

Emergent Equivariance in Deep Ensembles

Jan E. Gerken



UNIVERSITY OF
GOTHENBURG

WASP | WALLENBERG AI
AUTONOMOUS SYSTEMS
AND SOFTWARE PROGRAM

in collaboration with



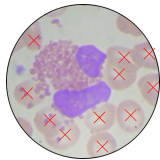
Pan Kessel



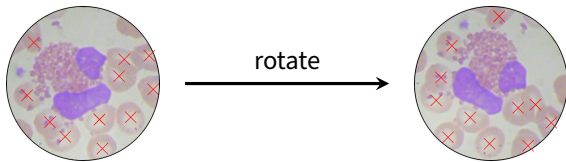
Philipp Misof

Symmetries in deep learning

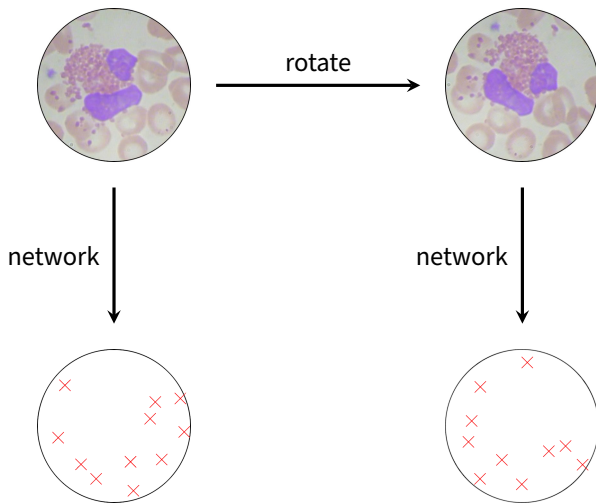
Symmetries in deep learning



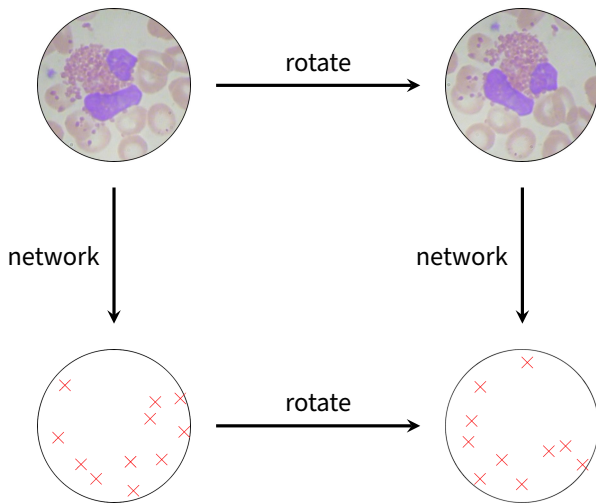
Symmetries in deep learning



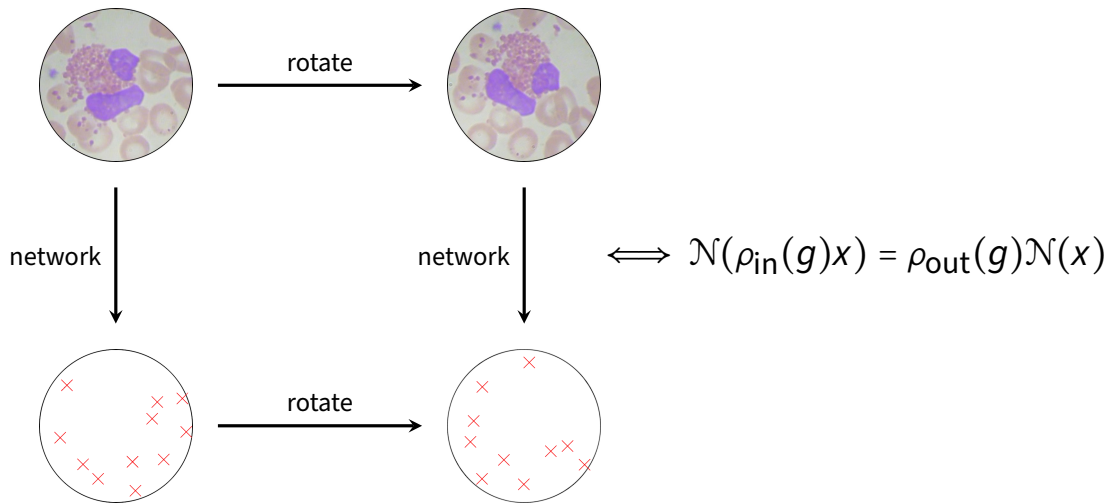
Symmetries in deep learning



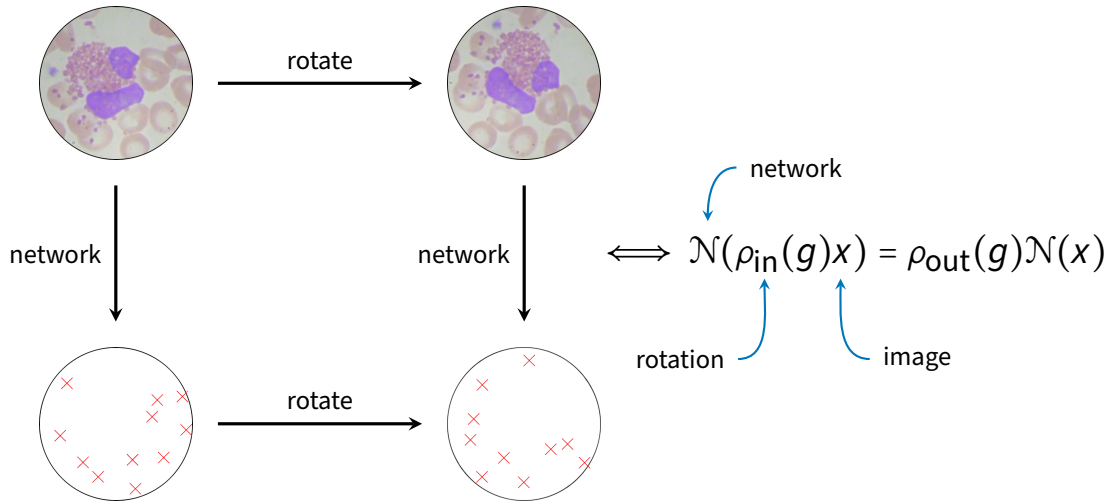
Symmetries in deep learning



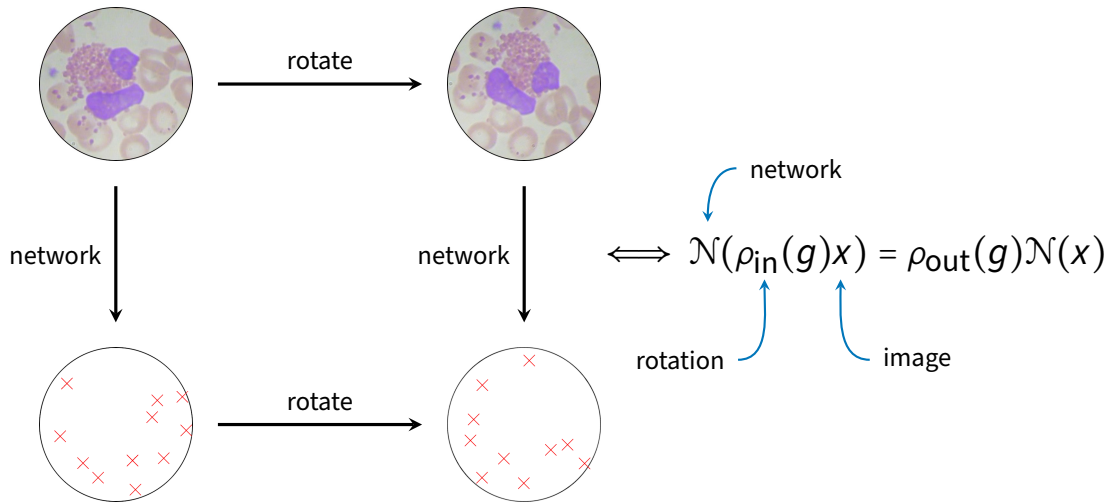
Symmetries in deep learning



Symmetries in deep learning



Equivariance



Equivariant neural networks

Equivariant neural networks

Group Equivariant Convolutional Networks

Taco S. Cohen

University of Amsterdam

Max Welling

University of Amsterdam

University of California Irvine

Canadian Institute for Advanced Research

T.S.COHEN@UVA.NL

M.WELLING@UVA.NL

Abstract

We introduce Group equivariant Convolutional Neural Networks (G-CNNs), a natural generalization of convolutional neural networks that reduces sample complexity by exploiting symme-

Convolution layers can be used effectively in a *deep* network because all the layers in such a network are *translation equivariant*: shifting the image and then feeding it through a *number* of layers is the same as feeding the original image through the same layers and then shifting the resulting feature maps (at least up to edge-effects). In

Equivariant neural networks

Group Equivariant Convolutional Networks

Taco S. Cohen
University of Amsterdam
Max Welling
University of Amsterdam
University of California Irvine
Canadian Institute for Advanced Research

T.S.COHEN@UVA.NL

M.WELLING@UVA.NL

Abstract

We introduce Group equivariant Convolutional Neural Networks (G-CNNs), a natural generalization of convolutional neural networks that reduces sample complexity by exploiting symme-

Convolution layers can be used effectively in a deep network because all the layers in such a network are *translation equivariant*: shifting the image and then feeding it through a number of layers is the same as feeding the original image through the same layers and then shifting the resulting feature maps (at least up to edge-effects). In

Equivariant Transformer Networks

Kai Sheng Tai¹ Peter Bailis¹ Gregory Valiant¹

Abstract

How can prior knowledge on the transformation invariances of a domain be incorporated into the architecture of a neural network? We propose Equivariant Transformers (ETs), a family of differentiable image-to-image mappings that improve the robustness of models towards pre-defined continuous transformation groups. Through the use of specially-derived canonical coordinate systems, ETs incorporate functions that

scaling to each training image). While data augmentation typically helps reduce the test error of CNN-based models, there is no guarantee that transformation invariance will be enforced for data not seen during training.

In contrast to training time approaches like data augmentation, recent work on group equivariant CNNs (Cohen & Welling, 2016; Dieleman et al., 2016; Marcos et al., 2017; Worrall et al., 2017; Henriques & Veličković, 2017; Cohen et al., 2018) has explored new CNN architectures that are *invariant to certain modifiable to particular transforms*.

Equivariant neural networks

Group Equivariant Convolutional Networks

Taco S. Cohen
University of Amsterdam
Max Welling
University of Amsterdam
University of California Irvine
Canadian Institute for Advanced Research

T.S.COHN@UVA.NL

M.WELLING@UVA.NL

Abstract

We introduce Group equivariant Convolutional Neural Networks (G-CNNs), a natural generalization of convolutional neural networks that reduces sample complexity by exploiting symme-

Convolution layers can be used effectively in a deep network because all the layers in such a network are *translation equivariant*: shifting the image and then feeding it through a number of layers is the same as feeding the original image through the same layers and then shifting the resulting feature maps (at least up to edge-effects). In

Equivariant Transformer Networks

Kai Sheng Tai¹ Peter Bailis¹ Gregory Valiant¹

Abstract

How can prior knowledge on the transformation invariances of a domain be incorporated into the architecture of a neural network? We propose Equivariant Transformers (ETs), a family of differentiable image-to-image mappings that improve the robustness of models towards pre-defined continuous transformation groups. Through the use of specially-derived canonical coordinate systems, ETs incorporate functions that

scaling to each training image). While data augmentation typically helps reduce the test error of CNN-based models, there is no guarantee that transformation invariance will be enforced for data not seen during training.

In contrast to training time approaches like data augmentation, recent work on group equivariant CNNs (Cohen & Welling, 2016; Dieleman et al., 2016; Marcos et al., 2017; Worrall et al., 2017; Henriques & Veličković, 2017; Cohen et al., 2018) has explored new CNN architectures that are

Theory for Equivariant Quantum Neural Networks

Quynh T. Nguyen,^{1,2} Louis Schatzki,^{3,4} Paolo Bucci,^{1,5} Michael Ragone,^{1,6} Patrick J. Cules,¹ Frédéric Sauvage,¹ Martin Laroche,^{1,7} and M. Cerezo¹

¹Theoretical Division, Los Alamos National Laboratory, Los Alamos, New Mexico 87545, USA
²School of Engineering and Applied Sciences, Harvard University, Cambridge, Massachusetts 02138, USA
³Information Sciences, Los Alamos National Laboratory, Los Alamos, New Mexico 87545, USA
⁴Department of Electrical and Computer Engineering, University of Illinois at Urbana-Champaign, Urbana, Illinois 61801, USA
⁵Dipartimento di Fisica e Astronomia, Università di Firenze, Sesto Fiorentino (FI), 50019, Italy
⁶Department of Mathematics, University of California Davis, Davis, California 95616, USA
⁷Center for Nonlinear Studies, Los Alamos National Laboratory, Los Alamos, New Mexico 87545, USA

Quantum neural network architectures that have little-to-no inductive biases are known to face trainability and generalization issues. Inspired by a similar problem, recent breakthroughs in machine learning address this challenge by creating models encoding the symmetries of the learning task. This is materialized through the usage of equivariant neural networks whose action commutes with that of the symmetries. In this work, we present these ideas in the quantum realm by

Equivariant neural networks

Group Equivariant Convolutional Networks

Taco S. Cohen
University of Amsterdam

T.S.COHEM@UVA.NL

Max Welling
University of Amsterdam
University of California Irvine
Canadian Institute for Advanced Research

M.WELLING@UVA.NL

Abstract

We introduce Group equivariant Convolutional Neural Networks (G-CNNs), a natural generalization of convolutional neural networks that reduces sample complexity by exploiting symme-

Convolution layers can be used effectively in a deep network because all the layers in such a network are *translation equivariant*: shifting the image and then feeding it through a number of layers is the same as feeding the original image through the same layers and then shifting the resulting feature maps (at least up to edge-effects). In

Equivariant Transformer Networks

Kai Sheng Tai¹ Peter Bailis¹ Gregory Valiant¹

Abstract

How can prior knowledge on the transformation invariances of a domain be incorporated into the architecture of a neural network? We propose Equivariant Transformers (ETs), a family of differentiable image-to-image mappings that improve the robustness of models towards pre-defined continuous transformation groups. Through the use of specially-derived canonical coordinate systems, ETs incorporate functions that

scaling to each training image). While data augmentation typically helps reduce the test error of CNN-based models, there is no guarantee that transformation invariance will be enforced for data not seen during training.

In contrast to training time approaches like data augmentation, recent work on group equivariant CNNs (Cohen & Welling, 2016; Dieleman et al., 2016; Marcos et al., 2017; Worrall et al., 2017; Henriques & Veličković, 2017; Cohen et al., 2018) has explored new CNN architectures that are invariant to various non-linear transformations.

Theory for Equivariant Quantum Neural Networks

Quynh T. Nguyen,^{1,2} Louis Schatzki,^{3,4} Paolo Braccia,^{1,5} Michael Ragone,^{1,6} Patrick J. Cules,¹ Frédéric Sauvage,¹ Martin Laroche,^{1,7} and M. Cerezo¹

¹Theoretical Division, Los Alamos National Laboratory, Los Alamos, New Mexico 87545, USA
²School of Engineering and Applied Sciences, Harvard University, Cambridge, Massachusetts 02138, USA
³Information Sciences, Los Alamos National Laboratory, Los Alamos, New Mexico 87545, USA
⁴Department of Electrical and Computer Engineering, University of Illinois at Urbana-Champaign, Urbana, Illinois 61801, USA
⁵Dipartimento di Fisica e Astronomia, Università di Firenze, Sesto Fiorentino (FI), 50019, Italy
⁶Department of Mathematics, University of California Davis, Davis, California 95616, USA
⁷Center for Nonlinear Studies, Los Alamos National Laboratory, Los Alamos, New Mexico 87545, USA

Quantum neural network architectures that have little-to-no inductive biases are known to face trainability and generalization issues. Inspired by a similar problem, recent breakthroughs in machine learning address this challenge by creating models encoding the symmetries of the learning task. This is materialized through the usage of equivariant neural networks whose action commutes with that of the symmetries. In this work, we present these ideas in the quantum realm by

An Efficient Lorentz Equivariant Graph Neural Network for Jet Tagging

Shiqi Gong^{a,b,c,1} Qi Meng^b Jue Zhang^b Hulin Qu^c Congqiao Li² Sitian Qian^d Weitao Du^a Zhi-Ming Ma^a Tie-Yan Liu^b

^aAcademy of Mathematics and Systems Science, Chinese Academy of Sciences, Zhongguancun East Road, Beijing 100190, China

^bMicrosoft Research Asia, Danling Street, Beijing 100080, China

^cCERN, EP Department, CH-1211 Geneva 23, Switzerland

^dSchool of Physics, Peking University, Chenmefu Road, Beizina 100871, China

Equivariant neural networks

Group Equivariant Convolutional Networks

Taco S. Cohen
University of Amsterdam

T.S.COHEM@UVA.NL

Max Welling
University of Amsterdam
University of California Irvine
Canadian Institute for Advanced Research

M.WELLING@UVA.NL

Abstract

We introduce Group equivariant Convolutional Neural Networks (G-CNNs), a natural generalization of convolutional neural networks that reduces sample complexity by exploiting symme-

Convolution layers can be used effectively in a deep network because all the layers in such a network are *translation equivariant*: shifting the image and then feeding it through a number of layers is the same as feeding the original image through the same layers and then shifting the resulting feature maps (at least up to edge-effects). In

Equivariant Transformer Networks

Kal Sheng Tai¹ Peter Bailis¹ Gregory Valiant¹

Abstract

How can prior knowledge on the transformation invariances of a domain be incorporated into the architecture of a neural network? We propose Equivariant Transformers (ETs), a family of differentiable image-to-image mappings that improve the robustness of models towards pre-defined continuous transformation groups. Through the use of specially-derived canonical coordinate systems, ETs incorporate functions that

scaling to each training image). While data augmentation typically helps reduce the test error of CNN-based models, there is no guarantee that transformation invariance will be enforced for data not seen during training.

In contrast to training time approaches like data augmentation, recent work on group equivariant CNNs (Cohen & Welling, 2016; Dieleman et al., 2016; Marcos et al., 2017; Worrall et al., 2017; Henriques & Veličković, 2017; Cohen et al., 2018) has explored new CNN architectures that are

Theory for Equivariant Quantum Neural Networks

Quynh T. Nguyen,^{1,2} Louis Schatzki,^{3,4} Paolo Braccia,^{1,5} Michael Ragone,^{1,6} Patrick J. Cules,¹ Frédéric Sauvage,¹ Martin Larooca,^{1,7} and M. Cerezo¹

¹Theoretical Division, Los Alamos National Laboratory, Los Alamos, New Mexico 87545, USA
²School of Engineering and Applied Sciences, Harvard University, Cambridge, Massachusetts 02138, USA
³Information Sciences, Los Alamos National Laboratory, Los Alamos, New Mexico 87545, USA
⁴Department of Electrical and Computer Engineering, University of Illinois at Urbana-Champaign, Urbana, Illinois 61801, USA
⁵Dipartimento di Fisica e Astronomia, Università di Firenze, Sesto Fiorentino (FI), 50019, Italy
⁶Department of Mathematics, University of California Davis, Davis, California 95616, USA
⁷Center for Nonlinear Studies, Los Alamos National Laboratory, Los Alamos, New Mexico 87545, USA

Quantum neural network architectures that have little-to-no inductive biases are known to face trainability and generalization issues. Inspired by a similar problem, recent breakthroughs in machine learning address this challenge by creating models encoding the symmetries of the learning task. This is materialized through the usage of equivariant neural networks whose action commutes with that of the symmetries. In this work, we present these ideas in the quantum realm by

An Efficient Lorentz Equivariant Graph Neural Network for Jet Tagging

Shiqi Gong^{a,b,c} Qi Meng^b Jue Zhang^b Hulin Qu^c Congqiao Li^c Sitian Qian^d Weitao Du^a Zhi-Ming Ma^a Tie-Yan Liu^b

^aAcademy of Mathematics and Systems Science, Chinese Academy of Sciences, Zhongguancun East Road, Beijing 100190, China

^bMicrosoft Research Asia, Danling Street, Beijing 100080, China

^cCERN, EP Department, CH-1211 Geneva 23, Switzerland

^dSchool of Physics, Peking University, Cheneiya Road, Beizhina 100871, China

E(3)-Equivariant Graph Neural Networks for Data-Efficient and Accurate Interatomic Potentials

Simon Batzner^{a,1} Albert Muscelian,¹ Litxin Sun,¹ Mario Geiger,² Jonathan P. Mailon,³ Mordechai Korbath,² Nicola Molinari,¹ Tess E. Smidt,^{4,5} and Boris Kozinsky^{a,1,2}

¹John A. Pashon School of Engineering and Applied Sciences,

Harvard University, Cambridge, MA 02138, USA

²École Polytechnique Fédérale de Lausanne, 1015 Lausanne, Switzerland

³Robert Bosch Research and Technology Center, Cambridge, MA 02139, USA

⁴Computational Research Division and Center for Advanced Mathematics for Energy Research Applications,

Lawrence Berkeley National Laboratory, Berkeley, CA 94720, USA

⁵Massachusetts Institute of Technology, Department of Electrical

Engineering and Computer Science, Cambridge, MA 02142, USA

This work presents Neural Equivariant Interatomic Potentials (NequIP), an E(3)-equivariant neural network approach for learning interatomic potentials from *ab-initio* calculations for molecular dynamics simulations. While most contemporary symmetry-aware models use invariant convolutions and only act on scalars, NequIP employs E(3)-equivariant convolutions for interactions of geometric tensors, resulting in a more informative-rich and faithful representation of atomic environments. The method achieves state-of-the-art accuracy on a challenging and diverse set of molecules and

Equivariant neural networks

Group Equivariant Convolutional Networks

Taco S. Cohen
University of Amsterdam

T.S.COHEN@UVA.NL

Max Welling
University of Amsterdam
University of California Irvine
Canadian Institute for Advanced Research

M.WELLING@UVA.NL

Abstract

We introduce Group equivariant Convolutional Neural Networks (G-CNNs), a natural generalization of convolutional neural networks that reduces sample complexity by exploiting symme-

Convolution layers can be used effectively in a deep network because all the layers in such a network are *translation equivariant*: shifting the image and then feeding it through a number of layers is the same as feeding the original image through the same layers and then shifting the resulting feature maps (at least up to edge-effects). In

Equivariant Transformer Networks

Kal Sheng Tai¹ Peter Bailis¹ Gregory Valiant¹

Abstract

How can prior knowledge on the transformation invariances of a domain be incorporated into the architecture of a neural network? We propose Equivariant Transformers (ETs), a family of differentiable image-to-image mappings that improve the robustness of models towards pre-defined continuous transformation groups. Through the use of specially-derived canonical coordinate systems, ETs incorporate functions that

scaling to each training image). While data augmentation typically helps reduce the test error of CNN-based models, there is no guarantee that transformation invariance will be enforced for data not seen during training.

In contrast to training time approaches like data augmentation, recent work on group equivariant CNNs (Cohen & Welling, 2016; Dieleman et al., 2016; Marcos et al., 2017; Worrall et al., 2017; Henriques & Veličković, 2017; Cohen et al., 2018) has explored new CNN architectures that are

Theory for Equivariant Quantum Neural Networks

Quynh T. Nguyen,^{1,2} Louis Schatzki,^{3,4} Paolo Braccia,^{1,5} Michael Ragone,^{1,6} Patrick J. Cules,¹ Frédéric Sauvage,¹ Martin Lucocca,^{1,7} and M. Cerezo¹

¹Theoretical Division, Los Alamos National Laboratory, Los Alamos, New Mexico 87545, USA
²School of Engineering and Applied Sciences, Harvard University, Cambridge, Massachusetts 02138, USA
³Information Sciences, Los Alamos National Laboratory, Los Alamos, New Mexico 87545, USA
⁴Department of Electrical and Computer Engineering, University of Illinois at Urbana-Champaign, Urbana, Illinois 61801, USA
⁵Dipartimento di Fisica e Astronomia, Università di Firenze, Sesto Fiorentino (FI), 50019, Italy
⁶Department of Mathematics, University of California Davis, Davis, California 95616, USA
⁷Center for Nonlinear Studies, Los Alamos National Laboratory, Los Alamos, New Mexico 87545, USA

Quantum neural network architectures that have little-to-no inductive biases are known to face trainability and generalization issues. Inspired by a similar problem, recent breakthroughs in machine learning address this challenge by creating models encoding the symmetries of the learning task. This is materialized through the usage of equivariant neural networks whose action commutes with that of the symmetries. In this work, we present these ideas in the quantum realm by

An Efficient Lorentz Equivariant Graph Neural Network for Jet Tagging

Shiqi Gong^{a,b,1} Qi Meng^b Jue Zhang^b Hulin Qu^b Gonggao Li² Sitian Qian¹ Weitao Du^b Zhi-Ming Ma^a Tie-Yan Liu^b

^aAcademy of Mathematics and Systems Science, Chinese Academy of Sciences, Zhongguancun East Road, Beijing 100190, China

^bMicrosoft Research Asia, Danling Street, Beijing 100080, China

^cCERN, EP Department, CH-1211 Geneva 23, Switzerland

^dSchool of Physics, Peking University, Cheneiya Road, Beizhen 100871, China

E(3)-Equivariant Graph Neural Networks for Data-Efficient and Accurate Interatomic Potentials

Simon Batzner^{a,1} Albert Muscelian,¹ Lixin Sun,¹ Mario Geiger,² Jonathan P. Mailon,³ Mordechai Kohnluth,² Nicola Molinari,¹ Tess E. Smidt,^{4,5} and Boris Kozinsky^{a,1,2}

¹John A. Paulson School of Engineering and Applied Sciences,

Harvard University, Cambridge, MA 02138, USA

²École Polytechnique Fédérale de Lausanne, 1015 Lausanne, Switzerland

³Robert Bosch Research and Technology Center, Cambridge, MA 02139, USA

⁴Computational Research Division and Center for Advanced Mathematics for Energy Research Applications,

Lawrence Berkeley National Laboratory, Berkeley, CA 94720, USA

⁵Massachusetts Institute of Technology, Department of Electrical

Engineering and Computer Science, Cambridge, MA 02138, USA

This work presents Neural Equivariant Interatomic Potentials (NequIP), an E(3)-equivariant neural network approach for learning interatomic potentials from *ab-initio* calculations for molecular dynamics simulations. While most contemporary symmetry-aware models use invariant convolutions and only act on scalars, NequIP employs E(3)-equivariant convolutions for interactions of geometric tensors, resulting in a more information-rich and faithful representation of atomic environments. The method achieves state-of-the-art accuracy on a challenging and diverse set of molecules and

HIERARCHICAL, ROTATION-EQUIVARIANT NEURAL NETWORKS TO SELECT STRUCTURAL MODELS OF PROTEIN COMPLEXES

Stephan Eismann^a
Department of Applied Physics
Stanford University
seismann@stanford.edu

Raphael J.L. Townsend^b
Department of Computer Science
Stanford University
raphael@cs.stanford.edu

Nathaniel Thomas^b
Department of Physics
Stanford University
nthomas103@gmail.com

Milind Jagota
Department of Electrical Engineering
Stanford University
mijagota@stanford.edu

Bowen Jing
Department of Computer Science
Stanford University
bjing@stanford.edu

Ron O. Dror
Department of Computer Science
Stanford University
rondror@cs.stanford.edu

ABSTRACT

Predicting the structure of multi-protein complexes is a grand challenge in biochemistry, with major implications for basic science and drug discovery. Computational structure prediction methods generally leverage pre-defined structural features to distinguish accurate structural models from less accurate ones. This raises the question of whether it is possible to learn characteristics of accurate models directly from atomic coordinates of protein complexes, with no prior assumptions. Here we introduce a machine learning method that learns directly from the 3D positions of all atoms to

Equivariant neural networks

Group Equivariant Convolutional Networks

Taco S. Cohen
University of Amsterdam
Max Welling
University of Amsterdam
University of California Irvine
Canadian Institute for Advanced Research

T.S.COHEN@UVA.NL

M.WELLING@UVA.

Abstract

We introduce Group equivariant Convolutional Neural Networks (G-CNNs), a natural generalization of convolutional neural networks that reduces sample complexity by exploiting symme-

Convolution layers can be used effectively in a deep work because all the layers in such a network are *invariant equivariant*: shifting the image and then feeding it through a number of layers is the same as feeding original image through the same layers and then shift the resulting feature maps (at least up to edge-effects

An Efficient Lorentz Equivariant Graph Neural Network for Jet Tagging

Shiqi Gong^{a,b,1} Qi Meng^b Jue Zhang^b Hulin Qu^c Congqiao Li² Sitian Qian^d Weitao Du^b Zhi-Ming Ma^a Tie-Yan Liu^b

^aAcademy of Mathematics and Systems Science, Chinese Academy of Sciences, Zhongguancun East Road, Beijing 100190, China

^bMicrosoft Research Asia, Danling Street, Beijing 100080, China

^cCERN, EP Department, CH-1211 Geneva 23, Switzerland

^dSchool of Physics, Peking University, Cheneiya Road, Beizhina 100871, China

Equivariant Transformer Networks

Karl Schaefer^{1,2,3} Boris Reif^{1,2} Gerasimos Voulas¹

Geometric Deep Learning and Equivariant Neural Networks

JAN E. GERKEN¹, JIMMY ARONSSON^{1*}, OSCAR CARLSSON^{1*}, HAMPUS LINANDER²,
FREDRIK OHLSSON³, CHRISTOFFER PETERSSON^{1,4} AND DANIEL PERSSON¹

¹ Chalmers University of Technology, Department of Mathematical Sciences
SE-412 96 Gothenburg, Sweden

² Gothenburg University, Department of Physics
SE-412 96, Gothenburg, Sweden

³ Umeå University, Department of Mathematics and Mathematical Statistics
SE-901 87, Umeå, Sweden

⁴ Zenseact
SE-417 56, Gothenburg, Sweden

* equal contribution

Academy of Mathematics and Systems Science, Chinese Academy of Sciences,
Zhongguancun East Road, Beijing 100190, China

This work presents Neural Equivariant Interaction Potentials (NeqIP), an E(3)-equivariant neural network approach for learning interatomic potentials from *ab-initio* calculations for molecular dynamics simulations. While most contemporary symmetry-aware models use invariant convolutions and only act on scalars, NeqIP employs E(3)-equivariant convolutions for interactions of geometric tensors, resulting in a more information-rich and faithful representation of atomic environments. The method achieves state-of-the-art accuracy on a challenging and diverse set of molecules and

Theory for Equivariant Quantum Neural Networks

Quynh T. Nguyen,^{1,2} Louis Schatzki,^{3,4} Paolo Braccia,^{1,5} Michael Ragone,^{1,6}
Patrick J. Cules,¹ Frédéric Sauvage,¹ Martin Lucero,^{1,7} and M. Cerezo¹

¹Theoretical Division, Los Alamos National Laboratory, Los Alamos, New Mexico 87545, USA
²School of Engineering and Applied Sciences, Harvard University, Cambridge, Massachusetts 02138, USA

³Information Science, Los Alamos National Laboratory, Los Alamos, New Mexico 87545, USA
⁴Department of Electrical and Computer Engineering,

University of Illinois at Urbana-Champaign, Urbana, Illinois 61801, USA
⁵Dipartimento di Fisica e Astronomia, Università di Firenze, Sesto Fiorentino (FI), 50019, Italy

⁶Department of Mathematics, University of California Davis, Davis, California 95616, USA
⁷Center for Nonlinear Studies, Los Alamos National Laboratory, Los Alamos, New Mexico 87545, USA

Quantum neural network architectures that have little-to-no inductive biases are known to face trainability and generalization issues. Inspired by a similar problem, recent breakthroughs in machine learning address this challenge by creating models encoding the symmetries of the learning task. This is materialized through the usage of equivariant neural networks whose action commutes with that of the symmetries. In this work, we present these ideas to the quantum realm by

Hi-C Hierarchy, Rotation-Equivariant Neural Networks to Select Structural Models of Protein Complexes

Stephan Eismann*
Department of Applied Physics
Stanford University
seismann@stanford.edu

Raphael J.L. Townshend*
Department of Computer Science
Stanford University
raphael@cs.stanford.edu

Nathaniel Thomas*
Department of Physics
Stanford University
nthomas103@gmail.com

Milind Jagota
Department of Electrical Engineering
Stanford University
mijagota@stanford.edu

Bowen Jing
Department of Computer Science
Stanford University
bjing@stanford.edu

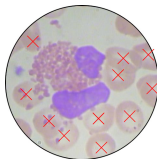
Ron O. Dror
Department of Computer Science
Stanford University
rondror@cs.stanford.edu

ABSTRACT

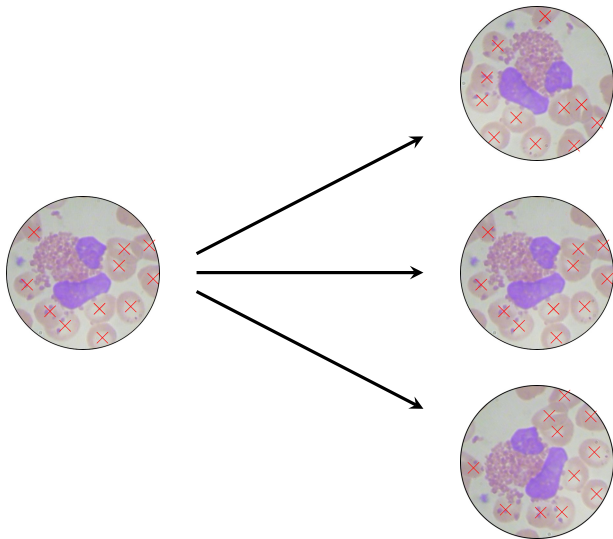
Predicting the structure of multi-protein complexes is a grand challenge in biochemistry, with major implications for basic science and drug discovery. Computational structure prediction methods generally leverage pre-defined structural features to distinguish accurate structural models from less accurate ones. This raises the question of whether it is possible to learn characteristics of accurate models directly from atomic coordinates of protein complexes, with no prior assumptions. Here we introduce a machine learning method that learns directly from the 3D positions of all atoms to

Data augmentation

Data augmentation



Data augmentation



Data augmentation

Data augmentation

Article

Accurate structure prediction of biomolecular interactions with AlphaFold 3

<https://doi.org/10.1038/s41586-024-03487-y>

Received: 18 December 2022

Accepted: 28 April 2024

Published online: 8 May 2024

Open access

 Check for updates

Josh Abramson^{1,2}, Anna Adler^{1,2}, Jack Dongus^{1,2}, Richard Evans^{1,2}, Tim Green^{1,2}, Alexander Pritzel^{1,2}, Olaf Ronneberger^{1,2}, Lindsay Wilkerson^{1,2}, Andrew J. Ballard¹, Joshua Bonbrink¹, Sebastian W. Bowers^{1,2}, David A. Evans¹, Chao-Chun Hung¹, Michael O'Neill¹, David Reissner¹, Kathryn Tarnopolska^{1,2}, Zachary Wu¹, Abhishek Baghel^{1,2}, Erni Arvaniti¹, Charles Beattie¹, Ottavia Bertolli¹, Alan Bridgland¹, Alessio Chignone¹, Miles Congreve¹, Alexander I. Cowen-Rivers¹, Andrew Cowie¹, Michael Figurnov¹, Fabian B. Fuchs¹, Harshad Guadagni¹, Rohan Jain¹, Yusuf A. Khan^{1,2}, Caroline M. M. Lee¹, Kuba Peltin¹, Anna Potapenko¹, Pascal Savay¹, Sukhdeep Singh¹, Adrian Stecula¹, Ashutk Thilakavandana¹, Catherine Tong¹, Sergei Yelensky¹, Ellen D. Zhang^{1,2}, Michal Zaidman¹, Augustin Eden¹, Victor Bapst^{1,2}, Pauliusius Kubilius^{1,2}, Max Jaderberg^{1,2}, Demis Hassabis^{1,2,3} & John M. Jumper^{1,2}

The introduction of AlphaFold 2¹ has spurred a revolution in modelling the structure of proteins and their interactions, enabling a huge range of applications in protein modelling and design^{2–5}. Here we describe our AlphaFold 3 model with a substantially updated diffusion-based architecture that is capable of predicting the joint structure of complexes including proteins, nucleic acids, small molecules, ions and modified residues. The new AlphaFold model demonstrates substantially improved accuracy

Data augmentation

Article

Accurate structure prediction of biomolecular interactions with AlphaFold 3

<https://doi.org/10.1038/s41586-024-03847-y>

Received: 18 December 2022

Accepted: 29 April 2024

Published online: 8 May 2024

Open access

 Check for updates

Josh Abramson^{1,2}, Anna Adler^{1,2}, Jack Dongus^{1,2}, Richard Evans^{1,2}, Tim Green^{1,2}, Alexander Holm^{1,2}, Otil Ronneberger^{1,2}, Lindsay Wilkerson^{1,2}, Andrew J. Ballard¹, Joshua Bonbright¹, Sebastian W. Boderstein¹, David A. Evans¹, Chao-Chun Hung¹, Michael O'Neill¹, David Reissner¹, Kathryn Tarnopolsky^{1,2}, Zachary Wu¹, Abhishek Bera^{1,2}, Evin Arvanis¹, Charles Beattie¹, Ottavia Bertolli¹, Alan Bridgland¹, Alessio Chiarapane¹, Miles Congreve¹, Alexander I. Cowen-Rivers¹, Andrew Cowie¹, Michael Figarova¹, Fabian B. Fuchs¹, Hannah Gaudin¹, Rishabh Jain¹, Yousaf A. Khan^{1,2}, Caroline M. K. Lee¹, Kuba Perlis¹, Anna Potapenko¹, Pascal Savry¹, Sukhdeep Singh¹, Adrian Stecula¹, Ashok Thirumangalakudi¹, Catherine Tong¹, Sergei Yalovoy¹, Ellen D. Zhang^{1,2}, Michael Zielinski¹, Augustin Eden¹, Victor Bapst^{1,2}, Pauliusius Kuba^{1,2}, Max Jaderberg^{1,2}, Denis Hassabis^{1,2,3} & John M. Jumper^{1,2}

The introduction of AlphaFold 2¹ has spurred a revolution in modelling the structure of proteins and their interactions, enabling a huge range of applications in protein modelling and design^{2–5}. Here we describe our AlphaFold 3 model with a substantially updated diffusion-based architecture that is capable of predicting the joint structure of complexes including proteins, nucleic acids, small molecules, ions and modified residues. The new AlphaFold model demonstrates substantial accuracy

The Importance of Being Scalable: Improving the Speed and Accuracy of Neural Network Interatomic Potentials Across Chemical Domains

Eric Qu
UC Berkeley
ericqu@berkeley.edu

Aditi S. Krishnapriyan
UC Berkeley, LBNL
aditik1@berkeley.edu

Abstract

Scaling has been a critical factor in improving model performance and generalization across various fields of machine learning. It involves how a model's performance changes with increases in model size or input data, as well as how efficiently computational resources are utilized to support this growth. Despite successes in scaling other types of machine learning models, the study of scaling in Neural Network Interatomic Potentials (NNIPs) remains limited. NNIPs act as

Data augmentation

Article

Accurate structure prediction of biomolecular interactions with AlphaFold 3

<https://doi.org/10.1038/s41586-024-03847-9>

Received: 18 December 2022

Accepted: 29 April 2024

Published online: 8 May 2024

Open access

 Check for updates

Josh Abramson^{1,2}, Jonas Adler^{3,4}, Jack Angus^{5,6}, Richard Evans^{7,8}, Tim Green⁹, Alexander Pritzel¹⁰, Olaf Ronneberger¹¹, Lindsay Wilkerson¹², Andrew J. Ballard¹³, Joshua Bonbright¹⁴, Sebastian W. Bowers¹⁵, David A. Evans¹⁶, Chao-Chun Hung¹⁷, Michael O'Neill¹⁸, David Reissner¹⁹, Kathryn Tarmann-Grosholz²⁰, Zachary Wu²¹, Abella Zarraga-Pol²², Eric Arvaniti²³, Charles Beattie²⁴, Ottavia Bertolli²⁵, Alex Bridgland²⁶, Alessio Cingolano²⁷, Miles Congreve²⁸, Alexander I. Cowen-Rivers²⁹, Andrew Cowie³⁰, Michael Figurnov³¹, Fabian D. Fuchs³², Harsh Gaudin³³, Rishabh Jain³⁴, Yousaf A. Khan³⁵, Caroline M. L. Lee³⁶, Kuba Peltin³⁷, Anna Potapenko³⁸, Pascal Savy³⁹, Sukhdeep Singh⁴⁰, Adrian Stecula⁴¹, Ashok Thirumangalakudi⁴², Catherine Tong⁴³, Sergei Yelensky⁴⁴, Ellen D. Zhang⁴⁵, Michael Zaltewski⁴⁶, Augustin Zidek⁴⁷, Victor Zinger⁴⁸, Pauliusius Zukas⁴⁹, Max Jaderberg^{49,50}, Demis Hassabis^{49,51} & John M. Jumper^{49,52}

The introduction of AlphaFold 2¹ has spurred a revolution in modelling the structure of proteins and their interactions, enabling a huge range of applications in protein modelling and design^{2–5}. Here we describe our AlphaFold 3 model with a substantially updated diffusion-based architecture that is capable of predicting the joint structure of complexes including proteins, nucleic acids, small molecules, ions and modified residues. The new AlphaFold model demonstrates substantial improved accuracy

The Importance of Being Scalable: Improving the Speed and Accuracy of Neural Network Interatomic Potentials Across Chemical Domains

Eric Qu
UC Berkeley
ericqu@berkeley.edu

Aditi S. Krishnapriyan
UC Berkeley, LBNL
aditiki@berkeley.edu

Abstract

Scaling has been a critical factor in improving model performance and generalization across various fields of machine learning. It involves how a model's performance changes with increases in model size or input data, as well as how efficiently computational resources are utilized to support this growth. Despite successes in scaling other types of machine learning models, the study of scaling in Neural Network Interatomic Potentials (NNIPs) remains limited. NNIPs act as

Swallowing the Bitter Pill: Simplified Scalable Conformer Generation

Yiyang Wang¹, Ahmed A. Elhag^{1,2}, Navdeep Jaitly¹, Joshua M. Susskind¹, Miguel Ángel Bautista¹

Abstract

We present a novel way to predict molecular conformers through a simple formulation that sidesteps many of the heuristics of prior works and achieves state-of-the-art results by using the advantages of scale. By training a diffusion generative model directly on 3D atomic positions without making assumptions about the explicit structure of molecules (e.g. modeling torsional angles) we are able to radically simplify structure prediction and enable us to scale our model

is the vast complexity of the 3D structure space, encompassing factors such as bond lengths and torsional angles. Despite the molecular graph dictating potential 3D conformers through specific constraints, such as bond types and spatial arrangements determined by chiral centers, the conformational space experiences exponential growth with the expansion of the graph size and the number of rotatable bonds (Aschrod & Gomez-Bombarelli, 2022). This complicates brute force and exhaustive approaches, making them virtually unfeasible for even moderately small molecules. Systematic methods like OMPGA (Hawkins et al., 2010)

Data augmentation

Article

Accurate structure prediction of biomolecular interactions with AlphaFold 3

<https://doi.org/10.1038/s41586-024-03877-w>

Received: 18 December 2023

Accepted: 28 April 2024

Published online: 8 May 2024

Open access

 Check for updates

Josh Abramson^{1,2}, Jonas Adler^{1,2}, Jack Angus^{1,2}, Richard Evans^{1,2}, Tim Green^{1,2}, Alexander Grigorev^{1,2}, Olaf Ronneberger^{1,2}, Lindsay Wilkerson^{1,2}, Andrew J. Ballard¹, Joshua Bonbright¹, Sebastian W. Boderstein¹, David A. Evans¹, Chao-Chun Hung¹, Michael D. Heide¹, David Reissner¹, Kathryn Tarmann-Reissner¹, Zachary Wu¹, Abhishek Bera^{1,3,4}, Enzo Arcañes¹, Charles Beattie¹, Ottavia Bertolli¹, Alex Bridgland¹, Alessio Cincorin^{1,5}, Miles Congreve¹, Alexander I. Cowen-Rivers¹, Andrew Cowie¹, Michael Figurnov¹, Fabian B. Fuchs¹, Hannah Gaudin¹, Rishabh Jain¹, Yousuik Kim^{1,2}, Caroline M. Lee¹, Kuba Peltin¹, Anna Potapenko¹, Pascal Savary¹, Sukhdeep Singh¹, Adrian Stecula¹, Ashish Thakkar^{1,6,7}, Catherine Tong¹, Sergei Yeliseyev¹, Ellen D. Zhang^{1,2}, Michael Zaitsev^{1,2}, Augustin Zidek¹, Victor Zippori¹, Pauline Zito^{1,2}, Max Jaderberg^{1,2}, Denis Hassabis^{1,2,3,4} & John M. Jumper^{1,2}

The introduction of AlphaFold 2¹ has spurred a revolution in modelling the structure of proteins and their interactions, enabling a huge range of applications in protein modelling and design^{2–5}. Here we describe our AlphaFold 3 model with a substantially updated diffusion-based architecture that is capable of predicting the joint structure of complexes including proteins, nucleic acids, small molecules, ions and modified residues. The new AlphaFold 3 model demonstrates substantial improvements in accuracy

The Importance of Being Scalable: Improving the Speed and Accuracy of Neural Network Interatomic Potentials Across Chemical Domains

Eric Qu
UC Berkeley
ericqu@berkeley.edu

Aditi S. Krishnapriyan
UC Berkeley, LBNL
aditik1@berkeley.edu

Abstract

Scaling has been a critical factor in improving model performance and generalization across various fields of machine learning. It involves how a model's performance changes with increases in model size or input data, as well as how efficiently computational resources are utilized to support this growth. Despite successes in scaling other types of machine learning models, the study of scaling in Neural Network Interatomic Potentials (NNIPs) remains limited. NNIPs act as

Swallowing the Bitter Pill: Simplified Scalable Conformer Generation

Yiyang Wang¹, Ahmed A. Elhag^{1,2}, Navdeep Jaitly¹, Joshua M. Susskind¹, Miguel Ángel Bautista¹

Abstract

We present a novel way to predict molecular conformers through a simple formulation that sidesteps many of the heuristics of prior works and achieves state-of-the-art results by using the advantages of scale. By training a diffusion generative model directly on 3D atomic positions without making assumptions about the explicit structure of molecules (e.g. modeling torsional angles) we are able to radically simplify structure prediction and make it suitable for scale.

is the vast complexity of the 3D structure space, encompassing factors such as bond lengths and torsional angles. Despite the molecular graph dictating potential 3D conformers through specific constraints, such as bond types and spatial arrangements determined by chiral centers, the conformational space experiences exponential growth with the expansion of the graph size and the number of rotatable bonds (Aschrod & Gomez-Bombarelli, 2022). This complicates brute force and exhaustive approaches, making them virtually unfeasible for even moderately small molecules.

Systematic methods like OMEGA (Hawkins et al., 2010)

Probing the effects of broken symmetries in machine learning

Marcel F. Langer¹, Sergey N. Poudyvalov² and Michele Ceriotti¹

Laboratory of Computational Science and Modeling and National Centre for Computational Design and Discovery of Novel Materials MARVEL, Institute of Materials, École Polytechnique Fédérale de Lausanne, 1015 Lausanne, Switzerland

¹ Author to whom any correspondence should be addressed.

E-mail: michele.ceriotti@epfl.ch

Keywords: machine learning, symmetry-constrained models, atomistic modeling, molecular simulation

Supplementary material for this article is available online

Abstract

Symmetry is one of the most central concepts in physics, and it is no surprise that it has also been widely adopted as an inductive bias for machine-learning models applied to the physical sciences. This is especially true for models targeting the properties of matter at the atomic scale. Both established and state-of-the-art approaches, with almost no exceptions, are built to be exactly equivariant to translations, permutations, and rotations of the atoms. Incorporating symmetries—rotations in particular—constrains the model design space and implies more complicated architectures that are often also computationally demanding. There are indications

Data augmentation

Article

Accurate structure prediction of biomolecular interactions with AlphaFold 3

<https://doi.org/10.1038/s41586-024-03847-w>

Received: 18 December 2023

Accepted: 28 April 2024

Published online: 8 May 2024

Open access

 Check for updates

Josh Abramson^{1,2}, Jonas Adler³, Jack Angus⁴, Richard Evans⁵, Tim Green^{1,2}, Alexander Pritzel⁶, Olaf Ronneberger⁶, Lindsay Wilkerson⁷, Andrew J. Ballard⁸, Joshua Bonbright⁹, Sebastian W. Bowers¹⁰, David A. Evans¹, Chao-Chun Hung¹, Michael O'Neill¹, David Rea¹⁰, Kathryn Teyssie^{10,11}, Zachary Wu¹, Alexia Bergel^{10,11}, Enzo Anyan¹, Charles Beattie¹, Ottavia Bertolli¹, Alex Bridgland¹, Alessio Cingolani¹, Miles Congreve¹, Alexander I. Cowen-Rivers¹, Andrew Cowie¹, Michael Figurnov¹, Fabian D. Fuchs¹, Hannah Gussone¹, Rohan Jain¹, Yousaf A. Khan¹², Caroline M. Lee¹, Kuba Peltin¹, Anna Potapenko¹, Pascal Savay¹, Bultheep Singh¹, Adrian Stecula¹, Ashwin Thirumangalakudi¹, Catherine Tong¹, Sergei Yelensky¹, Ellen D. Zhang¹³, Michal Zaidman¹⁴, Augustin Zidek¹⁵, Victor Zippori¹⁶, Paulius Zukas¹⁷, Max Jaderberg^{18,19}, Demis Hassabis^{1,20} & John M. Jumper^{1,21}

The introduction of AlphaFold 2¹ has spurred a revolution in modelling the structure of proteins and their interactions, enabling a huge range of applications in protein modelling and design^{2–5}. Here we describe our AlphaFold 3 model with a substantially updated diffusion-based architecture that is capable of predicting the joint structure of complexes including proteins, nucleic acids, small molecules, ions and modified residues. The new AlphaFold 3 model demonstrates substantial improved accuracy

The Importance of Being Scalable: Improving the Speed and Accuracy of Neural Network Interatomic Potentials Across Chemical Domains

Eric Qu
UC Berkeley
ericqu@berkeley.edu

Aditi S. Krishnamurthy
UC Berkeley, LBNL
aditik1@berkeley.edu

Abstract

Scaling has been a critical factor in improving model performance and generalization across various fields of machine learning. It involves how a model's performance changes with increases in model size or input data, as well as how efficiently computational resources are utilized to support this growth. Despite successes in scaling other types of machine learning models, the study of scaling in Neural Network Interatomic Potentials (NNIPs) remains limited. NNIPs act as

Swallowing the Bitter Pill: Simplified Scalable Conformer Generation

Yiyang Wang¹, Ahmed A. Elhag^{1,2}, Navdeep Jaitly¹, Joshua M. Susskind¹, Miguel Ángel Bautista¹

Abstract

We present a novel way to predict molecular conformers through a simple formulation that sidesteps many of the heuristics of prior works and achieves state-of-the-art results by using the advantages of scale. By training a diffusion generative model directly on 3D atomic positions without making assumptions about the explicit structure of molecules (e.g. modeling torsional angles) we are able to radically simplify structure generation and make it suitable to scale on the

is the vast complexity of the 3D structure space, encompassing factors such as bond lengths and torsional angles. Despite the molecular graph dictating potential 3D conformers through specific constraints, such as bond types and spatial arrangements determined by chiral centers, the conformational space experiences exponential growth with the expansion of the graph size and the number of rotatable bonds (Ashford & Gomez-Bombarelli, 2022). This complicates brute force and exhaustive approaches, making them virtually unfeasible for even moderately small molecules.

Systematic methods like OMFGA (Hawkins et al., 2010)

Probing the effects of broken symmetries in machine learning

Marcel F. Langer¹, Sergey N. Rudnitsky² & Michele Ceriotti¹

Laboratory of Computational Science and Modeling and National Centre for Computational Design and Discovery of Novel Materials MARVEL, Institute of Materials, Ecole Polytechnique Fédérale de Lausanne, 1015 Lausanne, Switzerland

* Author to whom any correspondence should be addressed.

E-mail: michele.ceriotti@epfl.ch

Keywords: machine learning, symmetry-constrained models, atomistic modeling, molecular simulation

Supplementary material for this article is available online

Abstract

Symmetry is one of the most central concepts in physics, and it is no surprise that it has also been widely adopted as an inductive bias for machine-learning models applied to the physical sciences. This is especially true for models targeting the properties of matter at the atomic scale. Both established and state-of-the-art approaches, with almost no exceptions, are built to be exactly equivariant to translations, permutations, and rotations of the atoms. Incorporating symmetries—rotations in particular—constrains the model design space and implies more complicated architectures that are often also computationally demanding. There are indications

Two for One: Diffusion Models and Force Fields for Coarse-Grained Molecular Dynamics

Marloes Arts,^{1,2,3,4} Victor Garcia Satorras,^{5,6,7} Chin-Wei Huang,⁸ Daniel Zügner,⁹ Marco Federici,^{1,2} Cecilia Clementi,^{3,2} Frank Noé,¹ Robert Pinsler,⁸ and Rianne van den Berg⁵

¹Work done during an internship at Microsoft Research (Amsterdam).

²University of Copenhagen, Department of Computer Science, Universitetsparken 1, Copenhagen, 2100, Denmark.

³AI4Science, Microsoft Research, Evert van de Beekstraat 354, Amsterdam, 1118 CZ, The Netherlands.

⁴AI4Science, Microsoft Research, Karl-Liebknecht-Straße 32, Berlin, 10178, Germany.

⁵University of Amsterdam, Informatics Institute, Science Park 904, Amsterdam, 1098 XH, The Netherlands.

⁶Freie Universität Berlin, Department of Physics, Arnimallee 12, Berlin, 14195, Germany.

⁷AI4Science, Microsoft Research, 21 Station Road, Cambridge, CB1 3FB, United Kingdom.

⁸Equal contribution.

*E-mail: ma@cs.ku.dk; victorgarcia@microsoft.com

Abstract

Coarse-grained (CG) molecular dynamics enables the study of biological processes at temporal and spatial scales that would be intractable at an atomistic resolution. However, accurately learning a CG force field remains a challenge. In this work, we leverage connections between score-based generative models, force fields and molecular

Data augmentation

Article

Accurate structure prediction of biomolecular interactions with AlphaFold 3

<https://doi.org/10.1038/s41586-024-03887-w>

Received: 18 December 2023

Accepted: 28 April 2024

Published online: 8 May 2024

Open access

Check for updates

Josh Abramson^{1,2}, Jonas Adler³, Jack Angus¹, Richard Evans¹, Tim Green¹, Alexander Pritzel¹, Olaf Ronneberger¹, Lindsay Wilkerson¹, Andrew J. Ballard¹, Joshua Bonbrink¹, Sebastian W. Bodenstein¹, David A. Evans¹, Chao-Chun Hung¹, Michael D. Heule¹, David Henderson¹, Kathryn Torrance-Hunt¹, Zachary Wu¹, Abhishek Bera¹, Enzo Anyan¹, Charles Beattie¹, Ottavia Bertolli¹, Alex Bridgland¹, Alessio Cimpanaru¹, Miles Congreve¹, Alexander I. Cowen-Rivers¹, Andrew Cowie¹, Michael Figurnov¹, Fabian B. Fuchs¹, Harsh Gaudin¹, Rohan Joshi¹, Yousuik Kim¹, Caroline M. Lee¹, Kuba Peltin¹, Anna Potapenko¹, Pascal Savay¹, Sukhdeep Singh¹, Adrian Stecula¹, Ashwin Thillaisandaran¹, Catherine Tong¹, Sergei Yalovoy¹, Ellen D. Zhang¹, Michal Zaidman¹, Augustin Zidek¹, Victor Zippori¹, Paulius Zukas¹, Max Jaderberg^{4,5}, Demis Hassabis^{1,6,7} & John M. Jumper^{4,5}

The introduction of AlphaFold 2¹ has spurred a revolution in modelling the structure of proteins and their interactions, enabling a huge range of applications in protein modelling and design². Here we describe our AlphaFold 3 model with a substantially updated diffusion-based architecture that is capable of predicting the joint structure of complexes including proteins, nucleic acids, small molecules, ions and modified residues. The new AlphaFold 3 model demonstrates substantial improvements in accuracy

The Importance of Being Scalable: Improving the Speed and Accuracy of Neural Network Interatomic Potentials Across Chemical Domains

Eric Qu
UC Berkeley
ericqu@berkeley.edu

Aditi S. Krishnamurthy
UC Berkeley, LBNL
aditiks@berkeley.edu

Abstract

Scaling has been a critical factor in improving model performance and generalization across various fields of machine learning. It involves how a model's performance changes with increases in model size or input data, as well as how efficiently computational resources are utilized to support this growth. Despite successes in scaling other types of machine learning models, the study of scaling in Neural Network Interatomic Potentials (NNIPs) remains limited. NNIPs act as

Two for One: Diffusion Models and Force Fields for Coarse-Grained Molecular Dynamics

Marloes Arts,^{1,2,3,4} Victor Garcia Satorras,^{5,6,7} Chin-Wei Huang,⁸ Daniel Zügner,⁹ Marco Federici,^{1,2} Cecilia Clementi,^{3,4} Frank Noé,⁵ Robert Pinsler,⁶ and Rianne van den Berg⁸

¹Work done during an internship at Microsoft Research (Amsterdam).
²University of Copenhagen, Department of Computer Science, Universitetsparken 1, Copenhagen, 2100, Denmark.
³AI4Science, Microsoft Research, Evert van de Beekstraat 354, Amsterdam, 1118 CZ, The Netherlands.
⁴AI4Science, Microsoft Research, Karl-Liebknecht-Straße 32, Berlin, 10178, Germany.
⁵University of Amsterdam, Informatics Institute, Science Park 904, Amsterdam, 1098 XH, The Netherlands.
⁶Freie Universität Berlin, Department of Physics, Arnimallee 19, Berlin, 14195, Germany.
⁷AI4Science, Microsoft Research, 21 Station Road, Cambridge, CB1 3FB, United Kingdom.
⁸Equal contribution.

*E-mail: ma@cs.ku.dk; victorga@microsoft.com

Abstract

Coarse-grained (CG) molecular dynamics enables the study of biological processes at temporal and spatial scales that would be intractable at an atomistic resolution. However, accurately learning a CG force field remains a challenge. In this work, we leverage connections between score-based generative models, force fields and molecular

Swallowing the Bitter Pill: Simplified Scalable Conformer Generation

Yiyang Wang¹, Ahmed A. Elhag^{1,2}, Navdeep Jaitly¹, Joshua M. Susskind¹, Miguel Ángel Bautista¹

Abstract

We present a novel way to predict molecular conformers through a simple formulation that sidesteps many of the heuristics of prior works and achieves state-of-the-art results by using the advantages of scale. By training a diffusion generative model directly on 3D atomic positions without making assumptions about the explicit structure of molecules (e.g. modeling torsional angles) we are able to radically simplify structure prediction and make it suitable for scale.

is the vast complexity of the 3D structure space, encompassing factors such as bond lengths and torsional angles. Despite the molecular graph dictating potential 3D conformers through specific constraints, such as bond types and spatial arrangements determined by chiral centers, the conformational space experiences exponential growth with the expansion of the graph size and the number of rotatable bonds (Aschrod & Gomez-Bombarelli, 2022). This complicates brute force and exhaustive approaches, making them virtually unfeasible for even moderately small molecules. Scalability methods like OMGA (Hawkins et al., 2010)

Probing the effects of broken symmetries in machine learning

Marcel F. Langer¹, Sergey N. Vudnyakov² and Michele Corietti¹

¹Laboratoire Computational Science and Modeling and National Centre for Computational Design and Discovery of Novel Materials MARVEL, Institute of Materials, Ecole Polytechnique Fédérale de Lausanne, 1015 Lausanne, Switzerland

* Author to whom any correspondence should be addressed.

E-mail: michele.corietti@epfl.ch

Keywords: machine learning, symmetry-constrained models, atomistic modeling, molecular simulation
Supplementary material for this article is available online

Abstract

Symmetry is one of the most central concepts in physics, and it is no surprise that it has also been widely adopted as an inductive bias for machine-learning models applied to the physical sciences. This is especially true for models targeting the properties of matter at the atomic scale. Both established and state-of-the-art approaches, with almost no exceptions, are built to be exactly equivariant to translations, permutations, and rotations of the atoms. Incorporating symmetries—rotations in particular—constrains the model design space and implies more complicated architectures that are often also computationally demanding. There are indications

DOES EQUIVARIANCE MATTER AT SCALE?

Johann Brechmer¹, Sönke Behrendt², Pin de Haan³, Theo Cohen⁴
Qualcomm AI Research¹
naill1@qualcomm.com

ABSTRACT

Given large data sets and sufficient compute, is it beneficial to design neural architectures for the structure and symmetries of each problem? Or is it more efficient to learn them from data? We study empirically how equivariant and non-equivariant networks scale with compute and training samples. Focusing on a benchmark problem of rigid-body interactions and on general-purpose transformer architectures, we perform a series of experiments, varying the model size, training steps, and dataset size. We find evidence for three conclusions. First, equivariance improves data efficiency, but training non-equivariant models with data augmentation can close this gap given sufficient epochs. Second, scaling with compute follows a power law, with equivariant models outperforming non-equivariant ones at each tested compute budget. Finally, the optimal allocation of a compute budget onto model size and training duration differs between equivariant and non-equivariant models.

Data augmentation

- 👍 Easy to implement
- 👍 No specialized architecture necessary

Data augmentation

- 👍 Easy to implement
- 👍 No specialized architecture necessary
- 👎 No exact equivariance

Data augmentation

- 👍 Easy to implement
- 👍 No specialized architecture necessary
- 👎 No exact equivariance

Can we understand data augmentation theoretically?

Empirical NTK

- Consider continuous gradient descent

$$\frac{d\theta_\mu}{dt} = -\eta \frac{\partial \mathcal{L}(\mathcal{N}_\theta, \mathcal{D})}{\partial \theta_\mu} = -\frac{\eta}{N} \sum_{i=1}^N \frac{\partial L(\mathcal{N}_\theta(x_i), y_i)}{\partial \theta_\mu}$$

Empirical NTK

- Consider continuous gradient descent

$$\frac{d\theta_\mu}{dt} = -\eta \frac{\partial \mathcal{L}(\mathcal{N}_\theta, \mathcal{D})}{\partial \theta_\mu} = -\frac{\eta}{N} \sum_{i=1}^N \frac{\partial L(\mathcal{N}_\theta(x_i), y_i)}{\partial \theta_\mu}$$

- Then, the network evolves according to

$$\frac{d\mathcal{N}_\theta(x)}{dt} = \sum_{\mu} \frac{\partial \mathcal{N}_\theta(x)}{\partial \theta_\mu} \frac{d\theta_\mu}{dt}$$

Empirical NTK

- Consider continuous gradient descent

$$\frac{d\theta_\mu}{dt} = -\eta \frac{\partial \mathcal{L}(\mathcal{N}_\theta, \mathcal{D})}{\partial \theta_\mu} = -\frac{\eta}{N} \sum_{i=1}^N \frac{\partial L(\mathcal{N}_\theta(x_i), y_i)}{\partial \theta_\mu}$$

- Then, the network evolves according to

$$\frac{d\mathcal{N}_\theta(x)}{dt} = \sum_{\mu} \frac{\partial \mathcal{N}_\theta(x)}{\partial \theta_\mu} \frac{d\theta_\mu}{dt} = -\frac{\eta}{N} \sum_{i=1}^N \sum_{\mu} \frac{\partial \mathcal{N}(x)}{\partial \theta_\mu} \frac{\partial \mathcal{N}(x_i)}{\partial \theta_\mu} \frac{\partial L(\mathcal{N}_\theta(x_i), y_i)}{\partial \mathcal{N}(x_i)}$$

Empirical NTK

- Consider continuous gradient descent

$$\frac{d\theta_\mu}{dt} = -\eta \frac{\partial \mathcal{L}(\mathcal{N}_\theta, \mathcal{D})}{\partial \theta_\mu} = -\frac{\eta}{N} \sum_{i=1}^N \frac{\partial L(\mathcal{N}_\theta(x_i), y_i)}{\partial \theta_\mu}$$

- Then, the network evolves according to

$$\frac{d\mathcal{N}_\theta(x)}{dt} = \sum_{\mu} \frac{\partial \mathcal{N}_\theta(x)}{\partial \theta_\mu} \frac{d\theta_\mu}{dt} = -\frac{\eta}{N} \sum_{i=1}^N \sum_{\mu} \frac{\partial \mathcal{N}(x)}{\partial \theta_\mu} \frac{\partial \mathcal{N}(x_i)}{\partial \theta_\mu} \frac{\partial L(\mathcal{N}_\theta(x_i), y_i)}{\partial \mathcal{N}(x_i)}$$

With the **empirical neural tangent kernel (NTK)**

$$\Theta_{ij}^\theta(x, x') = \sum_{\mu} \frac{\partial \mathcal{N}_i(x)}{\partial \theta_\mu} \frac{\partial \mathcal{N}_j(x')}{\partial \theta_\mu}$$

Empirical NTK

Training dynamics under continuous gradient descent:

$$\frac{d\mathcal{N}_{\theta}(x)}{dt} = -\frac{\eta}{N} \sum_{i=1}^N \Theta_{\theta}(x, x_i) \frac{\partial L}{\partial \mathcal{N}(x_i)}$$

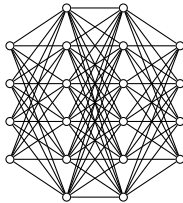
learning rate η

loss L

training sample x_i

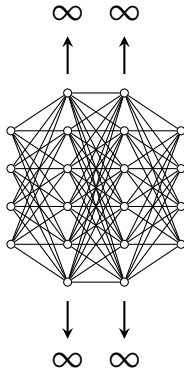
Infinite width limit

[Jacot et al. 2018]



Infinite width limit

[Jacot et al. 2018]



Infinite width limit

Consider an MLP in NTK parametrization

$$z^{(\ell)} = \frac{1}{\sqrt{n_{\ell-1}}} W^{(\ell)} \sigma(z^{(\ell-1)}(x)), \quad W^{(\ell)} \in \mathbb{R}^{n_{\ell} \times n_{\ell-1}}, \quad W_{ij}^{(\ell)} \sim \mathcal{N}(0, 1)$$

Infinite width limit

Consider an MLP in NTK parametrization

$$z^{(\ell)} = \frac{1}{\sqrt{n_{\ell-1}}} W^{(\ell)} \sigma(z^{(\ell-1)}(x)), \quad W^{(\ell)} \in \mathbb{R}^{n_{\ell} \times n_{\ell-1}}, \quad W_{ij}^{(\ell)} \sim \mathcal{N}(0, 1)$$

At infinite width

$$z_i^{(\ell)}(x) = \sqrt{n_{\ell-1}} \frac{1}{n_{\ell-1}} \sum_{j=1}^{n_{\ell-1}} W_{ij}^{(\ell)} \sigma(z_j^{(\ell-1)}(x))$$

Infinite width limit

Consider an MLP in NTK parametrization

$$z^{(\ell)} = \frac{1}{\sqrt{n_{\ell-1}}} W^{(\ell)} \sigma(z^{(\ell-1)}(x)), \quad W^{(\ell)} \in \mathbb{R}^{n_{\ell} \times n_{\ell-1}}, \quad W_{ij}^{(\ell)} \sim \mathcal{N}(0, 1)$$

At infinite width

$$z_i^{(\ell)}(x) = \sqrt{n_{\ell-1}} \underbrace{\frac{1}{n_{\ell-1}} \sum_{j=1}^{n_{\ell-1}}}_{\text{mean}} \underbrace{W_{ij}^{(\ell)} \sigma(z_j^{(\ell-1)}(x))}_{\text{i.i.d.}}$$

Infinite width limit

Consider an MLP in NTK parametrization

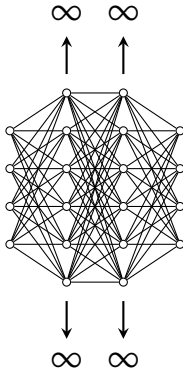
$$z^{(\ell)} = \frac{1}{\sqrt{n_{\ell-1}}} W^{(\ell)} \sigma(z^{(\ell-1)}(x)), \quad W^{(\ell)} \in \mathbb{R}^{n_{\ell} \times n_{\ell-1}}, \quad W_{ij}^{(\ell)} \sim \mathcal{N}(0, 1)$$

At infinite width

$$z_i^{(\ell)}(x) = \sqrt{n_{\ell-1}} \underbrace{\frac{1}{n_{\ell-1}} \sum_{j=1}^{n_{\ell-1}}}_{\text{mean}} \underbrace{W_{ij}^{(\ell)} \sigma(z_j^{(\ell-1)}(x))}_{\text{i.i.d.}} \sim \mathcal{N}(0, \underbrace{\text{Cov}(z_i^{(\ell)}, z_j^{(\ell)})}_{\text{NNGP}})$$

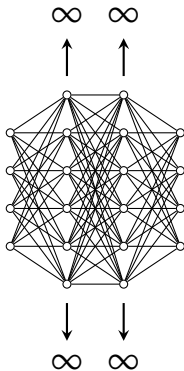
Infinite width limit

[Jacot et al. 2018]



Infinite width limit

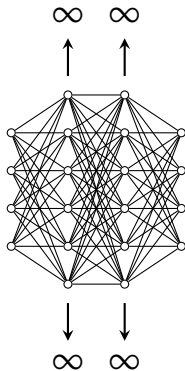
[Jacot et al. 2018]



👍 NTK becomes independent of initialization

Infinite width limit

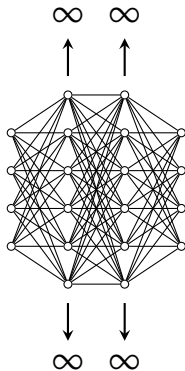
[Jacot et al. 2018]



- 👍 NTK becomes independent of initialization
- 👍 NTK becomes constant in training

Infinite width limit

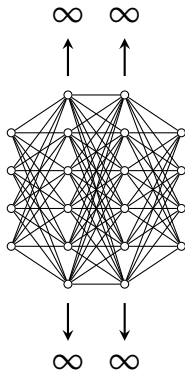
[Jacot et al. 2018]



- 👍 NTK becomes independent of initialization
- 👍 NTK becomes constant in training
- 👍 NTK can be computed for most networks

Infinite width limit

[Jacot et al. 2018]



- 👍 NTK becomes independent of initialization
- 👍 NTK becomes constant in training
- 👍 NTK can be computed for most networks
- ✓ Training dynamics can be solved

NTK at initialization

[Jacot et al. 2020]

When taking the layer widths to infinity sequentially, the empirical NTK $\Theta_{ij}^{\theta}(x, x')$ at initialization converges in probability to a deterministic kernel $\Theta(x, x')\delta_{ij}$

NTK at initialization

[Jacot et al. 2020]

When taking the layer widths to infinity sequentially, the empirical NTK $\Theta_{ij}^{\theta}(x, x')$ at initialization converges in probability to a deterministic kernel $\Theta(x, x')\delta_{ij}$

- The deterministic kernel is given in terms of a recursion over layers

NTK at initialization

[Jacot et al. 2020]

When taking the layer widths to infinity sequentially, the empirical NTK $\Theta_{ij}^{\theta}(x, x')$ at initialization converges in probability to a deterministic kernel $\Theta(x, x')\delta_{ij}$

- The deterministic kernel is given in terms of a recursion over layers
- For most common architectures, this recursion can be performed explicitly,
e.g. using neural-tangents Python package

[Novak et al. 2020]

Frozen NTK

[Jacot et al. 2020]

For a nonlinearity which is Lipschitz, twice differentiable and has bounded second derivative,

$$\Theta_{ij}^{\theta_t}(x, x') \rightarrow \Theta(x, x') \delta_{ij}$$

uniformly in t as the layer widths go to infinity sequentially.

Frozen NTK

[Jacot et al. 2020]

For a nonlinearity which is Lipschitz, twice differentiable and has bounded second derivative,

$$\Theta_{ij}^{\theta_t}(x, x') \rightarrow \Theta(x, x') \delta_{ij}$$

uniformly in t as the layer widths go to infinity sequentially.

- Intuitively, this happens because the weight updates vanish in the limit $n \rightarrow \infty$

Frozen NTK

[Jacot et al. 2020]

For a nonlinearity which is Lipschitz, twice differentiable and has bounded second derivative,

$$\Theta_{ij}^{\theta_t}(x, x') \rightarrow \Theta(x, x') \delta_{ij}$$

uniformly in t as the layer widths go to infinity sequentially.

- Intuitively, this happens because the weight updates vanish in the limit $n \rightarrow \infty$
- However, the network still learns

Mean prediction from NTK

[Jacot et al. 2018]

① At infinite width, the mean prediction is given by

$$\mu_t(x) = \Theta(x, X) \Theta(X, X)^{-1} (\mathbb{I} - e^{-\eta \Theta(X, X) t}) \gamma$$

Mean prediction from NTK

[Jacot et al. 2018]

① At infinite width, the mean prediction is given by



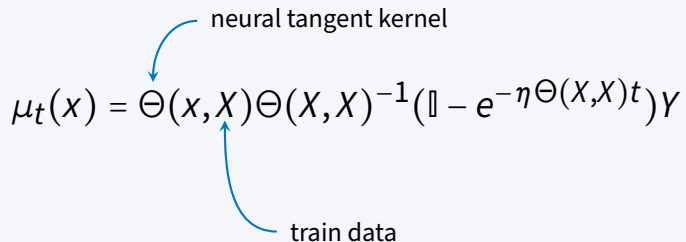
neural tangent kernel

$$\mu_t(x) = \Theta(x, X) \Theta(X, X)^{-1} (\mathbb{I} - e^{-\eta \Theta(X, X) t}) \gamma$$

Mean prediction from NTK

[Jacot et al. 2018]

① At infinite width, the mean prediction is given by



The diagram shows the equation $\mu_t(x) = \Theta(x, X) \Theta(X, X)^{-1} (\mathbb{I} - e^{-\eta \Theta(X, X) t}) \gamma$. A blue curved arrow points from the text "neural tangent kernel" to the $\Theta(x, X)$ term. Another blue curved arrow points from the text "train data" to the X in the $\Theta(X, X)$ term.

$$\mu_t(x) = \Theta(x, X) \Theta(X, X)^{-1} (\mathbb{I} - e^{-\eta \Theta(X, X) t}) \gamma$$

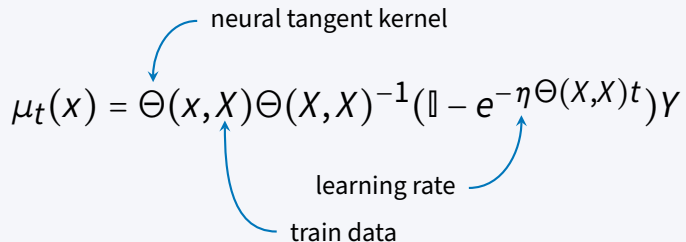
neural tangent kernel

train data

Mean prediction from NTK

[Jacot et al. 2018]

① At infinite width, the mean prediction is given by



The diagram shows the formula $\mu_t(x) = \Theta(x, X) \Theta(X, X)^{-1} (\mathbb{I} - e^{-\eta \Theta(X, X) t}) \gamma$ with three blue arrows pointing to specific parts: one from 'neural tangent kernel' to $\Theta(x, X)$, one from 'train data' to γ , and one from 'learning rate' to η .

$$\mu_t(x) = \Theta(x, X) \Theta(X, X)^{-1} (\mathbb{I} - e^{-\eta \Theta(X, X) t}) \gamma$$

neural tangent kernel

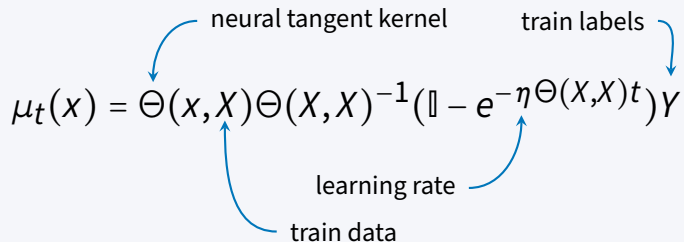
learning rate

train data

Mean prediction from NTK

[Jacot et al. 2018]

① At infinite width, the mean prediction is given by



The diagram shows the formula $\mu_t(x) = \Theta(x, X) \Theta(X, X)^{-1} (\mathbb{I} - e^{-\eta \Theta(X, X) t}) Y$ with four blue arrows pointing to its components: 'neural tangent kernel' points to $\Theta(x, X)$, 'train labels' points to Y , 'learning rate' points to η , and 'train data' points to X in the $\Theta(X, X)$ term.

$$\mu_t(x) = \Theta(x, X) \Theta(X, X)^{-1} (\mathbb{I} - e^{-\eta \Theta(X, X) t}) Y$$

Data augmentation

Data augmentation at infinite width

$$\mu_t(x) = \Theta(x, X) \Theta(X, X)^{-1} (\mathbb{I} - e^{-\eta \Theta(X, X)t}) \gamma$$

Data augmentation at infinite width

$$\mu_t(x) = \Theta(x, X) \Theta(X, X)^{-1} (\mathbb{I} - e^{-\eta \Theta(X, X) t}) \gamma$$

The diagram illustrates the components of the equation $\mu_t(x) = \Theta(x, X) \Theta(X, X)^{-1} (\mathbb{I} - e^{-\eta \Theta(X, X) t}) \gamma$. Blue arrows originate from the labels 'augmented data' and 'augmented labels' and point to specific parts of the equation:

- 'augmented data' points to x in $\Theta(x, X)$, X in $\Theta(X, X)$, and X in $\Theta(X, X)^{-1}$.
- 'augmented labels' points to γ and X in $\Theta(X, X)$.

Data augmentation at infinite width

group transformation

$$\mu_t(\rho(g)x) = \Theta(\rho(g)x, X) \Theta(X, X)^{-1} (\mathbb{I} - e^{-\eta \Theta(X, X)t}) Y$$

augmented data

augmented labels

The diagram illustrates the equation for data augmentation at infinite width. A blue arrow points from the text 'group transformation' to the term $\rho(g)$ in the expression $\rho(g)x$. Below the equation, the text 'augmented data' has four blue arrows pointing to the terms $\rho(g)x$, X , X , and X in the expression $\Theta(\rho(g)x, X) \Theta(X, X)^{-1}$. The text 'augmented labels' has two blue arrows pointing to the terms X and X in the expression $\Theta(X, X)^{-1}$. A final blue arrow points from 'augmented labels' to the Y term in the expression $(\mathbb{I} - e^{-\eta \Theta(X, X)t}) Y$.

Kernel transformation

The neural tangent kernel Θ as well as the NNGP kernel K transform according to

$$\begin{aligned}\Theta(\rho(g)x, \rho(g)x') &= \rho_K(g) \Theta(x, x') \rho_K^\top(g), \\ K(\rho(g)x, \rho(g)x') &= \rho_K(g) K(x, x') \rho_K^\top(g),\end{aligned}$$

for all $g \in G$ and $x, x' \in X$.

Kernel transformation

The neural tangent kernel Θ as well as the NNGP kernel K transform according to

$$\begin{aligned}\Theta(\rho(g)x, \rho(g)x') &= \rho_K(g) \Theta(x, x') \rho_K^\top(g), \\ K(\rho(g)x, \rho(g)x') &= \rho_K(g) K(x, x') \rho_K^\top(g),\end{aligned}$$

for all $g \in G$ and $x, x' \in X$.

Hence, for MLPs,

$$\Theta(\rho(g)x, \rho(g)x') = \Theta(x, x') \quad \Rightarrow \quad \Theta(\rho(g)x, x') = \Theta(x, \rho^{-1}(g)x')$$

Permutation shift

- On the training data, group transformations permute the samples

$$\rho(g)x_i = x_{\pi_g(i)}, \quad \pi_g \in S_N$$

Permutation shift

- On the training data, group transformations permute the samples

$$\rho(g)x_i = x_{\pi_g(i)}, \quad \pi_g \in S_N$$

- Therefore, for a permutation of training samples associate to g

$$\begin{aligned}\Pi(g)\Theta(X,X) &= \Theta(\rho(g)X,X) \\ &= \Theta(X,\rho^{-1}(g)X) \\ &= \Theta(X,X)(\Pi^{-1}(g))^\top \\ &= \Theta(X,X)\Pi(g)\end{aligned}$$

Data augmentation at infinite width

group transformation

$$\mu_t(\rho(g)x) = \Theta(\rho(g)x, X) \Theta(X, X)^{-1} (\mathbb{I} - e^{-\eta \Theta(X, X)t}) Y$$

augmented data

augmented labels

The diagram illustrates the equation for data augmentation at infinite width. A blue arrow points from the text 'group transformation' to the term $\rho(g)$ in the expression $\rho(g)x$. Below the equation, the text 'augmented data' has four blue arrows pointing to the terms $\rho(g)x$, X , X , and X in the expression $\Theta(\rho(g)x, X) \Theta(X, X)^{-1}$. The text 'augmented labels' has two blue arrows pointing to the terms X and X in the expression $\Theta(X, X)^{-1}$. A final blue arrow points from 'augmented labels' to the term Y in the expression $(\mathbb{I} - e^{-\eta \Theta(X, X)t}) Y$.

Data augmentation at infinite width

group transformation for augmented data

$$\mu_t(\rho(g)x) = \Theta(\rho(g)x, X) \Theta(X, X)^{-1} (\mathbb{I} - e^{-\eta \Theta(X, X)t}) Y$$

augmented data augmented labels

Data augmentation at infinite width

group transformation

$$\mu_t(\rho(g)x) = \Theta(x, X)\Theta(X, X)^{-1}(\mathbb{I} - e^{-\eta\Theta(X, X)t})\rho(g)Y$$

augmented data

augmented labels

The diagram illustrates the equation for data augmentation at infinite width. A blue arrow points from the text 'group transformation' to the $\rho(g)$ term in the equation. Another blue arrow points from the text 'augmented data' to the x term. A third blue arrow points from the text 'augmented labels' to the Y term. Additionally, there are three blue arrows pointing from the 'augmented data' label to the $\Theta(x, X)$ and $\Theta(X, X)^{-1}$ terms, and two blue arrows pointing from the 'augmented labels' label to the $\Theta(X, X)$ and $\Theta(X, X)^{-1}$ terms.

Data augmentation at infinite width

group transformation

augmented labels

$$\mu_t(\rho(g)x) = \Theta(x, X)\Theta(X, X)^{-1}(\mathbb{I} - e^{-\eta\Theta(X, X)t})\underbrace{\rho(g)Y}_{=Y}$$

for invariance

Data augmentation at infinite width

group transformation

$$\mu_t(\rho(g)x) = \Theta(x, X)\Theta(X, X)^{-1}(\mathbb{I} - e^{-\eta\Theta(X, X)t})\underbrace{\rho(g)Y}_{=Y}$$

for invariance

Mean prediction

$$\mu_t(x)$$

Mean prediction

$$\mu_t(x) = \mathbb{E}_{\theta_0 \sim \text{initializations}}[\mathcal{N}_{\theta_t}(x)]$$

Mean prediction

$$\mu_t(x) = \mathbb{E}_{\theta_0 \sim \text{initializations}}[\mathcal{N}_{\theta_t}(x)] = \lim_{n \rightarrow \infty} \frac{1}{n} \sum_{\theta_0 = \text{init}_1}^{\text{init}_n} \mathcal{N}_{\theta_t}(x)$$

Mean prediction

$$\mu_t(x) = \mathbb{E}_{\theta_0 \sim \text{initializations}}[\mathcal{N}_{\theta_t}(x)] = \lim_{n \rightarrow \infty} \underbrace{\frac{1}{n} \sum_{\theta_0 = \text{init}_1}^{\text{init}_n} \mathcal{N}_{\theta_t}(x)}_{\text{mean prediction of deep ensemble}}$$

Main conclusion

Deep ensembles trained with data augmentation are equivariant.

Main conclusion

Deep ensembles trained with data augmentation are equivariant.

- ✓ Proof of exact equivariance for
 - full data augmentation
 - infinite ensembles

Main conclusion

Deep ensembles trained with data augmentation are equivariant.

- ✓ Proof of exact equivariance for
 - full data augmentation
 - infinite ensembles
- ✓ Equivariance holds for all training times

Main conclusion

Deep ensembles trained with data augmentation are equivariant.

- ✓ Proof of exact equivariance for
 - full data augmentation
 - infinite ensembles
- ✓ Equivariance holds for all training times
- ✓ Equivariance holds away from the training data

Main conclusion

Deep ensembles trained with data augmentation are equivariant.

- ✓ Proof of exact equivariance for
 - full data augmentation
 - infinite ensembles
- ✓ Equivariance holds for all training times
- ✓ Equivariance holds away from the training data
- ✓ Holds also for finite-width networks

[Nordenfors, Flinth 2024]

Intuitive explanation

- ✓ Equivariance holds for all training times
- ✓ Equivariance holds away from the training data

Intuitive explanation

- ✓ Equivariance holds for all training times
- ✓ Equivariance holds away from the training data

ⓘ At infinite width, the mean output at initialization is zero everywhere.

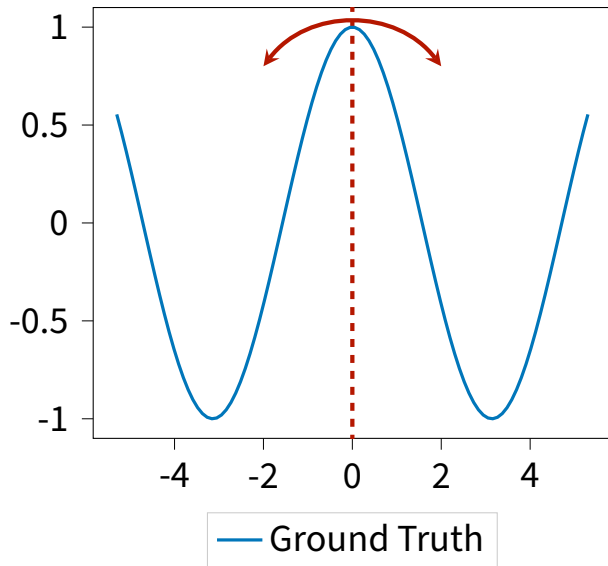
Intuitive explanation

- ✓ Equivariance holds for all training times
- ✓ Equivariance holds away from the training data

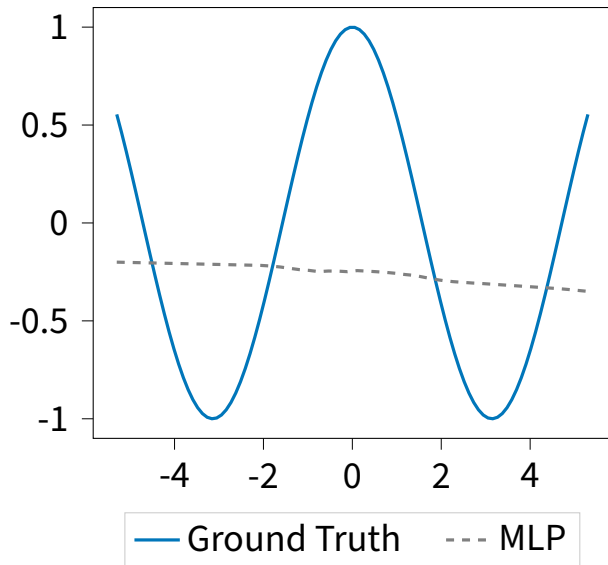
ⓘ At infinite width, the mean output at initialization is zero everywhere.

⇒ Training with full data augmentation leads to an equivariant function.

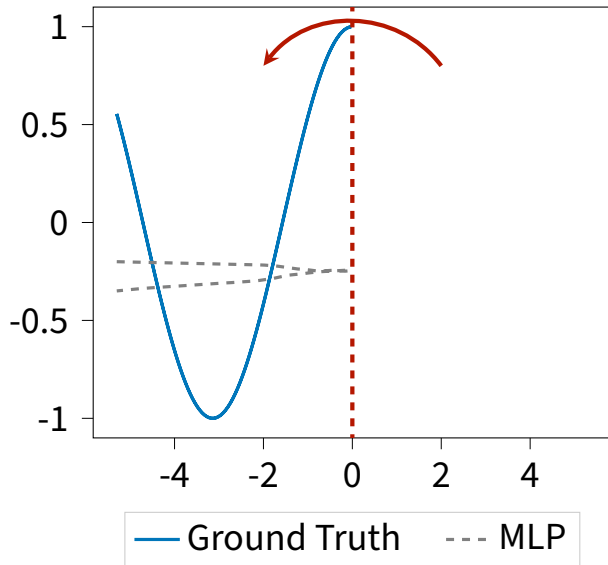
Toy example



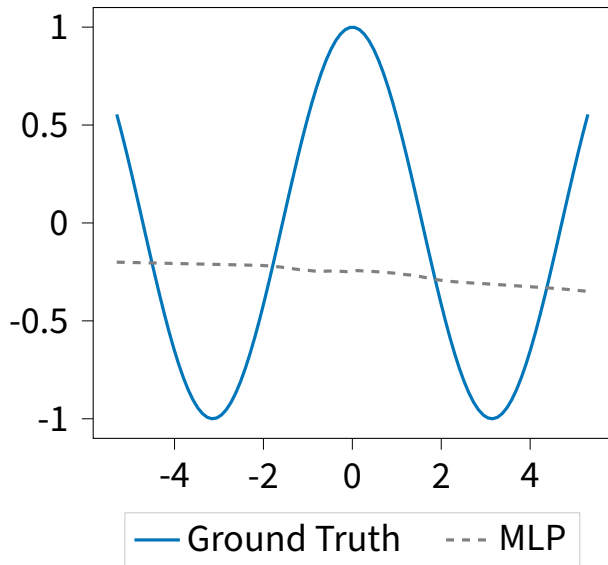
Initialization



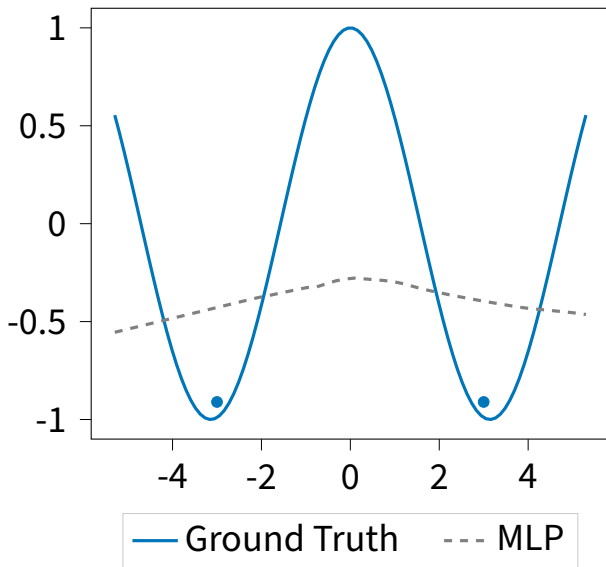
Initialization



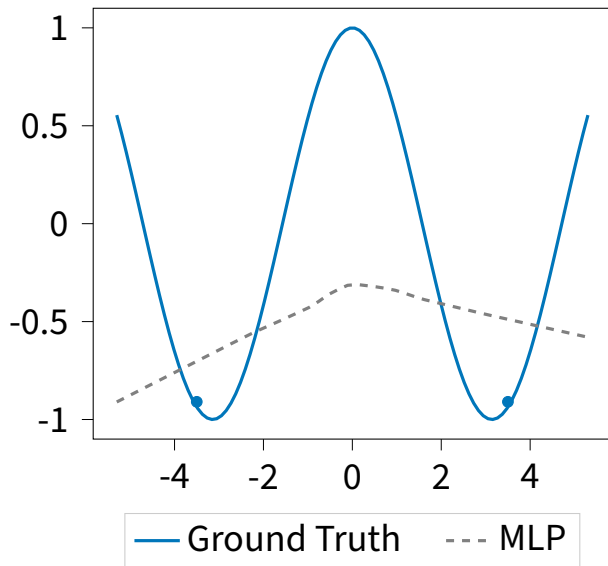
Initialization



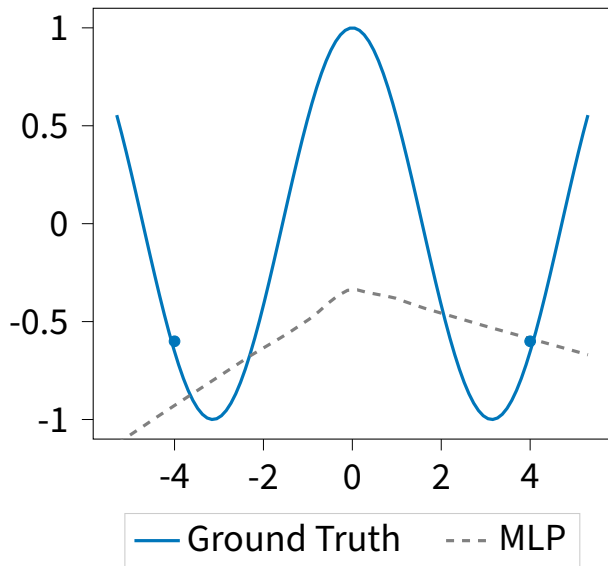
After 1 Training Step



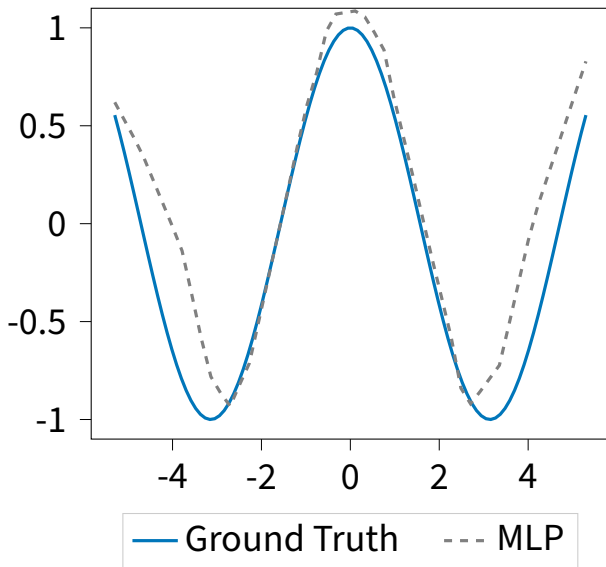
After 2 Training Steps



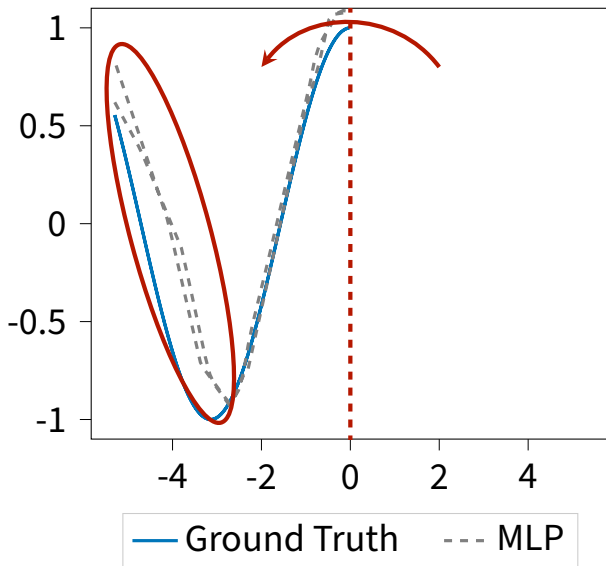
After 3 Training Steps



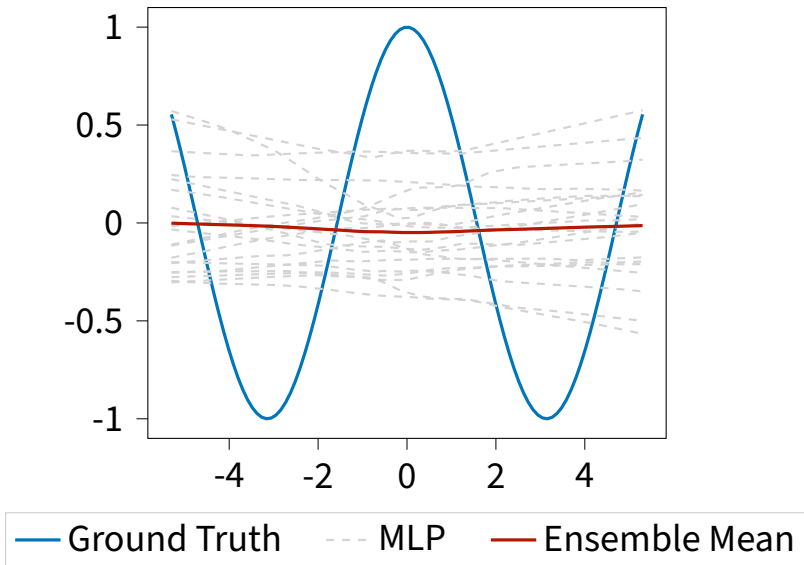
After 2000 Training Steps



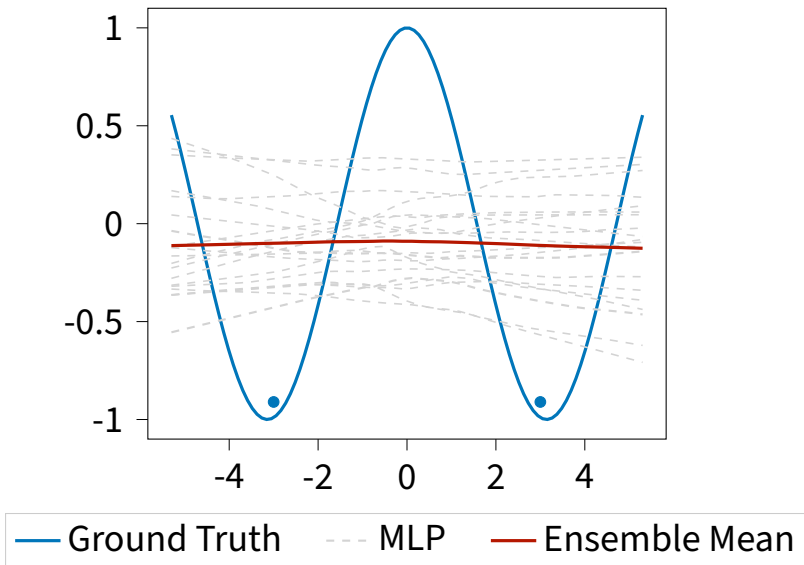
After 2000 Training Steps



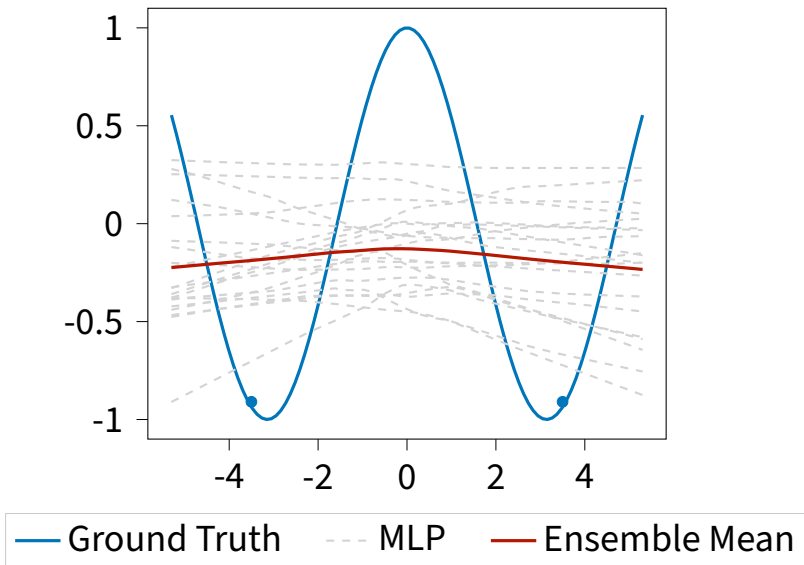
Initialization



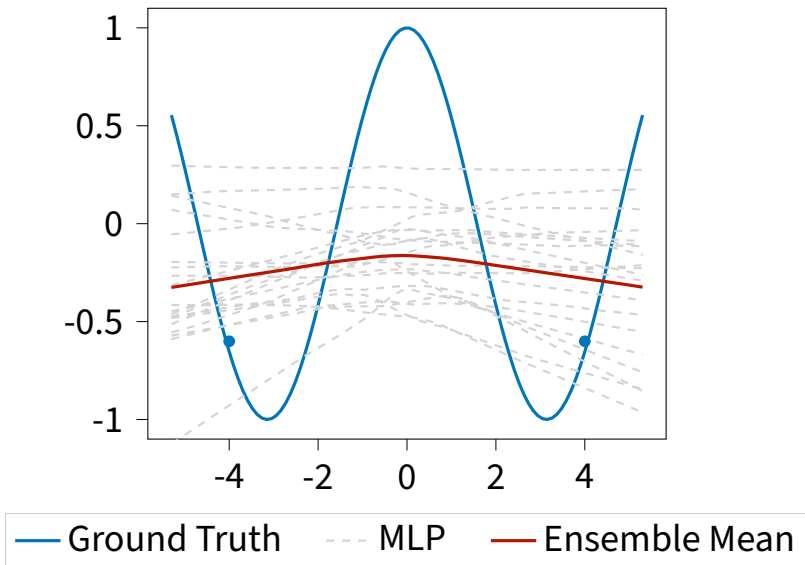
After 1 Training Step



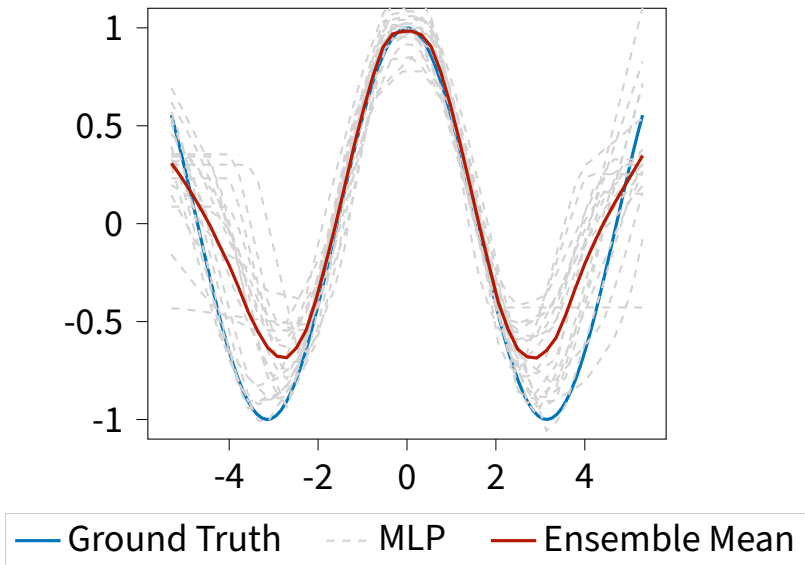
After 2 Training Steps



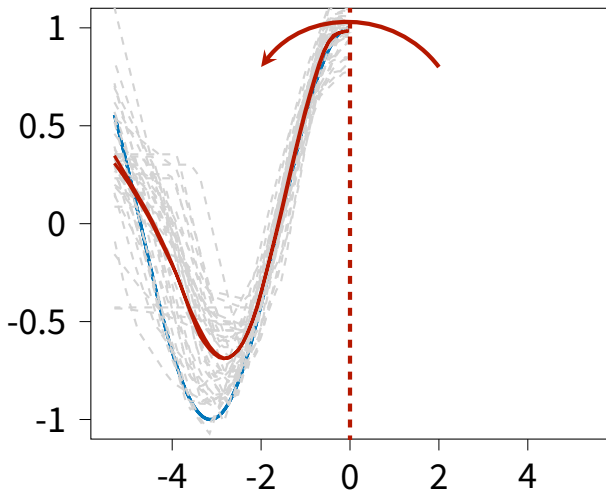
After 3 Training Steps



After 2000 Training Steps



After 2000 Training Steps

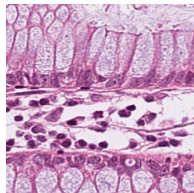
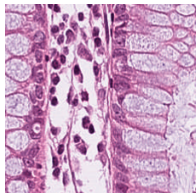


— Ground Truth - - - MLP — Ensemble Mean

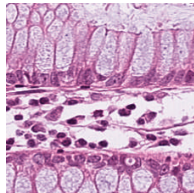
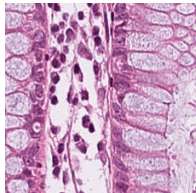
What Does An Augmented Ensemble Converge To?

Rotating images

Rotating images



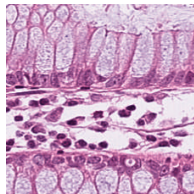
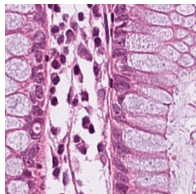
Rotating images



$f(x)$

$f : \text{pixels} \rightarrow \text{colors}$

Rotating images



$f(x)$
 $f : \text{pixels} \rightarrow \text{colors}$



$f(\rho(g^{-1})x)$
 $= [\rho_{\text{reg}}(g)f](x)$

Data augmentation and NTKs

Data augmentation and NTKs

Consider two ensembles:

trained without data augmentation

trained with data augmentation

Data augmentation and NTKs

Consider two ensembles:

trained without data augmentation

trained with data augmentation

If


$$\Theta^{\text{non-aug}}(f, f') = \frac{1}{|G|} \sum_{g \in G} \Theta^{\text{aug}}(f, \rho_{\text{reg}}(g)f')$$

Data augmentation and NTKs

Consider two ensembles:

trained without data augmentation

trained with data augmentation

If


$$\Theta^{\text{non-aug}}(f, f') = \frac{1}{|G|} \sum_{g \in G} \Theta^{\text{aug}}(f, \rho_{\text{reg}}(g)f')$$

Then

$$\mu_t^{\text{non-aug}}(x) = \mu_t^{\text{aug}}(x)$$

at infinite width.

Data augmentation and NTKs

Consider two ensembles:

trained without data augmentation

trained with data augmentation

If


$$\Theta^{\text{non-aug}}(f, f') = \frac{1}{|G|} \sum_{g \in G} \Theta^{\text{aug}}(f, \rho_{\text{reg}}(g)f')$$

Then

$$\mu_t^{\text{non-aug}}(x) = \mu_t^{\text{aug}}(x) \quad \forall t$$

at infinite width.

Data augmentation and NTKs

Consider two ensembles:

trained without data augmentation

trained with data augmentation

If


$$\Theta^{\text{non-aug}}(f, f') = \frac{1}{|G|} \sum_{g \in G} \Theta^{\text{aug}}(f, \rho_{\text{reg}}(g)f')$$

Then

$$\mu_t^{\text{non-aug}}(x) = \mu_t^{\text{aug}}(x) \quad \forall t \quad \forall x$$

at infinite width.

Data augmentation and NTKs

$$\Theta^{\text{non-aug}}(f, f') = \frac{1}{|G|} \sum_{g \in G} \Theta^{\text{aug}}(f, \rho_{\text{reg}}(g)f')$$

Data augmentation and NTKs

$$\Theta^{\text{non-aug}}(f, f') = \frac{1}{|G|} \sum_{g \in G} \Theta^{\text{aug}}(f, \rho_{\text{reg}}(g)f')$$

- ① Given an architecture with NTK Θ^{aug} ,
find an architecture with NTK $\Theta^{\text{non-aug}}$

Group convolutions

[Cohen, Welling 2016]

Group convolutions

[Cohen, Welling 2016]

Group conv's are the (unique) linear layers equivariant wrt ρ_{reg}

Group convolutions

[Cohen, Welling 2016]

Group conv's are the (unique) linear layers equivariant wrt ρ_{reg}

- Ordinary convolutions

$$f'(y) = \int_X dx \, \kappa(x - y) f(x)$$

Group convolutions

[Cohen, Welling 2016]

Group conv's are the (unique) linear layers equivariant wrt ρ_{reg}

- Ordinary convolutions

$$f'(y) = \int_X dx \, \kappa(x - y) f(x)$$

- Group convolutions

$$f'(g) = \int_X dx \, \kappa(\rho(g^{-1})x) f(x)$$

lifting

Group convolutions

[Cohen, Welling 2016]

Group conv's are the (unique) linear layers equivariant wrt ρ_{reg}

- Ordinary convolutions

$$f'(y) = \int_X dx \, \kappa(x - y) f(x)$$

- Group convolutions

$$f'(g) = \int_X dx \, \kappa(\rho(g^{-1})x) f(x) \quad \text{lifting}$$

$$f'(g) = \int_G dg \, \kappa(g^{-1}h) f(h) \quad \text{group convolution}$$

Group convolutions

[Cohen, Welling 2016]

Group conv's are the (unique) linear layers equivariant wrt ρ_{reg}

- Ordinary convolutions

$$f'(y) = \int_X dx \, \kappa(x - y) f(x)$$

- Group convolutions

$$f'(g) = \int_X dx \, \kappa(\rho(g^{-1})x) f(x) \quad \text{lifting}$$

$$f'(g) = \int_G dg \, \kappa(g^{-1}h) f(h) \quad \text{group convolution}$$

$$f' = \frac{1}{\text{vol}(G)} \int_G dg \, f(g) \quad \text{group pooling}$$

GCNNs

Stack GConv-layers to obtain an invariant network

GCNNs

Stack GConv-layers to obtain an invariant network



GCNNs

Stack GConv-layers to obtain an invariant network



→ lifting

GCNNs

Stack GConv-layers to obtain an invariant network



→ lifting → σ

GCNNs

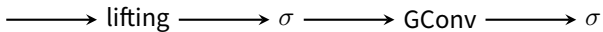
Stack GConv-layers to obtain an invariant network



→ lifting → σ → GConv

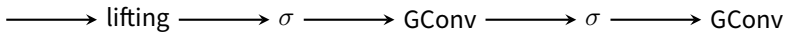
GCNNs

Stack GConv-layers to obtain an invariant network



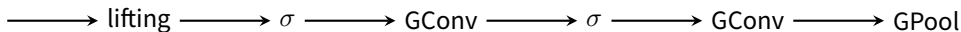
GCNNs

Stack GConv-layers to obtain an invariant network



GCNNs

Stack GConv-layers to obtain an invariant network



NTKs for GCNNs

For GCNN-layers, define the NNGP and NTK via

$$K_{g,g'}^{(\ell)}(f, f') = \mathbb{E} \left[[\mathcal{N}^{(\ell)}(f)](g) \left([\mathcal{N}^{(\ell)}(f')](g') \right)^\top \right]$$

NTKs for GCNNs

For GCNN-layers, define the NNGP and NTK via

$$K_{g,g'}^{(\ell)}(f, f') = \mathbb{E} \left[[\mathcal{N}^{(\ell)}(f)](g) \left([\mathcal{N}^{(\ell)}(f')](g') \right)^\top \right]$$
$$\Theta_{g,g'}^{(\ell)}(f, f') = \mathbb{E} \left[\sum_{\ell'=1}^{\ell} \frac{\partial [\mathcal{N}^{(\ell)}(f)](g)}{\partial \theta^{(\ell')}} \left(\frac{\partial [\mathcal{N}^{(\ell)}(f')](g')}{\partial \theta^{(\ell')}} \right)^\top \right]$$

NTKs for GCNNs

$$[\mathcal{N}^{(\ell)}(f)](g) = \int_G dg \kappa(g^{-1}h) [\mathcal{N}^{(\ell-1)}(f)](h)$$

The layer-recursion for a GCNN-layer is given by

$$K_{g,g'}^{(\ell+1)}(f,f') = \frac{1}{|S_K|} \int_{S_K} dh K_{gh,g'h}^{(\ell)}(f,f')$$

NTKs for GCNNs

$$[\mathcal{N}^{(\ell)}(f)](g) = \int_G dg \kappa(g^{-1}h) [\mathcal{N}^{(\ell-1)}(f)](h)$$

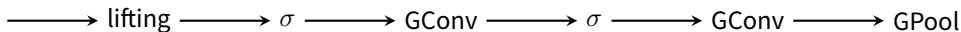
The layer-recursion for a GCNN-layer is given by

$$K_{g,g'}^{(\ell+1)}(f, f') = \frac{1}{|S_K|} \int_{S_K} dh K_{gh,g'h}^{(\ell)}(f, f')$$

$$\Theta_{g,g'}^{(\ell+1)}(f, f') = K_{g,g'}^{(\ell+1)}(f, f') + \frac{1}{|S_K|} \int_{S_K} dh \Theta_{gh,g'h}^{(\ell)}(f, f')$$

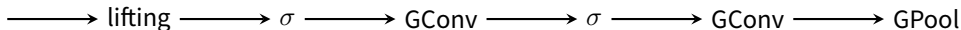
GCNNs

Stack GConv-layers to obtain an invariant network



NTKs for GCNNs

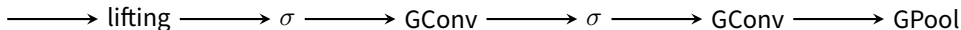
Stack GConv-layers to obtain an invariant network



Compute NTK with layer-wise recursion

NTKs for GCNNs

Stack GConv-layers to obtain an invariant network

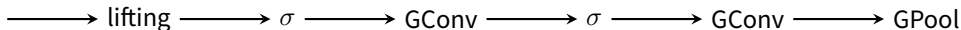


Compute NTK with layer-wise recursion

0

NTKs for GCNNs

Stack GConv-layers to obtain an invariant network

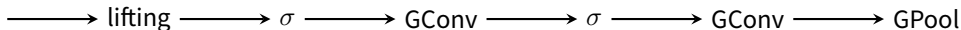


Compute NTK with layer-wise recursion

$$0 \longrightarrow \Theta_{g,g'}^{(1)}(f, f')$$

NTKs for GCNNs

Stack GConv-layers to obtain an invariant network

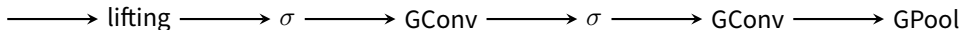


Compute NTK with layer-wise recursion

$$0 \longrightarrow \Theta_{g,g'}^{(1)}(f,f') \longrightarrow \Theta_{g,g'}^{(2)}(f,f')$$

NTKs for GCNNs

Stack GConv-layers to obtain an invariant network

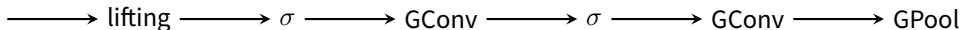


Compute NTK with layer-wise recursion

$$0 \longrightarrow \Theta_{g,g'}^{(1)}(f,f') \longrightarrow \Theta_{g,g'}^{(2)}(f,f') \longrightarrow \Theta_{g,g'}^{(3)}(f,f')$$

NTKs for GCNNs

Stack GConv-layers to obtain an invariant network

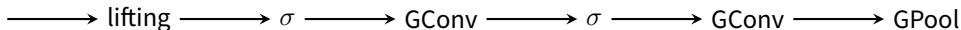


Compute NTK with layer-wise recursion

$$0 \longrightarrow \Theta_{g,g'}^{(1)}(f,f') \longrightarrow \Theta_{g,g'}^{(2)}(f,f') \longrightarrow \Theta_{g,g'}^{(3)}(f,f') \longrightarrow \Theta_{g,g'}^{(4)}(f,f')$$

NTKs for GCNNs

Stack GConv-layers to obtain an invariant network

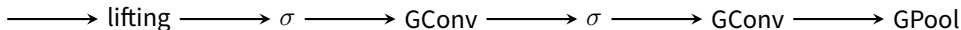


Compute NTK with layer-wise recursion

$$0 \longrightarrow \Theta_{g,g'}^{(1)}(f,f') \longrightarrow \Theta_{g,g'}^{(2)}(f,f') \longrightarrow \Theta_{g,g'}^{(3)}(f,f') \longrightarrow \Theta_{g,g'}^{(4)}(f,f') \longrightarrow \Theta_{g,g'}^{(5)}(f,f')$$

NTKs for GCNNs

Stack GConv-layers to obtain an invariant network



Compute NTK with layer-wise recursion

$$0 \longrightarrow \Theta_{g,g'}^{(1)}(f,f') \longrightarrow \Theta_{g,g'}^{(2)}(f,f') \longrightarrow \Theta_{g,g'}^{(3)}(f,f') \longrightarrow \Theta_{g,g'}^{(4)}(f,f') \longrightarrow \Theta_{g,g'}^{(5)}(f,f') \longrightarrow \Theta(f,f')$$

NTKs of MLPs and GCNNs

NTKs of MLPs and GCNNs

- Consider two neural networks

NTKs of MLPs and GCNNs

- Consider two neural networks

An MLP



NTKs of MLPs and GCNNs

- Consider two neural networks

An MLP



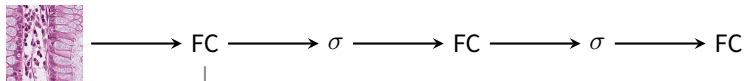
A GCNN



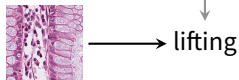
NTKs of MLPs and GCNNs

- Consider two neural networks

An MLP



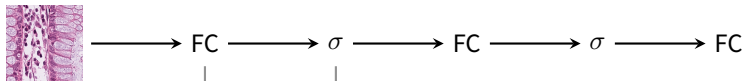
A GCNN



NTKs of MLPs and GCNNs

- Consider two neural networks

An MLP



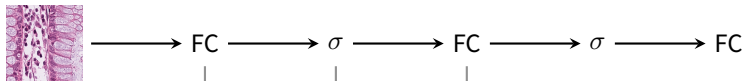
A GCNN



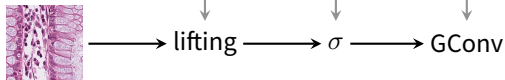
NTKs of MLPs and GCNNs

- Consider two neural networks

An MLP



A GCNN



NTKs of MLPs and GCNNs

- Consider two neural networks

An MLP

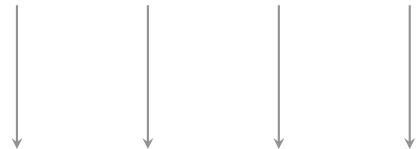


→ FC → σ → FC → σ → FC

A GCNN



→ lifting → σ → GConv → σ



NTKs of MLPs and GCNNs

- Consider two neural networks

An MLP



→ FC → σ → FC → σ → FC

A GCNN



→ lifting → σ → GConv → σ → GConv



NTKs of MLPs and GCNNs

- Consider two neural networks

An MLP



→ FC → σ → FC → σ → FC

A GCNN



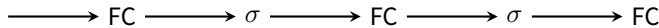
→ lifting → σ → GConv → σ → GConv → GPool



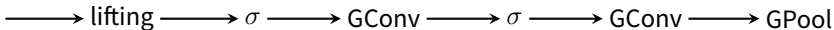
NTKs of MLPs and GCNNs

- Consider two neural networks

An MLP



A GCNN



- Then

$$\Theta^{\text{GCNN}}(f, f') = \frac{1}{|G|} \sum_{g \in G} \Theta^{\text{MLP}}(f, \rho_{\text{reg}}(g)f')$$

Data augmentation of MLPs

$$\Theta^{\text{GCNN}}(f, f') = \frac{1}{|G|} \sum_{g \in G} \Theta^{\text{MLP}}(f, \rho_{\text{reg}}(g)f')$$

Data augmentation of MLPs

before: non-aug


$$\Theta^{\text{GCNN}}(f, f') = \frac{1}{|G|} \sum_{g \in G} \Theta^{\text{MLP}}(f, \rho_{\text{reg}}(g)f')$$

Data augmentation of MLPs

before: non-aug

$$\Theta^{\text{GCNN}}(f, f') = \frac{1}{|G|} \sum_{g \in G} \Theta^{\text{MLP}}(f, \rho_{\text{reg}}(g)f')$$

before: aug

Data augmentation of MLPs

before: non-aug

$$\Theta^{\text{GCNN}}(f, f') = \frac{1}{|G|} \sum_{g \in G} \Theta^{\text{MLP}}(f, \rho_{\text{reg}}(g)f')$$

before: aug

⇒ training the MLP on
G-augmented data

Data augmentation of MLPs

before: non-aug

$$\Theta^{\text{GCNN}}(f, f') = \frac{1}{|G|} \sum_{g \in G} \Theta^{\text{MLP}}(f, \rho_{\text{reg}}(g)f')$$

before: aug

⇒ training the MLP on G-augmented data = training the GCNN on unaugmented data

Data augmentation of MLPs

before: non-aug

before: aug

$$\Theta^{\text{GCNN}}(f, f') = \frac{1}{|G|} \sum_{g \in G} \Theta^{\text{MLP}}(f, \rho_{\text{reg}}(g)f')$$



training the MLP on
G-augmented data

=

training the GCNN on
unaugmented data

in the ensemble mean

Data augmentation of MLPs

before: non-aug

before: aug

$$\Theta^{\text{GCNN}}(f, f') = \frac{1}{|G|} \sum_{g \in G} \Theta^{\text{MLP}}(f, \rho_{\text{reg}}(g)f')$$



training the MLP on
G-augmented data

=

training the GCNN on
unaugmented data



in the ensemble mean, $\forall t, \forall x$

Data augmentation of CNNs

Data augmentation of CNNs

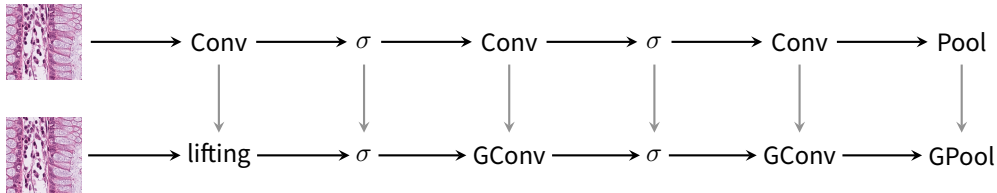
- Consider a CNN



→ Conv → σ → Conv → σ → Conv → Pool

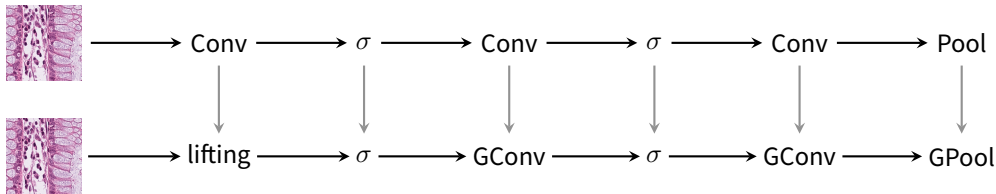
Data augmentation of CNNs

- Consider a CNN and a GCNN invariant wrt. roto-translations



Data augmentation of CNNs

- Consider a CNN and a GCNN invariant wrt. roto-translations

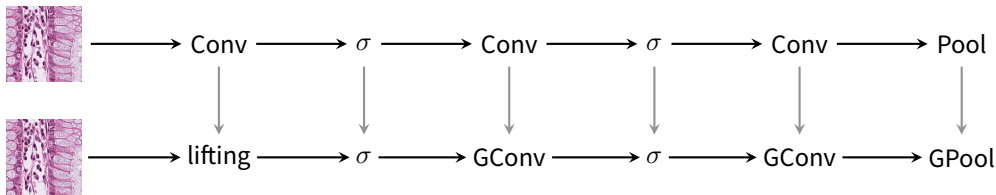


- Then

$$\Theta^{\text{GCNN}}(f, f') = \frac{1}{n} \sum_{r \in C_n} \Theta^{\text{CNN}}(f, \rho_{\text{reg}}(r)f')$$

Data augmentation of CNNs

- Consider a CNN and a GCNN invariant wrt. roto-translations



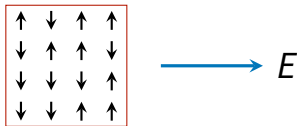
- Then

$$\Theta^{\text{GCNN}}(f, f') = \frac{1}{n} \sum_{r \in C_n} \Theta^{\text{CNN}}(f, \rho_{\text{reg}}(r)f')$$

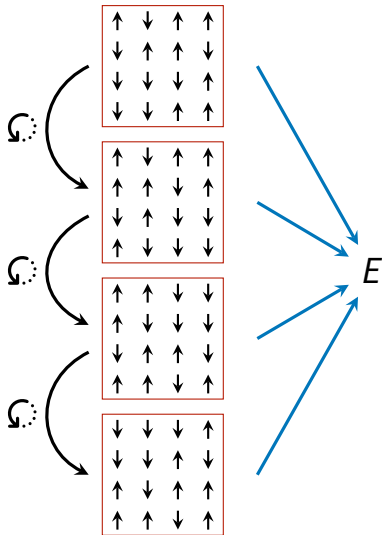
⇒ By training the CNN on rotated images, one obtains a roto-translation invariant GCNN

Experiments

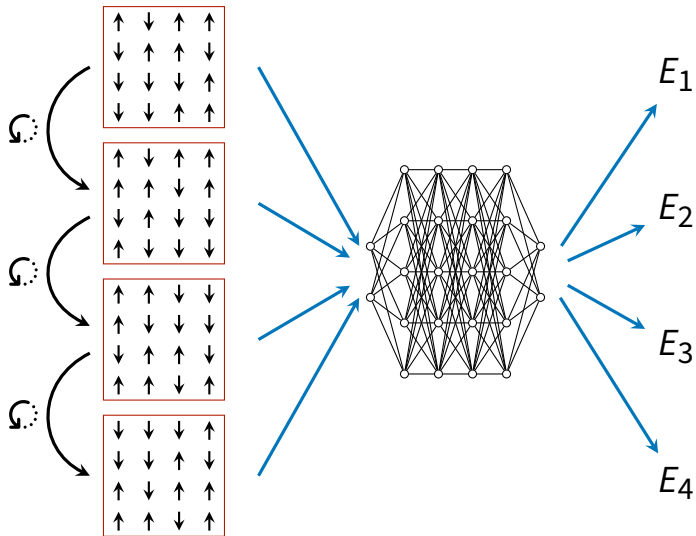
Ising model



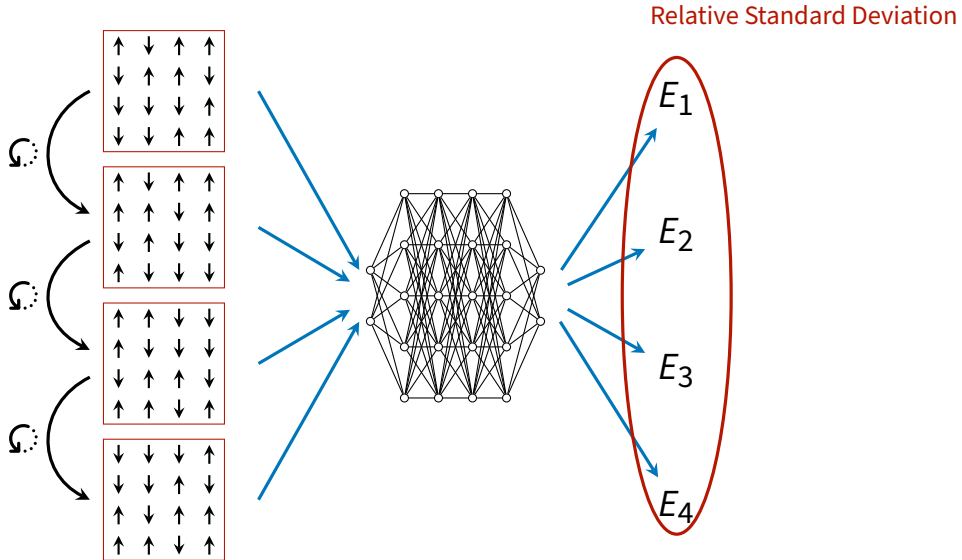
Ising model

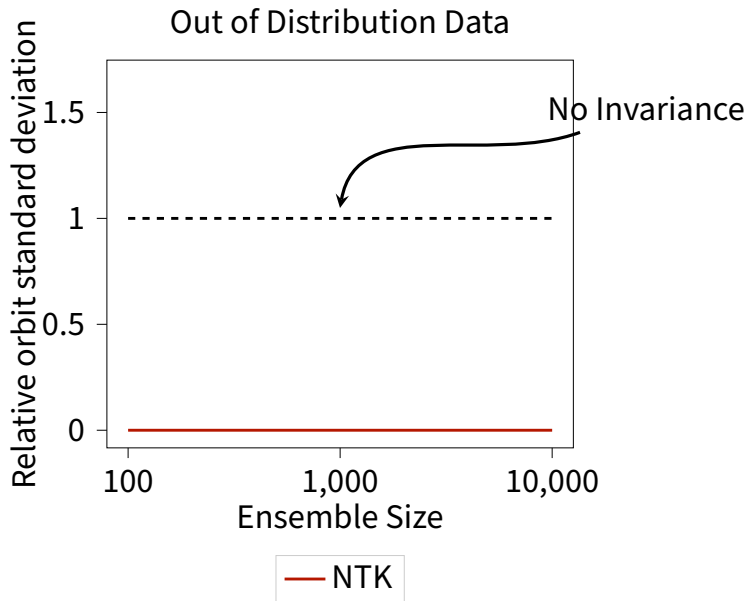


Ising model

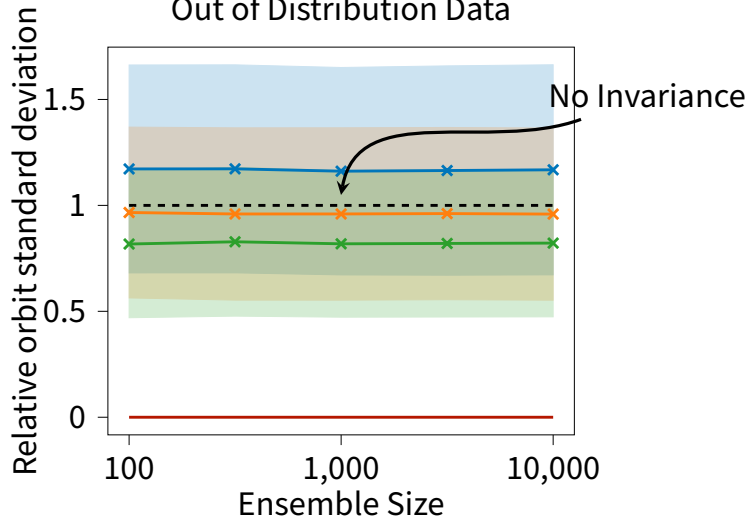


Ising model



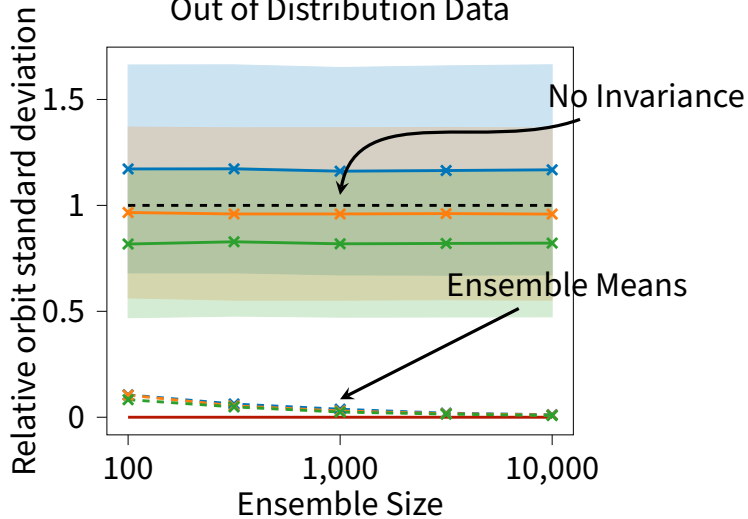


Out of Distribution Data

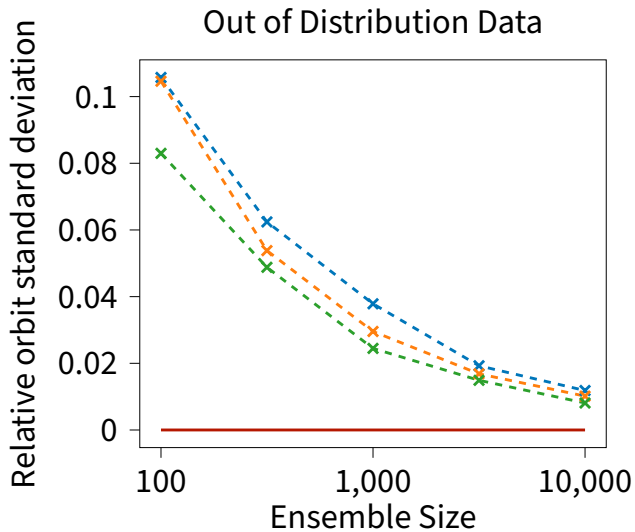


— NTK × Width 512 × Width 1024 × Width 2048

Out of Distribution Data



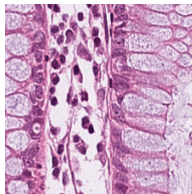
— NTK × Width 512 × Width 1024 × Width 2048



— NTK -x- Width 512 -x- Width 1024 -x- Width 2048

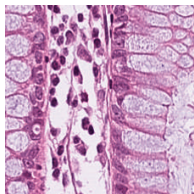
Histological slices

[Kather et al. 2018]



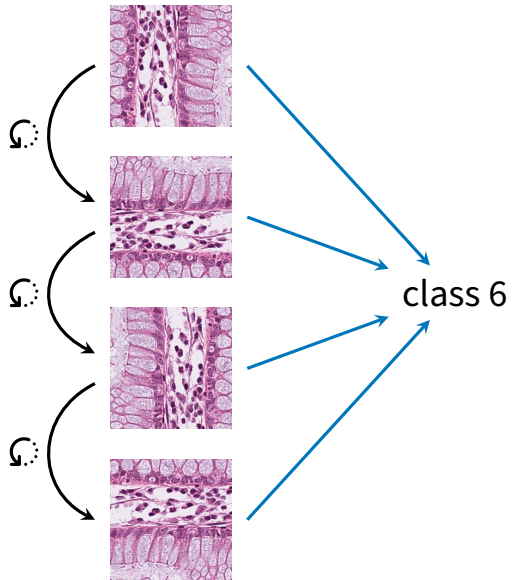
Histological slices

[Kather et al. 2018]

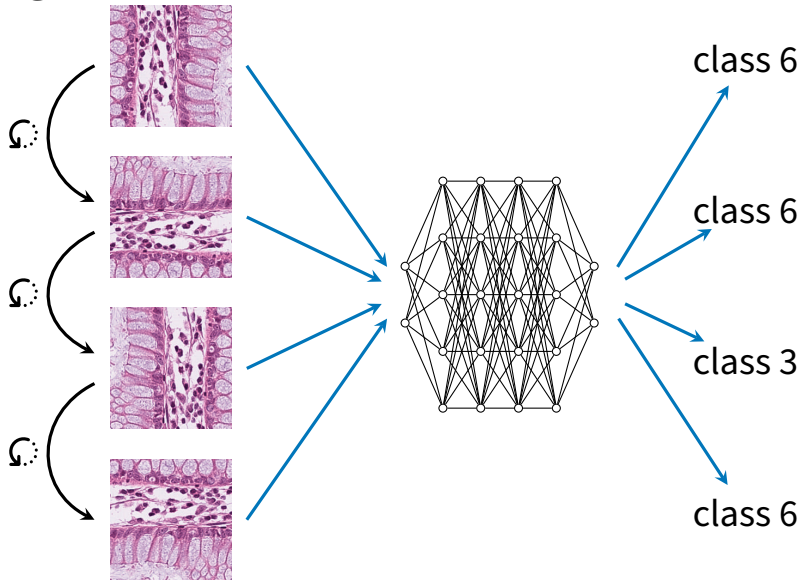


→ class 6

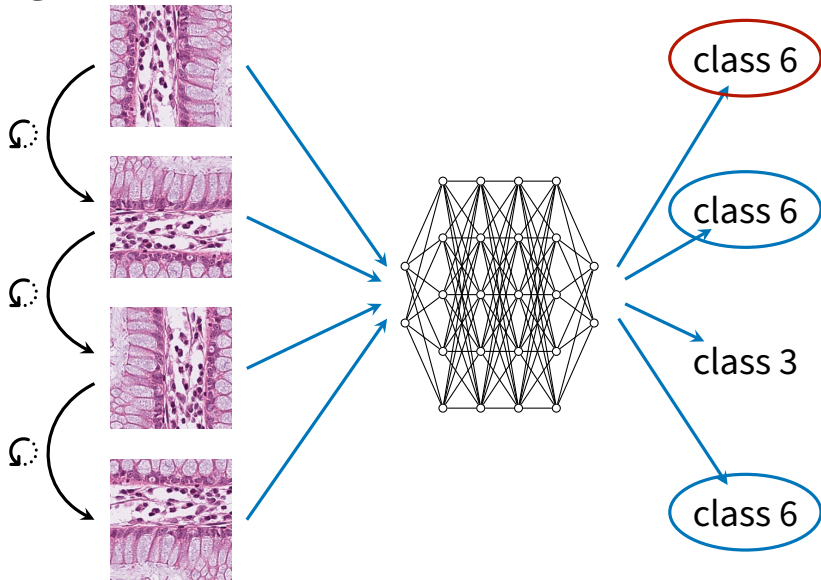
Histological slices



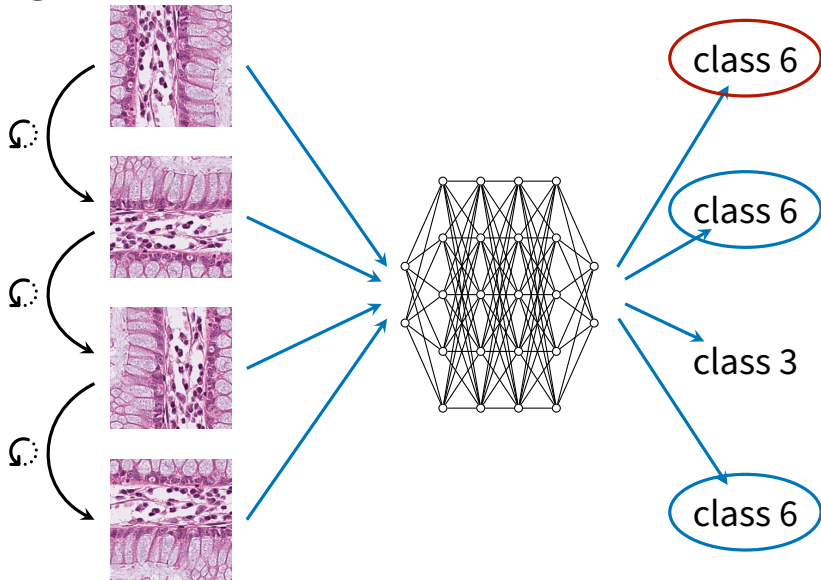
Histological slices



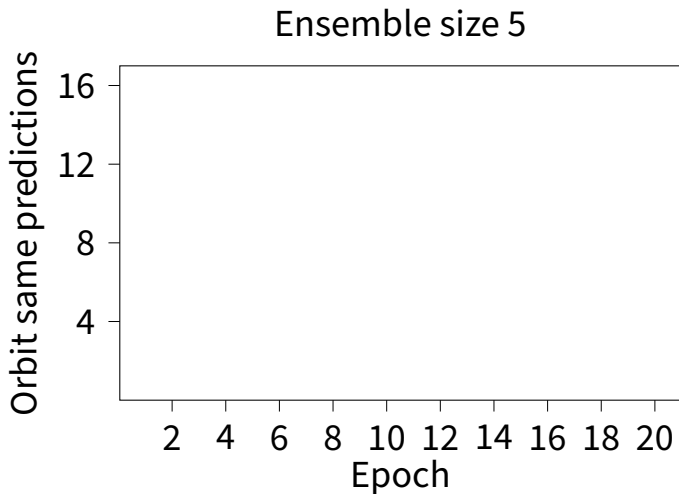
Histological slices



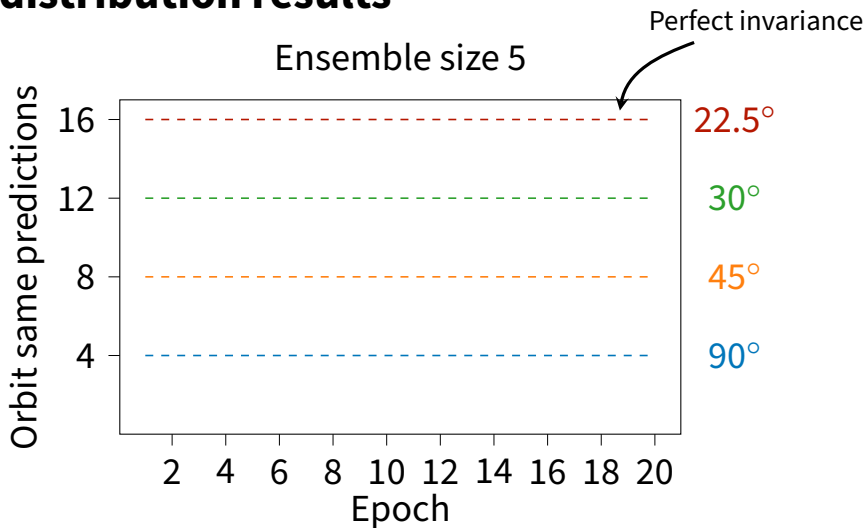
Histological slices



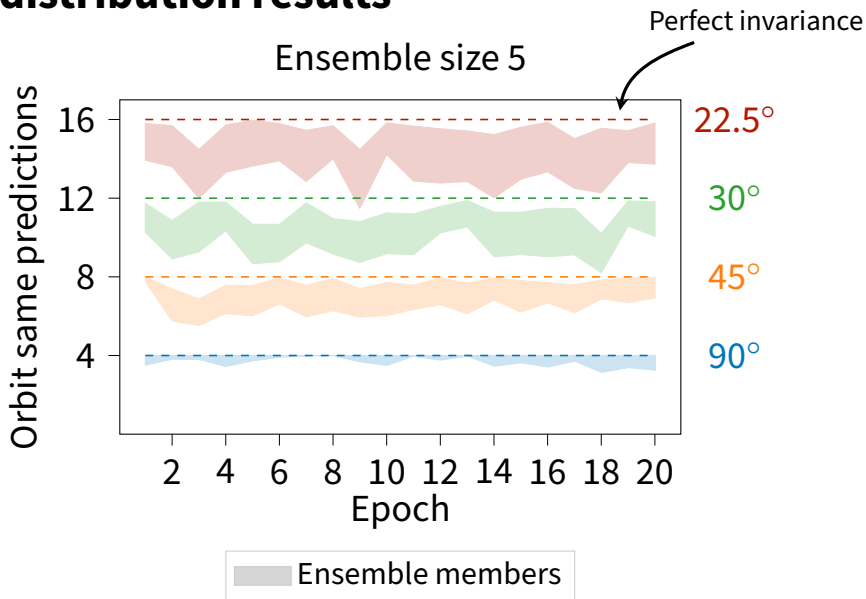
Out of distribution results



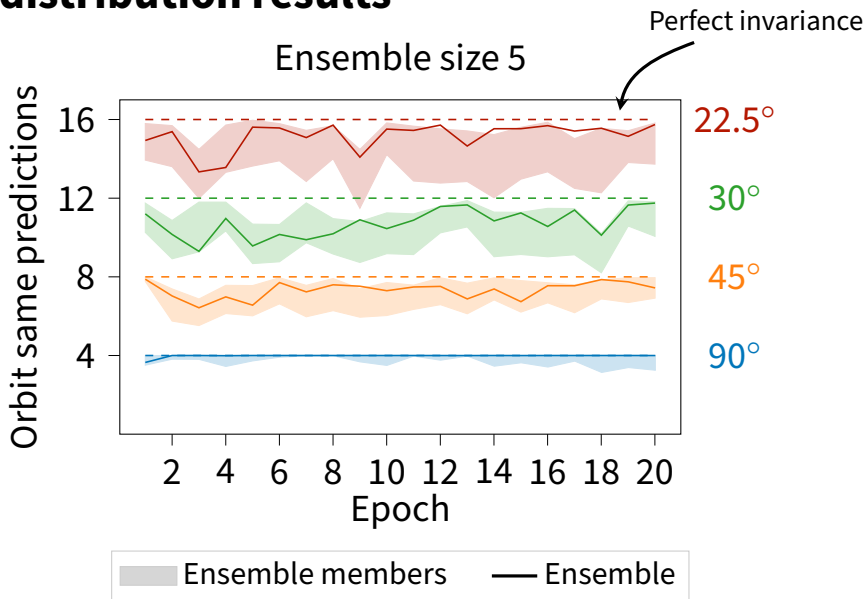
Out of distribution results



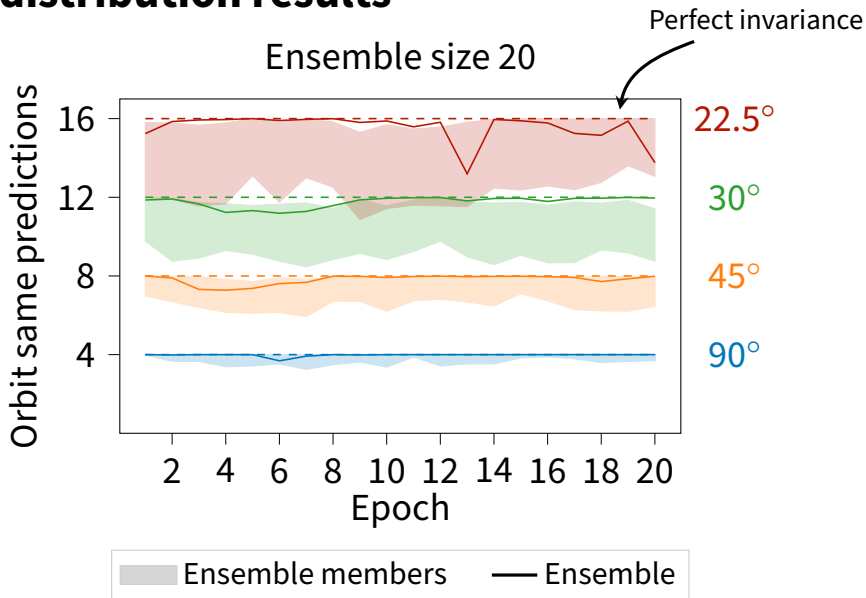
Out of distribution results



Out of distribution results



Out of distribution results



Further experimental results

Further experimental results

- ✓ Emergent invariance for rotated FashionMNIST

Further experimental results

- ✓ Emergent invariance for rotated FashionMNIST
- ✓ Partial augmentation for continuous symmetries

Further experimental results

- ✓ Emergent invariance for rotated FashionMNIST
- ✓ Partial augmentation for continuous symmetries
- ✓ Emergent equivariance (as opposed to invariance)

Comparison to other methods

Comparison to other methods

⇒ Models trained on rotated FashionMNIST

Comparison to other methods

⇒ Models trained on rotated FashionMNIST

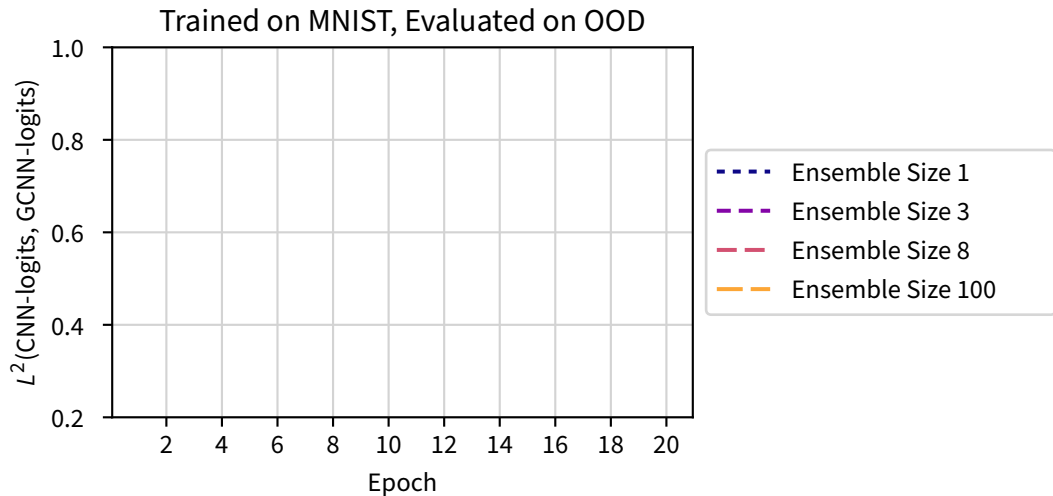
Orbit same predictions out of distribution:

	C_4	C_8	C_{16}
DeepEns+DA	3.85 ± 0.12	7.72 ± 0.34	15.24 ± 0.69
only DA	3.41 ± 0.18	6.73 ± 0.24	12.77 ± 0.71
E2CNN ¹	4 ± 0.0	7.71 ± 0.21	15.08 ± 0.34
Canon ²	4 ± 0.0	7.45 ± 0.14	12.41 ± 0.85

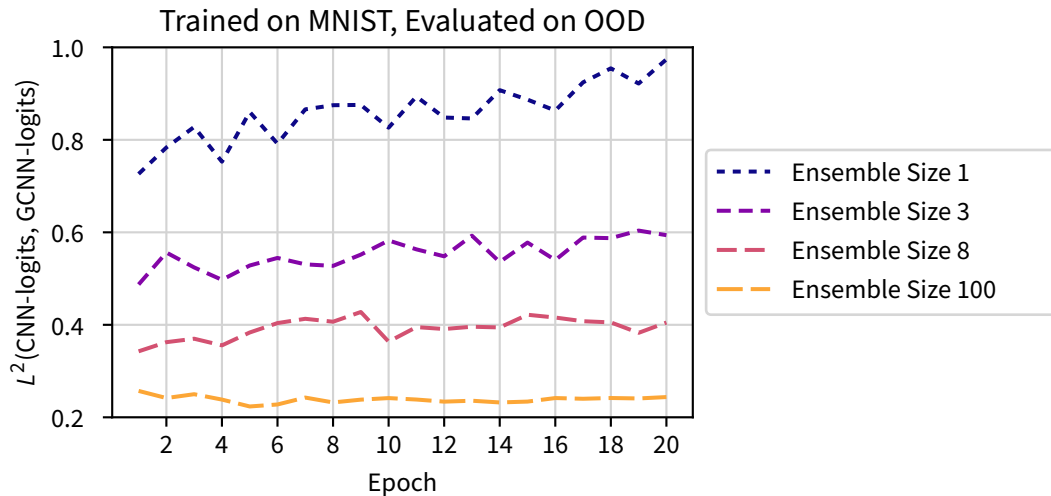
¹[Weiler et al. 2019], ²[Kaba et al. 2022]

Convergence of augmented CNNs to GCNNs

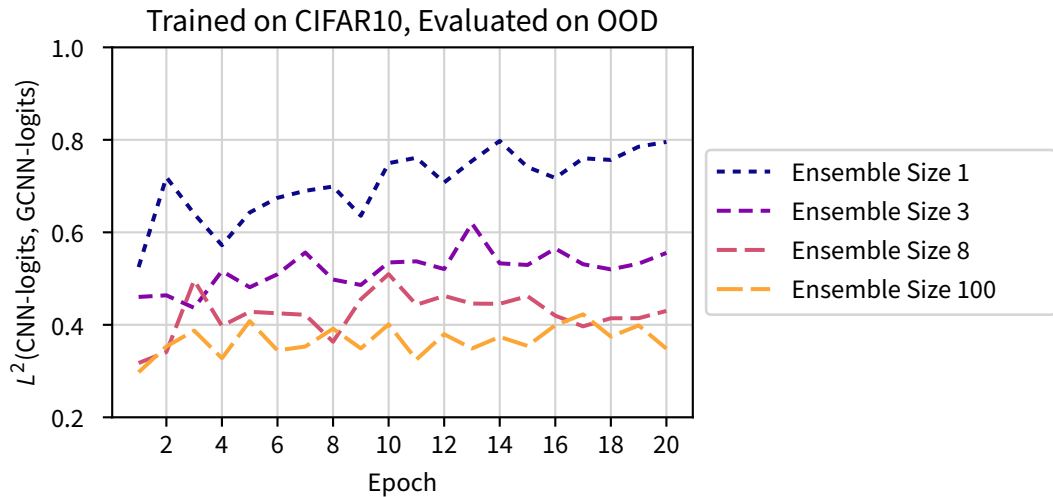
Convergence of augmented CNNs to GCNNs



Convergence of augmented CNNs to GCNNs



Convergence of augmented CNNs to GCNNs



Key takeaways

Key takeaways

If you need ensembles

- 👍 use data augmentation to obtain an equivariant model.

Key takeaways

If you need ensembles

👍 use data augmentation to obtain an equivariant model.

If you need data augmentation

👍 use an ensemble to boost the equivariance.

Papers

- **Emergent Equivariance in Deep Ensembles**

Jan E. Gerken*, Pan Kessel*

ICML 2024 (Oral)

* Equal contribution

- **Equivariant Neural Tangent Kernels**

Philipp Misof, Pan Kessel, Jan E. Gerken

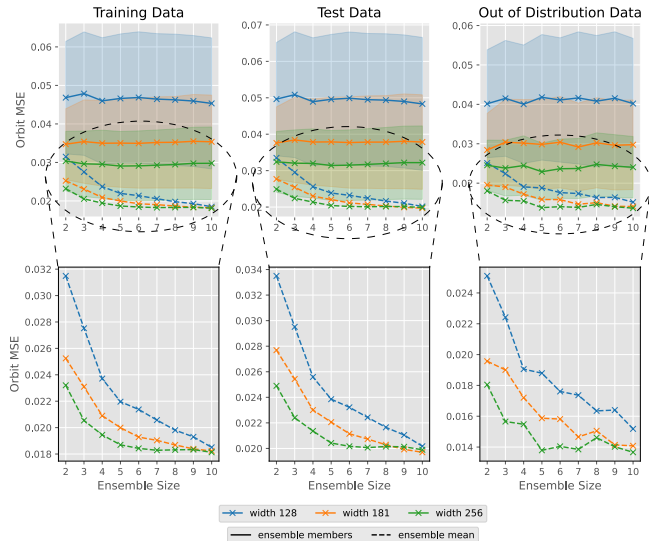
ICML 2025



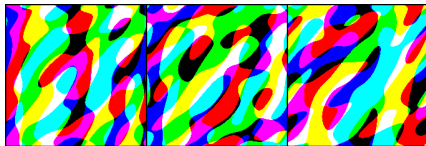
Thank you

Backup

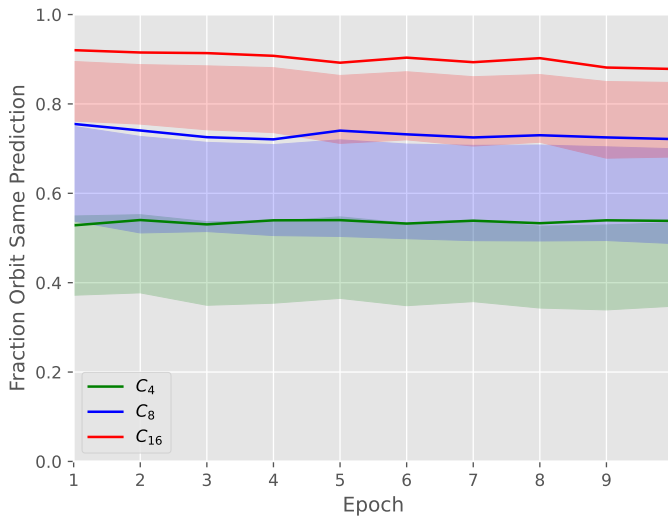
Emergent equivariance of cross products



Histological Data – OOD samples

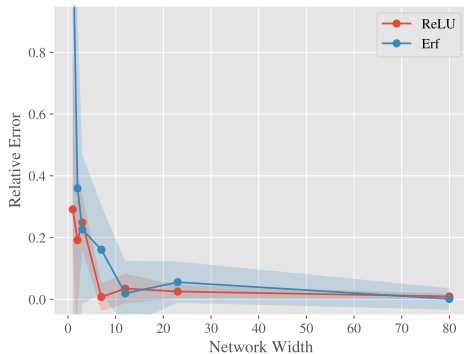


Emergent continuous symmetry on FashionMNIST

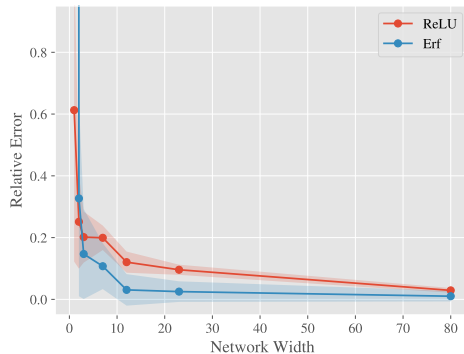


Kernel convergence

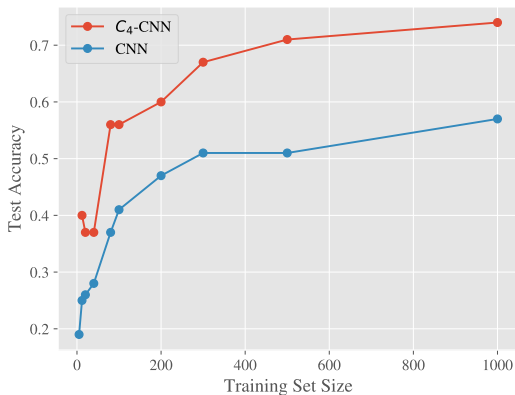
NNGP



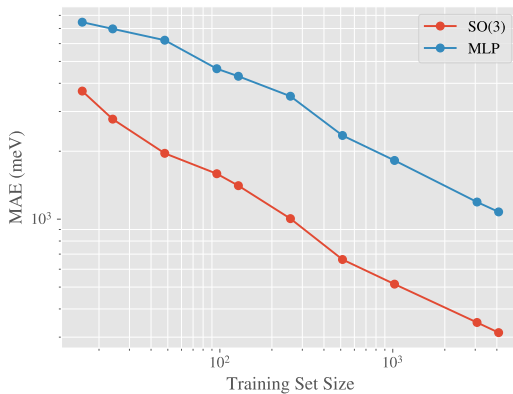
NTK



Equivariant NTKs for medical image classification

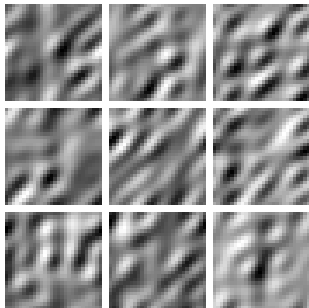


Equivariant NTKs for molecular property regression



OOD samples for CNN to GCNN convergence

MNIST



CIFAR10

