

# An Exact Kernel Equivalence for Finite Classification Models

Brian Bell<sup>\*12</sup> Michael Geyer<sup>\*13</sup> David Glickenstein<sup>2</sup> Amanda Fernandez<sup>3</sup> Juston Moore<sup>1</sup>

## Abstract

We explore the equivalence between neural networks and kernel methods by deriving the first exact representation of any finite-size parametric classification model trained with gradient descent as a kernel machine. We compare our exact representation to the well-known Neural Tangent Kernel (NTK) and discuss approximation error relative to the NTK and other non-exact path kernel formulations. We experimentally demonstrate that the kernel can be computed for realistic networks up to machine precision. We use this exact kernel to show that our theoretical contribution can provide useful insights into the predictions made by neural networks, particularly the way in which they generalize.

## 1. Introduction

This study investigates the relationship between kernel methods and finite parametric models. To date, interpreting the predictions of complex models, like neural networks, has proven to be challenging. Prior work has shown that the inference-time predictions of a neural network can be exactly written as a sum of independent predictions computed with respect to each training point. We formally show that classification models trained with cross-entropy loss can be exactly formulated as a kernel machine. It is our hope that these new theoretical results will open new research directions in the interpretation of neural network behavior.

There has recently been a surge of interest in the connection between neural networks and kernel methods (Bietti & Mairal, 2019; Du et al., 2019; Tancik et al., 2020; Abdar et al., 2021; Geifman et al., 2020; Chen et al., 2020; Alemohammad et al., 2021). Much of this work has been motivated by the the neural tangent kernel (NTK), which

<sup>\*</sup>Equal contribution <sup>1</sup>Los Alamos National Lab <sup>2</sup>University of Arizona <sup>3</sup>University of Texas San Antonio. Correspondence to: Brian Bell <bwbell@math.arizona.edu>, Michael Geyer <mgeyer@lanl.gov>.

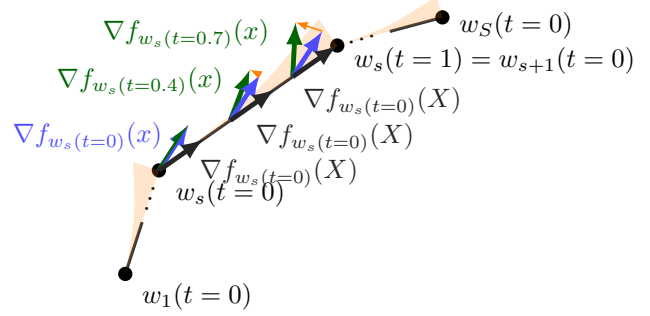


Figure 1: Comparison of test gradients used by Discrete Path Kernel (DPK) from prior work (Blue) and the Exact Path Kernel (EPK) proposed in this work (green) versus total training vectors (black) used for both kernel formulations along a discrete training path with  $S$  steps. Orange shading indicates cosine error of DPK test gradients versus EPK test gradients shown in practice in Fig. 2.

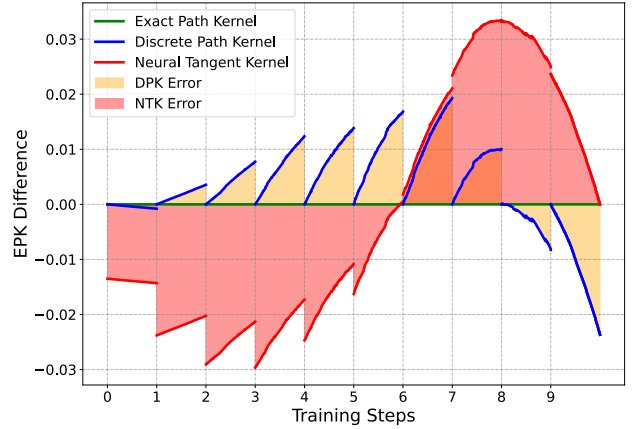


Figure 2: Measurement of gradient alignment on test points across the training path. The EPK is used as a frame of reference. The y-axis is exactly the difference between the EPK and other representations. For example  $EPK - DPK = \langle \phi_{s,t}(X), \phi_{s,t}(x) - \phi_{s,0}(x) \rangle$  (See Definition 3.4). Shaded regions indicate total accumulated error. Note: this is measuring an angle of error in weight space; therefore, equivalent positive and negative error will not result in zero error.

describes the training dynamics of neural networks in the infinite limit of network width (Jacot et al., 2018). We argue that many intriguing behaviors arise in the *finite* parameter regime (Bubeck & Sellke, 2021). All prior works, to the best of our knowledge, appeal to discrete approximations of the kernel corresponding to a neural network. Specifically, prior approaches are derived under the assumption that training step size is small enough to guarantee close approximation of a gradient flow (Ghojogh et al., 2021; Shawe-Taylor et al., 2004; Zhao & Grishman, 2005).

In this work, we show that the simplifying assumptions used in prior works (i.e. infinite network width and infinitesimal gradient descent steps) are not necessary. Our **Exact Path Kernel (EPK)** provides the first, exact method to study the behavior of finite-sized neural networks used for classification. Previous results are limited in application (Incudini et al., 2022) due to dependence of the kernel on test data unless strong conditions are imposed on the training process as by (Chen et al., 2021). We show, however, that the training step sizes used in practice do not closely follow this gradient flow, introducing significant error into all prior approaches (Figure 2).

Our experimental results build on prior studies attempting to evaluate empirical properties of the kernels corresponding to finite neural networks (Lee et al., 2018; Chen et al., 2021). While the properties of infinite neural networks are fairly well understood (Neal & Neal, 1996), we find that the kernels learned by finite neural networks have non-intuitive properties that may explain the failures of modern neural networks on important tasks such as robust classification and calibration on out-of-distribution data.

This paper makes the following significant theoretical and experimental contributions:

1. We prove that finite-sized neural networks trained with finite-sized gradient descent steps and cross-entropy loss can be exactly represented as kernel machines using the EPK. Our derivation incorporates a previously-proposed path kernel, but extends this method to account for practical training procedures (Domingos, 2020b; Chen et al., 2021).
2. We demonstrate that it is computationally tractable to estimate the kernel underlying a neural network classifier, including for small convolutional computer vision models.
3. We compute Gram matrices using the EPK and use them to illuminate prior theory of neural networks and their understanding of uncertainty.
4. We employ Gaussian processes to compute the covariance of a neural network’s logits and show that this

reiterates previously observed shortcomings of neural network generalization.

## 2. Related Work

Fundamentally, the neural tangent kernel (NTK) is rooted in the concept that all information necessary to represent a parametric model is stored in the Hilbert space occupied by the model’s weight gradients up to a constant factor. This is very well supported in infinite width (Jacot et al., 2018). In this setting, it has been shown that neural networks are equivalent to support vector machines, drawing a connection to maximum margin classifiers (Chen et al., 2021; Chizat & Bach, 2020). Similarly, Shah et al. demonstrate that this maximum margin classifier exists in Wasserstien space; however, they also show that model gradients may not contain the required information to represent this (Shah et al., 2021).

The correspondence between kernel machines and parametric models trained by gradient descent has been previously developed in the case of a continuous training path (i.e. the limit as gradient descent step size  $\varepsilon \rightarrow 0$ ) (Domingos, 2020a). We will refer to the previous finite approximation of this kernel as the Discrete Path Kernel (DPK). However, a limitation of this formulation is its reliance on a continuous integration over a gradient flow, which differs from the discrete forward Euler steps employed in real-world model training. This discrepancy raises concerns regarding the applicability of the continuous path kernel to practical scenarios (Incudini et al., 2022). Moreover, the formulation of the sample weights and bias term in the DPK depends on its test points. Chen et al. propose that this can be addressed, in part, by imposing restrictions on the loss function used for training, but did not entirely disentangle the kernel formulation from sample importance weights on training points (Chen et al., 2021).

We address the limitations of Domingos (2020a) and Chen et al. (2021) in Subsection 3.5. By default, their approach produces a system which can be viewed as an ensemble of kernel machines, but without a single aggregated kernel which can be analyzed directly. Chen et al. (2021) propose that the resulting sum over kernel machines can be formulated as a kernel machine so long as the sign of the gradient of the loss stays constant through training; however, we show that this is not necessarily a sufficient restriction. Instead, their formulation leads to one of several non-symmetric functions which can serve as a surrogate to replicate a given models behavior, but without retaining properties of a kernel.

### 3. Theoretical Results

Our goal is to show an equivalence between any given finite parametric model trained with gradient descent  $f_w(x)$  (e.g. neural networks) and a kernel based prediction that we construct. We define this equivalence in terms of the output of the parametric model  $f_w(x)$  and our kernel method in the sense that they form identical maps from input to output. In the specific case of neural network classification models, we consider the mapping  $f_w(x)$  to include all layers of the neural network up to and including the log-softmax activation function. Formally:

**Definition 3.1.** A kernel is a function of two variables which is symmetric and positive semi-definite.

**Definition 3.2.** Given a Hilbert space  $X$ , a test point  $x \in X$ , and a training set  $X_T = \{x_1, x_2, \dots, x_n\} \subset X$  indexed by  $I$ , a *Kernel Machine* is a model characterized by

$$K(x) = b + \sum_{i \in I} a_i k(x, x_i) \quad (1)$$

where the  $a_i \in \mathbb{R}$  do not depend on  $x$ ,  $b \in \mathbb{R}$  is a constant, and  $k$  is a kernel. (Rasmussen et al., 2006)

By Mercer’s theorem (Ghojogh et al., 2021) a kernel can be produced by composing an inner product on a Hilbert space with a mapping  $\phi$  from the space of data into the chosen Hilbert space. We use this property to construct a kernel machine of the following form.

$$K(x) = b + \sum_{i \in I} a_i \langle \phi(x), \phi(x_i) \rangle \quad (2)$$

Where  $\phi$  is a function mapping input data into the weight space via gradients. Our  $\phi$  will additionally differentiate between test and training points to resolve a discontinuity that arises under discrete training.

#### 3.1. Exact Path Kernels

We first derive a kernel which is an exact representation of the change in model output over one training step, and then compose our final representation by summing along the finitely many steps. Models trained by gradient descent can be characterized by a discrete set of intermediate states in the space of their parameters. These discrete states are often considered to be an estimation of the gradient flow, however in practical settings where  $\epsilon \not\rightarrow 0$  these discrete states differ from the true gradient flow. Our primary theoretical contribution is an algorithm which accounts for this difference by observing the true path the model followed during training. Here we consider the training dynamics of practical gradient descent steps by integrating a discrete path for weights whose states differ from the gradient flow induced by the training set.

**Gradient Along Training Path vs Gradient Field:** In order to compute the EPK, gradients on training data must serve two purposes. First, they are the reference points for comparison (via inner product) with test points. Second, they determine the path of the model in weight space. In practice, the path followed during gradient descent does not match the gradient field exactly. Instead, the gradient used to move the state of the model forward during training is only computed for finitely many discrete weight states of the model. In order to produce a path kernel, we must *continuously* compare the model’s gradient at test points with *fixed* training gradients along each discrete training step  $s$  whose weights we interpolate linearly by  $w_s(t) = w_s - t(w_s - w_{s+1})$ . We will do this by integrating across the gradient field induced by test points, but holding each training gradient fixed along the entire discrete step taken. This creates an asymmetry, where test gradients are being measured continuously but the training gradients are being measured discretely (see Figure 1).

To account for this asymmetry in representation, we will redefine our data using an indicator to separate training points from all other points in the input space.

**Definition 3.3.** Let  $X$  be two copies of a Hilbert space  $H$  with indices 0 and 1 so that  $X = H \times \{0, 1\}$ . We will write  $x \in H \times \{0, 1\}$  so that  $x = (x_H, x_I)$  (For brevity, we will omit writing  $H$  and assume each of the following functions defined on  $H$  will use  $x_H$  and  $x_I$  will be a hidden indicator). Let  $f_w$  be a differentiable function on  $H$  parameterized by  $w \in \mathbb{R}^d$ . Let  $X_T = \{(x_i, 1)\}_{i=1}^M$  be a finite subset of  $X$  of size  $M$  with corresponding observations  $Y_T = \{y_{x_i}\}_{i=1}^M$  with initial parameters  $w_0$  so that there is a constant  $b \in \mathbb{R}$  such that for all  $x$ ,  $f_{w_0}(x) = b$ . Let  $L$  be a differentiable loss function of two values which maps  $(f(x), y_x)$  into the positive real numbers. Starting with  $f_{w_0}$ , let  $\{w_s\}$  be the sequence of points attained by  $N$  forward Euler steps of fixed size  $\epsilon$  so that  $w_{s+1} = w_s - \epsilon \nabla L(f(X_T), Y_T)$ . Let  $x \in H \times \{0\}$  be arbitrary and within the domain of  $f_w$  for every  $w$ . Then  $f_{w_s(t)}$  is a *finite parametric gradient model (FPGM)*.

**Definition 3.4.** Let  $f_{w_s(t)}$  be an FPGM with all corresponding assumptions. Then, for a given training step  $s$ , the *exact path kernel (EPK)* can be written

$$K_{\text{EPK}}(x, x', s) = \int_0^1 \langle \phi_{s,t}(x), \phi_{s,t}(x') \rangle dt \quad (3)$$

where

$$\phi_{s,t}(x) = \nabla_w f_{w_s(t),x}(x) \quad (4)$$

$$w_s(t) = w_s - t(w_s - w_{s+1}) \quad (5)$$

$$w_s(t, x) = \begin{cases} w_s(0), & \text{if } x_I = 1 \\ w_s(t), & \text{if } x_I = 0 \end{cases} \quad (6)$$

**Note:**  $\phi$  is deciding whether to select a continuously or discrete gradient based on whether the data is from the training or testing copy of the Hilbert space  $H$ . This is due to the inherent asymmetry that is apparent from the derivation of this kernel (see Appendix section A.2). This choice avoids potential discontinuity in the kernel output when a test set happens to contain training points.

**Lemma 3.5.** *The exact path kernel (EPK) is a kernel.*

**Theorem 3.6** (Exact Kernel Ensemble Representation). *A model  $f_{w_N}$  trained using discrete steps matching the conditions of the exact path kernel has the following exact representation as an ensemble of  $N$  kernel machines:*

$$f_{w_N} = KE(x) := \sum_{s=1}^N \sum_{i=1}^M a_{i,s} K_{EPK}(x, x', s) + b \quad (7)$$

where

$$a_{i,s} = -\varepsilon \frac{dL(f_{w_s(0)}(x_i), y_i)}{df_{w_s(0)}(x_i)} \quad (8)$$

$$b = f_{w_0}(x) \quad (9)$$

*Proof Sketch.* Assuming the theorem hypothesis, we'll measure the change in model output as we interpolate across each training step  $s$  by measuring the change in model state along a linear parametrization  $w_s(t) = w_s - t(w_s - w_{s+1})$ . We will let  $d$  denote the number of parameters of  $f_w$ . For brevity, we define  $L(x_i, y_i) = l(f_{w_s(0)}(x_i), y_i)$  where  $l$  is the loss function used to train the model.

$$\begin{aligned} \frac{d\hat{y}}{dt} &= \sum_{j=1}^d \frac{d\hat{y}}{dw_j} \frac{dw_j}{dt} \\ &= \sum_{j=1}^d \frac{df_{w_s(t)}(x)}{dw_j} \left( -\varepsilon \sum_{i=1}^M \frac{\partial L(x_i, y_i)}{\partial f_{w_s(0)}(x_i)} \frac{\partial f_{w_s(0)}(x_i)}{\partial w_j} \right) \end{aligned} \quad (10)$$

$$(11)$$

We use fundamental theorem of calculus to integrate this equation from step  $s$  to step  $s+1$  and then add up across all steps. See Appendix A.2 for the full proof.  $\square$

**Remark 1** Note that in this formulation,  $b$  depends on the test point  $x$ . In order to ensure information is not being leaked from the kernel into this bias term the model  $f$  must have constant output for all input. When relaxing this property, to allow for models that have a non-constant starting output, but still requiring  $b$  to remain constant, we note that this representation ceases to be exact for all  $x$ . The resulting approximate representation has logit error bounded by its initial bias which can be chosen as  $b = \text{mean}(f_{w_0(0)}(X_T))$ . Starting bias can be minimized by starting with small parameter values which will be out-weighted by contributions

from training. In practice, we sidestep this issue by initializing all weights in the final layer to 0, resulting in  $b = \log(\text{softmax}(0))$ , thus removing  $b$ 's dependence on  $x$ .

**Remark 2** The exactness of this proof hinges on the *separate* measurement of how the model's parameters change. The gradients on training data, which are fixed from one step to the next, measure how the parameters are changing. This is opposed to the gradients on test data, which are *not* fixed and vary with time. These measure a continuous gradient field for a given point. We are using interpolation as a way to measure the difference between the step-wise linear training path and the continuous loss gradient field.

**Theorem 3.7** (Exact Kernel Machine Reduction). *Let  $\nabla L(f(w_s(x), y)$  be constant across steps  $s$ ,  $(a_{i,s}) = (a_{i,0})$ . Let the kernel across all  $N$  steps be defined as  $K_{NEPK}(x, x') = \sum_{s=1}^N a_{i,0} K_{EPK}(x, x', s)$  Then the exact kernel ensemble representation for  $f_{w_N}$  can be reduced exactly to the kernel machine representation:*

$$f_{w_N}(x) = KM(x) := b + \sum_{i=1}^M a_{i,0} K_{NEPK}(x, x') \quad (12)$$

See Appendix A.3 for full proof. By combining theorems 3.6 and 3.7, we can construct an exact kernel machine representation for any arbitrary parameterized model trained by gradient descent which satisfies the additional property of having constant loss across training steps (e.g. any ANN using categorical cross-entropy loss (CCE) for classification). This representation will produce exactly identical output to the model across the model's entire domain. This establishes exact kernel-neural equivalence for classification ANNs. Furthermore, Theorem 3.6 establishes an exact kernel ensemble representation without limitation to models using loss functions with constant derivatives across steps. It remains an open problem to determine other conditions under which this ensemble may be reduced to a single kernel representation.

### 3.2. Discussion

$\phi_{s,t}(x)$  depends on both  $s$  and  $t$ , which is non-standard but valid, however an important consequence of this mapping is that the output of this representation is not guaranteed to be continuous. This discontinuity is exactly measuring the error between the model along the exact path compared with the gradient flow for each step.

We can write another function  $k'$  which is continuous but not symmetric, yet still produces an exact representation:

$$k'(x, x') = \langle \nabla_w f_{w_s(t)}(x), \nabla_w f_{w_s(0)}(x') \rangle \quad (13)$$

The resulting function is a valid kernel if and only if for

**Algorithm 1** Exact Path Kernel: Given a training set  $(X, Y)$  with  $M$  data points, a testing point  $x$  and  $N$  weight states  $\{w_0, w_1 \dots w_N\}$ , the kernel machine corresponding to the exact path kernel can be calculated for a model with  $W$  weights and  $K$  outputs. We estimate the integral across test points by calculating the Riemann sum with sufficient steps ( $T$ ) to achieve machine precision. For loss functions that do not have constant gradient values throughout training, this algorithm produces an ensemble of kernel machines.

---

```

b =  $f(w_0, x)$ 
for  $s = 0$  to  $N$  do
     $J^X = \nabla_w f_{w_s(0)}(X)$                                 {Jacobian of training point outputs w.r.t model weights  $[M \times K \times W]$  }
    for  $t$  from  $0$  to  $1$  with step  $1/T$  do
         $w_s(t) = w_s + t(w_{s+1} - w_s)$ 
         $J^x += \frac{1}{T} \nabla_w f_{w_s(t)}(x)$     {Jacobian of testing point output w.r.t model weights averaged across  $T$  steps  $[K \times W]$  }
    end for
     $G_{ijk} = \sum_w J_{ijw}^X J_{kw}^x$                                 {Inner product on the weight space, this is the kernel value  $[M \times K \times K]$  }
     $L' = \nabla_f L(f_{w_s(0)}(X), Y)$                             {Jacobian of loss w.r.t model output of training points  $[M \times K]$  }
     $P_{ik}^s = \sum_j L'_{ij} G_{ijk}$                                 {Inner product of kernel value scaled by loss gradients  $[M \times K]$  }
end for
 $\mathcal{P}_{sik} = \{P^0, P^1, \dots, P^N\}$                             {Stack values across all training steps  $[N \times M \times K]$  }
 $\hat{p} = -\varepsilon \frac{1}{M} \sum_s \sum_i \mathcal{P}_{sik} + b$     {Sum across training steps and average across training points for final prediction  $[K]$  }
    
```

---

every  $s$  and every  $x$ ,

$$\int_0^1 \nabla_w f_{w_s(t)}(x) dt = \nabla_w f_{w_s(0)}(x) \quad (14)$$

We note that since  $f$  is being trained using forward Euler, we can write:

$$\frac{\partial w_s(t)}{\partial t} = -\varepsilon \nabla_w L(f_{w_s(0)}(x_i), y_i) \quad (15)$$

In other words, our parameterization of this step depends on the step size  $\varepsilon$  and as  $\varepsilon \rightarrow 0$ , we have

$$\int_0^1 \nabla_w f_{w_s(t)}(x) dt \approx \nabla_w f_{w_s(0)}(x) \quad (16)$$

In particular, given a model  $f$  that admits a Lipschitz constant  $K$  this approximation has error bounded by  $\varepsilon K$  and a proof of this convergence is direct. This demonstrates that the asymmetry of this function is exactly measuring the disagreement between the discrete steps taken during training with the gradient field. This function is one of several subjects for further study, particularly in the context of Gaussian processes whereby the asymmetric Gram matrix corresponding with this function can stand in for a covariance matrix. It may be that the not-symmetric analogue of the covariance in this case has physical meaning relative to uncertainty.

### 3.3. Independence from Optimization Scheme

We can see that by changing equation 15 we can produce an exact representation for any first order discrete optimization

scheme that can be written in terms of model gradients aggregated across subsets of training data. This could include backward Euler, leapfrog, and any variation of adaptive step sizes. This includes stochastic gradient descent, and other forms of subsampling (for which the training sums need only be taken over each sample). One caveat is adversarial training, whereby the  $a_i$  are now sampling a measure over the continuum of adversarial images. We can write this exactly, however computation will require approximation across the measure. Modification of this kernel for higher order optimization schemes remains an open problem.

### 3.4. Ensemble Reduction

In order to reduce the ensemble representation of Equation (7) to the kernel representation of Equation (12), we require that the sum over steps still retain the properties of the kernel (symmetry and positive semi-definiteness). In particular we require that for every subset of the training data  $x_i$  and arbitrary  $\alpha_i$  and  $\alpha_j$ , we have

$$\sum_{i=1}^n \sum_{j=1}^n \sum_{l=1}^M \sum_{s=1}^N \alpha_i \alpha_j a_{l,s} \int_0^1 K_{\text{EPK}}(x_i, x_j) dt \geq 0 \quad (17)$$

A sufficient condition for this reduction is that the gradient of the loss function does not change throughout training. This is the case for categorical cross-entropy where labels are in  $\{0, 1\}$ . In fact, in this specific context the gradient of the loss function does not depend on  $f(x)$ , and are fully determined by the ground truth label, making the gradient of the cross-entropy loss a constant value throughout training (See Appendix section A.3). Showing the positive-definiteness of more general loss functions (e.g.

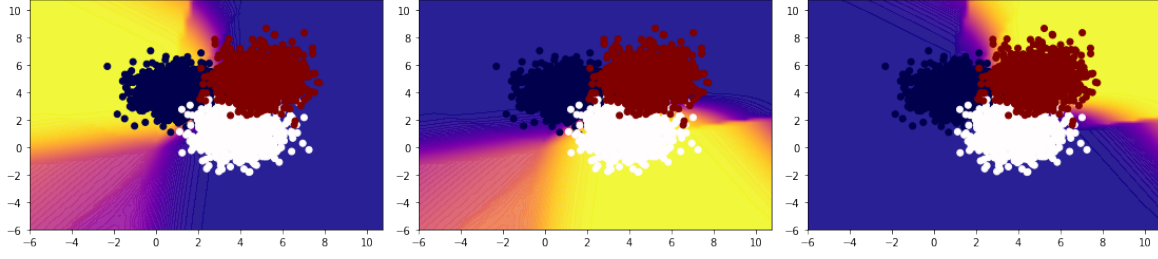


Figure 3: Updated predictions with kernel  $a_i$  updated via gradient descent with training data overlaid for classes 1 (left), 2 (middle), and 3 (right). The high prediction confidence in regions far from training points demonstrates that the learned kernel is non-stationary.

mean squared error loss) will likely require additional regularity conditions on the training path, and is left as future work.

### 3.5. Prior Work

Constant sign loss functions have been previously studied by Chen et al. (Chen et al., 2021), however the kernel that they derive for a finite-width case is of the form

$$K(x, x_i) = \int_0^T |\nabla_f L(f_t(x_i), y_i)| \langle \nabla_w f_t(x), \nabla_w f_t(x_i) \rangle dt \quad (18)$$

The summation across these terms satisfies the positive semi-definite requirement of a kernel, however the weight  $|\nabla L(f_t(x_i), y_i)|$  depends on  $x_i$  which is one of the two inputs. This makes the resulting function  $K(x, x_i)$  asymmetric and therefore not a kernel.

### 3.6. Uniqueness

Uniqueness of this kernel is not guaranteed. The mapping from paths in gradient space to kernels is in fact a function, meaning that each finite continuous path has a unique exact kernel representation of the form described above. However, this function is not necessarily onto the set of all possible kernels. This is evident from the existence of kernels for which representation by a finite parametric function is impossible. Nor is this function necessarily one-to-one since there is a continuous manifold of equivalent parameter configurations for neural networks. For a given training path, we can pick another path of equivalent configurations whose gradients will be separated by some constant  $\delta > 0$ . The resulting kernel evaluation along this alternate path will be exactly equivalent to the first, despite being a unique path. We also note that the linear path  $l_2$  interpolation is not the only valid path between two discrete points in weight space. Following the changes in model weights along a path defined by Manhattan Distance is equally valid and will produce a kernel machine with equivalent outputs. It remains an open problem to compute paths from two different

starting points which both satisfy the constant bias condition from Definition (3.4) which both converge to the same final parameter configuration and define different kernels.

## 4. Experimental Results

Our first experiments test the kernel formulation on a dataset which can be visualized in 2d. These experiments serve as a sanity check and provide an interpretable representation of what the kernel is learning.

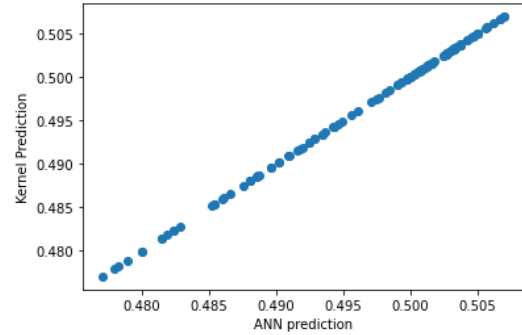


Figure 4: Class 1 EPK Kernel Prediction (Y) versus neural network prediction (X) for 100 test points, demonstrating extremely close agreement.

### 4.1. Evaluating The Kernel

A small test data set within 100 dimensions is created by generating 1000 random samples with means  $(1, 4, 0, \dots)$ ,  $(4, 1, 0, \dots)$  and  $(5, 5, 0, \dots)$  and standard deviation 1.0. These points are labeled according to the mean of the Gaussian used to generate them, providing 1000 points each from 3 classes. A fully connected ReLU network with 1 hidden layer is trained using categorical cross-entropy (CCE) and gradient descent with gradients aggregated across the entire training set for each step. We then compute the EPK for this network, approximating the integral from Equation 3 with 100 steps which replicates the output from the ReLU

network within machine precision. The EPK (Kernel) outputs are compared with neural network predictions in Fig. 4 for class 1. Having established this kernel, and its corresponding kernel machine, one natural extension is to allow the kernel weights  $a_i$  to be retrained. We perform this updating of the kernel weights using a SVM and present its predictions for each of three classes in Fig. 3.

#### 4.2. Kernel Analysis

Having established the efficacy of this kernel for model representation, the next step is to analyze this kernel to understand how it may inform us about the properties of the corresponding model. In practice, it becomes immediately apparent that this kernel lacks typical properties preferred when humans select kernels. Fig. 3 show that the weights of this kernel are non-stationary on our toy problem, with very stable model predictions far away from training data. Next, we use this kernel to estimate uncertainty. Consistent with many other research works on Gaussian processes for classification (Rasmussen et al., 2006) we use a GP to regress to logits. We then use Monte-Carlo to estimate posteriors with respect to probabilities (post-soft-max) for each prediction across a grid spanning the training points of our toy problem. The result is shown on the right-hand column of Fig. 5. We can see that the kernel values are more confident (lower standard deviation) and more stable (higher kernel values) the farther they get from the training data in most directions.

In order to further understand how these strange kernel properties come about, we exercise another advantage of a kernel by analyzing the points that are contributing to the kernel value for a variety of test points. In Fig. 6 we examine the kernel values for each of the training points during evaluation of three points chosen as the mean of the generating distribution for each class. The most striking property of these kernel point values is the fact that they are not proportional to the euclidean distance from the test point. This appears to indicate a set of basis vectors relative to each test point learned by the model based on the training data which are used to spatially transform the data in preparation for classification. This may relate to the correspondence between neural networks and maximum margin classifiers discussed in related work (Chizat & Bach, 2020) (Shah et al., 2021)). Another more subtle property is that some individual data points, mostly close to decision boundaries are slightly over-weighted compared to the other points in their class. This latter property points to the fact that during the latter period of training, once the network has already achieved high accuracy, only the few points which continue to receive incorrect predictions, i.e. caught on the wrong side of a decision boundary, will continue contributing to the training gradient and therefore to the kernel value.

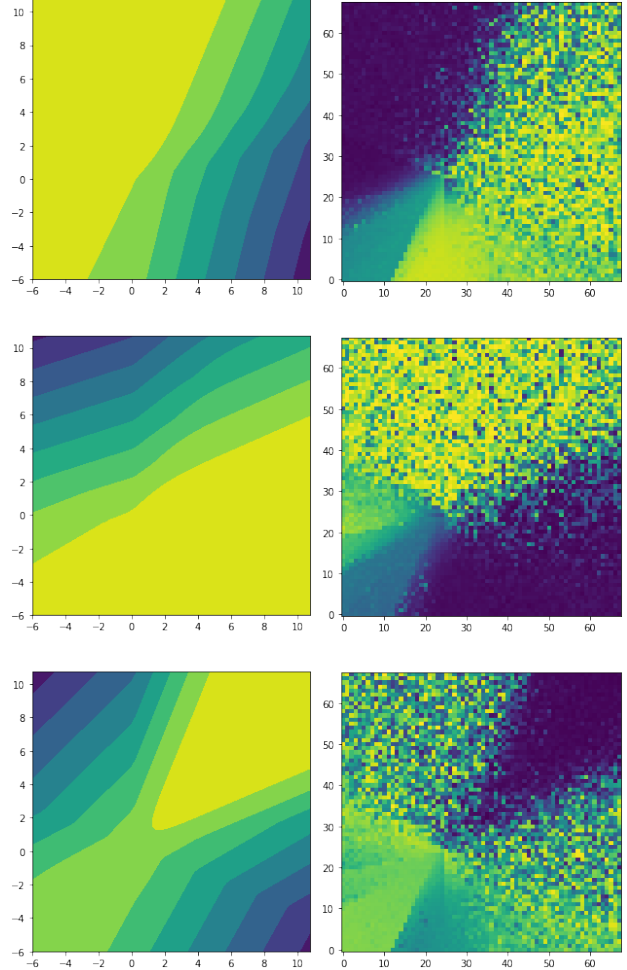


Figure 5: (left) Kernel values measured on a grid around the training set for our 2D problem. Bright yellow means high kernel value (right) Monte-Carlo estimated standard deviation based on gram matrices generated using our kernel for the same grid as the kernel values. Yellow means high standard deviation, blue means low standard deviation.

#### 4.3. Extending To Image Data

We perform experiments on MNIST to demonstrate the applicability to image data. This kernel representation was generated for convolutional ReLU Network with the categorical cross-entropy loss function, using Pytorch (Paszke et al., 2019). The model was trained using forward Euler (gradient descent) using gradients generated as a sum over all training data for each step. The state of the model was saved for every training step. In order to compute the per-training-point gradients needed for the kernel representation, the per-input jacobians are computed at execution time in the representation by loading the model for each training step  $i$ , computing the jacobians for each training input to compute  $\nabla_w f_{w_s(0)}(x_i)$ , and then repeating this procedure

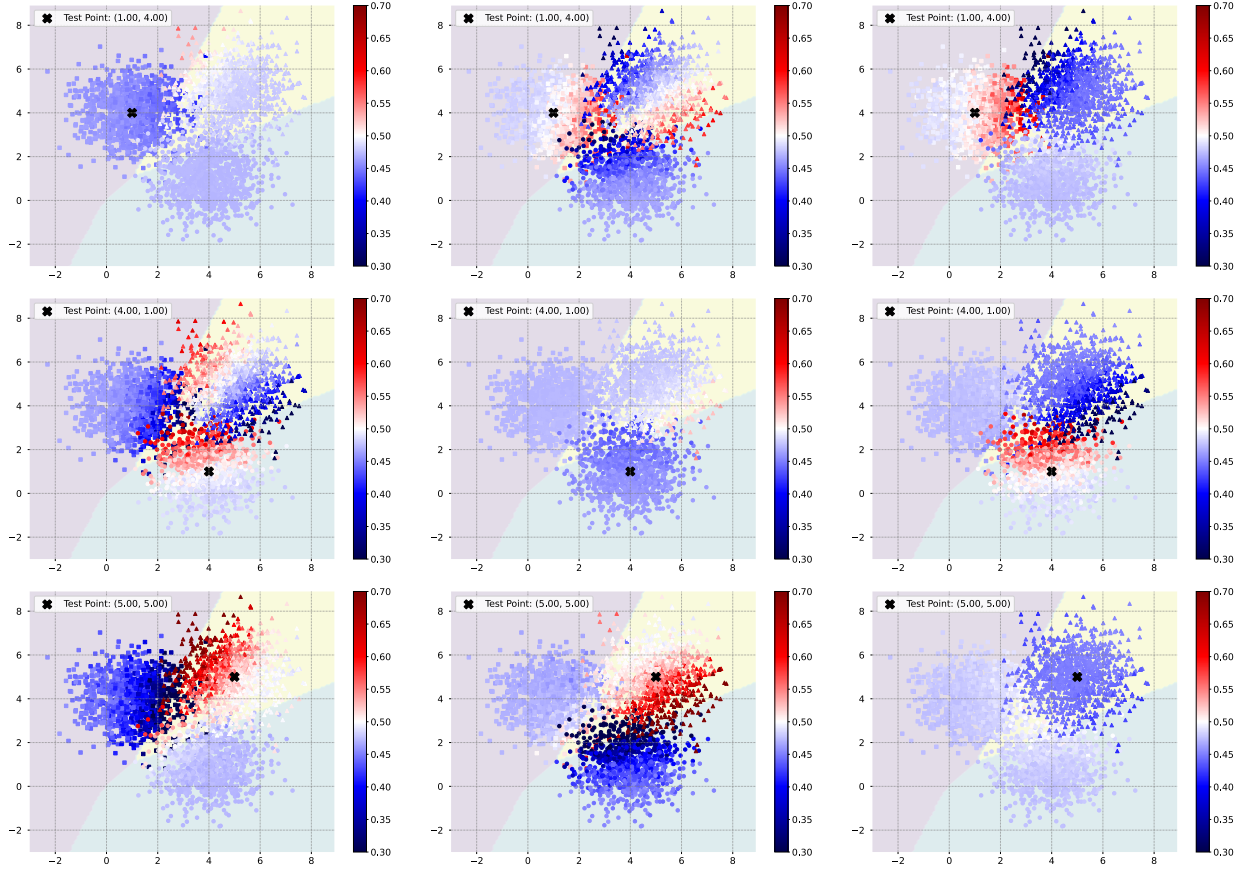


Figure 6: Plots showing kernel values for each training point relative to a test point. Because our kernel is replicating the output of a network, there are three kernel values per sample on a three class problem. This plot shows kernel values for all three classes across three different test points selected as the mean of the generating distribution. Figures on the diagonal show kernel values of the predicted class. Background shading is the neural network decision boundary.

for 200  $t$  values between 0 and 1 in order to approximate  $\int_0^1 f_{w_\theta(t)}(x)$ . For MNIST, the resulting prediction is very sensitive to the accuracy of this integral approximation, as shown in Fig. 7. The top plot shows approximation of the above integral with only one step, which corresponds to the DPK from previous work ((Chen et al., 2021), (Domingos, 2020a), (Incudini et al., 2022)) and as we can see, careful approximation of this integral is necessary to achieve an accurate match between the model and kernel.

## 5. Conclusion and Outlook

The implications of a practical and finite kernel representation for the study of neural networks are profound and yet importantly limited by the networks that they are built from. For most gradient trained models, there is a disconnect between the input space (e.g. images) and the parameter space of a network. Parameters are intrinsically difficult to interpret and much work has been spent building approximate mappings that convert model understanding back into the

input space in order to interpret features, sample importance, and other details (Simonyan et al., 2013; Lundberg & Lee, 2017; Selvaraju et al., 2019). The EPK is composed of a direct mapping from the input space into parameter space. This mapping allows for a much deeper understanding of gradient trained models because the internal state of the method has an exact representation mapped from the input space. As we have shown in Fig. 6, kernel values derived from gradient methods tell an odd story. We have observed a kernel that picks inputs near decision boundaries to emphasize and derives a spatial transform whose basis vectors depend neither uniformly nor continuously on training points. Although kernel values are linked to sample importance, we have shown that most contributions to the kernel’s prediction for a given point are measuring an overall change in the network’s internal representation. This supports the notion that most of what a network is doing is fitting a spatial transform based on a wide aggregation of data, and only doing a trivial calculation to the data once this spatial transform has been determined (Chizat & Bach,

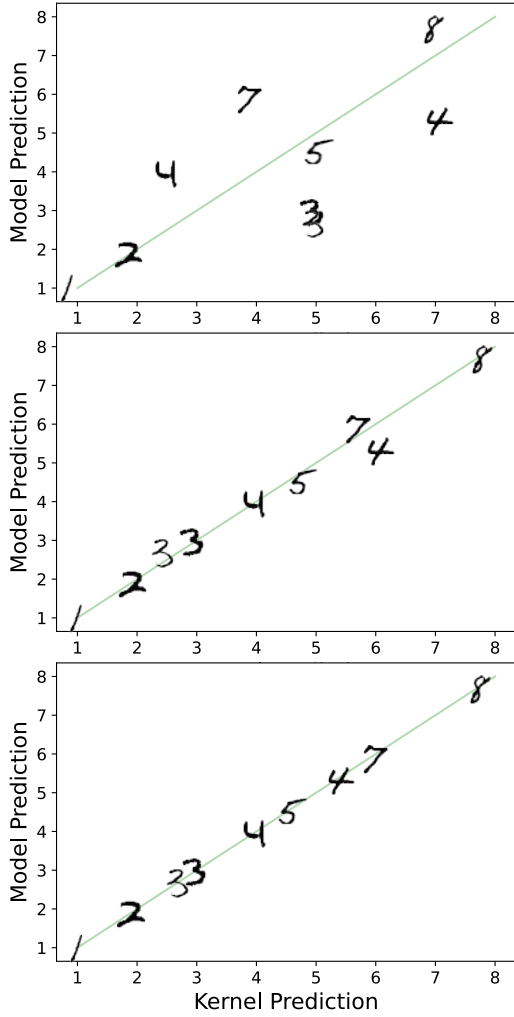


Figure 7: Experiment demonstrating the relationship between model predictions and kernel predictions for varying precision of the integrated path kernel. The top figure shows the integral estimated using only a single step. This is equivalent to the discrete path kernel (DPK) of previous work (Domingos, 2020b; Chen et al., 2021). The middle figure shows the kernel evaluated using 10 integral steps. The final figure shows the path kernel evaluated using 200 integral steps.

2020). As stated in previous work (Domingos, 2020a), this representation has strong implications about the structure of gradient trained models and how they can understand the problems that they solve. Since the kernel weights in this representation are fixed derivatives with respect to the loss function  $L$ ,  $a_{i,s} = -\varepsilon \frac{\partial L(f_{w_s(0)}(x_i), y_i)}{\partial f_i}$ , nearly all of the information used by the network is represented by the kernel mapping function and inner product. Inner products

are not just measures of distance, they also measure angle. In fact, figure 8 shows that for a typical training example, the  $L_2$  norm of the weights changes monotonically by only 20-30% during training. This means that the "learning" of a gradient trained model is dominated by change in angle, which is predicted for kernel methods in high dimensions (Härdle et al., 2004).

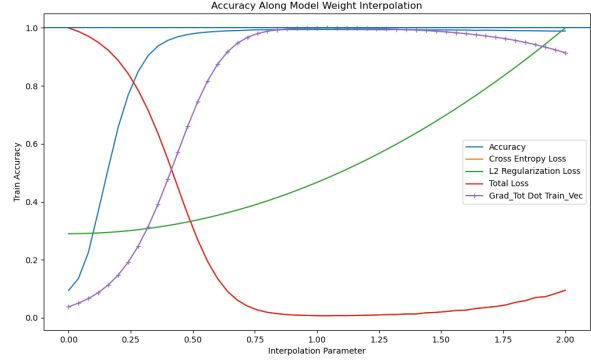


Figure 8: This plot shows a linear interpolation  $w(t) = w_0 + t(w_1 - w_0)$  of model parameters  $w$  from their starting random state  $w_0$  to their ending trained state  $w_1$ . The hatched purple line shows the dot product of the sum of the gradient over the training data  $X$ ,  $\langle \nabla_w f_{w(t)}(X), (w_1 - w_0)/|w_1 - w_0| \rangle$ . The other lines indicate accuracy (blue), total loss (red decreasing), and L2 Regularization (green increasing)

For kernel methods, our result also represents a new direction. Despite their firm mathematical foundations, kernel methods have lost ground since the early 2000s because the features implicitly learned by deep neural networks yield better accuracy than any known hand-crafted kernels for complex high-dimensional problems (Bengio et al., 2005). We're hopeful about the scalability of learned kernels based on recent results in scaling kernel methods (Snelson & Ghahramani, 2005). Exact kernel equivalence could allow the use of neural networks to implicitly construct a kernel. This could allow kernel based classifiers to approach the performance of neural networks on complex data. Kernels built in this way may be used with Gaussian processes to allow meaningful direct uncertainty measurement. This would allow for much more significant analysis for out-of-distribution samples including adversarial attacks (Szegedy et al., 2013; Ilyas et al., 2019). There is significant work to be done in improving the properties of the kernels learned by neural networks for these tools to be used in practice. We are confident that this direct connection between practical neural networks and kernels is a strong first step towards achieving this goal.

## References

- Abdar, M., Pourpanah, F., Hussain, S., Rezazadegan, D., Liu, L., Ghavamzadeh, M., Fieguth, P., Cao, X., Khosravi, A., Acharya, U. R., Makarencov, V., and Nahavandi, S. A review of uncertainty quantification in deep learning: Techniques, applications and challenges. *Information Fusion*, 76:243–297, 2021. ISSN 1566-2535. doi: <https://doi.org/10.1016/j.inffus.2021.05.008>. URL <https://www.sciencedirect.com/science/article/pii/S1566253521001081>.
- Alemohammad, S., Wang, Z., Balestriero, R., and Baraniuk, R. The Recurrent Neural Tangent Kernel, June 2021. URL <http://arxiv.org/abs/2006.10246>. arXiv:2006.10246 [cs, stat].
- Bengio, Y., Delalleau, O., and Roux, N. The curse of highly variable functions for local kernel machines. In Weiss, Y., Schölkopf, B., and Platt, J. (eds.), *Advances in Neural Information Processing Systems*, volume 18. MIT Press, 2005. URL <https://proceedings.neurips.cc/paper/2005/file/663772ea088360f95bac3dc7ffb841be-Paper.pdf>.
- Bietti, A. and Mairal, J. On the inductive bias of neural tangent kernels. In Wallach, H., Larochelle, H., Beygelzimer, A., d’Alché-Buc, F., Fox, E., and Garnett, R. (eds.), *Advances in Neural Information Processing Systems*, volume 32. Curran Associates, Inc., 2019. URL [https://proceedings.neurips.cc/paper\\_files/paper/2019/file/c4ef9c39b300931b69a36fb3dbb8d60e-Paper.pdf](https://proceedings.neurips.cc/paper_files/paper/2019/file/c4ef9c39b300931b69a36fb3dbb8d60e-Paper.pdf).
- Bubeck, S. and Sellke, M. A universal law of robustness via isoperimetry. In Ranzato, M., Beygelzimer, A., Dauphin, Y. N., Liang, P., and Vaughan, J. W. (eds.), *Advances in Neural Information Processing Systems 34: Annual Conference on Neural Information Processing Systems 2021, NeurIPS 2021, December 6-14, 2021, virtual*, pp. 28811–28822, 2021. URL <https://proceedings.neurips.cc/paper/2021/hash/f197002b9a0853eca5e046d9ca4663d5-Abstract.html>.
- Chen, Y., Huang, W., Nguyen, L., and Weng, T.-W. On the equivalence between neural network and support vector machine. *Advances in Neural Information Processing Systems*, 34:23478–23490, 2021.
- Chen, Z., Cao, Y., Gu, Q., and Zhang, T. Mean-field analysis of two-layer neural networks: Non-asymptotic rates and generalization bounds. *CoRR*, abs/2002.04026, 2020. URL <https://arxiv.org/abs/2002.04026>.
- Chizat, L. and Bach, F. Implicit bias of gradient descent for wide two-layer neural networks trained with the logistic loss. In Abernethy, J. and Agarwal, S. (eds.), *Proceedings of Thirty Third Conference on Learning Theory*, volume 125 of *Proceedings of Machine Learning Research*, pp. 1305–1338. PMLR, 09–12 Jul 2020. URL <https://proceedings.mlr.press/v125/chizat20a.html>.
- Domingos, P. Every model learned by gradient descent is approximately a kernel machine. *CoRR*, abs/2012.00152, 2020a. URL <https://arxiv.org/abs/2012.00152>.
- Domingos, P. Every model learned by gradient descent is approximately a kernel machine. *arXiv preprint arXiv:2012.00152*, 2020b.
- Du, S. S., Hou, K., Salakhutdinov, R. R., Poczos, B., Wang, R., and Xu, K. Graph neural tangent kernel: Fusing graph neural networks with graph kernels. In Wallach, H., Larochelle, H., Beygelzimer, A., d’Alché-Buc, F., Fox, E., and Garnett, R. (eds.), *Advances in Neural Information Processing Systems*, volume 32. Curran Associates, Inc., 2019. URL [https://proceedings.neurips.cc/paper\\_files/paper/2019/file/663fd3c5144fd10bd5ca6611a9a5b92d-Paper.pdf](https://proceedings.neurips.cc/paper_files/paper/2019/file/663fd3c5144fd10bd5ca6611a9a5b92d-Paper.pdf).
- Geifman, A., Yadav, A., Kasten, Y., Galun, M., Jacobs, D., and Ronen, B. On the Similarity between the Laplace and Neural Tangent Kernels. In *Advances in Neural Information Processing Systems*, volume 33, pp. 1451–1461. Curran Associates, Inc., 2020. URL <https://proceedings.neurips.cc/paper/2020/hash/1006ff12c465532f8c574aeaa4461b16-Abstract.html>.
- Ghojogh, B., Ghodsi, A., Karray, F., and Crowley, M. Reproducing kernel hilbert space, mercer’s theorem, eigenfunctions, nyström method, and use of kernels in machine learning: Tutorial and survey, 2021. URL <https://arxiv.org/abs/2106.08443>.
- Härdle, W., Müller, M., Sperlich, S., Werwatz, A., et al. *Nonparametric and semiparametric models*, volume 1. Springer, 2004.
- Ilyas, A., Santurkar, S., Tsipras, D., Engstrom, L., Tran, B., and Madry, A. Adversarial examples are not bugs, they are features. *Advances in neural information processing systems*, 32, 2019.
- Includini, M., Grossi, M., Mandarino, A., Vallecorsa, S., Pierro, A. D., and Windridge, D. The quantum path kernel: a generalized quantum neural tangent kernel for deep quantum machine learning, 2022.

- Jacot, A., Gabriel, F., and Hongler, C. Neural tangent kernel: Convergence and generalization in neural networks. *Advances in neural information processing systems*, 31, 2018.
- Lee, J., Bahri, Y., Novak, R., Schoenholz, S. S., Pennington, J., and Sohl-Dickstein, J. Deep neural networks as gaussian processes. In *6th International Conference on Learning Representations, ICLR 2018, Vancouver, BC, Canada, April 30 - May 3, 2018, Conference Track Proceedings*. OpenReview.net, 2018. URL <https://openreview.net/forum?id=B1EA-M-0Z>.
- Lundberg, S. and Lee, S.-I. A unified approach to interpreting model predictions, 2017. URL <https://arxiv.org/abs/1705.07874>.
- Neal, R. M. and Neal, R. M. Priors for infinite networks. *Bayesian learning for neural networks*, pp. 29–53, 1996.
- Paszke, A., Gross, S., Massa, F., Lerer, A., Bradbury, J., Chanan, G., Killeen, T., Lin, Z., Gimelshein, N., Antiga, L., Desmaison, A., Kopf, A., Yang, E., DeVito, Z., Raison, M., Tejani, A., Chilamkurthy, S., Steiner, B., Fang, L., Bai, J., and Chintala, S. Pytorch: An imperative style, high-performance deep learning library. In *Advances in Neural Information Processing Systems 32*, pp. 8024–8035. Curran Associates, Inc., 2019.
- Rasmussen, C. E., Williams, C. K., et al. *Gaussian processes for machine learning*, volume 1. Springer, 2006.
- Selvaraju, R. R., Cogswell, M., Das, A., Vedantam, R., Parikh, D., and Batra, D. Grad-CAM: Visual explanations from deep networks via gradient-based localization. *International Journal of Computer Vision*, 128(2):336–359, oct 2019. doi: 10.1007/s11263-019-01228-7. URL <https://doi.org/10.1007%2Fs11263-019-01228-7>.
- Shah, H., Jain, P., and Netrapalli, P. Do input gradients highlight discriminative features? *Advances in Neural Information Processing Systems*, 34:2046–2059, 2021.
- Shawe-Taylor, J., Cristianini, N., et al. *Kernel methods for pattern analysis*. Cambridge university press, 2004.
- Simonyan, K., Vedaldi, A., and Zisserman, A. Deep inside convolutional networks: Visualising image classification models and saliency maps. *arXiv preprint arXiv:1312.6034*, 2013.
- Snelson, E. and Ghahramani, Z. Sparse gaussian processes using pseudo-inputs. *Advances in neural information processing systems*, 18, 2005.
- Szegedy, C., Zaremba, W., Sutskever, I., Bruna, J., Erhan, D., Goodfellow, I., and Fergus, R. Intriguing properties of neural networks. *arXiv preprint arXiv:1312.6199*, 2013.
- Tancik, M., Srinivasan, P., Mildenhall, B., Fridovich-Keil, S., Raghavan, N., Singhal, U., Ramamoorthi, R., Barron, J., and Ng, R. Fourier features let networks learn high frequency functions in low dimensional domains. In Larochelle, H., Ranzato, M., Hadsell, R., Balcan, M., and Lin, H. (eds.), *Advances in Neural Information Processing Systems*, volume 33, pp. 7537–7547. Curran Associates, Inc., 2020. URL [https://proceedings.neurips.cc/paper\\_files/paper/2020/file/55053683268957697aa39fba6f231c68-Paper.pdf](https://proceedings.neurips.cc/paper_files/paper/2020/file/55053683268957697aa39fba6f231c68-Paper.pdf).
- Zhao, S. and Grishman, R. Extracting relations with integrated information using kernel methods. In *Proceedings of the 43rd annual meeting of the association for computational linguistics (acl’05)*, pp. 419–426, 2005.