

Lattice Rules
Kernel Methods
Deep Neural Networks
and how to connect them

Frances Kuo

`f.kuo@unsw.edu.au`

School of Mathematics and Statistics, UNSW Sydney, Australia

Outline of my talk

- **Lattice rule** for high dimensional integration

14 minutes

- **Lattice-based kernel interpolant** for high dimensional function approximation

1 minute

- SIDE STEP: **Uncertainty Quantification (UQ)** and modeling random fields

10 minutes

- **Lattice-based training for deep neural network** as function approximation

15 minutes

Joint work with **Alexander Keller** (NVIDIA), **Dirk Nuyens** (KU Leuven), **Ian Sloan** (UNSW Sydney)

Thanks to my collaborators



NZ & AU

Austria

Belgium

Germany

Poland

Switzerland

UK

USA

Thanks to my collaborators



AU

Austria

Belgium

Ecuador

Finland

France

Germany

Switzerland

UK

USA

Shameless advertising

- **Uncertainty quantification (UQ) using periodic random variables** SINUM (2020)
Vesa Kaarnioja (FU Berlin), Ian Sloan (UNSW Sydney)
- **Kernel method for function approximation in UQ** Numer. Math. (2021)
Yoshihito Kazashi (Strathclyde), Fabio Nobile (EPFL), Vesa Kaarnioja (FU Berlin), Ian Sloan
- **Lattice algorithms for multivariate integration and approximation**
Ronald Cools & Dirk Nuyens (KU Leuven), Ian Sloan MCOM75 (2020); MCOM (2021)
Laurence Wilkes & Weiwen Mo (KU Leuven) J. Complexity (2023); Constr. App. (2024)
- **Smoothing by preintegration; CDF/PDF estimation; option pricing; log-normal PDE**
Alexander Gilbert & Ian Sloan & Abirami Srikumar (UNSW Sydney)
MCOM (2022); Springer (2022); SINUM (2022); SINUM (2025)
- **Parabolic PDE-constrained optimal control under uncertainty**
Philipp Guth (JKU Linz), Vesa Kaarnioja & Claudia Schillings (FU Berlin), Ian Sloan
Numer. Math. (2024); JUQ (2021)
- **Poisson equation on random domain** Numer. Math. (2024)
Harri Hakula (Aalto), Helmut Harbrecht (Basel), Vesa Kaarnioja (FU Berlin), Ian Sloan
- **Helmholtz equation in random medium** submitted (2025)
Ivan Graham & Euan Spence (Bath), Dirk Nuyens (KU Leuven), Ian Sloan
- **Regularity and tailored regularization of DNNs with application to parametric PDEs**
Alexander Keller (Nvidia), Dirk Nuyens (KU Leuven), Ian Sloan submitted (2025)

Lattice Rules

Kernel Methods

(SIDE STEP: Parametric PDEs in UQ)

Deep Neural Networks

Monte Carlo v.s. Quasi-Monte Carlo methods

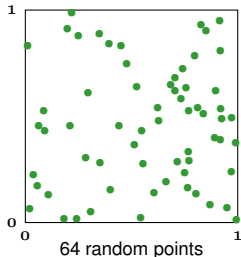
$$\int_{[0,1]^s} F(\mathbf{y}) \, d\mathbf{y} \approx \frac{1}{N} \sum_{k=0}^{N-1} F(\mathbf{t}_k)$$

Monte Carlo method (MC)

$\mathbf{t}_k$ random uniform

$N^{-1/2}$ convergence

Order of variables irrelevant



Quasi-Monte Carlo methods (QMC)

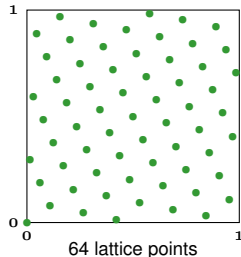
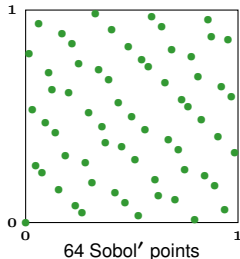
$\mathbf{t}_k$ deterministic

Close to N^{-1} convergence or better

Better for earlier variables and lower-order projections

Order of variables very important

Good for integrands with *low effective dimension*



Quasi-Monte Carlo methods (QMC)

$$\int_{[0,1]^s} F(\mathbf{y}) \, d\mathbf{y} \approx \frac{1}{N} \sum_{k=0}^{N-1} F(\mathbf{t}_k)$$

- How do we choose good QMC points $\mathbf{t}_k$?
- How would the error behave w.r.t. N ?
- How would the error behave w.r.t. s ?

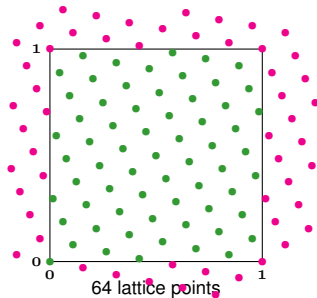
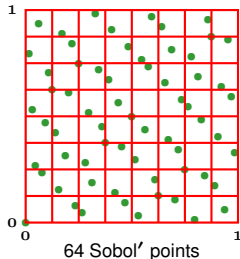
Niederreiter (1992)

Sloan, Joe (1994)

Dick, Pillichshammer (2010)

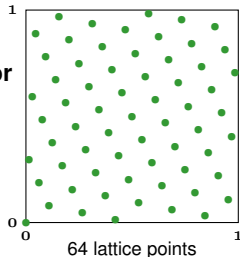
Dick, Kuo, Sloan (2013)

Dick, Kritzer, Pillichshammer (2022)



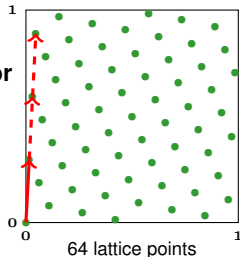
Lifting the curse of dimensionality

- How would the error behave with respect to the dimension s ?
 - Want **strong tractability**, i.e., error bounded independently of s
 - Work in **weighted** function space setting (“effective dimension”)
Sloan, Woźniakowski (1998), Novak, Woźniakowski (2008, 2010, 2012)
- How would the error behave with respect to the number of points N ?
 - Want **higher order** convergence rate when the integrand is smooth
- How do we choose good QMC points t_k ?
 - Use **fast component-by-component construction**
 - For lattice rules, we need just one **generating vector**
- How do we estimate the error in practice?
 - Use **randomization**
 - For lattice rules, we use **random shifts**



Lifting the curse of dimensionality

- How would the error behave with respect to the dimension s ?
 - Want **strong tractability**, i.e., error bounded independently of s
 - Work in **weighted** function space setting (“effective dimension”)
Sloan, Woźniakowski (1998), Novak, Woźniakowski (2008, 2010, 2012)
- How would the error behave with respect to the number of points N ?
 - Want **higher order** convergence rate when the integrand is smooth
- How do we choose good QMC points t_k ?
 - Use **fast component-by-component construction**
 - For lattice rules, we need just one **generating vector**
- How do we estimate the error in practice?
 - Use **randomization**
 - For lattice rules, we use **random shifts**

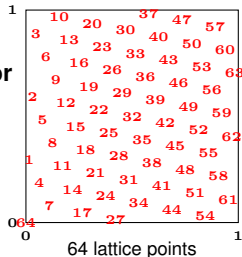


Lifting the curse of dimensionality

- How would the error behave with respect to the dimension s ?
 - Want **strong tractability**, i.e., error bounded independently of s
 - Work in **weighted** function space setting (“effective dimension”)
- How would the error behave with respect to the number of points N ?
 - Want **higher order** convergence rate when the integrand is smooth

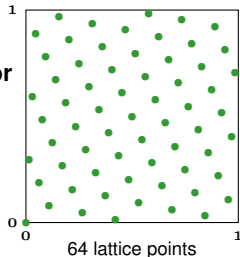
Sloan, Woźniakowski (1998), Novak, Woźniakowski (2008, 2010, 2012)

- How do we choose good QMC points t_k ?
 - Use **fast component-by-component construction**
 - For lattice rules, we need just one **generating vector**
- How do we estimate the error in practice?
 - Use **randomization**
 - For lattice rules, we use **random shifts**



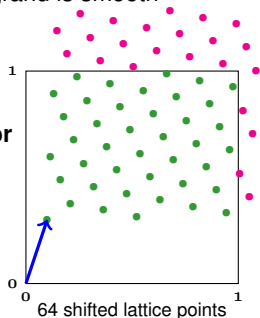
Lifting the curse of dimensionality

- How would the error behave with respect to the dimension s ?
 - Want **strong tractability**, i.e., error bounded independently of s
 - Work in **weighted** function space setting (“effective dimension”)
Sloan, Woźniakowski (1998), Novak, Woźniakowski (2008, 2010, 2012)
- How would the error behave with respect to the number of points N ?
 - Want **higher order** convergence rate when the integrand is smooth
- How do we choose good QMC points t_k ?
 - Use **fast component-by-component construction**
 - For lattice rules, we need just one **generating vector**
- How do we estimate the error in practice?
 - Use **randomization**
 - For lattice rules, we use **random shifts**



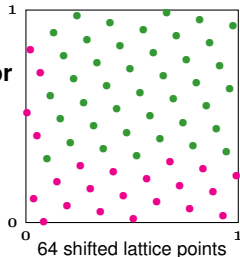
Lifting the curse of dimensionality

- How would the error behave with respect to the dimension s ?
 - Want **strong tractability**, i.e., error bounded independently of s
 - Work in **weighted** function space setting (“effective dimension”)
Sloan, Woźniakowski (1998), Novak, Woźniakowski (2008, 2010, 2012)
- How would the error behave with respect to the number of points N ?
 - Want **higher order** convergence rate when the integrand is smooth
- How do we choose good QMC points t_k ?
 - Use **fast component-by-component construction**
 - For lattice rules, we need just one **generating vector**
- How do we estimate the error in practice?
 - Use **randomization**
 - For lattice rules, we use **random shifts**



Lifting the curse of dimensionality

- How would the error behave with respect to the dimension s ?
 - Want **strong tractability**, i.e., error bounded independently of s
 - Work in **weighted** function space setting (“effective dimension”)
Sloan, Woźniakowski (1998), Novak, Woźniakowski (2008, 2010, 2012)
- How would the error behave with respect to the number of points N ?
 - Want **higher order** convergence rate when the integrand is smooth
- How do we choose good QMC points t_k ?
 - Use **fast component-by-component construction**
 - For lattice rules, we need just one **generating vector**
- How do we estimate the error in practice?
 - Use **randomization**
 - For lattice rules, we use **random shifts**



Component-by-component construction

- Want to find $\mathbf{z}$ for which some error criterion is as small as possible?

- Exhaustive search is practically impossible - too many choices!

- Component-by-component construction (CBC)**

Korobov (1959),
Sloan, Reztsov (2002),
Sloan, Kuo, Joe (2022), ...

- Set $z_1 = 1$
- With z_1 fixed, choose z_2 to minimize the error criterion in 2 dimensions
- With z_1, z_2 fixed, choose z_3 to minimize the error criterion in 3 dimensions
- etc.

Kuo (2003), Dick (2004), ...

- Optimal rate of convergence $\mathcal{O}(N^{-1+\delta})$ in “weighted Sobolev space”,
 $\mathcal{O}(N^{-\alpha+\delta})$ in “weighted Korobov space”,
independently of s under an appropriate condition on the weights

- Averaging argument: there is always one choice as good as average!

- Cost of algorithm for “product weights” is $\mathcal{O}(s N \log N)$ using FFT

Nuyens, Cools (2006), ...

- Extensible/embedded variants

Hickernell, Hong, L'Ecuyer, Lemieux (2000), Hickernell, Niederreiter (2003),
Cools, Kuo, Nuyens (2006), Dick, Pillichshammer, Waterhouse (2007)

How to apply QMC theory (most important slide)

- Common features among “modern” QMC theoretical settings:

- Separation in error bound

$$(\text{rms}) \text{ QMC error} \leq (\text{rms}) \text{ worst case error}_\gamma \times \text{norm of integrand}_\gamma$$

- Weighted spaces
- Pairing of QMC rule with function space
- Optimal rate of convergence
- Rate and constant independent of dimension
- Fast component-by-component (CBC) construction
- Extensible or embedded rules

How to apply QMC theory (most important slide)

- Common features among “modern” QMC theoretical settings:

- Separation in error bound

$$(\text{rms}) \text{ QMC error} \leq (\text{rms}) \text{ worst case error}_\gamma \times \text{norm of integrand}_\gamma$$

- Weighted spaces
- Pairing of QMC rule with function space
- Optimal rate of convergence
- Rate and constant independent of dimension
- Fast component-by-component (CBC) construction
- Extensible or embedded rules

- Application of QMC theory

- Estimate the norm (critical step)
- Choose the weights
- Weights as input to the CBC construction

Two (out of many) theoretical settings for lattice rules

$$(\text{rms}) \text{ QMC error} \leq (\text{rms}) \text{ worst case error}_{\gamma} \times \text{norm of integrand}_{\gamma}$$

(a) Randomly-shifted lattice rules + weighted Sobolev space of smoothness 1

$$\|F\|_{\mathcal{W}_{1,\gamma}}^2 := \sum_{u \subseteq \{1:s\}} \frac{1}{\gamma_u} \int_{[0,1]^{|u|}} \left| \int_{[0,1]^{s-|u|}} \partial^{\mathbf{1}_u} F(\mathbf{y}) d\mathbf{y}_{-u} \right|^2 d\mathbf{y}_u$$

2^s subsets Small "weight" γ_u means that F depends weakly on the variables $\mathbf{y}_u$ Mixed first derivatives are square integrable

(b) Lattice rules + (periodic) weighted Korobov space of smoothness α

$$\|F\|_{\mathcal{W}_{\alpha,\gamma}}^2 := \sum_{u \subseteq \{1:s\}} \frac{1}{\gamma_u} \int_{[0,1]^{|u|}} \left| \int_{[0,1]^{s-|u|}} \partial^{\alpha_u} F(\mathbf{y}) d\mathbf{y}_{-u} \right|^2 d\mathbf{y}_u$$

Two (out of many) theoretical settings for lattice rules

$$(\text{rms}) \text{ QMC error} \leq (\text{rms}) \text{ worst case error}_{\gamma} \times \text{norm of integrand}_{\gamma}$$

(a) Randomly-shifted lattice rules + weighted Sobolev space of smoothness 1

$$\|F\|_{\mathcal{W}_{1,\gamma}}^2 := \sum_{u \subseteq \{1:s\}} \frac{1}{\gamma_u} \int_{[0,1]^{|u|}} \left| \int_{[0,1]^{s-|u|}} \partial^{\mathbf{1}_u} F(\mathbf{y}) d\mathbf{y}_{-u} \right|^2 d\mathbf{y}_u \leq \sum_{u \subseteq \{1:s\}} \frac{B_u}{\gamma_u}$$

2^s subsets $\uparrow$ γ_u $\nwarrow$ Small "weight" γ_u means that F depends weakly on the variables $\mathbf{y}_u$

$\partial^{\mathbf{1}_u}$ $\nwarrow$ Mixed first derivatives are square integrable

(b) Lattice rules + (periodic) weighted Korobov space of smoothness α

$$\|F\|_{\mathcal{W}_{\alpha,\gamma}}^2 := \sum_{u \subseteq \{1:s\}} \frac{1}{\gamma_u} \int_{[0,1]^{|u|}} \left| \int_{[0,1]^{s-|u|}} \partial^{\alpha_u} F(\mathbf{y}) d\mathbf{y}_{-u} \right|^2 d\mathbf{y}_u \leq \sum_{u \subseteq \{1:s\}} \frac{B_u}{\gamma_u}$$

Two (out of many) theoretical settings for lattice rules

$$(\text{rms}) \text{ QMC error} \leq (\text{rms}) \text{ worst case error}_{\gamma} \times \text{norm of integrand}_{\gamma}$$

(a) Randomly-shifted lattice rules + weighted Sobolev space of smoothness 1

$$\|F\|_{\mathcal{W}_{1,\gamma}}^2 := \sum_{u \subseteq \{1:s\}} \frac{1}{\gamma_u} \int_{[0,1]^{|u|}} \left| \int_{[0,1]^{s-|u|}} \partial^{\mathbf{1}_u} F(\mathbf{y}) d\mathbf{y}_{-u} \right|^2 d\mathbf{y}_u \leq \sum_{u \subseteq \{1:s\}} \frac{B_u}{\gamma_u}$$

2^s subsets $\uparrow$ γ_u Small "weight" γ_u means that F depends weakly on the variables $\mathbf{y}_u$ $\nwarrow$ Mixed first derivatives are square integrable

$$\text{rms } e^{\text{wor-int}} \leq \left(\frac{2}{N} \sum_{u \subseteq \{1:s\}} \gamma_u^\lambda A_u \right)^{\frac{1}{2\lambda}} \quad \forall \lambda \in (\frac{1}{2}, 1]$$

close to $\mathcal{O}(N^{-1})$

(b) Lattice rules + (periodic) weighted Korobov space of smoothness α

$$\|F\|_{\mathcal{W}_{\alpha,\gamma}}^2 := \sum_{u \subseteq \{1:s\}} \frac{1}{\gamma_u} \int_{[0,1]^{|u|}} \left| \int_{[0,1]^{s-|u|}} \partial^{\alpha \mathbf{1}_u} F(\mathbf{y}) d\mathbf{y}_{-u} \right|^2 d\mathbf{y}_u \leq \sum_{u \subseteq \{1:s\}} \frac{B_u}{\gamma_u}$$

$$e^{\text{wor-int}} \leq \left(\frac{2}{N} \sum_{u \subseteq \{1:s\}} \gamma_u^\lambda A_u \right)^{\frac{1}{2\lambda}} \quad \forall \lambda \in (\frac{1}{2\alpha}, 1]$$

close to $\mathcal{O}(N^{-\alpha})$

Two (out of many) theoretical settings for lattice rules

$$(\text{rms}) \text{ QMC error} \leq (\text{rms}) \text{ worst case error}_{\gamma} \times \text{norm of integrand}_{\gamma}$$

(a) Randomly-shifted lattice rules + weighted Sobolev space of smoothness 1

$$\|F\|_{\mathcal{W}_{1,\gamma}}^2 := \sum_{u \subseteq \{1:s\}} \frac{1}{\gamma_u} \int_{[0,1]^{|u|}} \left| \int_{[0,1]^{s-|u|}} \partial^{\mathbf{1}_u} F(\mathbf{y}) d\mathbf{y}_{-u} \right|^2 d\mathbf{y}_u \leq \sum_{u \subseteq \{1:s\}} \frac{B_u}{\gamma_u}$$

2^s subsets $\uparrow$ Small "weight" γ_u means that F depends weakly on the variables $\mathbf{y}_u$ $\nwarrow$ Mixed first derivatives are square integrable

$$\text{rms } e^{\text{wor-int}} \leq \left(\frac{2}{N} \sum_{u \subseteq \{1:s\}} \gamma_u^\lambda A_u \right)^{\frac{1}{2\lambda}} \quad \forall \lambda \in (\frac{1}{2}, 1]$$

close to $\mathcal{O}(N^{-1})$

(b) Lattice rules + (periodic) weighted Korobov space of smoothness α

$$\|F\|_{\mathcal{W}_{\alpha,\gamma}}^2 := \sum_{u \subseteq \{1:s\}} \frac{1}{\gamma_u} \int_{[0,1]^{|u|}} \left| \int_{[0,1]^{s-|u|}} \partial^{\alpha_u} F(\mathbf{y}) d\mathbf{y}_{-u} \right|^2 d\mathbf{y}_u \leq \sum_{u \subseteq \{1:s\}} \frac{B_u}{\gamma_u}$$

$$e^{\text{wor-int}} \leq \left(\frac{2}{N} \sum_{u \subseteq \{1:s\}} \gamma_u^\lambda A_u \right)^{\frac{1}{2\lambda}} \quad \forall \lambda \in (\frac{1}{2\alpha}, 1]$$

close to $\mathcal{O}(N^{-\alpha})$

Choose weights for CBC construction

$$\text{minimize} \quad \left(\sum_{u \subseteq \{1:s\}} \gamma_u^\lambda A_u \right)^{\frac{1}{2\lambda}} \left(\sum_{u \subseteq \{1:s\}} \frac{B_u}{\gamma_u} \right)^{\frac{1}{2}} \quad \Rightarrow \quad \gamma_u = \left(\frac{B_u}{A_u} \right)^{\frac{1}{1+\lambda}}$$

Lattice Rules

Kernel Methods

(SIDE STEP: Parametric PDEs in UQ)

Deep Neural Networks

From integration to function approximation

● Integration over the unit cube

$$\int_{[0,1]^s} F(\mathbf{y}) \, d\mathbf{y} \approx \frac{1}{N} \sum_{k=0}^{N-1} F(\mathbf{t}_k)$$

● Integration over the Euclidean space

$$\dots = \int_{\mathbb{R}^s} F(\mathbf{y}) \prod_{i=1}^s \phi(y_i) \, d\mathbf{y} = \int_{[0,1]^s} F(\Phi^{-1}(\mathbf{w})) \, d\mathbf{w} \approx \frac{1}{N} \sum_{k=0}^{N-1} F(\Phi^{-1}(\mathbf{t}_k))$$

From integration to function approximation

● Integration over the unit cube

$$\int_{[0,1]^s} F(\mathbf{y}) \, d\mathbf{y} \approx \frac{1}{N} \sum_{k=0}^{N-1} F(\mathbf{t}_k)$$

● Integration over the Euclidean space

$$\dots = \int_{\mathbb{R}^s} F(\mathbf{y}) \prod_{i=1}^s \phi(y_i) \, d\mathbf{y} = \int_{[0,1]^s} F(\Phi^{-1}(\mathbf{w})) \, d\mathbf{w} \approx \frac{1}{N} \sum_{k=0}^{N-1} F(\Phi^{-1}(\mathbf{t}_k))$$

● Function approximation over the unit cube

$$\begin{aligned} \int_{[0,1]^s} F(\mathbf{y}) e^{-2\pi i \mathbf{h} \cdot \mathbf{y}} \, d\mathbf{y} &\approx \frac{1}{N} \sum_{k=0}^{N-1} F(\mathbf{t}_k) e^{-2\pi i \mathbf{h} \cdot \mathbf{t}_k} \\ \text{● Truncated trig. series} \quad F(\mathbf{y}) &= \sum_{\mathbf{h} \in \mathbb{Z}^s} \hat{F}_{\mathbf{h}} e^{2\pi i \mathbf{h} \cdot \mathbf{y}} \approx \underbrace{\sum_{\mathbf{h} \in \mathcal{A}_s} \hat{F}_{\mathbf{h}}^a e^{2\pi i \mathbf{h} \cdot \mathbf{y}}}_{=: A_N^{\text{trig}}(F)(\mathbf{y})} \end{aligned}$$

From integration to function approximation

● Integration over the unit cube

$$\int_{[0,1]^s} F(\mathbf{y}) \, d\mathbf{y} \approx \frac{1}{N} \sum_{k=0}^{N-1} F(\mathbf{t}_k)$$

● Integration over the Euclidean space

$$\dots = \int_{\mathbb{R}^s} F(\mathbf{y}) \prod_{i=1}^s \phi(y_i) \, d\mathbf{y} = \int_{[0,1]^s} F(\Phi^{-1}(\mathbf{w})) \, d\mathbf{w} \approx \frac{1}{N} \sum_{k=0}^{N-1} F(\Phi^{-1}(\mathbf{t}_k))$$

● Function approximation over the unit cube

$$\int_{[0,1]^s} F(\mathbf{y}) e^{-2\pi i \mathbf{h} \cdot \mathbf{y}} \, d\mathbf{y} \approx \frac{1}{N} \sum_{k=0}^{N-1} F(\mathbf{t}_k) e^{-2\pi i \mathbf{h} \cdot \mathbf{t}_k}$$

$$\text{● Truncated trig. series} \quad F(\mathbf{y}) = \sum_{\mathbf{h} \in \mathbb{Z}^s} \hat{F}_{\mathbf{h}} e^{2\pi i \mathbf{h} \cdot \mathbf{y}} \approx \underbrace{\sum_{\mathbf{h} \in \mathcal{A}_s} \hat{F}_{\mathbf{h}} e^{2\pi i \mathbf{h} \cdot \mathbf{y}}}_{=: A_N^{\text{trig}}(F)(\mathbf{y})}$$

$$\text{● Kernel method} \quad F(\mathbf{y}) \approx \underbrace{\sum_{k=0}^{N-1} a_k K(\mathbf{t}_k, \mathbf{y})}_{=: A_N^{\text{ker}}(F)(\mathbf{y})} \quad \text{so that} \quad F(\mathbf{t}_\ell) = A_N^{\text{ker}}(F)(\mathbf{t}_\ell) \quad \forall \ell$$

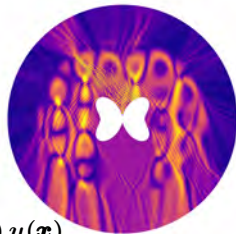
Lattice Rules

Kernel Methods

(SIDE STEP: **Parametric PDEs in UQ**)

Deep Neural Networks

Eg.1 Sound-soft plane-wave scattering



● Helmholtz equation

- Frequency domain: wavenumber k

$$F(\mathbf{x}, t) = \exp(i k t) g(\mathbf{x}) \implies U(\mathbf{x}, t) = \exp(i k t) u(\mathbf{x})$$

$$-\nabla \cdot (\mathbf{A} \nabla u) - k^2 a u = g$$

plus boundary conditions

● Uncertainty Quantification (UQ)

- Random coefficients: $\mathbf{A}(\mathbf{x}, \omega)$ and $a(\mathbf{x}, \omega)$
- Quantities of interest: expected values with respect to ω
- **Forward problem:** how do fluctuations in $\mathbf{A}$ and a affect u
- **Inverse problem:** given observations of u , infer distribution of $\mathbf{A}$ and a

Graham, Kuo, Nuyens, Spence, Sloan (submitted 2025)

Eg.2 Uncertainty in groundwater flow

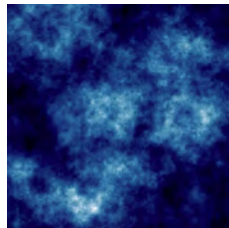
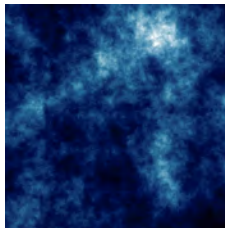
● Risk analysis of radwaste disposal or CO₂ sequestration

● Darcy's law

● Mass conservation law

$$\begin{cases} q(\mathbf{x}) + a(\mathbf{x}, \omega) \nabla p(\mathbf{x}) = f(\mathbf{x}) \\ \nabla \cdot q(\mathbf{x}) = 0 \end{cases}$$

in $D \subset \mathbb{R}^d$, $d = 1, 2, 3$



● Uncertainty in $a(\mathbf{x}, \omega)$ leads to uncertainty in $q(\mathbf{x}, \omega)$ and $p(\mathbf{x}, \omega)$

Graham, Kuo, Nuyens, Scheichl, Sloan (J. Comput. Physics 2011), ...

Eg.3 Criticality problem for nuclear reactors

● PDE eigenvalue problem

$$\nabla \cdot (\underbrace{a(\mathbf{x}, \omega)}_{\text{diffusion}} \nabla u(\mathbf{x})) + \underbrace{b(\mathbf{x}, \omega)}_{\text{absorption}} u(\mathbf{x}) = \lambda \underbrace{c(\mathbf{x}, \omega)}_{\text{fission}} u(\mathbf{x})$$

- Smallest eigenvalue λ_1 measures *criticality* of reactor
- Eigenfunction $u_1(\mathbf{x})$ is the *neutron flux* at the point $\mathbf{x}$

- $\lambda_1 \approx 1 \implies$ operating efficiently
- $\lambda_1 > 1 \implies$ not self-sustaining
- $\lambda_1 < 1 \implies$ supercritical



<https://www.youtube.com/watch?v=xIDytUCRtTA>

Gilbert, Graham, Kuo, Scheichl, Sloan (Numer. Math. 2019), ...

Eg.4 Optimal control under uncertainty

- PDE-constrained robust optimal control: steer u toward target g

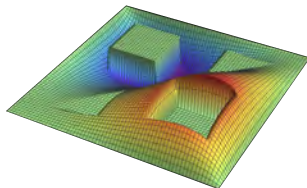
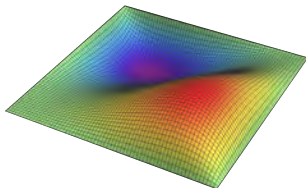
$$\min_{z \in \mathcal{Z}} J(u, z), \quad J(u, z) = \frac{1}{2} \int_U \|u(\cdot, \mathbf{y}) - g\|^2 d\mathbf{y} + \frac{\alpha}{2} \|z\|^2$$

subject to

$$-\nabla \cdot (\mathbf{a}(\mathbf{x}, \mathbf{y}) \nabla u(\mathbf{x}, \mathbf{y})) = z(\mathbf{x}) \quad \mathbf{x} \in D, \mathbf{y} \in U$$

$$u(\mathbf{x}, \mathbf{y}) = 0 \quad \mathbf{x} \in \partial D, \mathbf{y} \in U$$

$$z \in \mathcal{Z} = \{z \in L^2(D) : z_{\min}(\mathbf{x}) \leq z(\mathbf{x}) \leq z_{\max}(\mathbf{x}) \text{ a.e. in } D\}$$



Guth, Kaarnioja, Kuo, Schillings, Sloan (SIAM/ASA J. Uncertain. Quantif. 2021; Numer. Math. 2024), ...

Eg.5 Poisson equation on random domain

● Poisson equation on random domains

$$\Delta u(\mathbf{x}, \omega) = f(\mathbf{x}) \quad \mathbf{x} \in D(\omega)$$

$$u(\mathbf{x}, \omega) = 0 \quad \mathbf{x} \in \partial D(\omega)$$

● Domain mapping method

● Reference domain D_{ref} (Lipschitz)

● Admissible domain $D(\mathbf{y}) = V(D_{\text{ref}}, \mathbf{y})$

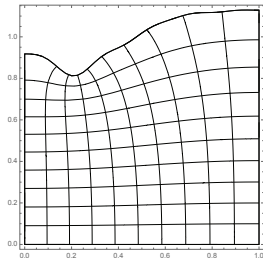
● Perturbation field $V(\mathbf{x}, \mathbf{y})$ and Jacobian matrix $J(\mathbf{x}, \mathbf{y})$

● Variational formulation

$$\int_{D(\mathbf{y})} \nabla u(\mathbf{x}, \mathbf{y}) \cdot \nabla v(\mathbf{x}, \mathbf{y}) \, d\mathbf{x} = \int_{D(\mathbf{y})} f(\mathbf{x}) v(\mathbf{x}, \mathbf{y}) \, d\mathbf{x} \quad \forall v \in H_0^1(D(\mathbf{y}))$$

$$\int_{D_{\text{ref}}} (\underbrace{A(\mathbf{x}, \mathbf{y})}_{\text{transported coefficient}} \nabla \hat{u}(\mathbf{x}, \mathbf{y})) \cdot \nabla \hat{v}(\mathbf{x}) \, d\mathbf{x} = \int_{D_{\text{ref}}} \underbrace{f_{\text{ref}}(\mathbf{x}, \mathbf{y})}_{\text{transported source term}} \hat{v}(\mathbf{x}) \, d\mathbf{x} \quad \forall \hat{v} \in H_0^1(D_{\text{ref}})$$

$$u(\cdot, \mathbf{y}) = \hat{u}(V^{-1}(\cdot, \mathbf{y}), \mathbf{y}) \iff \hat{u}(\cdot, \mathbf{y}) = u(V(\cdot, \mathbf{y}), \mathbf{y})$$



Harbrecht, Peters, Siebenmorgen (2016),

Hakula, Harbrecht, Kaarnioja, Kuo, Sloan (Numer. Maths. 2024), ...

Modeling random fields

● Affine

$$a(\mathbf{x}, \mathbf{y}) = \psi_0(\mathbf{x}) + \sum_{j \geq 1} \mathbf{y}_j \psi_j(\mathbf{x})$$

Cohen, DeVore, Schwab (2010)

$\mathbf{y}_j \in [-\frac{1}{2}, \frac{1}{2}]$ uniform

● Lognormal (Karhunen–Loève expansion, circulant embedding, H-matrix)

$$a(\mathbf{x}, \mathbf{y}) = \exp\left(\psi_0(\mathbf{x}) + \sum_{j \geq 1} \mathbf{y}_j \psi_j(\mathbf{x})\right) \quad \mathbf{y}_j \in \mathbb{R} \text{ normal}$$

Graham, Kuo, Nuyens, Scheichl, Sloan (2011, 2018)

Graham, Kuo, Nichols, Scheichl, Schwab, Sloan (2015)

Feischl, Kuo, Sloan (2018)

Modeling random fields

● Affine

$$a(\mathbf{x}, \mathbf{y}) = \psi_0(\mathbf{x}) + \sum_{j \geq 1} \mathbf{y}_j \psi_j(\mathbf{x})$$

Cohen, DeVore, Schwab (2010)

$\mathbf{y}_j \in [-\frac{1}{2}, \frac{1}{2}]$ uniform

or

$\mathbf{y}_j \in [-1, 1]$ Chebyshev

Adcock, Brugiapaglia, Webster (2022)

● Lognormal (Karhunen–Loève expansion, circulant embedding, H-matrix)

$$a(\mathbf{x}, \mathbf{y}) = \exp\left(\psi_0(\mathbf{x}) + \sum_{j \geq 1} \mathbf{y}_j \psi_j(\mathbf{x})\right) \quad \mathbf{y}_j \in \mathbb{R} \text{ normal}$$

Graham, Kuo, Nuyens, Scheichl, Sloan (2011, 2018)

Graham, Kuo, Nichols, Scheichl, Schwab, Sloan (2015)

Feischl, Kuo, Sloan (2018)

● Periodic

$$a(\mathbf{x}, \mathbf{y}) = \psi_0(\mathbf{x}) + \sum_{j \geq 1} \sin(2\pi \mathbf{y}_j) \psi_j(\mathbf{x})$$

$\mathbf{y}_j \in [0, 1]$ uniform

Kaarnioja, Kuo, Sloan (2020)

Key features of parametric PDE problems

- PDE solution $u(\mathbf{x}, \mathbf{y})$ is infinitely smooth w.r.t. $\mathbf{y}$ (“differentiate” the weak form)
- Truncate to s terms and solve by finite element method $\implies u_{s,h}(\mathbf{x}, \mathbf{y})$
- Interested in
 - **Integration** — expected value of a linear functional $\mathbb{E}_{\mathbf{y}}[\mathcal{G}(u_{s,h}(\cdot, \mathbf{y}))]$
 - **Function approximation** — solution $u_{s,h}(\mathbf{x}^\dagger, \mathbf{y})$ as a function of $\mathbf{y}$
- Assume “ p -summability”: $\sum_{j \geq 1} \|\psi_j\|_\infty^p < \infty, p \in (0, 1)$
Small $p \implies$ faster decay $\implies$ low effective dimension $\implies$ higher order convergence
- Freedom to choose the QMC function space setting for error analysis
 - The function space is not given
 - The smoothness parameter is not given
 - The function space weights are not given
 - The QMC rules are not given
- Goal: best convergence rate with constant independent of dimension s

“The” elliptic PDE example (regularity of the PDE solution)

● Elliptic PDE

$$\begin{aligned} -\nabla \cdot (a(\mathbf{x}, \mathbf{y}) \nabla u(\mathbf{x}, \mathbf{y})) &= f(\mathbf{x}) & \mathbf{x} \in D \subset \mathbb{R}^d, \quad d = 1, 2, 3 \\ u(\mathbf{x}, \mathbf{y}) &= 0 & \mathbf{x} \in \partial D \end{aligned}$$

(a) Affine uniform random field model

Cohen, DeVore, Schwab (2010)

$$\|\partial^\nu u(\cdot, \mathbf{y})\|_{H_0^1(D)} \leq \frac{\|f\|_{H^{-1}(D)}}{a_{\min}} |\nu|! \prod_{j \geq 1} b_j^{\nu_j}, \quad b_j = \frac{\|\psi_j\|_\infty}{a_{\min}}$$

(b) Periodic random field model

Kaarnioja, Kuo, Sloan (2020)

$$\|\partial^\nu u(\cdot, \mathbf{y})\|_{H_0^1(D)} \leq \frac{\|f\|_{H^{-1}(D)}}{a_{\min}} (2\pi)^{|\nu|} \sum_{\mathbf{m} \leq \nu} |\mathbf{m}|! \prod_{j \geq 1} (b_j^{m_j} \mathcal{S}(\nu_j, m_j))$$

Stirling numbers of the second kind

“The” elliptic PDE example (error analysis)

$$\int_{(-\frac{1}{2}, \frac{1}{2})^N} \mathcal{G}(u(\cdot, \mathbf{y})) \, d\mathbf{y} \approx \frac{1}{N} \sum_{k=0}^{N-1} \mathcal{G}(u_{\mathbf{s}, \mathbf{h}}(\cdot, \mathbf{t}_k - \frac{1}{2}))$$

(1) dimension truncation

(2) FE discretization

(3) QMC quadrature

$$(\text{rms}) \text{ error} \leq \underbrace{\text{truncation error}}_{s-2/p+1} + \underbrace{\text{FE error}}_{h^2} + \underbrace{(\text{rms}) \text{ QMC error}}_{n^{-r}}$$

$$r = \begin{cases} \min(\frac{1}{p} - \frac{1}{2}, 1 - \delta) & \text{affine uniform random field + “QMC Setting 1”} \\ & \text{Kuo, Schwab, Sloan (2012)} \\ \frac{1}{p} & \text{affine uniform random field + “QMC Setting 3”} \\ & \text{Dick, Kuo, Le Gia, Nuyens, Schwab, Sloan (2014)} \\ \frac{1}{p} & \text{periodic random field + “QMC Setting 4”} \\ & \text{Kaarnioja, Kuo, Sloan (2020)} \end{cases}$$

Key assumptions:

• $\sum_{j \geq 1} \|\psi_j\|_{\infty}^p < \infty$ for some $p \in (0, 1)$

• first order FEM; $f, \mathcal{G} \in L_2(D)$

• $(\text{rms}) \text{ QMC error} \leq (\text{rms}) \text{ worst case error}_{\gamma} \times \|\mathcal{G}(u_{\mathbf{s}, \mathbf{h}})\|_{\gamma}$

Lattice Rules

Kernel Methods

(SIDE STEP: Parametric PDEs in UQ)

Deep Neural Networks

Data-to-Observables map

- Given data $\mathbf{y} \in \mathbf{Y} \subset \mathbb{R}^s$, compute observable $G(\mathbf{y}) \in \mathbb{R}^{N_{\text{obs}}}$

- A vector of s input parameters to a parametric PDE
- A vector of N_{obs} observables obtained from the solution to a PDE

Examples

- Average of the PDE solution over a subset of the domain $\implies N_{\text{obs}} = 1$
- Vector of PDE solutions at finite element mesh $\implies N_{\text{obs}} = \text{\#FE nodes}$

- Deep Neural Network (DNN) as function approximation**

- $G(\mathbf{y}) \approx G_{\theta}^{[L]}(\mathbf{y})$
- Want to have a small $\|G - G_{\theta}^{[L]}\|_{L_2}$

Deep Neural Network (DNN)

Standard “feed-forward” DNN

$$G_{\theta}^{[L]}(\mathbf{y}) := W_L \sigma(\cdots W_2 \sigma(W_1 \sigma(W_0 \mathbf{y} + \mathbf{v}_0) + \mathbf{v}_1) + \mathbf{v}_2 \cdots) + \mathbf{v}_L$$

- Input: $\mathbf{y} \in Y := [0, 1]^s$, $d_0 = s$ is the dimension of the input vector $\mathbf{y}$
- Output: $G_{\theta}^{[L]}(\mathbf{y}) \in \mathbb{R}^{N_{\text{obs}}}$, $d_{L+1} = N_{\text{obs}}$ is the number of observables
- Unknowns to be “learned”: $\theta := \{(W_{\ell}, \mathbf{v}_{\ell})\}_{\ell=0}^L \in \Theta$ “layer” ℓ
 - W_{ℓ} is a $d_{\ell+1} \times d_{\ell}$ “weight” matrix
 - $\mathbf{v}_{\ell}$ is a $d_{\ell+1} \times 1$ “bias” vector
- Hyperparameters (network architecture): “depth” L , “width” $d_{\ell}, \dots$
- σ is a univariate **activation function**: $\sigma(x) = \max(x, 0)$ (rectified linear unit / ReLU),
 $\sigma(x) = \frac{1}{1+e^{-x}}$ (logistic sigmoid), $\sigma(x) = \tanh(x)$, $\sigma(x) = \frac{x}{1+e^{-x}}$ (swish)

Deep Neural Network (DNN)

(a) Non-periodic DNN

$$G_{\theta}^{[L]}(\mathbf{y}) := W_L \sigma(\cdots W_2 \sigma(W_1 \sigma(W_0 \mathbf{y} + \mathbf{v}_0) + \mathbf{v}_1) + \mathbf{v}_2 \cdots) + \mathbf{v}_L$$

(b) Periodic DNN

Keller, Kuo, Nuyens, Sloan (submitted 2025)

$$G_{\theta}^{[L]}(\mathbf{y}) := W_L \sigma(\cdots W_2 \sigma(W_1 \sigma(W_0 \sin(2\pi \mathbf{y}) + \mathbf{v}_0) + \mathbf{v}_1) + \mathbf{v}_2 \cdots) + \mathbf{v}_L$$

- Input: $\mathbf{y} \in \mathbf{Y} := [0, 1]^s$, $d_0 = s$ is the dimension of the input vector $\mathbf{y}$
- Output: $G_{\theta}^{[L]}(\mathbf{y}) \in \mathbb{R}^{N_{\text{obs}}}$, $d_{L+1} = N_{\text{obs}}$ is the number of observables
- Unknowns to be “learned”: $\theta := \{(W_{\ell}, \mathbf{v}_{\ell})\}_{\ell=0}^L \in \Theta$ “layer” ℓ
 - W_{ℓ} is a $d_{\ell+1} \times d_{\ell}$ “weight” matrix
 - $\mathbf{v}_{\ell}$ is a $d_{\ell+1} \times 1$ “bias” vector
- Hyperparameters (network architecture): “depth” L , “width” $d_{\ell}, \dots$
- σ is a univariate activation function: $\sigma(x) = \max(x, 0)$ (rectified linear unit / ReLU), $\sigma(x) = \frac{1}{1+e^{-x}}$ (logistic sigmoid), $\sigma(x) = \tanh(x)$, $\sigma(x) = \frac{x}{1+e^{-x}}$ (swish)

Error analysis for DNN

- “Training” – to “learn” network parameters

Mishra, Rusch (2021)

Longo, Mishra, Rusch, Schwab (2021)

$$\theta^* := \operatorname{argmin}_{\theta \in \Theta} \left(\mathcal{J}(\theta) + \lambda \mathcal{R}(\theta) \right), \quad \mathcal{J}(\theta) := \frac{1}{N} \sum_{k=0}^{N-1} |G(\mathbf{t}_k) - G_{\theta}^{[L]}(\mathbf{t}_k)|^2$$

Training points

- Generalization error

$$\mathcal{E}_G := \left(\int_{[0,1]^s} |G(\mathbf{y}) - G_{\theta}^{[L]}(\mathbf{y})|^2 d\mathbf{y} \right)^{1/2} = \|G - G_{\theta}^{[L]}\|_{L_2}$$

- Training error

$$\mathcal{E}_T := \left(\frac{1}{N} \sum_{k=0}^{N-1} |G(\mathbf{t}_k) - G_{\theta}^{[L]}(\mathbf{t}_k)|^2 \right)^{1/2}$$

Error analysis for DNN

- “Training” – to “learn” network parameters

Mishra, Rusch (2021)

Longo, Mishra, Rusch, Schwab (2021)

$$\theta^* := \operatorname{argmin}_{\theta \in \Theta} \left(\mathcal{J}(\theta) + \lambda \mathcal{R}(\theta) \right), \quad \mathcal{J}(\theta) := \frac{1}{N} \sum_{k=0}^{N-1} |G(\mathbf{t}_k) - G_{\theta}^{[L]}(\mathbf{t}_k)|^2$$

Training points

- Generalization error

$$\mathcal{E}_G := \left(\int_{[0,1]^s} |G(\mathbf{y}) - G_{\theta}^{[L]}(\mathbf{y})|^2 d\mathbf{y} \right)^{1/2} = \|G - G_{\theta}^{[L]}\|_{L_2}$$

- Training error

$$\mathcal{E}_T := \left(\frac{1}{N} \sum_{k=0}^{N-1} |G(\mathbf{t}_k) - G_{\theta}^{[L]}(\mathbf{t}_k)|^2 \right)^{1/2}$$

- Generalization gap

Quadrature error for $F(\mathbf{y}) := (G(\mathbf{y}) - G_{\theta}^{[L]}(\mathbf{y}))^2$

$$|\mathcal{E}_G - \mathcal{E}_T| \leq \sqrt{|\mathcal{E}_G^2 - \mathcal{E}_T^2|} = \left| \int_{[0,1]^s} (G(\mathbf{y}) - G_{\theta}^{[L]}(\mathbf{y}))^2 d\mathbf{y} - \frac{1}{N} \sum_{k=0}^{N-1} (G(\mathbf{t}_k) - G_{\theta}^{[L]}(\mathbf{t}_k))^2 \right|^{\frac{1}{2}}$$

Error analysis for DNN

- “Training” – to “learn” network parameters

Mishra, Rusch (2021)

Longo, Mishra, Rusch, Schwab (2021)

$$\theta^* := \operatorname{argmin}_{\theta \in \Theta} \left(\mathcal{J}(\theta) + \lambda \mathcal{R}(\theta) \right), \quad \mathcal{J}(\theta) := \frac{1}{N} \sum_{k=0}^{N-1} |G(\mathbf{t}_k) - G_{\theta}^{[L]}(\mathbf{t}_k)|^2$$

Training points

- Generalization error

$$\mathcal{E}_G := \left(\int_{[0,1]^s} |G(\mathbf{y}) - G_{\theta}^{[L]}(\mathbf{y})|^2 d\mathbf{y} \right)^{1/2} = \|G - G_{\theta}^{[L]}\|_{L_2}$$

- Training error

$$\mathcal{E}_T := \left(\frac{1}{N} \sum_{k=0}^{N-1} |G(\mathbf{t}_k) - G_{\theta}^{[L]}(\mathbf{t}_k)|^2 \right)^{1/2}$$

- Generalization gap

Quadrature error for $F(\mathbf{y}) := (G(\mathbf{y}) - G_{\theta}^{[L]}(\mathbf{y}))^2$

$$|\mathcal{E}_G - \mathcal{E}_T| \leq \sqrt{|\mathcal{E}_G^2 - \mathcal{E}_T^2|} = \left| \int_{[0,1]^s} (G(\mathbf{y}) - G_{\theta}^{[L]}(\mathbf{y}))^2 d\mathbf{y} - \frac{1}{N} \sum_{k=0}^{N-1} (G(\mathbf{t}_k) - G_{\theta}^{[L]}(\mathbf{t}_k))^2 \right|^{\frac{1}{2}}$$

- Theoretical bound for generalization error

Need to know the regularity of $(G_{\theta}^{[L]}(\mathbf{y}))^2$

$$\mathcal{E}_G \leq \mathcal{E}_T + |\mathcal{E}_G - \mathcal{E}_T| \leq \mathcal{E}_T + \left(e^{\text{wor-int}} \times \|(G - G_{\theta}^{[L]})^2\|_{\mathcal{W}_{\alpha,\gamma}} \right)^{1/2}$$

(a) **Non-periodic DNN**
$$\begin{cases} G_{\theta}^{[0]}(\mathbf{y}) := \mathbf{W}_0 \mathbf{y} + \mathbf{v}_0 \\ G_{\theta}^{[\ell]}(\mathbf{y}) := \mathbf{W}_{\ell} \sigma(G_{\theta}^{[\ell-1]}(\mathbf{y})) + \mathbf{v}_{\ell} \quad \text{for } \ell \geq 1 \end{cases}$$

(b) **Periodic DNN**
$$\begin{cases} G_{\theta}^{[0]}(\mathbf{y}) := \mathbf{W}_0 \sin(2\pi \mathbf{y}) + \mathbf{v}_0 \\ G_{\theta}^{[\ell]}(\mathbf{y}) := \mathbf{W}_{\ell} \sigma(G_{\theta}^{[\ell-1]}(\mathbf{y})) + \mathbf{v}_{\ell} \quad \text{for } \ell \geq 1 \end{cases}$$

- Assumption 1 The columns of the matrix $\mathbf{W}_0$ have bounded vector ∞ -norms
 $\|\mathbf{W}_{0,:j}\|_{\infty} \leq \beta_j \quad \text{for all } j \in \{1 : s\}$

Cf. Longo, Mishra, Rusch, Schwab (2021)

- Assumption 2 The matrices $\mathbf{W}_{\ell}$ for $\ell \geq 1$ have bounded matrix ∞ -norms
 $\|\mathbf{W}_{\ell}\|_{\infty} \leq R_{\ell} \quad \text{for all } \ell \geq 1$

Cf. Longo, Mishra, Rusch, Schwab (2021)

- Assumption 3 The activation function σ is smooth and satisfies
 $\|\sigma^{(n)}\|_{\infty} \leq A_n \quad \text{for all } n \geq 1$

Regularity bounds for DNNs

Theorem 1 (Keller, Kuo, Nuyens, Sloan) For any multiindex $\nu \neq \mathbf{0}$

(a) Non-periodic DNN

$$|\partial^\nu G_\theta^{[L]}(\mathbf{y})| \leq R_L \Gamma_{|\nu|}^{[L]} \prod_{j=1}^s \beta_j^{\nu_j}$$

(b) Periodic DNN

$$|\partial^\nu G_\theta^{[L]}(\mathbf{y})| \leq R_L \sum_{\mathbf{m} \leq \nu} \Gamma_{|\mathbf{m}|}^{[L]} \prod_{j=1}^s (\beta_j^{m_j} \mathcal{S}(\nu_j, m_j))$$

Stirling numbers of the second kind

• The sequence $\Gamma_n^{[\ell]}$ is defined recursively by

$$\textcircled{1} \Gamma_n^{[1]} := A_n \text{ for } n \geq 1, \quad \Gamma_n^{[\ell]} := \sum_{\lambda=1}^n A_\lambda R_{\ell-1}^\lambda \mathbb{B}_{n,\lambda}^{[\ell-1]} \text{ for } \ell \geq 2 \text{ and } n \geq 1$$

$$\textcircled{2} \mathbb{B}_{n,1}^{[\ell]} := \Gamma_n^{[\ell]} \text{ for } n \geq 1, \quad \mathbb{B}_{n,\lambda}^{[\ell]} := \sum_{i=\lambda-1}^{n-1} \binom{n-1}{i} \Gamma_{n-i}^{[\ell]} \mathbb{B}_{i,\lambda-1}^{[\ell]} \text{ for } \ell \geq 1 \text{ and } n \geq \lambda \geq 2$$

Regularity bounds for DNNs

Theorem 2 (Keller, Kuo, Nuyens, Sloan) For any multiindex ν and $A_n = \xi \tau^n n!$

(a) Non-periodic DNN

$$|\partial^\nu G_\theta^{[L]}(\mathbf{y})| \leq C_L |\nu|! \prod_{j=1}^s (S_L \beta_j)^{\nu_j}, \quad S_L := \tau \sum_{\ell=0}^{L-1} \prod_{k=1}^{\ell} (\xi \tau R_k)$$

(b) Periodic DNN

$$|\partial^\nu G_\theta^{[L]}(\mathbf{y})| \leq C_L \sum_{\mathbf{m} \leq \nu} |\mathbf{m}|! \prod_{j=1}^s ((S_L \beta_j)^{m_j} \mathcal{S}(\nu_j, m_j))$$

Common activation functions

$$\left\{ \begin{array}{ll} \text{sigmoid: } \sigma(x) = \frac{1}{1+e^{-x}} & \sigma^{(n)}(x) \leq A_n = n! \quad (\xi = 1, \tau = 1) \\ \text{tanh: } \sigma(x) = \frac{e^x - e^{-x}}{e^x + e^{-x}} & \sigma^{(n)}(x) \leq A_n = 2^n n! \quad (\xi = 1, \tau = 2) \\ \text{swish: } \sigma(x) = \frac{x}{1+e^{-x}} & \sigma^{(n)}(x) \leq A_n = 1.1 n! \quad (\xi = 1.1, \tau = 1) \\ \text{swish}_c: \sigma(x) = \frac{x}{1+e^{-cx}} & \sigma^{(n)}(x) \leq A_n = \frac{1.1}{c} c^n n! \quad (\xi = \frac{1.1}{c}, \tau = c) \end{array} \right.$$

$S_L = L$ for sigmoid, when $R_\ell \leq 1$ for all $\ell \in \{1 : L-1\}$

Norm bounds for DNNs

Theorem 3 (Keller, Kuo, Nuyens, Sloan)

(a) Non-periodic DNN

$$\|(G_{\theta}^{[L]})^2\|_{\mathcal{W}_1, \gamma}^2 \leq C_L^4 \sum_{u \subseteq \{1:s\}} \frac{1}{\gamma_u} \left((|u| + 1)! \prod_{j \in u} (S_L \beta_j) \right)^2$$

(b) Periodic DNN

$$\|(G_{\theta}^{[L]})^2\|_{\mathcal{W}_{\alpha}, \gamma}^2 \leq C_L^4 \sum_{u \subseteq \{1:s\}} \frac{(2\pi)^{2\alpha|u|}}{\gamma_u} \left(\sum_{\mathbf{m}_u \leq \alpha_u} (|\mathbf{m}_u| + 1)! \prod_{j \in u} ((S_L \beta_j)^{m_j} \mathcal{S}(\alpha, m_j)) \right)^2$$

Norm bounds for DNNs

Theorem 3 (Keller, Kuo, Nuyens, Sloan)

(a) Non-periodic DNN + PDE with affine uniform random field

$$\begin{aligned}\|(G_\theta^{[L]})^2\|_{\mathcal{W}_{1,\gamma}}^2 &\leq C_L^4 \sum_{u \subseteq \{1:s\}} \frac{1}{\gamma_u} \left((|u| + 1)! \prod_{j \in u} (S_L \beta_j) \right)^2 \\ \|G^2\|_{\mathcal{W}_{1,\gamma}}^2 &\leq C^4 \sum_{u \subseteq \{1:s\}} \frac{1}{\gamma_u} \left((|u| + 1)! \prod_{j \in u} b_j \right)^2, \quad b_j := \frac{\|\psi_j\|_\infty}{a_{\min}}\end{aligned}$$

(b) Periodic DNN + PDE with periodic random field

$$\begin{aligned}\|(G_\theta^{[L]})^2\|_{\mathcal{W}_{\alpha,\gamma}}^2 &\leq C_L^4 \sum_{u \subseteq \{1:s\}} \frac{(2\pi)^{2\alpha|u|}}{\gamma_u} \left(\sum_{\mathbf{m}_u \leq \alpha_u} (|\mathbf{m}_u| + 1)! \prod_{j \in u} ((S_L \beta_j)^{m_j} \mathcal{S}(\alpha, m_j)) \right)^2 \\ \|G^2\|_{\mathcal{W}_{\alpha,\gamma}}^2 &\leq C^4 \sum_{u \subseteq \{1:s\}} \frac{(2\pi)^{2\alpha|u|}}{\gamma_u} \left(\sum_{\mathbf{m}_u \leq \alpha_u} (|\mathbf{m}_u| + 1)! \prod_{j \in u} (b_j^{m_j} \mathcal{S}(\alpha, m_j)) \right)^2\end{aligned}$$

Norm bounds for the generalization gap

Theorem 4 (Keller, Kuo, Nuyens, Sloan)

- restrict the elements of $\mathbf{W}_\ell$ so that $R_\ell \leq \rho$ for all $\ell \in \{1 : L - 1\}$
- restrict the elements of $\mathbf{W}_0$ so that $\mathbf{S}_L \beta_j \leq b_j := \frac{\|\psi_j\|_\infty}{a_{\min}}$ for all $j \in \{1 : s\}$
- restrict the elements of $\mathbf{W}_L$ and $\mathbf{v}_L$ so that $C_L \leq C$

(a) Non-periodic DNN + PDE with affine uniform random field

$$\|(G - G_\theta^{[L]})^2\|_{\mathcal{W}_1, \gamma}^2 \leq 2 C^4 \sum_{u \subseteq \{1:s\}} \frac{1}{\gamma_u} \left((|u| + 1)! \prod_{j \in u} b_j \right)^2$$

(b) Periodic DNN + PDE with periodic random field

$$\|(G - G_\theta^{[L]})^2\|_{\mathcal{W}_\alpha, \gamma}^2 \leq 2 C^4 \sum_{u \subseteq \{1:s\}} \frac{(2\pi)^{2\alpha|u|}}{\gamma_u} \left(\sum_{\mathbf{m}_u \leq \alpha_u} (|\mathbf{m}_u| + 1)! \prod_{j \in u} (b_j^{m_j} \mathcal{S}(\alpha, m_j)) \right)^2$$

Norm bounds for the generalization gap

Theorem 4 (Keller, Kuo, Nuyens, Sloan)

- restrict the elements of $\mathbf{W}_\ell$ so that $R_\ell \leq \rho$ for all $\ell \in \{1 : L - 1\}$
- restrict the elements of $\mathbf{W}_0$ so that $\mathbf{S}_L \beta_j \leq b_j := \frac{\|\psi_j\|_\infty}{a_{\min}}$ for all $j \in \{1 : s\}$
- restrict the elements of $\mathbf{W}_L$ and $\mathbf{v}_L$ so that $C_L \leq C$

(a) Non-periodic DNN + PDE with affine uniform random field

$$\|(G - G_\theta^{[L]})^2\|_{\mathcal{W}_{1,\gamma}}^2 \leq 2C^4 \sum_{\mathbf{u} \subseteq \{1:s\}} \frac{1}{\gamma_{\mathbf{u}}} \left((|\mathbf{u}| + 1)! \prod_{j \in \mathbf{u}} b_j \right)^2 \quad e^{\text{wor-int}} = \mathcal{O}(N^{-1+\delta})$$

(b) Periodic DNN + PDE with periodic random field

$$\|(G - G_\theta^{[L]})^2\|_{\mathcal{W}_{\alpha,\gamma}}^2 \leq 2C^4 \sum_{\mathbf{u} \subseteq \{1:s\}} \frac{(2\pi)^{2\alpha|\mathbf{u}|}}{\gamma_{\mathbf{u}}} \left(\sum_{\mathbf{m}_{\mathbf{u}} \leq \alpha_{\mathbf{u}}} (|\mathbf{m}_{\mathbf{u}}| + 1)! \prod_{j \in \mathbf{u}} (b_j^{m_j} \mathcal{S}(\alpha, m_j)) \right)^2$$

$$e^{\text{wor-int}} = \mathcal{O}(N^{-\alpha+\delta})$$

● Theoretical bound for generalization error

$$\mathcal{E}_G \leq \mathcal{E}_T + |\mathcal{E}_G - \mathcal{E}_T| \leq \mathcal{E}_T + \left(e^{\text{wor-int}} \times \|(G - G_\theta^{[L]})^2\|_{\mathcal{W}_{\alpha,\gamma}} \right)^{1/2} \leq \text{tol} + \mathcal{O}(N^{-r/2})$$

Numerical experiments

- Algebraic equation $a(\mathbf{y}) G(\mathbf{y}) = 1$ mimicking parametric PDE, “decay” $q = 2.5$,

$$G(\mathbf{y}) = \frac{1}{a(\mathbf{y})}, \quad a(\mathbf{y}) = 1 + \sum_{j=1}^s \sin(2\pi \mathbf{y}_j) \psi_j, \quad \psi_j = \frac{0.5}{j^q}, \quad b_j = \frac{0.5}{(1 - 0.5 \zeta(q)) j^q}$$

- Hyperparameters: 1. $L = 3$, $N_{\text{obs}} = 1$, $d_\ell = 32$, $s = 50$ (3777 parameters)
2. $L = 12$, $N_{\text{obs}} = 1$, $d_\ell = 30$, $s = 50$ (11791 parameters)

- Tailored regularization** to “encourage” $\beta_j \leq \frac{b_j}{L}$ for all $j \in \{1 : s\}$

$$\theta^* = \operatorname{argmin}_{\theta \in \Theta} \left(\frac{1}{N} \sum_{k=0}^{N-1} (G(\mathbf{t}_k) - G_\theta^{[L]}(\mathbf{t}_k))^2 + \lambda \|\theta\|_2^2 + \lambda_1 \mathcal{R}_1(\theta) \right)$$

$$\mathcal{R}_1(\theta) = \frac{1}{s} \sum_{j=1}^s \frac{1}{d_1} \sum_{p=1}^{d_1} \left(w_{0,p,j}^2 \frac{L^2}{b_j^2} \right)^{m/2}, \quad m = 6$$

- PyTorch. Full-batch ADAM. Random Glorot initialization.
- Learning rate (step size) 10^{-4} . Regularization parameters $\lambda = 10^{-8}$, $\lambda_1 = 10^{-8}$.
- Stop when training error reaches $\text{tol} = 10^{-3}$ or maximum 40000 epochs (steps).
- Off-the-shelf embedded lattice point set with $N = 2^5, \dots, 2^{12}$ points.
- Different point set with $M = 2^{15}$ points to estimate the generalization error.

1. $L = 3, N_{\text{obs}} = 1, d_\ell = 32, s = 50$ (3777 parameters)

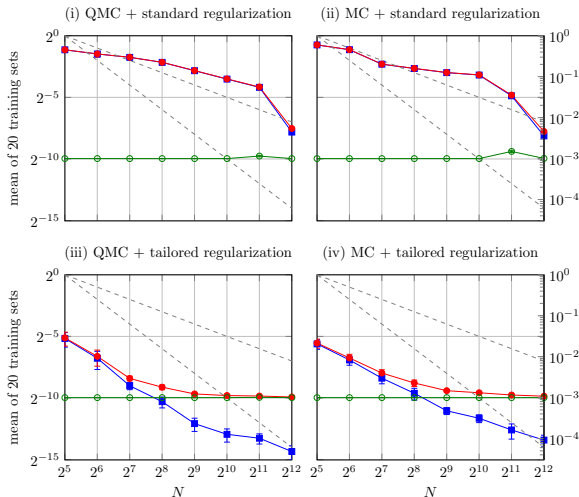
$$\tilde{\mathcal{E}}_G \leq \mathcal{E}_T + |\tilde{\mathcal{E}}_G - \mathcal{E}_T| \leq \text{tol} + \mathcal{O}(N^{-r/2})$$

$\mathcal{E}_T$ training error

$\tilde{\mathcal{E}}_G$ (estimated) generalization error

$|\tilde{\mathcal{E}}_G - \mathcal{E}_T|$ (estimated) generalization gap

sigmoid activation function



1. $L = 3, N_{\text{obs}} = 1, d_\ell = 32, s = 50$ (3777 parameters)

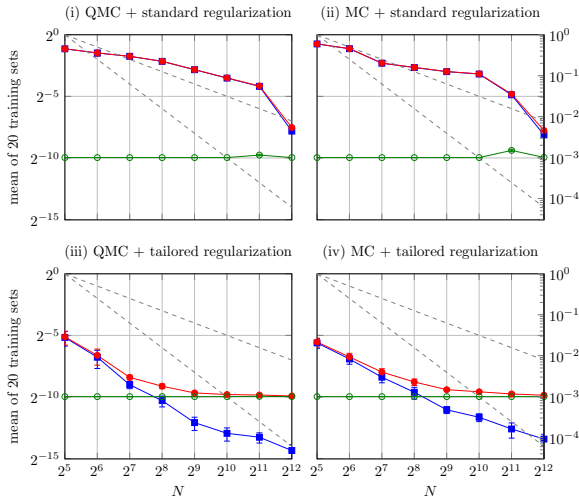
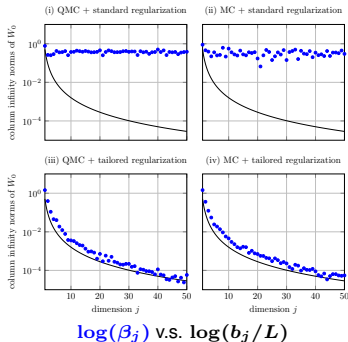
$$\tilde{\mathcal{E}}_G \leq \mathcal{E}_T + |\tilde{\mathcal{E}}_G - \mathcal{E}_T| \leq \text{tol} + \mathcal{O}(N^{-r/2})$$

$\mathcal{E}_T$ training error

$\tilde{\mathcal{E}}_G$ (estimated) generalization error

$|\tilde{\mathcal{E}}_G - \mathcal{E}_T|$ (estimated) generalization gap

sigmoid activation function



2. $L = 12, N_{\text{obs}} = 1, d_\ell = 30, s = 50$ (11791 parameters)

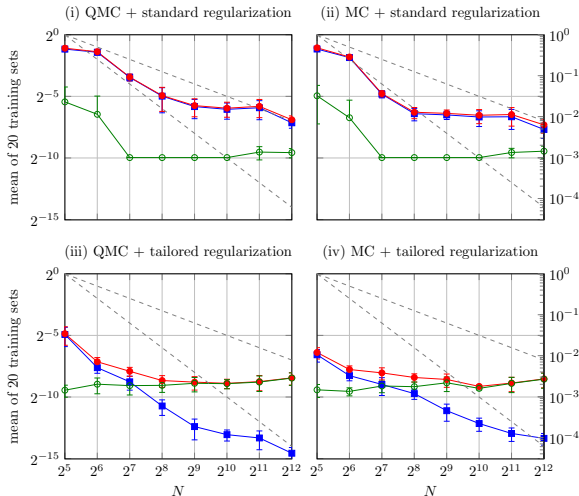
$$\tilde{\mathcal{E}}_G \leq \mathcal{E}_T + |\tilde{\mathcal{E}}_G - \mathcal{E}_T| \leq \text{tol} + \mathcal{O}(N^{-r/2})$$

$\mathcal{E}_T$ training error

$\tilde{\mathcal{E}}_G$ (estimated) generalization error

$|\tilde{\mathcal{E}}_G - \mathcal{E}_T|$ (estimated) generalization gap

sigmoid activation function



2. $L = 12$, $N_{\text{obs}} = 1$, $d_\ell = 30$, $s = 50$ (11791 parameters)

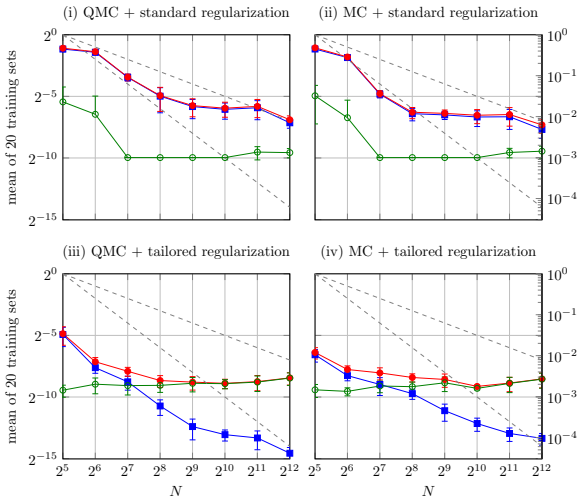
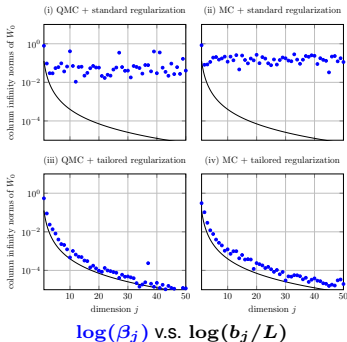
$$\tilde{\mathcal{E}}_G \leq \mathcal{E}_T + |\tilde{\mathcal{E}}_G - \mathcal{E}_T| \leq \text{tol} + \mathcal{O}(N^{-\tau/2})$$

$\mathcal{E}_T$ training error

$\tilde{\mathcal{E}}_G$ (estimated) generalization error

$|\tilde{\mathcal{E}}_G - \mathcal{E}_T|$ (estimated) generalization gap

sigmoid activation function

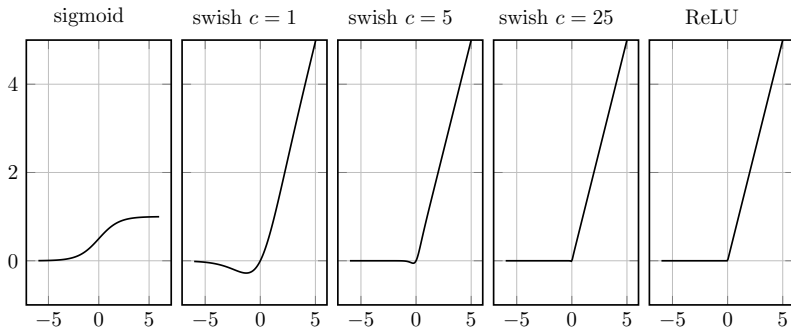


Different activation functions

$$\text{sigmoid}(x) = \frac{1}{1 + e^{-x}}$$

$$\text{swish}_{\textcolor{red}{c}}(x) = \frac{x}{1 + e^{-\textcolor{red}{c}x}}$$

$$\text{ReLU}(x) = \max(x, 0)$$



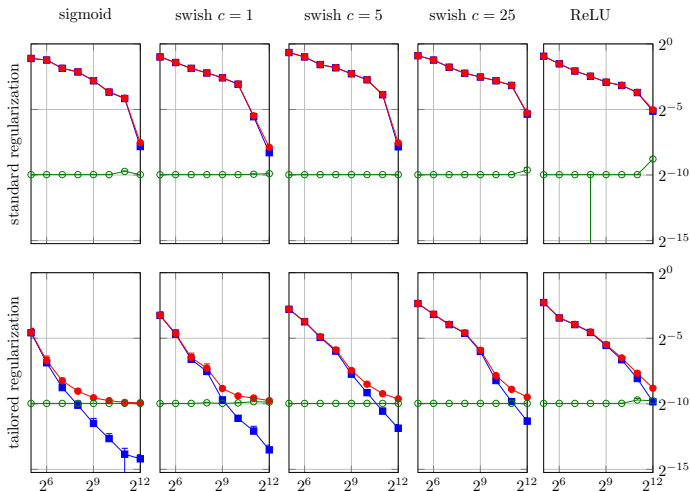
1. $L = 3, N_{\text{obs}} = 1, d_\ell = 32, s = 50$ (3777 parameters)

$$\tilde{\mathcal{E}}_G \leq \mathcal{E}_T + |\tilde{\mathcal{E}}_G - \mathcal{E}_T| \leq \text{tol} + \mathcal{O}(N^{-r/2})$$

$\mathcal{E}_T$ training error

$\tilde{\mathcal{E}}_G$ (estimated) generalization error

$|\tilde{\mathcal{E}}_G - \mathcal{E}_T|$ (estimated) generalization gap



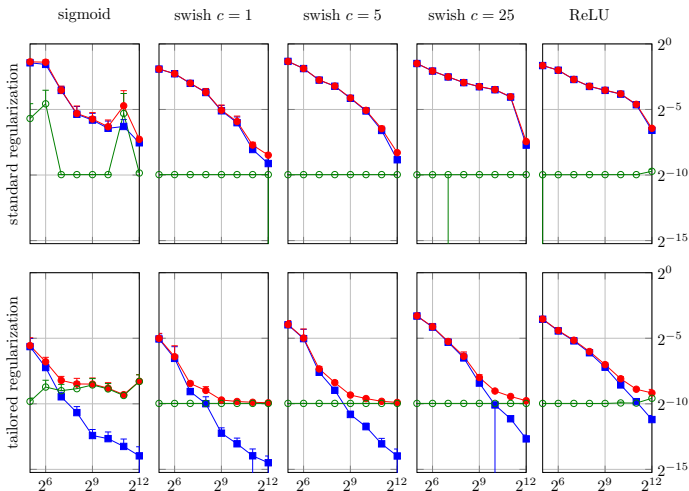
2. $L = 12, N_{\text{obs}} = 1, d_\ell = 30, s = 50$ (11791 parameters)

$$\tilde{\mathcal{E}}_G \leq \mathcal{E}_T + |\tilde{\mathcal{E}}_G - \mathcal{E}_T| \leq \text{tol} + \mathcal{O}(N^{-r/2})$$

$\mathcal{E}_T$ training error

$\tilde{\mathcal{E}}_G$ (estimated) generalization error

$|\tilde{\mathcal{E}}_G - \mathcal{E}_T|$ (estimated) generalization gap



Summary

- When can we use QMC?
 - High dimensional integration
 - Multivariate function approximation
 - Density estimation (for CDF and PDF, using “preintegration”)
 - Training points for DNNs
 - ...
- QMC have **higher convergence rates with respect to N** for smooth functions
- QMC error bounds can be **independent of dimension s** in “weighted” spaces
- QMC points can be **tailored to the applications** (fast CBC construction)

Summary

- When can we use QMC?
 - High dimensional integration
 - Multivariate function approximation
 - Density estimation (for CDF and PDF, using “preintegration”)
 - Training points for DNNs
 - ...
- QMC have **higher convergence rates with respect to N** for smooth functions
- QMC error bounds can be **independent of dimension s** in “weighted” spaces
- QMC points can be **tailored to the applications** (fast CBC construction)
- QMC for DNNs
 - **Explicit derivative bounds** on the non-periodic and periodic DNNs
 - **Tailor-constructed QMC** rules to ensure good convergence rates, with error bound independent of dimension
 - **Tailored regularization** during training to “encourage” network parameters to match the derivative features of the target function
 - Promising numerical results on simple test function