



Machine-Learning-Accelerated Multigrid Setup for Lattice QCD Dirac Solves

Simon Pfahler Daniel Knüttel Riccardo Costantini Christoph Lehner Tilo Wettig

Department of Physics, University of Regensburg



Daniel Knüttel



Riccardo
Costantini



Christoph Lehner



Tilo Wettig

Overview

- Gauge-field generation in Lattice QCD is expensive, mostly because of the required solves of the Dirac equation
- Algebraic multigrid¹ (MG) is the state-of-the-art preconditioner
- MG needs a costly setup → only limited benefit for gauge-field generation

→ Goal of this project:

Accelerate the MG setup by “moving the setup along the Markov-chain trajectory”

¹Osborn et al., PoS Lattice 2010

Preconditioning

- Linear equation $D\psi = \varphi$
with $\varphi \in \mathbb{C}^{V \times 4 \times 3}$
- Solve via iterative methods, e.g., GMRES
→ Convergence depends on spectrum
- Preconditioners can change the spectrum to obtain better convergence:

$$DM \underbrace{M^{-1}}_y \psi = \varphi$$

$$D = (m + 4) \mathbb{1}$$

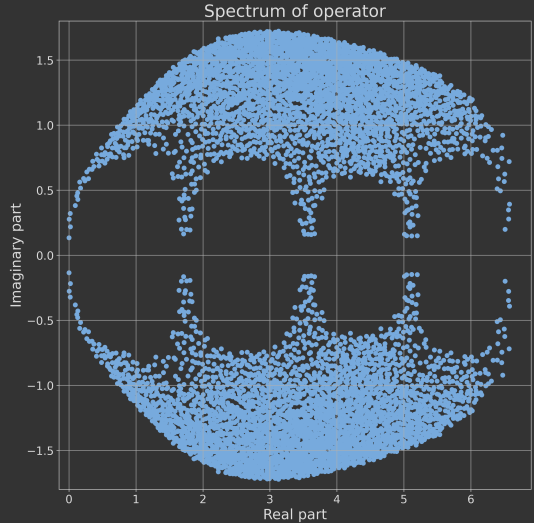
$$- \sum_{\mu} \frac{\mathbb{1} - \gamma_{\mu}}{2} H_{-\mu}^{(U)} - \frac{\mathbb{1} + \gamma_{\mu}}{2} H_{+\mu}^{(U)} - \frac{c_{\text{SW}}}{4} \sum_{\mu, \nu} \sigma_{\mu \nu} F_{\mu \nu}^{(U)}$$



Preconditioning

- Linear equation $D\psi = \varphi$ with $\varphi \in \mathbb{C}^{V \times 4 \times 3}$
- Solve via iterative methods, e.g., GMRES
 - Convergence depends on spectrum
- Preconditioners can change the spectrum to obtain better convergence:

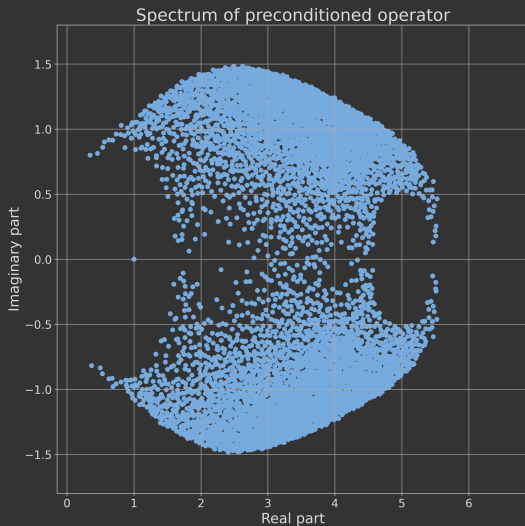
$$D \underbrace{M^{-1}\psi}_y = \varphi$$



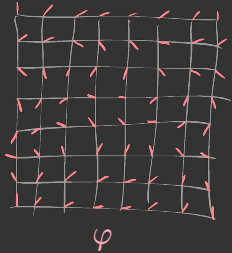
Preconditioning

- Linear equation $D\psi = \varphi$ with $\varphi \in \mathbb{C}^{V \times 4 \times 3}$
- Solve via iterative methods, e.g., GMRES
→ Convergence depends on spectrum
- Preconditioners can change the spectrum to obtain better convergence:

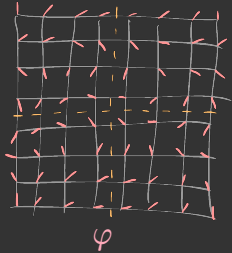
$$DM \underbrace{M^{-1}}_y \psi = \varphi$$



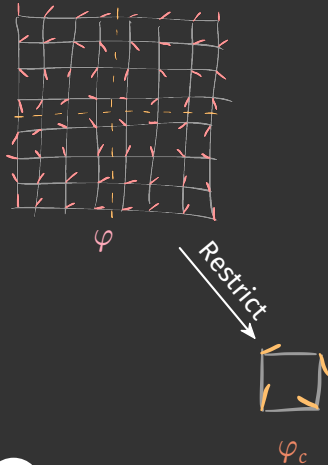
Multigrid - Overview



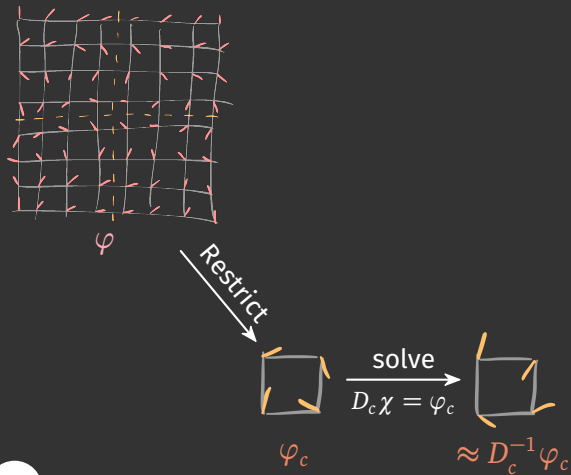
Multigrid - Overview



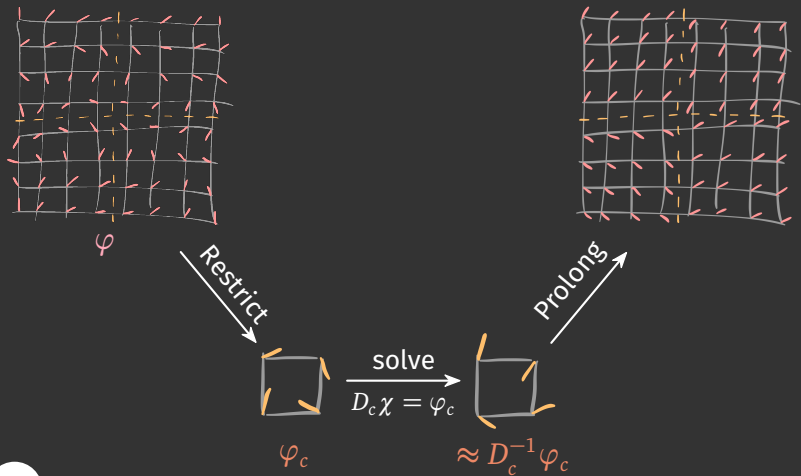
Multigrid - Overview



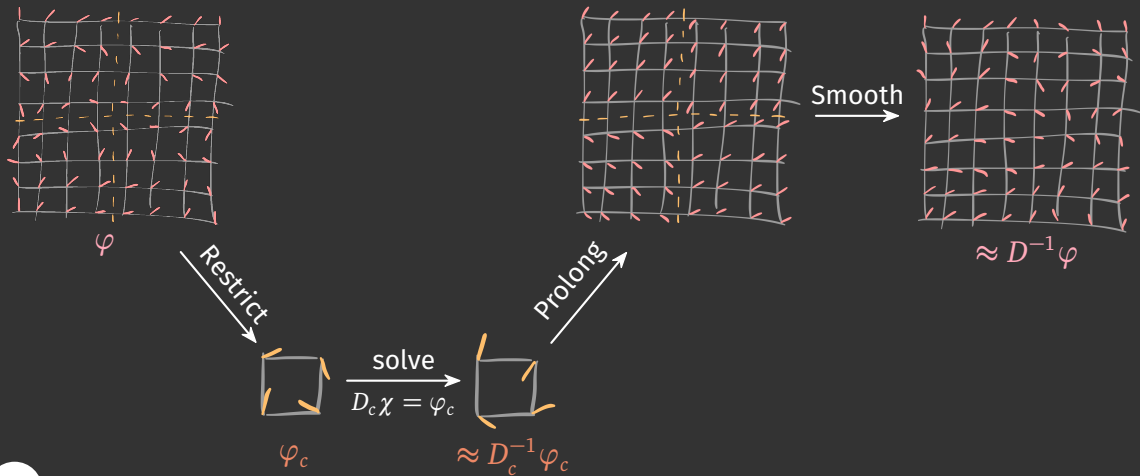
Multigrid - Overview



Multigrid - Overview



Multigrid - Overview



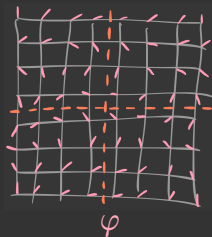
Multigrid - Restriction and Prolongation

How can we restrict φ to the coarse grid?²³

- Take N low modes ψ_i on the fine grid (e.g. via $D\psi_i \approx 0$)
- Block-orthonormalize ψ_i to get $\hat{\psi}_i|_B$ on each block B
- Take scalar products

Prolongation = Restriction[†]
→ Coarse-grid operator:

$$D_c = RDP$$



²Babich et al., PRL 105(2010)

³Lüscher, JHEP 07(2007)

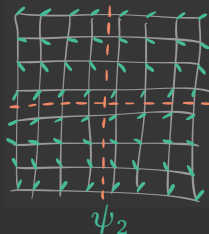
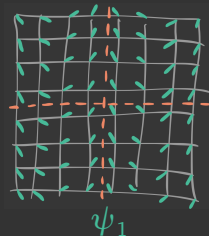
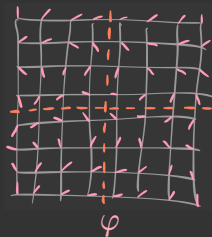
Multigrid - Restriction and Prolongation

How can we restrict φ to the coarse grid?²³

- Take N low modes ψ_i on the fine grid (e.g. via $D\psi_i \approx 0$)
- Block-orthonormalize ψ_i to get $\hat{\psi}_i|_B$ on each block B
- Take scalar products

Prolongation = Restriction[†]
→ Coarse-grid operator:

$$D_c = RDP$$



²Babich et al., PRL 105(2010)

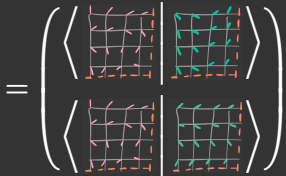
³Lüscher, JHEP 07(2007)

Multigrid - Restriction and Prolongation

How can we restrict φ to the coarse grid?²³

- Take N low modes ψ_i on the fine grid (e.g. via $D\psi_i \approx 0$)
- Block-orthonormalize ψ_i to get $\hat{\psi}_i|_B$ on each block B
- Take scalar products

Prolongation = Restriction[†]
→ Coarse-grid operator:

$$\varphi_c(B) = \begin{pmatrix} \langle \varphi|_B | \hat{\psi}_1|_B \rangle \\ \langle \varphi|_B | \hat{\psi}_2|_B \rangle \end{pmatrix} = \begin{pmatrix} \langle \text{grid} | \text{grid} \rangle \\ \langle \text{grid} | \text{grid} \rangle \end{pmatrix}$$


$$D_c = RDP$$

²Babich et al., PRL 105(2010)

³Lüscher, JHEP 07(2007)

Multigrid - Restriction and Prolongation

How can we restrict φ to the coarse grid?²³

- Take N low modes ψ_i on the fine grid (e.g. via $D\psi_i \approx 0$)
- Block-orthonormalize ψ_i to get $\hat{\psi}_i|_B$ on each block B
- Take scalar products

Prolongation = Restriction[†]
→ Coarse-grid operator:

$$\varphi_c(B) = \begin{pmatrix} \langle \varphi|_B | \hat{\psi}_1|_B \rangle \\ \langle \varphi|_B | \hat{\psi}_2|_B \rangle \end{pmatrix} = \begin{pmatrix} \langle \text{grid with red arrows} | \text{grid with green arrows} \rangle \\ \langle \text{grid with red arrows} | \text{grid with green arrows} \rangle \end{pmatrix}$$

$$D_c = RDP$$

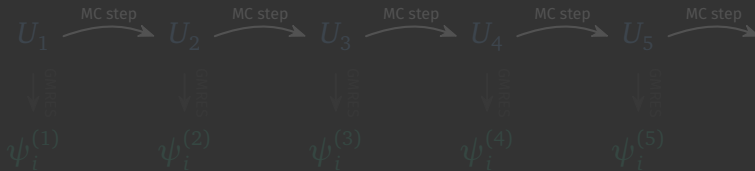
²Babich et al., PRL 105(2010)

³Lüscher, JHEP 07(2007)

MG setup transfer

- We can construct P and R for any set of approximate low modes ψ_i
- The solver speedup depends on the overlap of the chosen ψ_i with the low-mode space of D_U
- We typically generate configurations U in a Markov chain
- Generating the ψ_i anew for each configuration is expensive

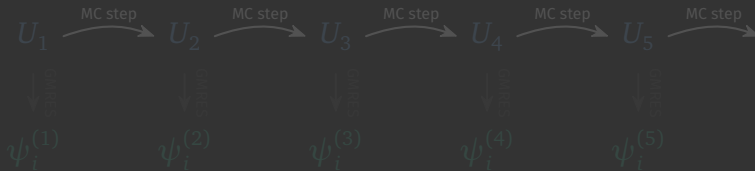
→ Move the ψ_i along the Markov chain!



MG setup transfer

- We can construct P and R for any set of approximate low modes ψ_i
- The solver speedup depends on the overlap of the chosen ψ_i with the low-mode space of D_U
- We typically generate configurations U in a Markov chain
- Generating the ψ_i anew for each configuration is expensive

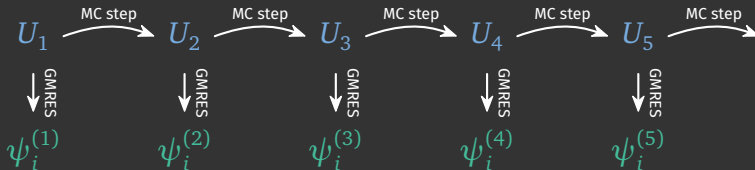
→ Move the ψ_i along the Markov chain!



MG setup transfer

- We can construct P and R for any set of approximate low modes ψ_i
- The solver speedup depends on the overlap of the chosen ψ_i with the low-mode space of D_U
- We typically generate configurations U in a Markov chain
- Generating the ψ_i anew for each configuration is expensive

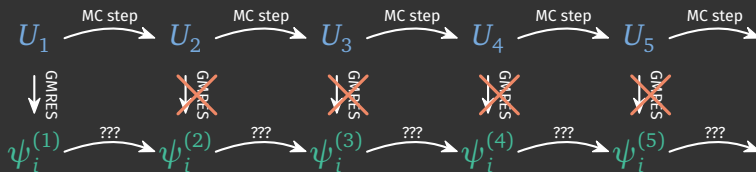
→ Move the ψ_i along the Markov chain!



MG setup transfer

- We can construct P and R for any set of approximate low modes ψ_i
- The solver speedup depends on the overlap of the chosen ψ_i with the low-mode space of D_U
- We typically generate configurations U in a Markov chain
- Generating the ψ_i anew for each configuration is expensive

→ Move the ψ_i along the Markov chain!



Machine Learning - Building blocks

We use 2 types of layers in our model, taken from⁴

Parallel-transport (PT) layer:

$$\psi_a(x) \stackrel{PT}{=} T_{p_a} \varphi_a(x)$$

for $a = 1, \dots, n$

Linear (L) layer:

$$\psi_b(x) \stackrel{L}{=} \sum_{a=1}^n W_{ba} \varphi_a(x)$$

for $a = 1, \dots, n$ and
 $b = 1, \dots, m$

Both are gauge-equivariant!

⁴Pfahler et al., PoS Lattice 2025

Machine Learning - Building blocks

We use 2 types of layers in our model, taken from⁴

Parallel-transport (PT) layer:

$$\psi_a(x) \stackrel{PT}{=} T_{p_a} \varphi_a(x)$$

for $a = 1, \dots, n$

Linear (L) layer:

$$\psi_b(x) \stackrel{L}{=} \sum_{a=1}^n W_{ba} \varphi_a(x)$$

for $a = 1, \dots, n$ and
 $b = 1, \dots, m$

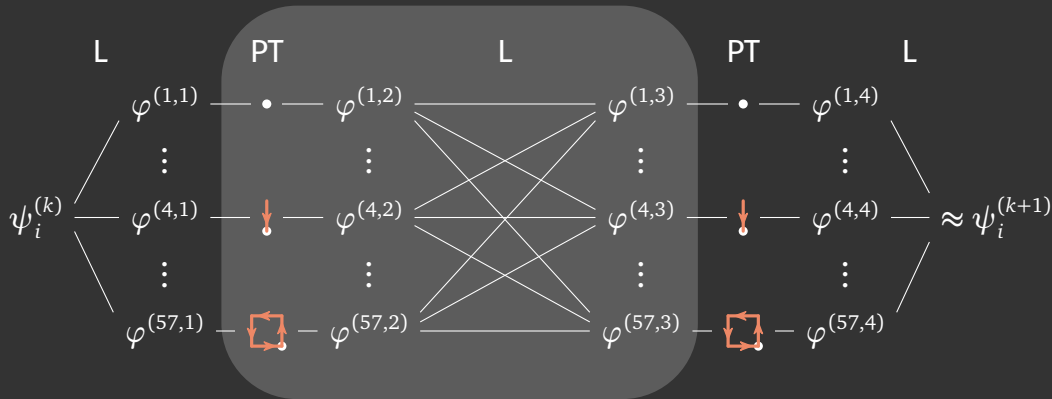
Both are gauge-equivariant!

⁴Pfahler et al., PoS Lattice 2025

Network architecture

All PT layers have the paths $\left\{ \bullet, \xrightarrow{\quad}, \xleftarrow{\quad}, \downarrow, \uparrow, \cdots, \begin{array}{c} \circlearrowleft \\ \square \end{array}, \begin{array}{c} \circlearrowright \\ \square \end{array}, \begin{array}{c} \circlearrowleft \\ \square \end{array}, \begin{array}{c} \circlearrowright \\ \square \end{array}, \cdots \right\}$

Network architecture



All PT layers have the paths $\left\{ \bullet, \rightarrow, \leftarrow, \downarrow, \uparrow, \dots, \square, \square, \square, \square, \dots \right\}$

Cost function for training

Aspects of our cost function:

- Average of a block-wise cost function
- On each block, calculate the overlap between the spaces spanned by a standard MG setup and our model output

$$C = 1 - \frac{1}{\#vectors \cdot \#blocks} \sum_{\text{blocks } B} \sum_i \sum_j \left| \left\langle \hat{m}_i|_B \left| \hat{\psi}_j^{(k+1)}|_B \right\rangle \right|^2$$

$\hat{\psi}_j^{(k+1)}|_B$ = standard MG setup vectors of U_{k+1} on block B

$\hat{m}_i|_B$ = orthonormalized model output on block B , with $\psi_i^{(k)}$ as model input

Cost function for training

Aspects of our cost function:

- Average of a block-wise cost function
- On each block, calculate the overlap between the spaces spanned by a standard MG setup and our model output

$$C = 1 - \frac{1}{\# \text{vectors} \cdot \# \text{blocks}} \sum_{\text{blocks } B} \sum_i \sum_j \left| \left\langle \hat{m}_i|_B \left| \hat{\psi}_j^{(k+1)}|_B \right\rangle \right|^2$$

$\hat{\psi}_j^{(k+1)}|_B$ = standard MG setup vectors of U_{k+1} on block B

$\hat{m}_i|_B$ = orthonormalized model output on block B , with $\psi_i^{(k)}$ as model input

Results - Overview

$8^3 \times 16$ ensemble ⁵:

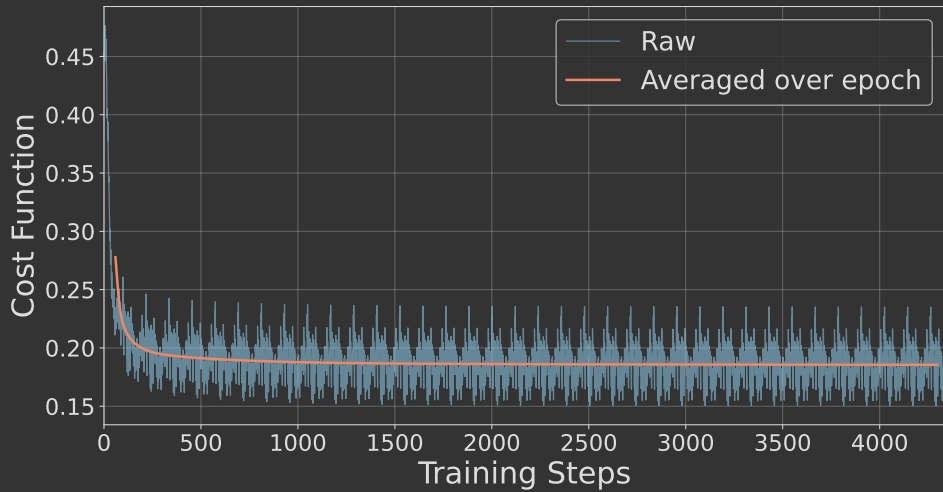
- Wilson action, $\beta = 6$
- quenched, heat bath
- Wilson-clover Dirac operator

$16^3 \times 32$ ensemble:

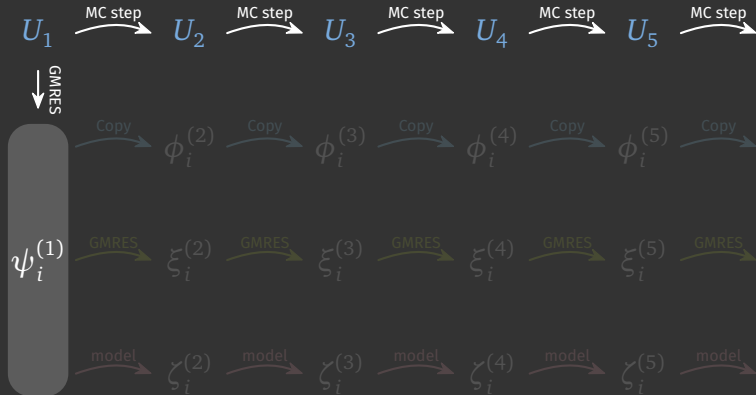
- Wilson action, $\beta = 6$
- quenched, HMC
- Wilson-clover Dirac operator

⁵Knüttel, doi.org/10.5281/zenodo.15018324

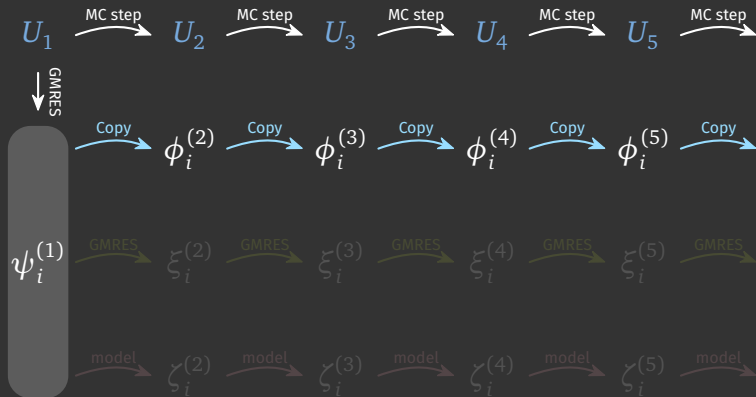
Training



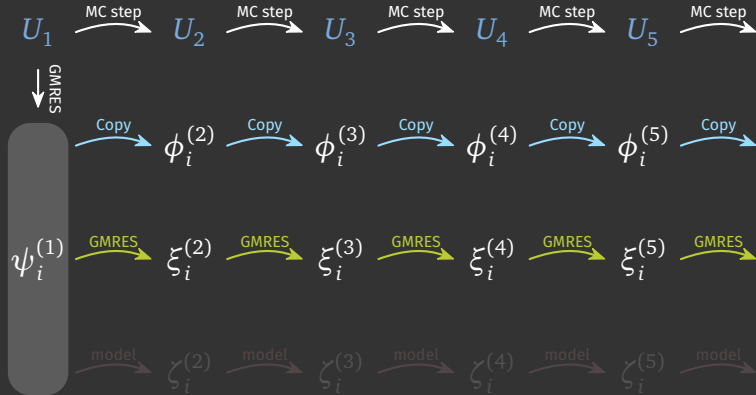
Evaluation



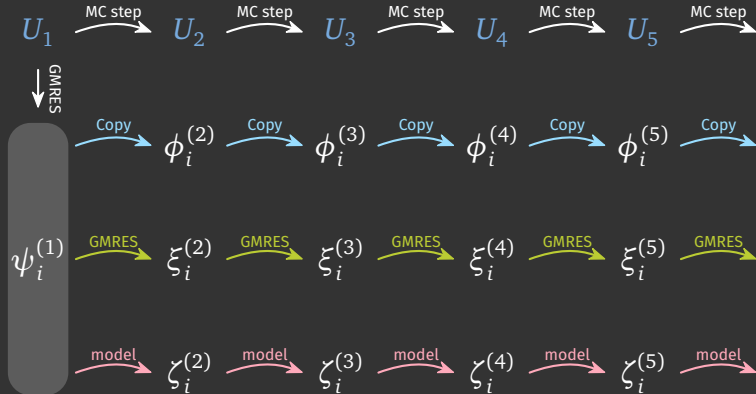
Evaluation



Evaluation

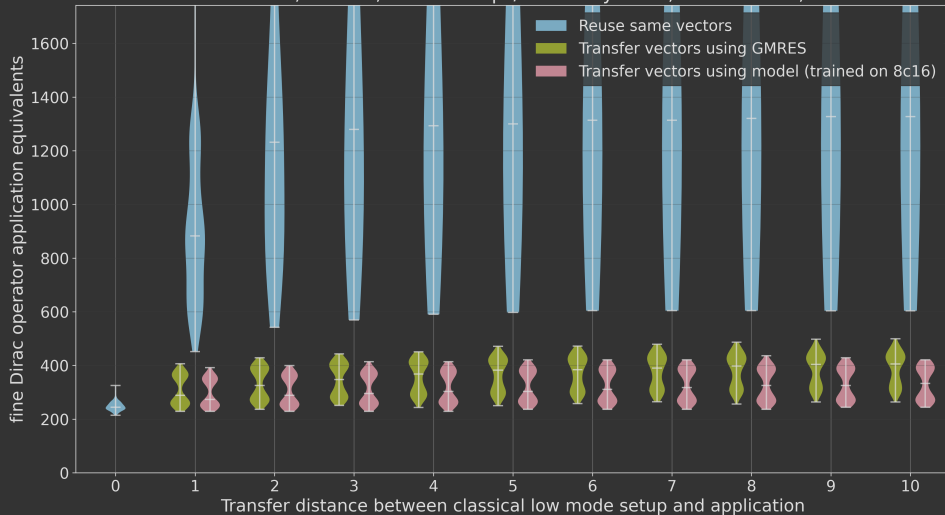


Evaluation

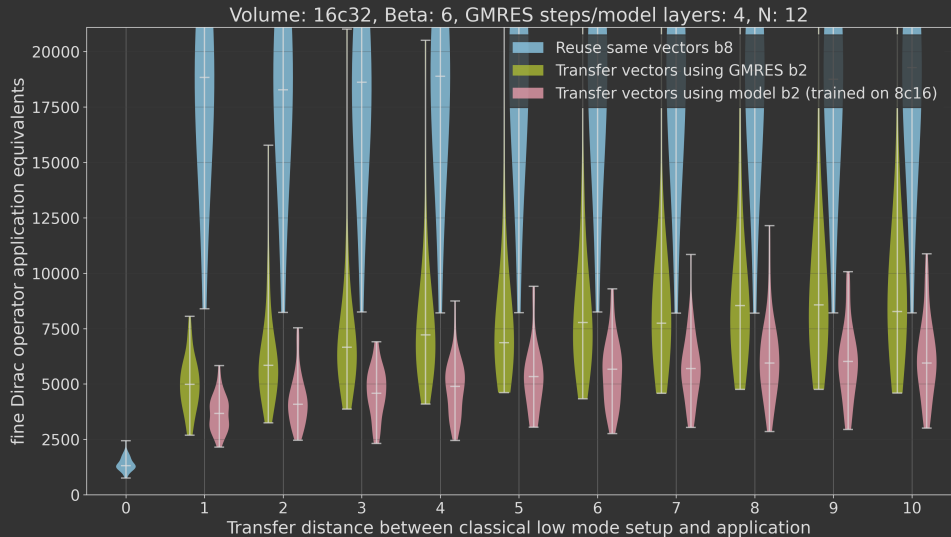


$8^3 \times 16$ volume

Volume: 8c16, Beta: 6, GMRES steps/model layers: 4, block size: 4, N: 12



Transfer to $16^3 \times 32$ volume



Retraining on $16^3 \times 32$ volume

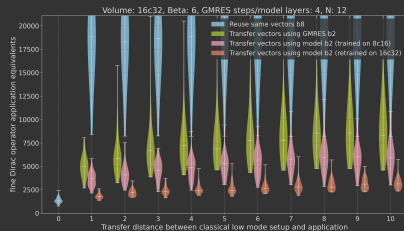
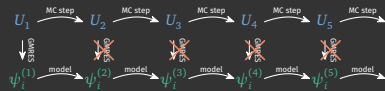


Summary

- MG setups can be transferred from one configuration to the next in a Markov chain
- Using Machine Learning, we can outperform standard methods like GMRES
- Outlook
 - Implement this for unquenched simulations
 - Investigate analytically what the models learn

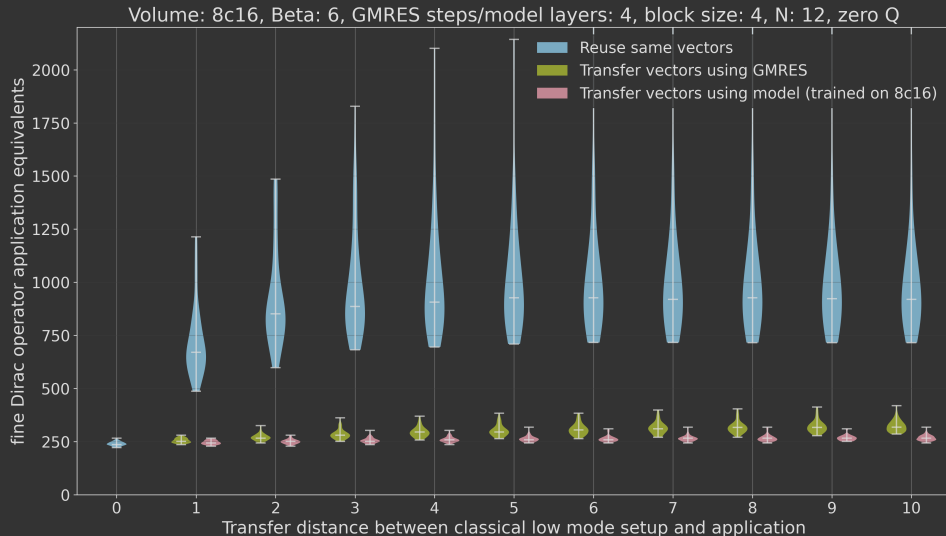


pfahler.online

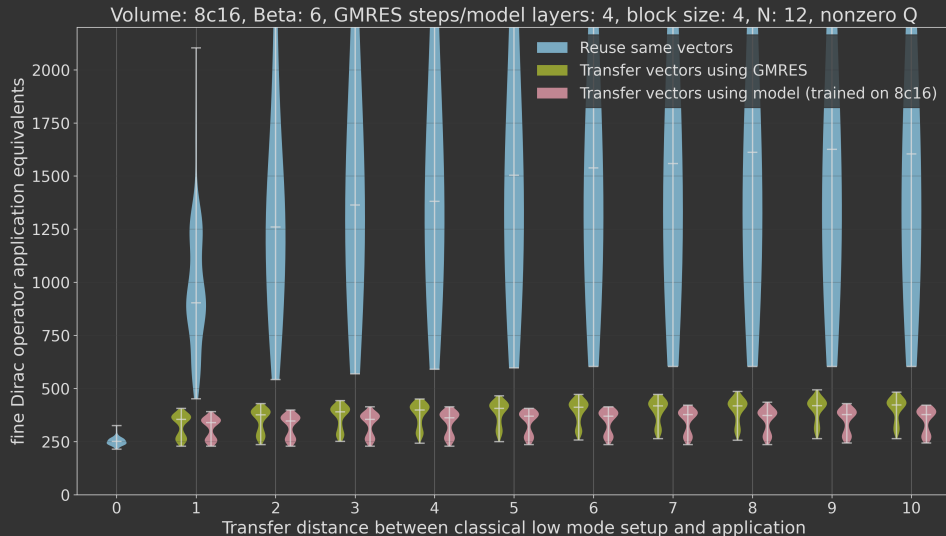


Bonus slides

Q-dependency



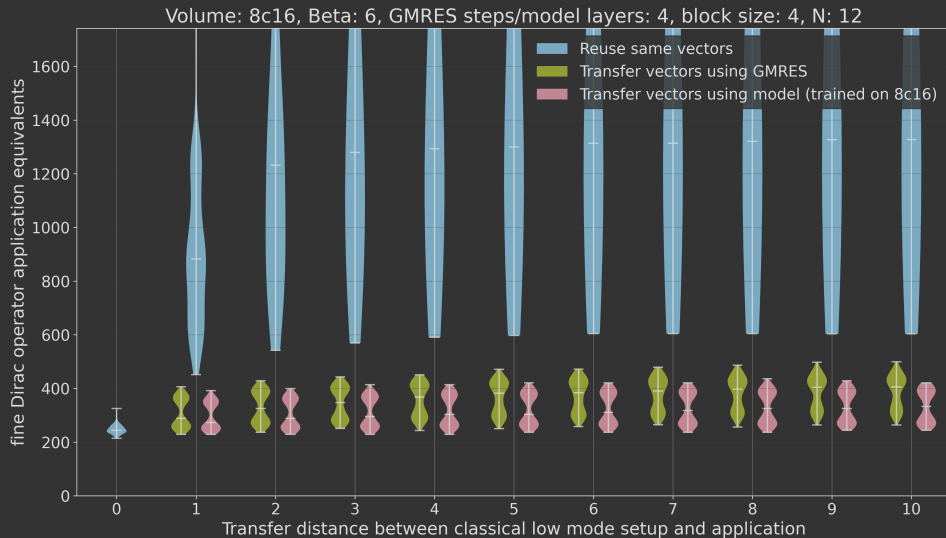
Q-dependency



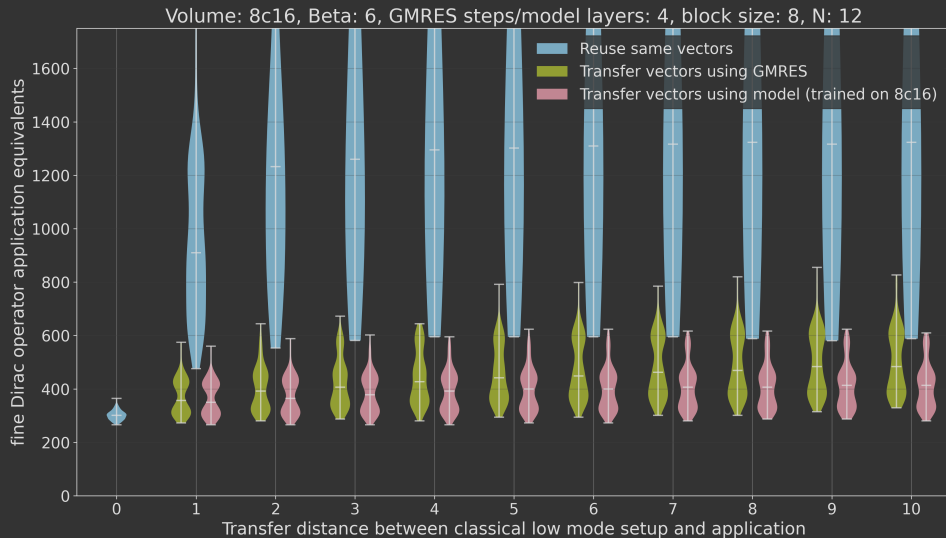
Dependency on block size - $8^3 \times 16$



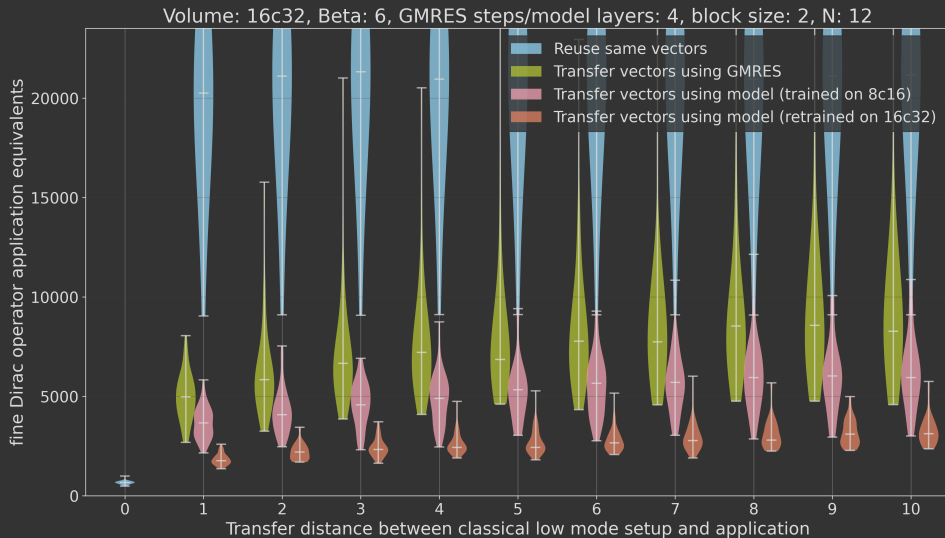
Dependency on block size - $8^3 \times 16$



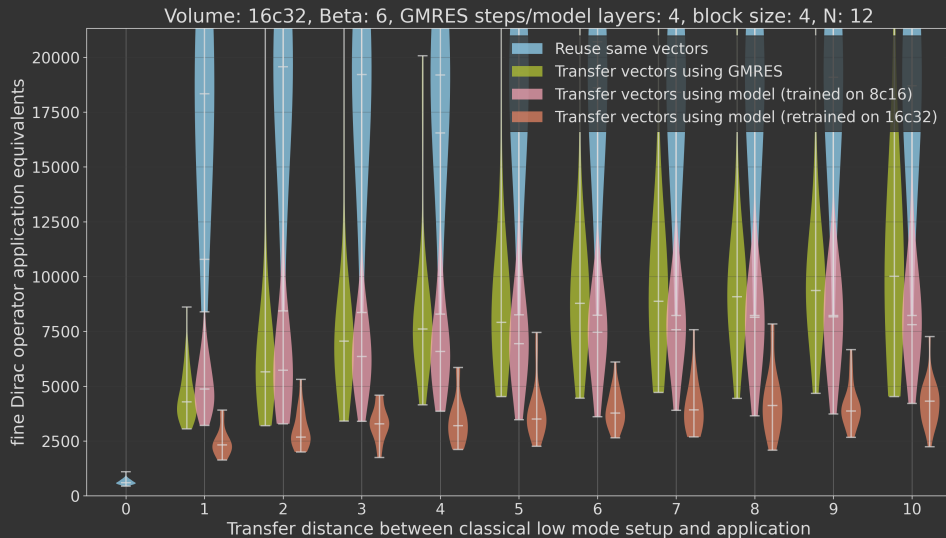
Dependency on block size - $8^3 \times 16$



Dependency on block size - $16^3 \times 32$



Dependency on block size - $16^3 \times 32$



Dependency on block size - $16^3 \times 32$

