4: {For any node  $i$ ,  $i.\text{sum\_g}$  denotes  $\sum_{l \in \mathcal{R} \mid i \in \mathcal{A}(l)} g(l)$ }
5: Traverse  $T_{\text{fractional}}$  in post order:
6:   Denote the current node by  $i$ 
7:   if  $i$  is a leaf then
8:      $\frac{\partial L}{\partial o_i} = \frac{\partial L}{\partial T} P(\{x \rightarrow i\})$ 
9:      $i.\text{sum\_g} = g(i)$ 
10:  else
11:     $a = \mu_1(x, i)$  ( $\text{left}(i).\text{sum\_g}$ )
12:     $b = \mu_2(x, i)$  ( $\text{right}(i).\text{sum\_g}$ )
13:     $\frac{\partial L}{\partial x} += w_i^T (a - b)$ 
14:     $\frac{\partial L}{\partial w_i} = x^T (a - b)$ 
15:     $i.\text{sum\_g} = \text{left}(i).\text{sum\_g} + \text{right}(i).\text{sum\_g}$ 
16:  end if
    
```

complexity can be significantly lower than that of a dense backward pass whose complexity is $\mathcal{O}(2^d p + 2^d k)$.

Memory Complexity: We store one scalar per node in the fractional tree (i.e., $i.\text{sum_g}$ for every node i in the fractional tree). Thus, the memory complexity is $\mathcal{O}(|\mathcal{F}| + U) = \mathcal{O}(U)$. If γ is chosen so that U is upper-bounded by a constant, then Algorithm 2 will require constant memory.

Connections to Backpropagation: An interesting observation in our approach is that the conditional backward pass generally has a better time complexity than the conditional forward pass. This is usually impossible in standard backpropagation for NNs, as the forward and backward passes traverse the same computational graph (Goodfellow et al., 2016). The improvement in complexity of the backward pass in our case is due to Algorithm 2 operating on the fractional tree, which can contain a significantly smaller number of nodes than the tree traversed by the forward pass. In the language of backpropagation, our fractional tree can be viewed as a “simplified” computational graph, where the simplifications are due to Theorem 1.

4. Experiments

We study the performance of TEL in terms of prediction, conditional computation, and compactness. We evaluate TEL as a standalone learner and as a layer in a NN, and compare to standard soft trees, GBDT, and dense layers.

Model Implementation: TEL is implemented in TensorFlow 2.0 using custom C++ kernels for forward and backward propagation, along with a Keras Python-accessible interface. The implementation is open source².

Datasets: We use a collection of 26 classification datasets

(binary and multiclass) from various domains (e.g., health-care, genetics, and image recognition). 23 of these are from the Penn Machine Learning Benchmarks (PMLB) (Olson et al., 2017), and the 3 remaining are CIFAR-10 (Krizhevsky et al., 2009), MNIST (LeCun et al., 1998), and Fashion MNIST (Xiao et al., 2017). Details are in the appendix.

Tuning, Toolkits, and Details: For all the experiments, we tune the hyperparameters using Hyperopt (Bergstra et al., 2013) with the Tree-structured Parzen Estimator (TPE). We optimize for either AUC or accuracy with stratified 5-fold cross-validation. NNs (including TEL) were trained using Keras with the TensorFlow backend, using Adam (Kingma & Ba, 2014) and cross-entropy loss. As discussed in Section 2, TEL is always preceded by a batch normalization layer. GBDT is from XGBoost (Chen & Guestrin, 2016), Logistic regression and CART are from Scikit-learn (Pedregosa et al., 2011). Additional details are in the appendix.

4.1. Soft Trees: Smooth-step vs. Logistic Activation

We compare the run time and performance of the smooth-step and logistic functions using 23 PMLB datasets.

Predictive Performance: We fix the TEL architecture to 10 trees of depth 4. We tune the learning rate, batch size, and number of epochs (ranges are in the appendix). We assume the following parametric form for the logistic function $f(t) = (1 + e^{-t/\alpha})^{-1}$, where α is a hyperparameter which we tune in the range $[10^{-4}, 10^4]$. The smooth-step’s parameter γ is tuned in the range $[10^{-4}, 1]$. Here we restrict the upper range of γ to 1 to enable conditional computation over the whole tuning range. While γ ’s larger than 1 can lead to slightly better predictive performance in some cases, they can slow down training significantly. For tuning, Hyperopt is run for 50 rounds with AUC as the metric. After tuning, models with the best hyperparameters are retrained. We repeat the training procedure 5 times using random weight initializations. The mean test AUC along with its standard error (SE) are in Table 2. The smooth-step outperforms the logistic function on 7 datasets (5 are statistically significant). The logistic function also wins on 7 datasets (4 are statistically significant). The two functions match on the rest of the datasets. The differences on the majority of the datasets are small (even when statistically significant), suggesting that using the smooth-step function does not hurt the predictive performance. However, as we will see next, the smooth-step has a significant edge in terms of computation time.

Training Time: We measure the training time over 50 epochs as a function of tree depth for both activation functions. We keep the same ensemble size (10) and use $\gamma = 1$ for the smooth-step as this corresponds to the worst-case training time (in the tuning range $[10^{-4}, 1]$), and we fix the optimization hyperparameters (batch size = 256 and learning rate = 0.1). We report the results for three of the datasets in

Table 2: Test AUC for the smooth-step and logistic functions (fixed TEL architecture). A * indicates statistical significance based on a paired two-sided t-test at a significance level of 0.05. Best results are in **bold**. AUCs on the 9 remaining datasets match and are hence omitted.

Dataset	Smooth-step	Logistic
ann-thyroid	0.997 $\pm$ 0.0001	0.996 $\pm$ 0.0006
breast-cancer-w.	0.992 $\pm$ 0.0015	0.994 $\pm$ 0.0002
churn	0.897 $\pm$ 0.0014	0.898 $\pm$ 0.0014
crx	0.916 $\pm$ 0.0025	0.929* $\pm$ 0.0021
diabetes	0.832* $\pm$ 0.0009	0.816 $\pm$ 0.0021
dna	0.993 $\pm$ 0.0004	0.994* $\pm$ 0.0
ecoli	0.97* $\pm$ 0.0004	0.952 $\pm$ 0.0038
flare	0.78 $\pm$ 0.0027	0.784 $\pm$ 0.0018
heart-c	0.936 $\pm$ 0.002	0.927 $\pm$ 0.0036
pima	0.828* $\pm$ 0.0005	0.82 $\pm$ 0.0003
satimage	0.988* $\pm$ 0.0002	0.987 $\pm$ 0.0002
solar-flare_2	0.926 $\pm$ 0.0002	0.927* $\pm$ 0.0007
vehicle	0.956 $\pm$ 0.0015	0.965* $\pm$ 0.0007
yeast	0.876* $\pm$ 0.0014	0.86 $\pm$ 0.0026
# wins	7	7

Figure 3; the results for the other datasets have very similar trends and are omitted due to space constraints. The results indicate a steep exponential increase in training time for the logistic activation after depth 6. In contrast, the smooth-step has a slow growth, achieving over 10x speed-up at depth 10.

4.2. TEL vs. Gradient Boosted Decision Trees

Predictive Performance: We compare the predictive performance of TEL and GBDT on the 23 PMLB datasets, and we include L2-regularized logistic regression (LR) and CART as baselines. For a fair comparison, we use TEL as a standalone layer. For TEL and GBDT, we tune over the # of trees, depth, learning rate, and L2 regularization. For TEL we also tune over the batch size, epochs, and $\gamma \in [10^{-4}, 1]$. For LR and CART, we tune the L2 regularization and depth, respectively. We use 50 tuning rounds in Hyperopt with AUC as the metric. We repeat the tuning/testing procedures on 15 random training/testing splits. The results are in Table 3.

As expected, no algorithm dominates on all the datasets. TEL outperforms GBDT on 9 datasets (5 are statistically significant). GBDT outperforms TEL on 8 datasets (7 of which are statistically significant). There were ties on the 6 remaining datasets; these typically correspond to easy tasks where an AUC of (almost) 1 can be attained. LR outperforms both TEL and GBDT on only 3 datasets with very marginal difference. Overall, the results indicate that TEL’s performance is competitive with GBDT. Moreover, adding feature representation layers before TEL can potentially improve its performance further, e.g., see Section 4.3.

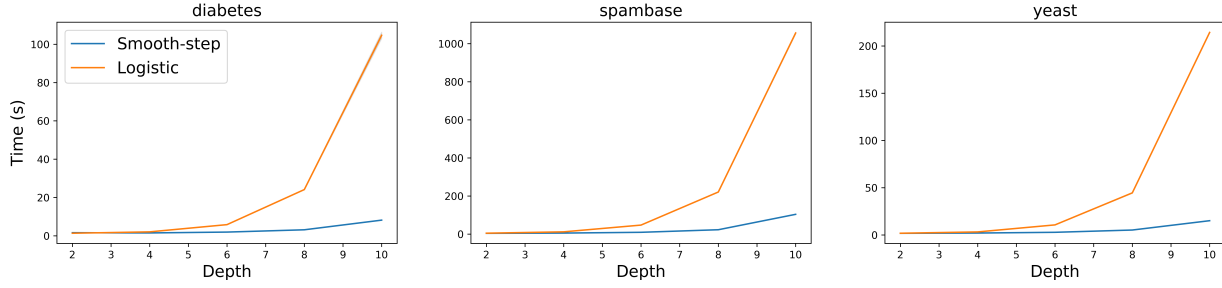


Figure 3: Training time (sec) vs. tree depth for the smooth-step and logistic functions, averaged over 5 repetitions.

 Table 3: Test AUC on 23 PMLB datasets. Averages over 15 random repetitions are reported along with the SE. A star (*) indicates statistical significance based on a paired two-sided t-test at a significance level of 0.05. Best results are in **bold**.

Dataset	TEL	GBDT	L2 Logistic Reg.	CART
ann-thyroid	0.996 $\pm$ 0.0	1.0* $\pm$ 0.0	0.92 $\pm$ 0.002	0.997 $\pm$ 0.0
breast-cancer-wisconsin	0.995* $\pm$ 0.001	0.992 $\pm$ 0.001	0.991 $\pm$ 0.001	0.929 $\pm$ 0.004
car-evaluation	1.0 $\pm$ 0.0	1.0 $\pm$ 0.0	0.985 $\pm$ 0.001	0.981 $\pm$ 0.001
churn	0.916 $\pm$ 0.004	0.92* $\pm$ 0.004	0.814 $\pm$ 0.003	0.885 $\pm$ 0.004
crx	0.911 $\pm$ 0.005	0.933* $\pm$ 0.004	0.916 $\pm$ 0.005	0.905 $\pm$ 0.005
dermatology	0.998 $\pm$ 0.001	0.998 $\pm$ 0.001	0.998 $\pm$ 0.001	0.962 $\pm$ 0.005
diabetes	0.831* $\pm$ 0.006	0.82 $\pm$ 0.006	0.824 $\pm$ 0.008	0.774 $\pm$ 0.008
dna	0.993 $\pm$ 0.0	0.994* $\pm$ 0.0	0.991 $\pm$ 0.0	0.964 $\pm$ 0.001
ecoli	0.97* $\pm$ 0.003	0.962 $\pm$ 0.003	0.972 $\pm$ 0.003	0.902 $\pm$ 0.007
flare	0.732 $\pm$ 0.009	0.738 $\pm$ 0.01	0.736 $\pm$ 0.009	0.717 $\pm$ 0.01
heart-c	0.903 $\pm$ 0.006	0.893 $\pm$ 0.008	0.908 $\pm$ 0.005	0.829 $\pm$ 0.012
hypothyroid	0.971 $\pm$ 0.003	0.987* $\pm$ 0.002	0.93 $\pm$ 0.005	0.926 $\pm$ 0.011
nursery	1.0 $\pm$ 0.0	1.0 $\pm$ 0.0	0.916 $\pm$ 0.001	0.996 $\pm$ 0.0
optdigits	1.0 $\pm$ 0.0	1.0 $\pm$ 0.0	0.998 $\pm$ 0.0	0.958 $\pm$ 0.001
pima	0.831 $\pm$ 0.008	0.825 $\pm$ 0.006	0.832 $\pm$ 0.008	0.758 $\pm$ 0.011
satimage	0.99 $\pm$ 0.0	0.99 $\pm$ 0.0	0.955 $\pm$ 0.001	0.949 $\pm$ 0.001
sleep	0.925 $\pm$ 0.0	0.927* $\pm$ 0.0	0.889 $\pm$ 0.0	0.876 $\pm$ 0.001
solar-flare_2	0.925 $\pm$ 0.002	0.924 $\pm$ 0.002	0.92 $\pm$ 0.002	0.907 $\pm$ 0.002
spambase	0.986 $\pm$ 0.001	0.989* $\pm$ 0.001	0.972 $\pm$ 0.001	0.926 $\pm$ 0.002
texture	1.0 $\pm$ 0.0	1.0 $\pm$ 0.0	1.0 $\pm$ 0.0	0.974 $\pm$ 0.001
twonorm	0.998* $\pm$ 0.0	0.997 $\pm$ 0.0	0.998 $\pm$ 0.0	0.865 $\pm$ 0.002
vehicle	0.953* $\pm$ 0.003	0.931 $\pm$ 0.002	0.941 $\pm$ 0.002	0.871 $\pm$ 0.004
yeast	0.861 $\pm$ 0.004	0.859 $\pm$ 0.004	0.852 $\pm$ 0.004	0.779 $\pm$ 0.005
# wins	12	14	6	0

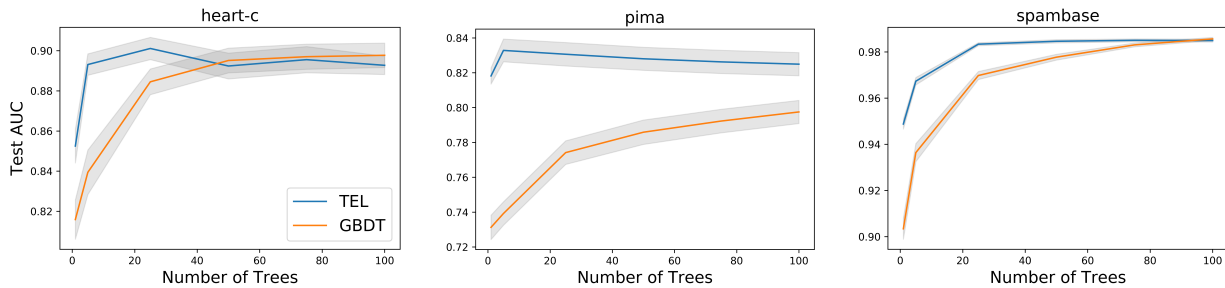


Figure 4: Mean test AUC vs # of trees (15 trials). SE is shaded. TEL and GBDT have (roughly) the same # of params/tree.

Table 4: Average and SE for test accuracy, loss and # of params for CNN-Dense and CNN-TEL over 5 random initializations. A star * indicates statistical significance based on a paired two-sided t-test at a level of 5%. Best values are in **bold**.

Dataset	CNN-Dense			CNN-TEL		
	Accuracy	Loss	# Params	Accuracy	Loss	# Params
CIFAR10	0.7278 ± 0.0047	1.673 ± 0.170	7, 548, 362	0.7296 ± 0.0109	$1.202^* \pm 0.011$	926, 465
MNIST	0.9926 ± 0.0002	0.03620 ± 0.00121	5, 830, 538	$0.9930 \pm 9e-5$	0.03379 ± 0.00093	699, 585
Fashion MNIST	0.9299 ± 0.0012	0.6930 ± 0.0291	5, 567, 882	0.9297 ± 0.0012	$0.3247^* \pm 0.0045$	699, 585

Compactness and Sensitivity: We compare the number of trees and sensitivity of TEL and GBDT on datasets from Table 3 where both models achieve comparable AUCs—namely, the heart-c, pima and spambase datasets. With similar predictive performance, compactness can be an important factor in choosing a model over the other. For TEL, we use the models trained in Table 3. As for GBDT, for each dataset, we fix the depth so that the number of parameters per tree in GBDT (roughly) matches that of TEL. We tune over the main parameters of GBDT (50 iterations of Hyperopt, under the same parameter ranges of Table 3). We plot the test AUC versus the number of trees in Figure 4. On all datasets, the test AUC of TEL peaks at a significantly smaller number of trees compared to GBDT. For example, on pima, TEL’s AUC peaks at 5 trees, whereas GBDT requires more than 100 trees to achieve a comparable performance—this is more than 20x reduction in the number of parameters. Moreover, the performance of TEL is less sensitive w.r.t. to changes in the number of trees. These observations can be attributed to the joint optimization performed in TEL, which can lead to more expressive ensembles compared to the stage-wise optimization in GBDT.

4.3. TEL vs. Dense Layers in CNNs

We study the potential benefits of replacing dense layers with TEL in CNNs, on the CIFAR-10, MNIST, and Fashion MNIST datasets. We consider 2 convolutional layers, followed by intermediate layers (max pooling, dropout, batch normalization), and finally dense layers; we refer to this as CNN-Dense. We also consider a similar architecture, where the final dense layers are replaced with a single dense layer followed by TEL; we refer to this model as CNN-TEL. We tune over the optimization hyperparameters, the number of filters in the convolutional layers, the number and width of the dense layers, and the different parameters of TEL (see appendix for details). We run Hyperopt for 25 iterations with classification accuracy as the target metric. After tuning, the models are trained using 5 random weight initializations.

The classification accuracy and loss on the test set and the total number of parameters are reported in Table 4. While the accuracies are comparable, CNN-TEL achieves a lower test loss on the three datasets, where the 28% and 53%

relative improvements on CIFAR and Fashion MNIST are statistically significant. Since we are using cross-entropy loss, this means that TEL gives higher scores on average, when it makes correct predictions. Moreover, the number of parameters in CNN-TEL is $\sim 8x$ smaller than CNN-Dense. This example also demonstrates how representation layers can be effectively leveraged by TEL—GBDT’s performance is significantly lower on MNIST and CIFAR-10, e.g., see the comparisons in Ponomareva et al. (2017).

5. Conclusion and Future Work

We introduced the tree ensemble layer (TEL) for neural networks. The layer is composed of an additive model of differentiable decision trees that can be trained end-to-end with the neural network, using first-order methods. Unlike differentiable trees in the literature, TEL supports conditional computation, i.e., each sample is routed through a small part of the tree’s architecture. This is achieved by using the smooth-step activation function for routing samples, along with specialized forward and backward passes for reducing the computational complexity. Our experiments indicate that TEL achieves competitive predictive performance compared to gradient boosted decision trees (GBDT) and dense layers, while leading to significantly more compact models. In addition, by effectively leveraging convolutional layers, TEL significantly outperforms GBDT on multiple image classification datasets.

One interesting direction for future work is to equip TEL with mechanisms for exploiting feature sparsity, which can further speed up computation. Promising works in this direction include feature bundling (Ke et al., 2017) and learning under hierarchical sparsity assumptions (Hazimeh & Mazumder, 2020). Moreover, it would be interesting to study whether the smooth-step function, along with specialized optimization methods, can be an effective alternative to the logistic function in other machine learning models.

Acknowledgements

We would like to thank Mehryar Mohri for the useful discussions. Part of this work was done when Hussein Hazimeh was at Google Research. At MIT, Hussein acknowledges research funding from the Office of Naval Research (ONR-N000141812298). Rahul Mazumder acknowledges research funding from the Office of Naval Research (ONR-N000141812298, Young Investigator Award), the National Science Foundation (NSF-IIS-1718258), and IBM.

References

- Bengio, E., Bacon, P., Pineau, J., and Precup, D. Conditional computation in neural networks for faster models. *CoRR*, abs/1511.06297, 2015. URL <http://arxiv.org/abs/1511.06297>.
- Bengio, Y., Courville, A., and Vincent, P. Representation learning: A review and new perspectives. *IEEE transactions on pattern analysis and machine intelligence*, 35(8): 1798–1828, 2013.
- Bergstra, J., Yamins, D., and Cox, D. D. Making a science of model search: Hyperparameter optimization in hundreds of dimensions for vision architectures. In *Proceedings of the 30th International Conference on International Conference on Machine Learning - Volume 28, ICML’13*, pp. I–115–I–123. JMLR.org, 2013.
- Biau, G., Scornet, E., and Welbl, J. Neural random forests. *Sankhya A*, 81(2):347–386, Dec 2019. ISSN 0976-8378. doi: 10.1007/s13171-018-0133-y. URL <https://doi.org/10.1007/s13171-018-0133-y>.
- Breiman, L., Friedman, J. H., Olshen, R. A., and Stone, C. J. Classification and regression trees. 1983.
- Chen, T. and Guestrin, C. Xgboost: A scalable tree boosting system. In *Proceedings of the 22nd acm sigkdd international conference on knowledge discovery and data mining*, pp. 785–794, 2016.
- Ebert, D. S., Musgrave, F. K., Peachey, D., Perlin, K., and Worley, S. *Texturing & modeling: a procedural approach*. Morgan Kaufmann, 2003.
- Friedman, J. H. Greedy function approximation: a gradient boosting machine. *Annals of statistics*, pp. 1189–1232, 2001.
- Frosst, N. and Hinton, G. E. Distilling a neural network into a soft decision tree. In *Proceedings of the First International Workshop on Comprehensibility and Explanation in AI and ML 2017*, CEUR Workshop Proceedings, 2017.
- Goodfellow, I., Bengio, Y., and Courville, A. *Deep Learning*. MIT Press, 2016. <http://www.deeplearningbook.org>.
- Hastie, T., Tibshirani, R., and Friedman, J. *The elements of statistical learning: data mining, inference, and prediction*. Springer Science & Business Media, 2009.
- Hazimeh, H. and Mazumder, R. Learning hierarchical interactions at scale: A convex optimization approach. In *International Conference on Artificial Intelligence and Statistics*, pp. 1833–1843, 2020.
- He, K., Zhang, X., Ren, S., and Sun, J. Delving deep into rectifiers: Surpassing human-level performance on imagenet classification. In *Proceedings of the IEEE international conference on computer vision*, pp. 1026–1034, 2015.
- Hehn, T. M., Kooij, J. F., and Hamprecht, F. A. End-to-end learning of decision trees and forests. *International Journal of Computer Vision*, pp. 1–15, 2019.
- Ioannou, Y., Robertson, D., Zikic, D., Kotschieder, P., Shotton, J., Brown, M., and Criminisi, A. Decision forests, convolutional networks and the models in-between, 2016.
- Ioffe, S. and Szegedy, C. Batch normalization: Accelerating deep network training by reducing internal covariate shift. In *Proceedings of the 32nd International Conference on International Conference on Machine Learning - Volume 37, ICML’15*, pp. 448–456. JMLR.org, 2015.
- Jernite, Y., Choromanska, A., and Sontag, D. Simultaneous learning of trees and representations for extreme classification and density estimation. In Precup, D. and Teh, Y. W. (eds.), *Proceedings of the 34th International Conference on Machine Learning*, volume 70 of *Proceedings of Machine Learning Research*, pp. 1665–1674, International Convention Centre, Sydney, Australia, 06–11 Aug 2017. PMLR. URL <http://proceedings.mlr.press/v70/jernite17a.html>.
- Johnson, R. and Zhang, T. Learning nonlinear functions using regularized greedy forest. *IEEE transactions on pattern analysis and machine intelligence*, 36(5):942–954, 2013.
- Jordan, M. I. and Jacobs, R. A. Hierarchical mixtures of experts and the em algorithm. *Neural computation*, 6(2): 181–214, 1994.
- Ke, G., Meng, Q., Finley, T., Wang, T., Chen, W., Ma, W., Ye, Q., and Liu, T.-Y. Lightgbm: A highly efficient gradient boosting decision tree. In *Advances in neural information processing systems*, pp. 3146–3154, 2017.
- Kingma, D. P. and Ba, J. Adam: A method for stochastic optimization. *arXiv preprint arXiv:1412.6980*, 2014.
- Kotschieder, P., Fiterau, M., Criminisi, A., and Bulò, S. R. Deep neural decision forests. In *2015 IEEE International*

- Conference on Computer Vision (ICCV)*, pp. 1467–1475, Dec 2015. doi: 10.1109/ICCV.2015.172.
- Krizhevsky, A., Hinton, G., et al. Learning multiple layers of features from tiny images. 2009.
- LeCun, Y., Bottou, L., Bengio, Y., and Haffner, P. Gradient-based learning applied to document recognition. *Proceedings of the IEEE*, 86(11):2278–2324, 1998.
- Murthy, S. K., Kasif, S., and Salzberg, S. A system for induction of oblique decision trees. *J. Artif. Int. Res.*, 2(1):1–32, August 1994. ISSN 1076-9757.
- Olson, R. S., La Cava, W., Orzechowski, P., Urbanowicz, R. J., and Moore, J. H. Pmlb: a large benchmark suite for machine learning evaluation and comparison. *BioData Mining*, 10(1):36, Dec 2017. ISSN 1756-0381. doi: 10.1186/s13040-017-0154-4. URL <https://doi.org/10.1186/s13040-017-0154-4>.
- Pedregosa, F., Varoquaux, G., Gramfort, A., Michel, V., Thirion, B., Grisel, O., Blondel, M., Prettenhofer, P., Weiss, R., Dubourg, V., Vanderplas, J., Passos, A., Cournapeau, D., Brucher, M., Perrot, M., and Duchesnay, E. Scikit-learn: Machine learning in Python. *Journal of Machine Learning Research*, 12:2825–2830, 2011.
- Ponomareva, N., Colthurst, T., Hendry, G., Haykal, S., and Radpour, S. Compact multi-class boosted trees. In *2017 IEEE International Conference on Big Data (Big Data)*, pp. 47–56. IEEE, 2017.
- Rost, R. J., Licea-Kane, B., Ginsburg, D., Kessenich, J., Lichtenbelt, B., Malan, H., and Weiblen, M. *OpenGL shading language*. Pearson Education, 2009.
- Shazeer, N., Mirhoseini, A., Maziarz, K., Davis, A., Le, Q. V., Hinton, G. E., and Dean, J. Outrageously large neural networks: The sparsely-gated mixture-of-experts layer. *CoRR*, abs/1701.06538, 2017. URL <http://arxiv.org/abs/1701.06538>.
- Tanno, R., Arulkumaran, K., Alexander, D. C., Criminisi, A., and Nori, A. Adaptive neural trees. *arXiv preprint arXiv:1807.06699*, 2018.
- Xiao, H., Rasul, K., and Vollgraf, R. Fashion-mnist: a novel image dataset for benchmarking machine learning algorithms. *arXiv preprint arXiv:1708.07747*, 2017.
- Yu, D. and Deng, L. *Automatic Speech Recognition*. Springer, 2016.
- Zoran, D., Lakshminarayanan, B., and Blundell, C. Learning deep nearest neighbor representations using differentiable boundary trees. *CoRR*, abs/1702.08833, 2017. URL <http://arxiv.org/abs/1702.08833>.