

ERG2026, University of Sussex

Solving Functional Renormalization Group Equations with Neural Networks

Phys. Rev. D **114**, 036026 (2026)

Yang-yang Tan



UTokyo



知の物理学 研究センター
Institute for Physics of Intelligence

IPI, The University of Tokyo



with **Wei-jie Fu**, **Lianyi He** and **Lingxiao Wang**

September 4, 2026

Introduction

Wetterich equation

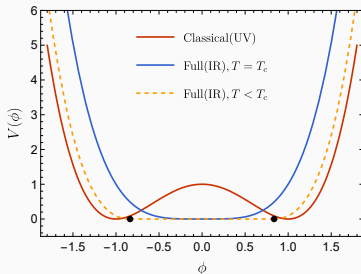
$$\partial_t \Gamma_k = \frac{1}{2} \text{Tr} \left[\partial_t R_k \left(\Gamma_k^{(2)} + R_k \right)^{-1} \right]$$

LPA flow of $V_k(\rho)$, a stiff PDE

Taylor, grids, Chebyshev, DG/FEM, FV

Bottleneck

stiffness, curse of dimensionality



$$\Gamma_{k=\Lambda} = S \longrightarrow \Gamma_{k \rightarrow 0} = \Gamma$$

Wetterich: PLB 301, 90 (1993). **Taylor:** Litim, NPB 631, 128 (2002). **Chebyshev:** Borchardt, Knorr, PRD 91, 105011 (2015); PRD 94, 025027 (2016); Chen, Wen, Fu, PRD 104, 054009 (2021).

DG / FEM: Grossi, Wink, SciPost Phys. Core 6, 071 (2023); Sattler, Pawłowski (DiFFRG), arXiv:2412.13043. **FV:** Zorbach, Koenigstein, Braun, arXiv:2412.16053.

Introduction

Wetterich equation

$$\partial_t \Gamma_k = \frac{1}{2} \text{Tr} \left[\partial_t R_k \left(\Gamma_k^{(2)} + R_k \right)^{-1} \right]$$

LPA flow of $V_k(\rho)$, a stiff PDE
Taylor, grids, Chebyshev, DG/FEM, FV

Bottleneck

stiffness, curse of dimensionality

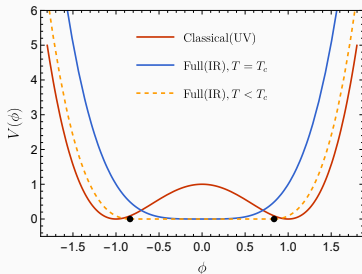
This talk

Neural Networks for the LPA flow

$V'_k(\rho)$ = a neural network

loss = flow-equation residual

learn only the **finite- N correction**



$$\Gamma_{k=\Lambda} = S \longrightarrow \Gamma_{k \rightarrow 0} = \Gamma$$

PINN for fRG on a lattice

Yokota, PRB 109, 214205 (2024)

Takeru's talk

Wetterich: PLB 301, 90 (1993). **Taylor:** Litim, NPB 631, 128 (2002). **Chebyshev:** Borchardt, Knorr, PRD 91, 105011 (2015); PRD 94, 025027 (2016); Chen, Wen, Fu, PRD 104, 054009 (2021).

DG / FEM: Grossi, Wink, SciPost Phys. Core 6, 071 (2023); Sattler, Pawłowski (DiFFRG), arXiv:2412.13043. **FV:** Zorbach, Koenigstein, Braun, arXiv:2412.16053.

Why a network instead of a grid?

Continuous and differentiable

smooth in $(t, \hat{\rho})$, no interpolation

Flexibility

boundary condition,
physics-driven constraints

Transfer learning

one model seeds a scan in T, μ, N

Physics-driven learning: Aarts, Fukushima, Hatsuda, Ipp, Shi, Wang, Zhou, Nat. Rev. Phys. 7, 154 (2025). **Vertex expansion in QCD:** Mitter, Pawłowski, Strodthoff, PRD 91, 054035 (2015); Fu, Huang, Pawłowski, Tan, Zhou, PRD 112, 054047 (2025). **Operator learning:** DeepONet, Lu et al., Nat. Mach. Intell. 3, 218 (2021); FNO, Li et al., arXiv:2010.08895. **Physics-informed kernels:** Ihssen, Pawłowski, Ann. Phys. 481, 170177 (2025); arXiv:2507.13011; Ihssen, Kapust, Pawłowski, arXiv:2510.26678; arXiv:2603.03159.

Why a network instead of a grid?

Continuous and differentiable

smooth in $(t, \hat{\rho})$, no interpolation

Flexibility

boundary condition,
physics-driven constraints

Transfer learning

one model seeds a scan in T, μ, N

No curse of dimensionality

grid cost $\sim N^{n_{\text{fields/momenta}}}$,
a network only adds inputs

Examples

multi-field V : $2 + 1$ flavour QCD
vertex expansion: $\Gamma^{(4)}(p_1, p_2, p_3)$,
 $32^6 \sim 10^9$ grid points

Physics-driven learning: Aarts, Fukushima, Hatsuda, Ipp, Shi, Wang, Zhou, Nat. Rev. Phys. 7, 154 (2025). **Vertex expansion in QCD:** Mitter, Pawłowski, Strodthoff, PRD 91, 054035 (2015); Fu, Huang, Pawłowski, Tan, Zhou, PRD 112, 054047 (2025). **Operator learning:** DeepONet, Lu et al., Nat. Mach. Intell. 3, 218 (2021); FNO, Li et al., arXiv:2010.08895. **Physics-informed kernels:** Ihssen, Pawłowski, Ann. Phys. 481, 170177 (2025); arXiv:2507.13011; Ihssen, Kapust, Pawłowski, arXiv:2510.26678; arXiv:2603.03159.

Why a network instead of a grid?

Continuous and differentiable

smooth in $(t, \hat{\rho})$, no interpolation

Flexibility

boundary condition,
physics-driven constraints

Transfer learning

one model seeds a scan in T, μ, N

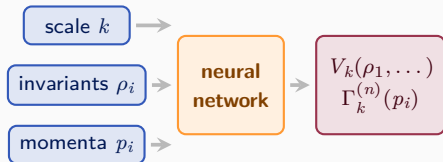
No curse of dimensionality

grid cost $\sim N^{n_{\text{fields/moments}}}$,
a network only adds inputs

Examples

multi-field V : 2 + 1 flavour QCD

vertex expansion: $\Gamma^{(4)}(p_1, p_2, p_3)$,
 $32^6 \sim 10^9$ grid points



universal approximators, strong representation power

ML for RG, RG for ML: **Jan's talk** physics-informed RG: **Friederike's talk**

Physics-driven learning: Aarts, Fukushima, Hatsuda, Ipp, Shi, Wang, Zhou, Nat. Rev. Phys. 7, 154 (2025). **Vertex expansion in QCD:** Mitter, Pawłowski, Strodthoff, PRD 91, 054035 (2015); Fu, Huang, Pawłowski, Tan, Zhou, PRD 112, 054047 (2025). **Operator learning:** DeepONet, Lu et al., Nat. Mach. Intell. 3, 218 (2021); FNO, Li et al., arXiv:2010.08895. **Physics-informed kernels:** Ihssen, Pawłowski, Ann. Phys. 481, 170177 (2025); arXiv:2507.13011; Ihssen, Kapust, Pawłowski, arXiv:2510.26678; arXiv:2603.03159.

The flow is stiff

$O(N)$ model, LPA (Model A)

$$\partial_t V_k = \mathcal{C} T k^d \left[\frac{1}{1 + \bar{m}_{\sigma,k}^2} + \frac{N-1}{1 + \bar{m}_{\pi,k}^2} \right]$$

$$\bar{m}_{\sigma,k}^2 = \frac{V'_k + 2\rho V''_k}{k^2}, \quad \bar{m}_{\pi,k}^2 = \frac{V'_k}{k^2}$$

Convexity restoration

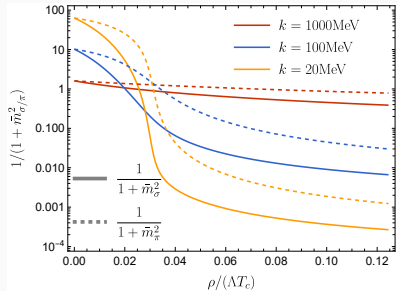
$$1 + \bar{m}_{\sigma/\pi}^2 \geq 0 \text{ at all } k$$

$$\rho < \rho_0: V'_k \rightarrow -k^2, \text{ so } 1 + \bar{m}^2 \rightarrow 0^+$$

Stiffness in a narrow field window

propagators span 10^{-4} – 10^2 ,

$\partial_t V'$ spans 10^2 – 10^7 ($k=20$ MeV)



$O(4)$, $T = 100$ MeV, broken phase.

$\rho = \frac{1}{2} \phi_a \phi_a$, $t = \ln(k/\Lambda)$, $\hat{\rho} = \rho/(\Lambda T_c)$. **Model A:** Hohenberg, Halperin, RMP 49, 435 (1977). **Flow equation from real-time fRG:** Tan, Chen, Fu, SciPost Phys. 12, 026 (2022); Chen, Tan, Fu, PRD 109, 094044 (2024).

Learn only the finite- N correction

Decomposition

$$\hat{V}'(t, \hat{\rho}) = \underbrace{\hat{V}'_{\text{LN}}(t, \hat{\rho})}_{\text{analytic}} + \underbrace{\Delta \hat{V}'_{\theta}(t, \hat{\rho})}_{\text{network}}$$

Large- N flow

$$\partial_t V_k = \mathcal{C} T k^d \frac{N-1}{1 + \bar{m}_{\pi,k}^2}$$

solvable by characteristics:

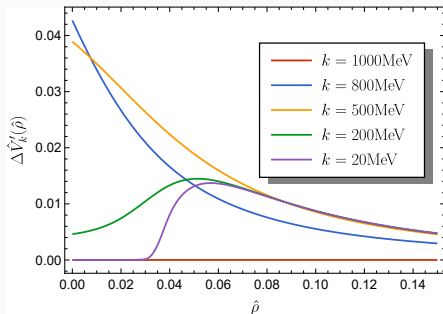
V' constant along them,

ρ_k in closed form (${}_2F_1$)

Why it helps

large N restores convexity;

the rest is **smooth and localized**



Learned $\Delta \hat{V}'_k(\hat{\rho})$, $O(4)$, $T = 100$ MeV.

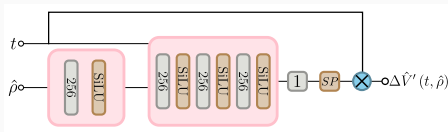
Zero at $k=\Lambda$, largest at $k \sim 0.2\text{--}0.6$ GeV.

Initial and boundary conditions

$$\Delta \hat{V}'_{\theta}(0, \hat{\rho}) = 0$$

$$\Delta \hat{V}'_{\theta}(t, 0), \Delta \hat{V}'_{\theta}(t, \hat{\rho}_{\text{max}}): \text{ not specified}$$

Solving the flow equation as an optimization problem

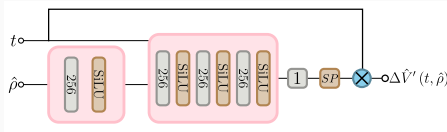


inputs $(t, \hat{\rho}) \rightarrow$ simple nonlinear layers $\rightarrow$ output $\Delta \hat{V}'_{\theta}$

- Residual** insert $\hat{V}'_{LN} + \Delta \hat{V}'_{\theta}$ into the flow equation on the grid:
 $\mathcal{R}_{ij}(\theta) = \partial_t \hat{V}'_{\theta} - \mathcal{F}[\hat{V}'_{\theta}, \hat{V}''_{\theta}, \hat{V}'''_{\theta}]|_{(\hat{\rho}_i, t_j)}$, zero for the exact solution

Details: 3-point finite differences on a 201×501 grid, 3–4 $\times$ less memory than autodiff; Adam; float32. **PINN**: Raissi, Perdikaris, Karniadakis, JCP 378, 686 (2019); Karniadakis et al., Nat. Rev. Phys. 3, 422 (2021). **PINNs for FRG / FDE / DSE**: Yokota, PRB 109, 214205 (2024); Miyagawa, Yokota, NeurIPS 2024; Terin, SciPost Phys. Core 8, 054 (2025). **FD derivatives in PINNs**: Chiu et al., CMAME 395, 114909 (2022).

Solving the flow equation as an optimization problem

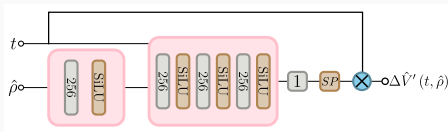


inputs $(t, \hat{\rho}) \rightarrow$ simple nonlinear layers $\rightarrow$ output $\Delta \hat{V}'_{\theta}$

- 1. Residual** insert $\hat{V}'_{LN} + \Delta \hat{V}'_{\theta}$ into the flow equation on the grid:
 $\mathcal{R}_{ij}(\theta) = \partial_t \hat{V}'_{\theta} - \mathcal{F}[\hat{V}'_{\theta}, \hat{V}''_{\theta}, \hat{V}'''_{\theta}]|_{(\hat{\rho}_i, t_j)}$, zero for the exact solution
- 2. Solve = minimize** $L(\theta) = \frac{1}{N_{\text{grid}}} \sum_{ij} |\mathcal{R}_{ij}(\theta)|^2$ over θ
gradient descent; the equation is the **only input**, no training data

Details: 3-point finite differences on a 201×501 grid, $3-4\times$ less memory than autograd; Adam; float32. **PINN**: Raissi, Perdikaris, Karniadakis, JCP 378, 686 (2019); Karniadakis et al., Nat. Rev. Phys. 3, 422 (2021). **PINNs for FRG / FDE / DSE**: Yokota, PRB 109, 214205 (2024); Miyagawa, Yokota, NeurIPS 2024; Terin, SciPost Phys. Core 8, 054 (2025). **FD derivatives in PINNs**: Chiu et al., CMAME 395, 114909 (2022).

Solving the flow equation as an optimization problem



inputs $(t, \hat{\rho}) \rightarrow$ simple nonlinear layers $\rightarrow$ output $\Delta \hat{V}'_{\theta}$

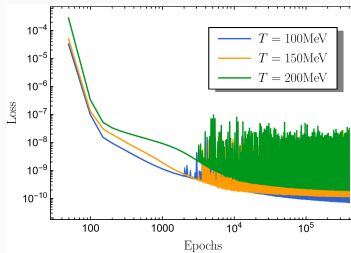
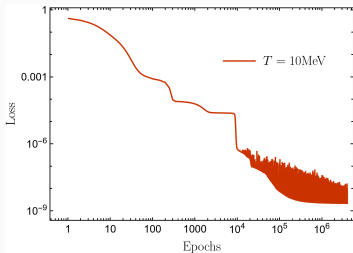
1. Residual insert $\hat{V}'_{LN} + \Delta \hat{V}'_{\theta}$ into the flow equation on the grid:
 $\mathcal{R}_{ij}(\theta) = \partial_t \hat{V}'_{\theta} - \mathcal{F}[\hat{V}'_{\theta}, \hat{V}''_{\theta}, \hat{V}'''_{\theta}]|_{(\hat{\rho}_i, t_j)}$, zero for the exact solution

2. Solve = minimize $L(\theta) = \frac{1}{N_{\text{grid}}} \sum_{ij} |\mathcal{R}_{ij}(\theta)|^2$ over θ
gradient descent; the equation is the **only input**, no training data

Built in initial condition (output $\propto t$); no field boundary condition

Details: 3-point finite differences on a 201×501 grid, 3–4 $\times$ less memory than autodiff; Adam; float32. **PINN**: Raissi, Perdikaris, Karniadakis, JCP 378, 686 (2019); Karniadakis et al., Nat. Rev. Phys. 3, 422 (2021). **PINNs for FRG / FDE / DSE**: Yokota, PRB 109, 214205 (2024); Miyagawa, Yokota, NeurIPS 2024; Terin, SciPost Phys. Core 8, 054 (2025). **FD derivatives in PINNs**: Chiu et al., CMAME 395, 114909 (2022).

Training



From scratch, $T=10 \text{ MeV}$, 4×10^6 epochs

Transfer from $T=120 \text{ MeV}$, 5×10^5 epochs

Optimizer

Adam, lr 5×10^{-4}
float32, then float64

Domain

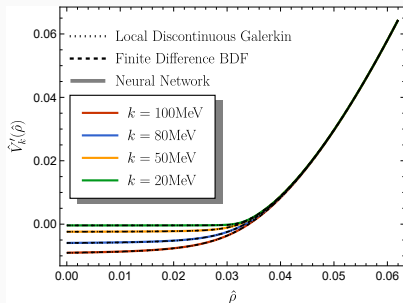
$k \in [20 \text{ MeV}, 1 \text{ GeV}]$
 $\hat{\rho} \in [0, 0.15]$
grid 201×501

Convergence

staircase descent,
plateau at $\sim 10^{-9}$

Grid 2001×501 and float64 refinement for the $T = 10 \text{ MeV}$ run. Denominators regularized by $\epsilon = 10^{-13}$. PyTorch, PhysicsNeMo, Lightning.

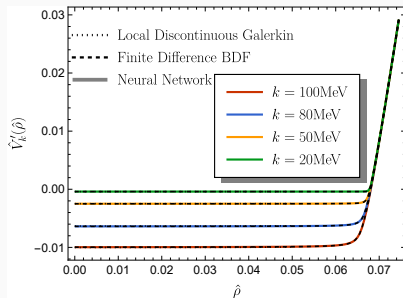
Broken phase: $\hat{V}'_k(\hat{\rho})$ vs. LDG and finite differences



$T = 100$ MeV, $k = 100, 80, 50, 20$ MeV

Benchmarks

LDG and FD-BDF



$T = 10$ MeV, deeply broken, same k

Accuracy

absolute error

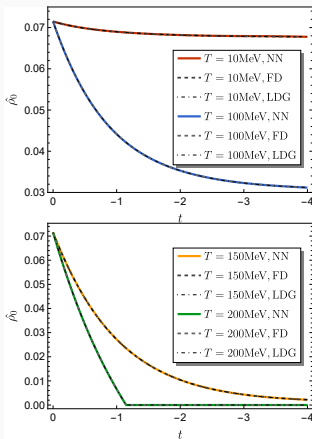
$$< 4 \times 10^{-5}$$

at all k and $\hat{\rho}$

Stiff region

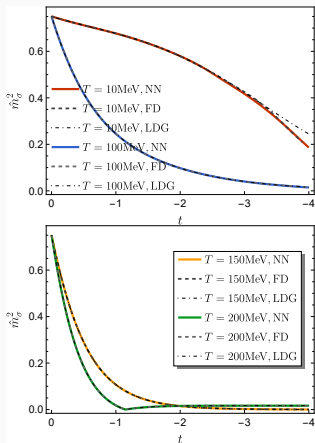
no oscillations at the
crossover $\hat{\rho}_0(k)$

Observables across the phase transition



Order parameter $\hat{V}'(\hat{\rho}_0) = 0$

$T < T_c$: $\hat{\rho}_0$ finite, $T \geq T_c$: $\hat{\rho}_0 \rightarrow 0$



Curvature mass

$$\hat{m}_\sigma^2 = \hat{V}'(\hat{\rho}_0) + 2\hat{\rho}_0 \hat{V}''(\hat{\rho}_0)$$

Same idea for the Wilson–Fisher fixed point

Fixed-point equation, $d = 3$, LPA

$$0 = -3u + \bar{\rho}u' + \mathcal{L} \left[\frac{1}{1 + u' + 2\bar{\rho}u''} + \frac{N-1}{1 + u'} \right]$$

Shooting is delicate

fine tuning to isolate WF from
Gaussian and multicritical branches

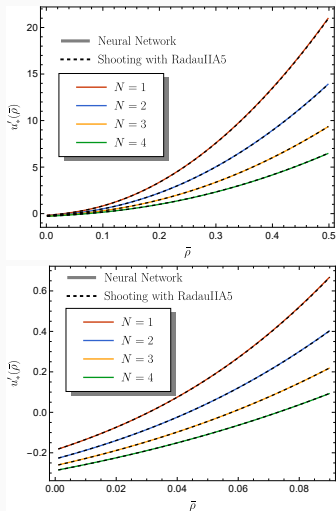
Network

$$u' = u'_{\text{LN}} + \Delta u'_\theta$$

u'_{LN} in closed form (${}_2F_1$)

4-layer MLP, ODE residual loss

no shooting, no Gaussian collapse



Shooting benchmark: Tan, Huang, Chen, Fu, EPJC 84, 897 (2024).

The large- N limit is not essential

Baselines from the equation itself

large field: $u'_* \rightarrow A\bar{\rho}^2$, $A = 3\gamma \simeq 84$

$$u' = A\bar{\rho}^2 + \Delta u'_\theta$$

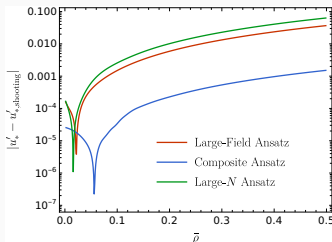
Composite ansatz ($O(1)$, $d = 3$)

$$u'_{\text{base}} = u'_{\text{poly}} [1 - w] + A\bar{\rho}^2 w$$

$$w = \frac{1}{2} \left\{ 1 + \tanh [(\bar{\rho} - 0.1)/0.02] \right\}$$

u'_{poly} : 10th-order Taylor series

$$u' = u'_{\text{base}} + \Delta u'_\theta$$



The large- N limit is not essential

Baselines from the equation itself

large field: $u'_* \rightarrow A\bar{\rho}^2$, $A = 3\gamma \simeq 84$

$$u' = A\bar{\rho}^2 + \Delta u'_\theta$$

Composite ansatz ($O(1)$, $d = 3$)

$$u'_{\text{base}} = u'_{\text{poly}} [1 - w] + A\bar{\rho}^2 w$$

$$w = \frac{1}{2} \{1 + \tanh [(\bar{\rho} - 0.1)/0.02]\}$$

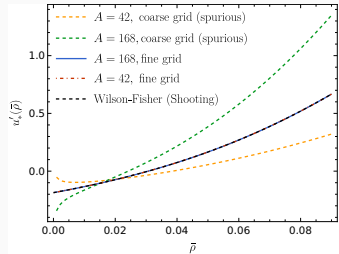
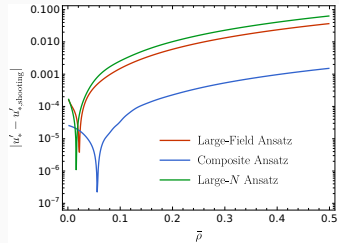
u'_{poly} : 10th-order Taylor series

$$u' = u'_{\text{base}} + \Delta u'_\theta$$

Mistuned amplitude $A/2$ or $2A$

coarse grid: a spurious solution

fine grid: **WF recovered**, A to 0.005%



Summary

Flow equation in the loss $\longrightarrow V_k(\rho)$ as a neural network

Finite- T $O(4)$ flows

all three regimes

symmetric, broken, critical

error $< 4 \times 10^{-5}$

Key ingredient

$$\hat{V}' = \hat{V}'_{\text{LN}} + \Delta \hat{V}'_{\theta}$$

stiffness absorbed

analytically

Wilson–Fisher, $N \leq 4$

no shooting

any baseline with the

right asymptotics

Outlook: multi-field potentials, vertex expansion,
operator learning, links to physics-informed kernels

Thank you!