

Richer Bayesian Last Layers with Subsampled NTK Features

ICML 2026

Presenter: Dongyoon Kim

26.07.16

Preliminaries — Notation

Two roles for $\mathbf{x}$ (paper's convention).

$$\underbrace{\mathbf{x}_i \in \mathbb{R}^d}_{\text{a single input point}}, \quad \underbrace{\mathbf{x} := (\mathbf{x}_1, \dots, \mathbf{x}_N)}_{\text{training set (used as subscript)}}, \quad \underbrace{\mathbf{x}' := (\mathbf{x}'_1, \dots, \mathbf{x}'_{N'})}_{\text{test set (used as subscript)}}.$$

Function arguments take single points; *matrix subscripts* take collections.

$\mathcal{D} = \{(\mathbf{x}_i, y_i)\}_{i=1}^N$ training data

$f_{\boldsymbol{\theta}} : \mathbb{R}^d \rightarrow \mathbb{R}$, $\boldsymbol{\theta} \in \mathbb{R}^P$ neural network with parameters $\boldsymbol{\theta}$

$\hat{\boldsymbol{\theta}}$ MAP estimate

$\phi^p(\mathbf{x}_i) := \nabla_{\boldsymbol{\theta}} f_{\hat{\boldsymbol{\theta}}}(\mathbf{x}_i) \in \mathbb{R}^p$ eNTK feature of *one* point

$\Phi_{\mathbf{x}}^p \in \mathbb{R}^{N \times p}$ feature matrix (rows are $\phi^p(\mathbf{x}_i)^\top$)

$k_{\mathbf{x}, \mathbf{x}'}^p := \Phi_{\mathbf{x}}^p \Phi_{\mathbf{x}'}^{p\top} \in \mathbb{R}^{N \times N'}$ eNTK Gram matrix (train $\times$ test)

σ^2 , $S_{\mathbf{x}', \mathbf{x}'}^{\text{ntk}} \in \mathbb{R}^{N' \times N'}$ noise variance; posterior-predictive cov.

$\phi^r(\mathbf{x}_i) \in \mathbb{R}^r$, $k_{\mathbf{x}, \mathbf{x}'}^r$, $S_{\mathbf{x}', \mathbf{x}'}^{\text{bll}}$ last-layer analogues (BLL), r = last-layer width

Three sizes to keep separate: d = input dim, $N/N' = \#$ train/test points, $p = m + r = \#$ params.

What is Neural Tangent Kernel?

The NTK *emerges naturally* from gradient-descent training dynamics. (Shown here for one point $(\mathbf{x}_i, y_i)$; sum over i for the full dataset.)

Step 1 — Continuous-time gradient flow. $\boldsymbol{\theta}_{k+1} = \boldsymbol{\theta}_k - \eta \nabla_{\boldsymbol{\theta}} \mathcal{L}(\boldsymbol{\theta}_k)$ becomes, as $\eta \rightarrow 0$,

$$\dot{\boldsymbol{\theta}}(t) = -\nabla_{\boldsymbol{\theta}} \mathcal{L}(\boldsymbol{\theta}(t)).$$

Step 2 — MSE loss. With $\mathcal{L}(\boldsymbol{\theta}) = \frac{1}{2}(f_{\boldsymbol{\theta}}(\mathbf{x}_i) - y_i)^2$,

$$\dot{\boldsymbol{\theta}} = -\nabla_{\boldsymbol{\theta}} f_{\boldsymbol{\theta}}(\mathbf{x}_i) (f_{\boldsymbol{\theta}}(\mathbf{x}_i) - y_i).$$

Step 3 — Chain rule: output evolution at a test point $\mathbf{x}'_j$.

$$\dot{f}_{\boldsymbol{\theta}}(\mathbf{x}'_j) = \nabla_{\boldsymbol{\theta}} f_{\boldsymbol{\theta}}(\mathbf{x}'_j)^{\top} \dot{\boldsymbol{\theta}} = - \underbrace{\nabla_{\boldsymbol{\theta}} f_{\boldsymbol{\theta}}(\mathbf{x}'_j)^{\top} \nabla_{\boldsymbol{\theta}} f_{\boldsymbol{\theta}}(\mathbf{x}_i)}_{=: H_{\boldsymbol{\theta}}(\mathbf{x}'_j, \mathbf{x}_i)} (f_{\boldsymbol{\theta}}(\mathbf{x}_i) - y_i).$$

Neural Tangent Kernel and eNTK

Neural Tangent Kernel (NTK).

In the *infinite-width* limit, training a neural network by gradient descent is equivalent to **kernel regression** with kernel

$$k_{\text{NTK}}(\mathbf{a}, \mathbf{b}) = \mathbb{E}_{\theta_0}[\langle \nabla_{\theta} f_{\theta}(\mathbf{a}), \nabla_{\theta} f_{\theta}(\mathbf{b}) \rangle], \quad \mathbf{a}, \mathbf{b} \in \mathbb{R}^d.$$

Empirical NTK (eNTK).

Finite-width, practical analogue: evaluate the tangent kernel at the trained (MAP) parameters $\hat{\theta}$,

$$k^p(\mathbf{a}, \mathbf{b}) := \langle \nabla_{\theta} f_{\hat{\theta}}(\mathbf{a}), \nabla_{\theta} f_{\hat{\theta}}(\mathbf{b}) \rangle = \phi^p(\mathbf{a})^{\top} \phi^p(\mathbf{b}).$$

Reading the notation. The kernel *function* $k^p(\mathbf{a}, \mathbf{b})$ takes two single points; stacking over training / test sets gives the Gram *matrix* $k_{\mathbf{x}, \mathbf{x}'}^p \in \mathbb{R}^{N \times N'}$.

From Neural Network to Bayesian Linear Model

Step 1 — Linearize f_{θ} around $\hat{\theta}$. For any input $\mathbf{a} \in \mathbb{R}^d$,

$$f_{\theta}(\mathbf{a}) \approx \underbrace{f_{\hat{\theta}}(\mathbf{a})}_{\text{fixed scalar}} + \underbrace{\phi^p(\mathbf{a})^{\top}}_{\text{fixed vector in } \mathbb{R}^p =: \Delta\theta, \text{ unknown in } \mathbb{R}^p} \underbrace{(\theta - \hat{\theta})}_{\text{}} .$$

Everything is computed once at $\hat{\theta}$; the only unknown is $\Delta\theta$.

Step 2 — Bayesian linear regression.

Gaussian prior on the deviation and Gaussian likelihood:

$$\Delta\theta \sim \mathcal{N}(\mathbf{0}, \mathbf{I}_p), \quad y_i \mid \mathbf{x}_i, \theta \sim \mathcal{N}(f_{\theta}(\mathbf{x}_i), \sigma^2) .$$

Weight-space $\leftrightarrow$ Function-space (GP view)

Compute mean and covariance directly. For any $\mathbf{a}, \mathbf{b} \in \mathbb{R}^d$, writing $f^{lin}(\mathbf{a}) = f_{\hat{\theta}}(\mathbf{a}) + \phi^p(\mathbf{a})^\top \Delta\theta$ with $\Delta\theta \sim \mathcal{N}(\mathbf{0}, \mathbf{I}_p)$:

$$\mathbb{E}[f^{lin}(\mathbf{a})] = f_{\hat{\theta}}(\mathbf{a}) + \phi^p(\mathbf{a})^\top \underbrace{\mathbb{E}[\Delta\theta]}_{=\mathbf{0}} = f_{\hat{\theta}}(\mathbf{a}),$$

$$\text{Cov}(f(\mathbf{a}), f(\mathbf{b})) = \phi^p(\mathbf{a})^\top \underbrace{\mathbb{E}[\Delta\theta \Delta\theta^\top]}_{=\mathbf{I}_p} \phi^p(\mathbf{b}) = k^p(\mathbf{a}, \mathbf{b}).$$

Joint Gaussianity. For any finite set of points, the corresponding output vector is a linear function of the Gaussian $\Delta\theta$, hence jointly Gaussian. This is the definition of a Gaussian process:

$$f \sim \mathcal{GP}(f_{\hat{\theta}}, k^p).$$

Preliminaries — Posterior Predictive

Step 3 — Condition on $\mathcal{D}$ (Gaussian conjugacy). Recall each kernel block is a feature outer product:

$k_{\mathbf{a},\mathbf{b}}^p = \Phi_{\mathbf{a}}^p \Phi_{\mathbf{b}}^{p\top}$. Under the GP prior, the joint distribution is Gaussian:

$$\begin{pmatrix} f^{\text{lin}}(\mathbf{x}') \\ \mathbf{y} \end{pmatrix}_{\substack{N' \times 1 \\ N \times 1}} \sim \mathcal{N} \left(\begin{pmatrix} f_{\hat{\theta}}(\mathbf{x}') \\ f_{\hat{\theta}}(\mathbf{x}) \end{pmatrix}, \begin{pmatrix} k_{\mathbf{x}',\mathbf{x}'}^p & k_{\mathbf{x}',\mathbf{x}}^p \\ k_{\mathbf{x},\mathbf{x}'}^p & k_{\mathbf{x},\mathbf{x}}^p + \sigma^2 \mathbf{I}_N \end{pmatrix}_{\substack{N' \times N' \\ N \times N}} \right).$$

Gaussian conditioning identity $\Rightarrow$ predictive covariance:

$$\begin{aligned} \text{Cov}[f^{\text{lin}}(\mathbf{x}') | \mathcal{D}] &= \underbrace{k_{\mathbf{x}',\mathbf{x}'}^p - k_{\mathbf{x}',\mathbf{x}}^p (k_{\mathbf{x},\mathbf{x}}^p + \sigma^2 \mathbf{I}_N)^{-1} k_{\mathbf{x},\mathbf{x}'}^p}_{= S_{\mathbf{x}',\mathbf{x}'}^{\text{ntk}} \in \mathbb{R}^{N' \times N'}}. \end{aligned}$$

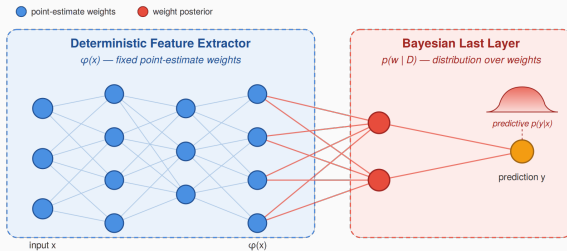
Bottleneck. $N \times N$ inverse above. **Woodbury identity** rewrites it as

$$(\sigma^2 \mathbf{I}_N + \Phi^p \Phi^{p\top})^{-1} = \sigma^{-2} \mathbf{I}_N - \sigma^{-2} \underbrace{\Phi^p (\sigma^2 \mathbf{I}_p + \Phi^{p\top} \Phi^p)^{-1} \Phi^{p\top}}_{p \times p},$$

swapping $N \times N$ for a $p \times p$ inverse. For full eNTK $p \gg N$ (millions of params) — Woodbury doesn't help.

Motivation

Bayesian Last Layers (BLLs) provide a convenient and computationally efficient way to estimate uncertainty in neural networks.



Key problem: BLL underestimates epistemic uncertainty because it applies a Bayesian treatment only to the final layer, ignoring uncertainty induced by earlier layers.

Methodology — Motivation

Two extremes so far.

	feature dim.	cost	uncertainty
Standard BLL	r	cheap	<i>underestimates</i>
Full eNTK / LLA	$p = m+r$	infeasible	richer

Goal. eNTK-quality uncertainty at BLL cost.

Key observation. The eNTK spectrum concentrates on few directions $\Rightarrow$ many of the p feature dimensions carry little information.

(*Spectrum* = eigenvalues of $\Phi_{\mathbf{x}}^p \Phi_{\mathbf{x}}^{p\top}$, i.e. squared singular values of $\Phi_{\mathbf{x}}^p$; rapid decay $\Rightarrow$ only a few singular directions matter.)

Idea. Approximate non-last-layer features by last-layer features:

$$\phi^m(\mathbf{x}_i) \approx \mathbf{A} \phi^r(\mathbf{x}_i), \quad \mathbf{A} \in \mathbb{R}^{m \times r}.$$

Inference lives in the r -dim last-layer space, richer than plain BLL.

Methodology — eNTK Feature Decomposition

Split parameters by layer. With a linear final layer:

$$\theta = \text{vec}(\theta^{\leq l}, \theta^{l+1}) \in \mathbb{R}^p, \quad \theta^{\leq l} \in \mathbb{R}^m, \quad \theta^{l+1} \in \mathbb{R}^r, \quad p = m + r.$$

Last-layer (NNGP) features. Gradient w.r.t. the last-layer weights and bias is exactly the last hidden layer's post-activation (augmented with a 1 for the bias):

$$\phi^r(\mathbf{x}_i) := \nabla_{\theta^{l+1}} f_{\hat{\theta}}(\mathbf{x}_i) = \text{vec}(\mathbf{a}_{\hat{\theta}}^l(\mathbf{x}_i), 1) \in \mathbb{R}^r.$$

Non-last-layer features. $\phi^m(\mathbf{x}_i) := \nabla_{\theta^{\leq l}} f_{\hat{\theta}}(\mathbf{x}_i) \in \mathbb{R}^m$.

Full eNTK features decompose as:

$$\phi^p(\mathbf{x}_i) = \text{vec}(\phi^m(\mathbf{x}_i), \phi^r(\mathbf{x}_i)) \in \mathbb{R}^{m+r}. \quad (6)$$

Standard **BLL** keeps only ϕ^r ; standard **eNTK/LLA** uses the full ϕ^p . We approximate the ϕ^m contribution using ϕ^r .

Methodology — Approximating ϕ^m by ϕ^r

Assumption. There exists $\mathbf{A} \in \mathbb{R}^{m \times r}$ such that $\phi^m(\mathbf{x}_i) \approx \mathbf{A} \phi^r(\mathbf{x}_i)$ for every training point. (7)

Fit $\mathbf{A}$ by least squares on the training data:

$$\min_{\mathbf{A}} \sum_{i=1}^N \|\phi^m(\mathbf{x}_i) - \mathbf{A} \phi^r(\mathbf{x}_i)\|_2^2. \quad (8)$$

Whenever $N \geq r$, $\Phi_{\mathbf{x}}^{r\top} \Phi_{\mathbf{x}}^r$ has full rank and

$$\mathbf{A} = \Phi_{\mathbf{x}}^{m\top} \Phi_{\mathbf{x}}^r (\Phi_{\mathbf{x}}^{r\top} \Phi_{\mathbf{x}}^r)^{-1} \in \mathbb{R}^{m \times r}. \quad (9)$$

Stack into a single lifting matrix.

$$\mathbf{B} := \begin{pmatrix} \mathbf{A} \\ \mathbf{I}_r \end{pmatrix} \in \mathbb{R}^{(m+r) \times r} \quad \implies \quad \phi^p(\mathbf{x}_i) \approx \mathbf{B} \phi^r(\mathbf{x}_i) =: \phi^B(\mathbf{x}_i). \quad (10, 11)$$

Equivalently $\Phi_{\mathbf{x}}^p \approx \Phi_{\mathbf{x}}^r \mathbf{B}^\top$.

Methodology — Modified Posterior Covariance

Approximate kernel matrix. $k_{\mathbf{x},\mathbf{x}'}^B := \Phi_{\mathbf{x}}^r \mathbf{B}^\top \mathbf{B} \Phi_{\mathbf{x}'}^{r\top} \in \mathbb{R}^{N \times N'}$. The $r \times r$ Gram $\mathbf{B}^\top \mathbf{B}$ acts as a data-driven metric on last-layer features, replacing the identity used by standard BLL.

From full eNTK to Rich-BLL: substitute $k^p \rightarrow k^B$.

$$\underbrace{S_{\mathbf{x}',\mathbf{x}'}^{\text{ntk}} = k_{\mathbf{x}',\mathbf{x}'}^p - k_{\mathbf{x}',\mathbf{x}}^p (k_{\mathbf{x},\mathbf{x}}^p + \sigma^2 \mathbf{I}_N)^{-1} k_{\mathbf{x},\mathbf{x}'}^p}_{\text{full eNTK}} \xrightarrow{k^p \rightarrow k^B} \underbrace{S_{\mathbf{x}',\mathbf{x}'}^B = k_{\mathbf{x}',\mathbf{x}'}^B - k_{\mathbf{x}',\mathbf{x}}^B (k_{\mathbf{x},\mathbf{x}}^B + \sigma^2 \mathbf{I}_N)^{-1} k_{\mathbf{x},\mathbf{x}'}^B}_{\text{Rich-BLL eNTK}}$$

Comparison at a glance.

	features	Gram/metric	effective dim.
Standard BLL	ϕ^r	$\mathbf{I}_r$	r
Full eNTK / LLA	ϕ^p	$\mathbf{I}_p$	$p = m+r$
Rich-BLL (ϕ^B)	$\mathbf{B}\phi^r$	$\mathbf{B}^\top \mathbf{B}$	r

Methodology — Cholesky Reformulation

Predictive covariance in the reduced feature space.

$$S_{\mathbf{x}', \mathbf{x}'}^B = \Phi_{\mathbf{x}'}^r \left(\frac{1}{\sigma^2} \Phi_{\mathbf{x}}^{r\top} \Phi_{\mathbf{x}}^r + (\mathbf{B}^\top \mathbf{B})^{-1} \right)^{-1} \Phi_{\mathbf{x}'}^{r\top}. \quad (12)$$

Cholesky factorization. Let

$$\mathbf{B}^\top \mathbf{B} = \mathbf{L} \mathbf{L}^\top, \quad \mathbf{L} \in \mathbb{R}^{r \times r},$$

where $\mathbf{L}$ is lower triangular, and define

$$\phi^L(\mathbf{x}) := \mathbf{L}^\top \phi^r(\mathbf{x}) \in \mathbb{R}^r. \quad (13)$$

Theorem 3.2. Using $\phi^L(\mathbf{x}) \in \mathbb{R}^r$ is equivalent to using $\phi^B(\mathbf{x}) \in \mathbb{R}^{m+r}$, and

$$S_{\mathbf{x}', \mathbf{x}'}^B = \Phi_{\mathbf{x}'}^L \left(\frac{1}{\sigma^2} \Phi_{\mathbf{x}}^{L\top} \Phi_{\mathbf{x}}^L + \mathbf{I}_r \right)^{-1} \Phi_{\mathbf{x}'}^{L\top} \quad (14)$$

$$= \Phi_{\mathbf{x}'}^r \mathbf{L} \left(\frac{1}{\sigma^2} \mathbf{L}^\top \Phi_{\mathbf{x}}^{r\top} \Phi_{\mathbf{x}}^r \mathbf{L} + \mathbf{I}_r \right)^{-1} \mathbf{L}^\top \Phi_{\mathbf{x}'}^{r\top}. \quad (15)$$

Methodology — Subsampling for Scalability

Remaining N -dependent bottleneck. Although the matrix inversion is only $r \times r$, constructing the feature covariance still requires all N training examples:

$$\Phi_{\mathbf{x}}^{L\top} \Phi_{\mathbf{x}}^L = \sum_{i=1}^N \phi^L(\mathbf{x}_i) \phi^L(\mathbf{x}_i)^\top =: N \hat{\Sigma}_N. \quad (17)$$

Uniform subsampling. Draw indices $s_1, \dots, s_k$ uniformly without replacement from $\{1, \dots, N\}$ and approximate

$$N \hat{\Sigma}_N \approx \frac{N}{k} \sum_{i=1}^k \phi^L(\mathbf{x}_{s_i}) \phi^L(\mathbf{x}_{s_i})^\top =: N \hat{\Sigma}_k. \quad (18)$$

Rich-BLL(S) predictive covariance. Replacing the full feature covariance with its subsampled estimate gives

$$S_{\mathbf{x}', \mathbf{x}'}^{B,k} := \Phi_{\mathbf{x}'}^L \left(\frac{1}{\sigma^2} N \hat{\Sigma}_k + \mathbf{I}_r \right)^{-1} \Phi_{\mathbf{x}'}^{L\top}. \quad (19)$$

Feature-covariance construction: $\mathcal{O}(Nr^2) \rightarrow \mathcal{O}(kr^2)$

Results — UCI Regression (NLL ↓)

Method	Boston	Concrete	Energy	Power	Wine
Rich-BLL	2.61±0.04	3.10±0.03	0.75±0.03	2.78±0.01	1.00±0.01
Rich-BLL (S)	2.62±0.04	3.10±0.04	0.78±0.03	2.78±0.01	1.01±0.01
NNGP (BLL)	2.78±0.06	3.39±0.06	0.92±0.01	2.82±0.01	1.03±0.01
VBLL	2.75±0.06	3.19±0.05	0.74±0.04	2.78±0.01	1.05±0.01
MAP	2.79±0.09	4.77±0.15	0.93±0.03	3.28±0.02	1.02±0.00
RBF GP	2.84±0.09	3.24±0.13	0.66 ±0.04	2.89±0.01	0.97 ±0.01
Dropout	2.62±0.05	3.19±0.02	1.80±0.03	2.81±0.01	1.01±0.01
Ensemble	2.48 ±0.09	3.15±0.03	0.94±0.02	2.75 ±0.02	0.98±0.01
SWAG	2.65±0.01	3.19±0.03	1.20±0.06	2.79±0.01	1.04±0.01
BBB	2.54±0.04	2.99 ±0.03	1.22±0.01	2.77±0.01	1.05±0.01

Results — Image Classification & OOD Detection

Method	NLL ↓	ECE ↓	AUROC ↑ (SVHN)	AUROC ↑ (CIFAR-100)
Rich-BLL (S)	0.56±0.00	0.029±0.01	0.91±0.02	0.65 ±0.01
NNGP (BLL)	0.58±0.00	0.046±0.01	0.88±0.01	0.62±0.02
LL Laplace	0.57±0.00	0.044±0.00	0.84±0.01	0.55±0.00
SNGP	0.89±0.00	0.027±0.001	0.85±0.01	0.52±0.00
LL Dropout	0.56±0.00	0.031±0.002	0.78±0.02	0.62±0.00
Dropout	0.31 ±0.05	0.013 ±0.003	0.92 ±0.001	0.61±0.02
Ensemble	0.47±0.00	0.030±0.002	0.80±0.20	0.65 ±0.01
BBB	1.39±0.02	0.220±0.007	0.89±0.01	0.54±0.00