
Learning Inter-Atomic Potentials without Explicit Equivariance

Ahmed A. Elhag^{1*}, Arun Raja^{1*}, Alex Morehead^{2*}, Samuel M. Blau²,
Garrett M. Morris¹, Michael M. Bronstein^{1,3}

¹University of Oxford, ²Lawrence Berkeley National Laboratory, ³AITHYRA

Abstract

Accurate and scalable machine-learned inter-atomic potentials (MLIPs) are essential for molecular simulations ranging from drug discovery to new material design. Current state-of-the-art models enforce roto-translational symmetries through equivariant neural network architectures, a hard-wired inductive bias that can often lead to reduced flexibility, computational efficiency, and scalability. In this work, we introduce **TransIP: Transformer-based Inter-Atomic Potentials**, a novel training paradigm for interatomic potentials achieving symmetry compliance without explicit architectural constraints. Our approach guides a generic non-equivariant Transformer-based model to learn $SO(3)$ -equivariance by optimizing its representations in the embedding space. Trained on the recent Open Molecules (OMol25) collection, a large and diverse molecular dataset built specifically for MLIPs and covering different types of molecules (including small organics, biomolecular fragments, and electrolyte-like species), TransIP effectively learns symmetry in its latent space, providing low equivariance error. Further, compared to a data augmentation baseline, TransIP achieves 40% to 60% improvement in performance across varying OMol25 dataset sizes. More broadly, our work shows that learned equivariance can be a powerful and efficient alternative to augmentation-based MLIP models.

1 Introduction

Atomistic simulations are a fundamental task in chemistry and materials science [1, 2], with Density Functional Theory (DFT) serving as a basis for accurately calculating interatomic forces and energies. However, the utility of DFT is severely restricted by its computational costs, which typically scale cubically with system size, rendering large-scale or long-timescale simulations intractable. This has motivated machine-learned interatomic potentials (MLIPs) to overcome this limitation by learning the potential energy surface from data, offering orders-of-magnitude speed-ups compared to DFT calculations [3–7].

Equivariant neural networks have become a central paradigm for MLIPs due to their ability to encode the three-dimensional structure of molecular graphs [8–11]. These architectures are designed to explicitly respect roto-translational symmetries ($SE(3)$ equivariance) by construction, often employing compute-intensive mechanisms like spherical harmonics or equivariant message passing [12–15]. However, due to the design difficulties and limited expressive power of these architectures [16, 17], a recent trend in predictive and generative modeling is to use unconstrained models when enough data is available [18–21].

In this paper, we introduce **TransIP (Transformer-based Interatomic Potentials)**, a training paradigm that achieves molecular symmetry for interatomic potentials *without* imposing architectural $SO(3)$

*Equal contribution

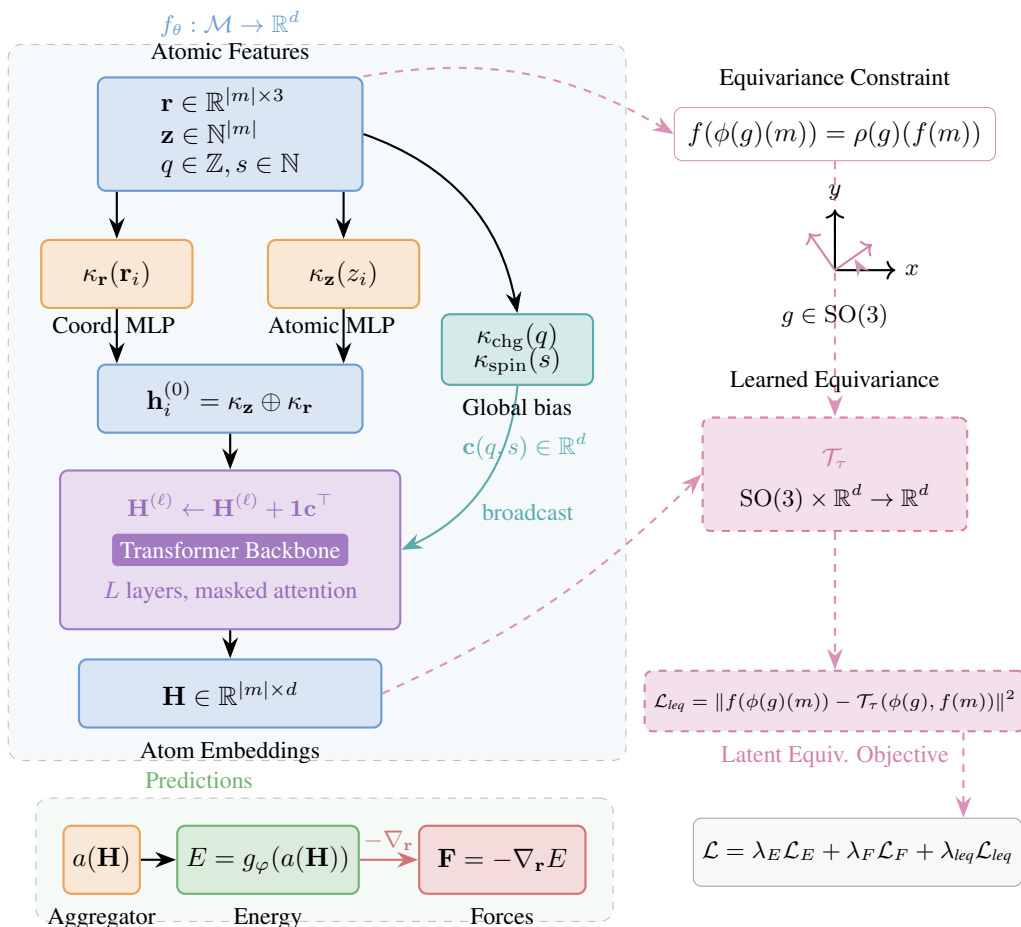


Figure 1: TransIP: Transformer-based Interatomic Potentials.

constraints. TransIP steers a standard transformer toward $\text{SO}(3)$ equivariance via an additional contrastive objective, allowing the model to retain the scalability and hardware efficiency of attention mechanisms while learning symmetry from data.

Our contributions are as follows:

- We propose an MLIP training pipeline with a general transformer-based model to obtain $\text{SO}(3)$ equivariance through training rather than hard-wired equivariant layers.
- We introduce an architecture-agnostic contrastive loss function that promotes $\text{SO}(3)$ equivariance in the embedding space of an unconstrained model. By aligning latent features across $\text{SO}(3)$ transformations in the model’s backbone, we show that TransIP scales better across different datasets and model sizes compared to traditional data augmentation techniques.
- On a diverse molecular benchmark, Open Molecules 25 [22] (that includes small organics, biomolecular fragments, electrolyte-like species), we show that TransIP outperforms data augmentation techniques often by a large margin, which might provide an alternative for the future of augmentation-based MLIP models.

2 Symmetry in Embedding Space

2.1 Problem Formulation

Molecular representations. Let $\mathcal{M}$ denote the space of molecular configurations. Each molecule $m \in \mathcal{M}$ is represented by atomic features $\mathbf{x} = (\mathbf{r}, \mathbf{z}, q, s)$, where $\mathbf{r} \in \mathbb{R}^{|m| \times 3}$ are atomic coordinates,

$\mathbf{z} \in \mathbb{N}^{|m|}$ are atomic numbers, $q \in \mathbb{Z}$ is the total molecular charge, and $s \in \mathbb{N}$ is the spin multiplicity, with $|m|$ denoting the number of atoms in molecule m .

Our goal is to learn an embedding function $f_\theta : \mathcal{M} \rightarrow \mathbb{R}^d$ that maps molecular configurations to a d -dimensional latent space, and a prediction function $g_\varphi : \mathbb{R}^d \rightarrow \mathbb{R}$ that acts in the embedding space $\mathbb{R}^d$ and outputs molecular properties (e.g., energy). Both f_θ and g_φ are neural networks parameterized by θ and φ , respectively.

Symmetry groups. We define a symmetry group G that acts on a set $\mathcal{X}$ as a group of bijective functions from $\mathcal{X}$ to itself, and the group operation is function composition. We say a function f is *equivariant* w.r.t. the group G if for every transformation $g \in G$ and every input $x \in \mathcal{X}$,

$$f(\phi(g)(x)) = \rho(g)(f(x)) \quad (1)$$

The group representations ϕ and ρ specify how we apply the elements of the group G on input and output data. As a concrete case, we can define G as a rotation group $\text{SO}(3)$ over molecular configurations $\mathcal{M}$, with $g \in \text{SO}(3)$ representing an element of G that acts on a molecule m by rotating the coordinates of each atom in 3D space. Formally, for a molecule $m = (\mathbf{r}, \mathbf{z}, q, s)$ with coordinates $\mathbf{r} = (\mathbf{r}_1, \dots, \mathbf{r}_{|m|})$, $\mathbf{r}_i \in \mathbb{R}^3$, the input action rotates each atom:

$$(\phi(g)m) = ((R\mathbf{r}_1, \dots, R\mathbf{r}_{|m|}), \mathbf{z}, q, s).$$

Here R is a 3×3 rotation matrix (orthogonal with $\det R = 1$); $\mathbf{z}, q, s$ are unchanged. An associated output representation rotates vector-valued quantities—e.g., for forces $\mathbf{F} = (\mathbf{F}_1, \dots, \mathbf{F}_{|m|})$, $\rho(g)\mathbf{F} = (R\mathbf{F}_1, \dots, R\mathbf{F}_{|m|})$ —while scalar outputs such as energies remain invariant, $\rho(g)E = E$.

2.2 Implicit Equivariance in Embedding Space

We seek an embedding function f that behaves equivariantly with respect to the symmetry group G , meaning there exists a transformation $\rho(g) : \mathbb{R}^d \rightarrow \mathbb{R}^d$ such that:

$$f(\phi(g)(m)) = \rho(g)(f(m)) \quad \forall g \in G, m \in \mathcal{M} \quad (2)$$

Common approaches enforce equivariance constraints through specialized architectures. Instead, we want the embedding function f to learn symmetry without equivariance constraints. However, with G being the rotation group $\text{SO}(3)$ on $\mathcal{M}$ and the output of f being a high-dimensional vector, there is no direct representation of $\rho(g)$ to act in the space of $\mathbb{R}^d$. Thus, rather than specifying $\rho(g)$ analytically, we propose to learn the group transformation on an embedding vector in $\mathbb{R}^d$ using a neural network $\mathcal{T}_\tau : \text{SO}(3) \times \mathbb{R}^d \rightarrow \mathbb{R}^d$ parameterized by τ . $\mathcal{T}$ can be understood as a non-linear function that learns the group action implicitly on a latent vector, by providing the group representation on the input data.

3 Learning Inter-Atomic Potentials without Explicit Equivariance

In this section, we introduce our training framework: TransIP (Transformer-based Inter-atomic Potentials), a new approach that achieves $\text{SO}(3)$ -equivariance through learned transformations in an embedding space without explicit equivariance constraints. Our method, illustrated in Figure 1, consists of three key components: (i) an unconstrained Transformer backbone that processes molecular configurations, (ii) a learned transformation network that performs group actions in the embedding space, and (iii) a contrastive objective that enforces latent equivariance (equiv.) during training.

3.1 TransIP: Transformer-based Interatomic Potentials

Atom as tokens. We model each molecule as a variable-length sequence of tokens, where each token represents an atom. Unlike conventional graph neural networks that construct edges based on distance cutoffs or neighbours’ atoms, we process all atoms within a molecule through self-attention, bounded by a maximum context length N_{ctx} . For batch processing, we use padding masks to prevent cross-molecule attention, ensuring each molecule is processed independently.

In addition, we apply rotary position embeddings (RoPE) [23] to the queries $\mathbf{q}_i \in \mathbb{R}^{d/h}$ and keys $\mathbf{k}_j \in \mathbb{R}^{d/h}$ of each attention head, where i, j denote the sequence positions of atoms within a

molecule, d is the model dimension, and h is the number of attention heads. The attention weights are computed as:

$$\begin{aligned}\tilde{\mathbf{q}}_i &= \text{RoPE}(\mathbf{q}_i, i), \quad \tilde{\mathbf{k}}_j = \text{RoPE}(\mathbf{k}_j, j) \\ \alpha_{ij} &= \text{softmax} \left(\frac{\tilde{\mathbf{q}}_i^\top \tilde{\mathbf{k}}_j}{\sqrt{d/h}} + m_{ij} \right)\end{aligned}$$

where $\text{RoPE}(\cdot, \cdot)$ is the rotary position encoding operator, and $m_{ij} \in \{0, -\infty\}$ is the attention mask that blocks padding tokens and enforces within-molecule attention. This approach eliminates the need for explicit distance cutoffs while maintaining flexibility in modeling molecular interactions.

Transformer Backbone. We implement the embedding function $f_\theta : \mathcal{M} \rightarrow \mathbb{R}^d$ as a Transformer encoder that processes atom-level tokens. Each atom i is initialized with a token representation:

$$\mathbf{h}_i^{(0)} = \kappa_{\mathbf{z}}(z_i) \oplus \kappa_{\mathbf{r}}(\mathbf{r}_i)$$

where $\kappa_{\mathbf{z}} : \mathbb{N} \rightarrow \mathbb{R}^d$ and $\kappa_{\mathbf{r}} : \mathbb{R}^3 \rightarrow \mathbb{R}^d$ are learnable MLPs that embed atomic numbers and centered coordinates (with $\mathbf{r}_i \leftarrow \mathbf{r}_i - \frac{1}{|m|} \sum_j \mathbf{r}_j$), and $\oplus$ denotes concatenation. These tokens are processed through L Transformer layers with masked self-attention within each molecule, producing final per-atom embeddings $\mathbf{H} = [\mathbf{h}_1, \dots, \mathbf{h}_{|m|}]^\top \in \mathbb{R}^{|m| \times d}$.

Global Molecular Properties. Following Levine et al. [22], we incorporate global molecular properties (total charge q and spin multiplicity s of a molecule m) through learnable embeddings, and form a graph-level bias:

$$\mathbf{c}(q, s) = \kappa_{\text{chg}}(q) + \kappa_{\text{spin}}(s) \in \mathbb{R}^d$$

where κ_{chg} and κ_{spin} are learnable embedding functions for charge and spin, respectively. This global bias is broadcast-added at each Transformer layer: $\mathbf{H}^{(\ell)} \leftarrow \mathbf{H}^{(\ell)} + \mathbf{1c}(q, s)^\top$.

Energy and Force Predictions. For molecular property prediction, we employ a permutation-invariant aggregator $a : \mathbb{R}^{|m| \times d} \rightarrow \mathbb{R}^d$ followed by an energy prediction head $g_\varphi : \mathbb{R}^d \rightarrow \mathbb{R}$:

$$E_\varphi(m) = g_\varphi(a(\mathbf{H}))$$

Forces are computed as conservative gradients of the energy with respect to atomic positions:

$$\mathbf{F}(m) = -\nabla_{\mathbf{r}} E_\varphi(m) \in \mathbb{R}^{|m| \times 3}$$

3.2 Learned Latent Equivariance

Transformation Network. We propose a transformation network $\mathcal{T}_\tau : \text{SO}(3) \times \mathbb{R}^d \rightarrow \mathbb{R}^d$ that learns how group actions (e.g., rotations) act on molecular embeddings. We implement $\mathcal{T}_\tau$ as a multilayer perceptron that takes as input the group representation in the input domain $\phi(g)$ and the molecular embedding $f(m)$. Formally,

$$\mathcal{T}_\tau(\phi(g), f(m)) = \text{MLP}_\tau([\phi(g), f(m)])$$

where $[\cdot, \cdot]$ denotes concatenation and MLP_τ is a multilayer perceptron with parameters τ .

Contrastive Objective for Latent Equivariance: To learn the molecular symmetry without architectural constraints, we define our latent equivariance loss as:

$$\mathcal{L}_{\text{leq}}(\phi(g), m, f, \mathcal{T}) = \|f(\phi(g)(m)) - \mathcal{T}_\tau(\phi(g), f(m))\|^2 \quad (3)$$

This loss encourages the embedding function f to behave equivariantly with respect to the symmetry group G , as mediated by the transformation network $\mathcal{T}_\tau$. During training, we sample a molecule m from the dataset and a rotation element g uniformly from $\text{SO}(3)$ and minimize the expected latent loss:

$$\min \mathbb{E}_{m \sim \mathcal{M}, g \sim \text{SO}(3)} [\mathcal{L}_{\text{leq}}(\phi(g), m, f, \mathcal{T})] \quad (4)$$

3.3 Training Objective

Our training objective combines three complementary losses for accurate prediction of energy and forces as well as implicitly learning molecular symmetry.

Prediction Losses. For energy and force predictions, we use:

$$\mathcal{L}_E = \frac{1}{|m|} |E_\varphi(m) - E^*| \quad (\text{per-atom mean absolute error (MAE)}) \quad (5)$$

$$\mathcal{L}_F = \frac{1}{3|m|} \|\mathbf{F}(m) - \mathbf{F}^*\|_F^2 \quad (\text{per-molecule mean squared error (MSE)}) \quad (6)$$

where E^* and $\mathbf{F}^*$ are ground-truth energies and forces, and $\|\cdot\|_F$ denotes the Frobenius norm. For energies, we use referenced targets as described by Levine et al. [22].

Combined Objective. Training combines three weighted terms: (i) the latent equivariance target $\mathcal{L}_{leq}$ defined in Eq. 3; (ii) energy loss $\mathcal{L}_E$; and (iii) force loss $\mathcal{L}_F$. The total objective is

$$\mathcal{L}_{\text{total}} = \lambda_E \mathcal{L}_E + \lambda_F \mathcal{L}_F + \lambda_{leq} \mathcal{L}_{leq} \quad (7)$$

where λ_E , λ_F , and λ_{leq} are hyperparameters for each loss. The optimal hyperparameters are given in Table 3 of Appendix B.

4 Related Work

ML Interatomic Potentials. Using machine learning (ML) methods to predict energies and forces of different molecular systems and materials has been an active area of research [24–29]. Due to the intricate 3D structures of atomistic systems, equivariant message-passing neural networks have been an essential backbone in this domain. For example, Gastegger et al. [30], Klicpera et al. [31] introduced equivariant directional message passing between pairs of atoms with a spherical harmonics representation. In contrast, Batzner et al. [4] developed equivariant convolution with tensor-products and Batatia et al. [5] built higher-order messages with equivariant graph neural networks [32]. Additionally, Passaro and Zitnick [33] reduced the computational complexity of SO(3) convolution and replaced it with SO(2) convolutions, which have been used as a backbone for MLIPs [11]. More recently, [34] presented Orb-v3 models with improved computational efficiency, built on Graph Network Simulators [35].

Unconstrained ML models. While current-state-of-the-art MLIP models primarily rely on equivariant GNNs, unconstrained models are actively used in other domains. For example, integrating data augmentation via image transformations has been used in different vision tasks, from classification [36–38] to segmentation [39, 40]. For geometric data, the use of unconstrained models and diffusion Transformers (without explicit equivariance constraints) has been a recent trend in generative tasks, e.g., AlphaFold 3 for biomolecular structure prediction [19] as well as molecular conformation and materials generation [18, 20, 21]. In contrast, several works have been introduced to overcome the limitations of strictly equivariant GNNs by enforcing symmetry via frame averaging over geometric inputs [41–45]; learning canonicalization functions that map inputs to a canonical orientation before prediction [46–49]; or learning equivariance through data augmentation with molecule-specific graph-based architectures [50, 51]. However, in this work, we demonstrate that an unconstrained general-purpose Transformer model can serve as a backbone for MLIPs, which replaces graph-based inductive biases with a scalable latent equivariance objective that implicitly learns equivariant features without explicit equivariance constraints.

5 Experimental Setup

Dataset. We train and evaluate our proposed method **TransIP** on the Open Molecules 2025 (OMol25) collection [22], a large-scale molecular DFT dataset for ML interatomic potentials. OMol25 covers 83 atomic elements and diverse chemistries (such as neutral organics, biomolecules, electrolytes, and metal complexes). Following Levine et al. [22], we use the official 4M training split (3,986,754) and the out-of-distribution composition validation split *Val-Comp* (2,762,021). *Val-Comp* consists of molecules gathered from various datasets and domains, such as biomolecules, neutral organics, and metal complexes.

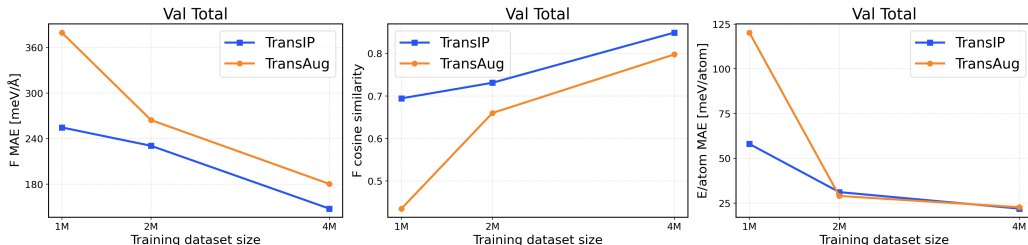


Figure 2: Val-Comp performance across dataset sizes (1M / 2M / 4M): force metrics (MAE, cosine similarity) and energy-per-atom MAE.

Model Configurations. We evaluate TransIP across three model scales: Small (14M parameters), Medium (85M parameters), and Large (302M parameters). All models use MLP-based coordinate embeddings and RoPE positional encodings. The transformation network $\mathcal{T}_\tau$ is a 2-layer MLP with GELU activations and 2d hidden dimension.

Training Setup. Using the standardized FAIRCHEM Python package [52], we train TransIP on the OMol25 dataset using an AdamW optimizer with learning rate 5×10^{-4} , weight decay 10^{-3} , and gradient norm clipping at 200. We use a cosine learning rate schedule with linear warmup over the first 1% of training, followed by cosine decay down to 1% of the initial 1τ . The loss weights are set to $\lambda_E = 5$ for energies and $\lambda_F = 15$ for forces. For the latent equivariance objective λ_{leq} , we sweep the values in $\{1, 5, 10, 100\}$ and selected $\lambda_{leq} = 5$ based on validation performance.

Scalability Experiments. We conduct two sets of experiments to assess TransIP’s scaling behavior:

- **Data scaling:** We train the Small (14M parameter) model on three dataset sizes (1M, 2M, 4M molecules) for 5 epochs using 8 NVIDIA 80GB GPUs, comparing TransIP with learned equivariance against an unconstrained Transformer version with SO(3) data augmentation (TransAug).
- **Model size scaling.** We compare TransIP and TransAug with different model sizes (Small/Medium/Large) trained on the same number of samples from the OMol25 4M dataset and report the evaluation metrics as a function of the processed number of atoms per second.
- **Extended training:** We train TransIP (Small) on the OMol25 4M dataset for 40 epochs using 64 NVIDIA 80GB GPUs to evaluate its performance against standardized equivariant baselines.

Baselines. We compare TransIP against: (i) an *unconstrained* TransIP variant trained with SO(3) rotation augmentation to assess the impact of learned latent equivariance versus data augmentations, and (ii) state-of-the-art equivariant models on OMol25: eSCN [11] in small/medium configurations with both direct and energy-conserving force variants as well as GemNet-OC [53].

Evaluation metrics. Following the OMol25 official benchmark, we report: Force MAE (eV/Å), Force cosine similarity, Energy per atom MAE (eV/atom), and Total energy MAE (eV). Detailed metric definitions are provided in Appendix B.4.

6 Results and Discussion

6.1 Scaling data size

To assess how performance scales with different training dataset sizes, we compare our latent equivariance-based model (TransIP) against an unconstrained baseline that uses SO(3) data augmentation (TransAug). Both models use a (small) 14M parameter Transformer architecture. Given our tight compute budget, we train on 1M, 2M, and 4M OMol25 molecules for 5 epochs and report validation (*Val-Comp*) results.

Performance in a limited data regime. Figure 2 shows that TransIP delivers large gains when trained on 1M samples and outperforms TransAug across all evaluation metrics with a large margin. The learned latent equivariance objective provides substantial improvements in force MAE (0.255 eV/Å vs

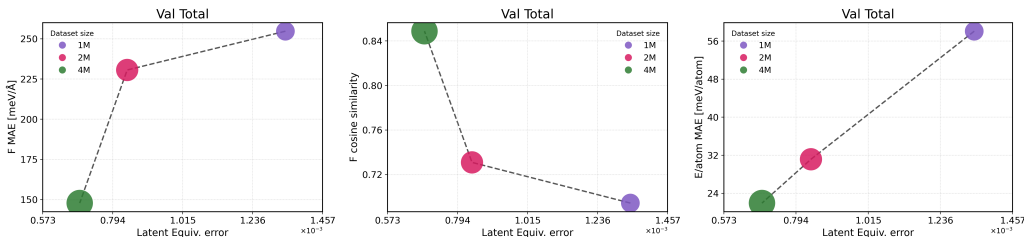


Figure 3: Latent equivariance (embedding) error vs validation performance: force metrics (MAE, cosine similarity) and energy-per-atom MAE.

0.6 eV/Å MAE) and directional consistency (0.7 vs 0.44 force cosine similarity). Energy predictions also benefit from the latent equivariance objective, with TransIP achieving 0.0581 eV/atom compared to TransAug’s 0.1203 eV/atom. These results suggest that learning equivariance in a latent space is a more effective scheme to incorporate molecular symmetry than data augmentation, particularly when training data is limited.

Performance in a larger data regime. As we scale to 2M and 4M molecules, both models (TransIP and TransAug) improve across the evaluation metrics. However, on larger datasets, TransIP still achieves better force MAEs and cosine similarity metrics compared to TransAug. This might indicate that the learned transformation network can successfully capture the geometric relationships necessary for accurate force predictions. Notably, energy prediction performance converges between the two at larger data scales, with both methods achieving comparable per-atom MAE values. This convergence suggests that while learned equivariance provides crucial benefits for force-related metrics in all data regimes, its advantages for energy prediction become less pronounced as the model can learn invariant energy representations from sufficient augmented data.

6.2 Learned latent equivariance

We investigate how learned equivariance affects the embedding space in relation to validation performance as the data scale increases. Figure 3 plots each metric against latent equiv. error for TransIP (Small) trained for 5 epochs on 1M, 2M, and 4M molecules (see Table 2 for a detailed definition of each model configuration).

Lower latent equivariance error leads to better accuracy. We found that the learned equiv. error serves as a strong predictor of model performance. Across all metrics, we observe a clear monotonic trend: lower equiv. error is associated with better performance (Figure 3). However, energy and force predictions respond differently to improvements in equivariance. Energy predictions show near-linear scaling with equiv. error, indicating that energy accuracy is directly limited by equivariance quality. This strong coupling aligns with energies being scalar invariants that depend primarily on learning correct symmetry-preserving features. In contrast, force predictions exhibit a two-regime behavior: initial improvements in equivariance (1M→2M) yield modest force improvements, while further tightening of equivariance (2M→4M) produces disproportionate gains. This might indicate that forces require both accurate equivariant features and sufficient data diversity to learn the energy landscape’s geometry.

These results demonstrate that implicitly learning equivariance through our learned transformation network provides an efficient inductive bias, accelerating learning. The 48% reduction in equiv. error from 1M to 4M training examples translates to 40-60% performance improvements, being more efficient than what would be expected from data scaling alone.

Learning equivariance leads to faster inference. To measure the inference efficiency of our method, we compare TransIP and TransAug with different model sizes (Small/Medium/Large) trained on 4M samples and report the evaluation metrics as a function of the processed number of atoms per second. However, due to limited compute, we compare models under a *fixed training budget* (i.e., with the same number of samples), which is 10k, 25k, and 100k steps for our Small, Medium, and Large models, respectively.

From the results in Figure 4, we see that TransIP scales smoothly with parameter count despite limited training: As model size grows, performance improves across all metrics. In contrast, TransAug

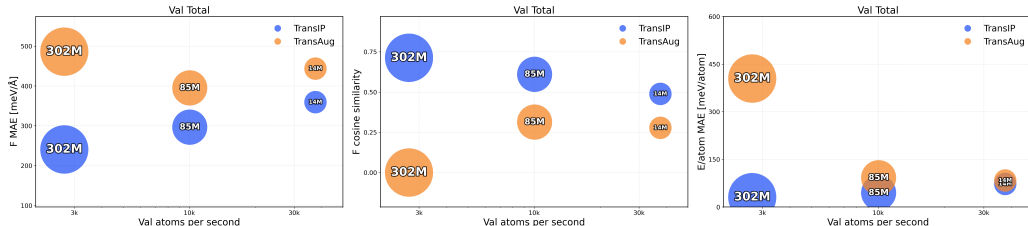


Figure 4: Validation total inference trade-off (atoms/s vs performance) for force metrics (MAE, cosine similarity) and energy-per-atom MAE.

Table 1: Comprehensive Val-Comp energy and force MAE results.

Model	Epochs	Biomolecules		Electrolytes		Metal Complexes		Neutral Organics		Total	
		Energy ↓	Forces ↓	Energy ↓	Forces ↓	Energy ↓	Forces ↓	Energy ↓	Forces ↓	Energy ↓	Forces ↓
eSEN-sm-d.	80	0.00088	0.00812	0.00193	0.01264	0.00337	0.04044	0.00216	0.02017	0.00219	0.01301
eSEN-sm-cons.	80	0.00086	0.00617	0.00161	0.01116	0.00272	0.03533	0.00150	0.01692	0.00189	0.01110
eSEN-md-d.	80	0.00047	0.00338	0.00118	0.00651	0.00253	0.02731	0.00121	0.00926	0.00132	0.00678
GemNet-OC-r6	80	0.00040	0.00584	0.00139	0.00937	0.00274	0.03360	0.00188	0.01655	0.00141	0.00983
GemNet-OC	80	0.00025	0.00520	0.00104	0.00842	0.00266	0.03276	0.00164	0.01559	0.00113	0.00898
TransAug	5	0.0166	0.2193	0.0175	0.1619	0.0207	0.1506	0.0289	0.2188	0.0235	0.1803
TransIP	5	0.0173	0.1811	0.0159	0.1296	0.0185	0.1325	0.0235	0.1650	0.0223	0.1466
TransIP	40	0.0138	0.1215	0.0127	0.0940	0.0152	0.1056	0.0185	0.1254	0.0179	0.1038
TransIP (in progress)	80	—	—	—	—	—	—	—	—	—	—

exhibits poorer scaling—larger models perform worse than smaller ones, with the Large model configuration yielding the lowest performance. This might indicate that augmentation alone does not provide a sufficiently informative and stable inductive bias for large-capacity models trained for molecular force field prediction.

6.3 Architectural equivariance versus learned equivariance

Table 1 compares the energy and force prediction performance of TransIP (Small) against TransAug (Small) as well as several well-known equivariant baselines for the OMol 2M Val-Comp evaluation dataset. The results of this comparison demonstrate that TransIP outperforms TransAug (trained for 5 epochs) in all but one evaluation metric, particularly differentiating itself in terms of force prediction. We further report the performance of TransIP trained for 40 epochs and, due to limited compute, we are currently training the model for a full 80 epochs (for fair comparison to each equivariant baseline). Results with TransIP after 40 training epochs suggest steady improvement is likely to be observed during the remainder of the model’s training epochs.

7 Conclusion

In this work, we introduced TransIP for modeling interatomic potentials with a modern Transformer-based architecture and a scalable latent equivariance objective. Empirical results across a variety of chemical systems as well as model and dataset scales suggest that TransIP’s latent equivariance objective enables better performance scaling than popular data augmentation-based alternatives to learning geometric equivariance. Further, we find that improvements in learning latent equivariance are strongly related to improved modeling of interatomic potentials, suggesting a complementary nature between the two prediction objectives. With sufficient compute, future work could involve studying the performance of TransIP in larger data, modeling, and runtime regimes in addition to the behavior of TransIP in a context amenable to the double-descent phenomenon [54].

While equivariant models for molecular machine learning have recently gained much research interest, with the large amount of data being generated and the need for larger model sizes, it is also important that models used for interatomic potentials be highly scalable. Through our work, we have shown that the generic Transformer is capable of modeling molecules accurately but is also able to learn equivariance effectively through our novel latent objective, all while being highly scalable. By making our code openly available to the research community, we hope that our work inspires future research that explores ways to leverage the simpler and more scalable Transformer architecture to better model equivariant molecular properties through learned equivariance.

Acknowledgments. This research is partially supported by EPSRC Turing AI World-Leading Research Fellowship No. EP/X040062/1, EPSRC AI Hub on Mathematical Foundations of Intelligence: An “Erlangen Programme” for AI No. EP/Y028872/1. Further, this research used resources of the National Energy Research Scientific Computing Center (NERSC), a U.S. Department of Energy (DOE) User Facility, using an AI4Sci@NERSC award (DDR-ERCAP 0034574) awarded to AM. S.M.B. acknowledges support from the Center for High Precision Patterning Science (CHiPPS), an Energy Frontier Research Center funded by the U.S. DOE, Office of Science, Basic Energy Sciences (BES). AR’s PhD is supported by the Agency for Science Technology and Research and the SABS R3 CDT program via the Engineering and Physical Sciences Research Council. AR also received compute resources from the DSO National Laboratories - AI Singapore (AISG) programme and the Lawrence Livermore National Laboratory. We would like to thank them for their resources, which played a significant role in this research. We would also like to thank Santiago Vargas, Chaitanya Joshi, and Chen Lin for their fruitful discussions.

References

- [1] Linfeng Zhang, Jiequn Han, Han Wang, Roberto Car, and Weinan E. Deep potential molecular dynamics: A scalable model with the accuracy of quantum mechanics. *Physical Review Letters*, 120(14), April 2018. ISSN 1079-7114. doi: 10.1103/physrevlett.120.143001. URL <http://dx.doi.org/10.1103/PhysRevLett.120.143001>.
- [2] Volker L. Deringer, Miguel A. Caro, and Gábor Csányi. Machine learning interatomic potentials as emerging tools for materials science. *Advanced Materials*, 2019.
- [3] Frank Noé, Alexandre Tkatchenko, Klaus-Robert Müller, and Cecilia Clementi. Machine learning for molecular simulation. *Annual Review of Physical Chemistry*, 71(1):361–390, April 2020. ISSN 1545-1593. doi: 10.1146/annurev-physchem-042018-052331. URL <http://dx.doi.org/10.1146/annurev-physchem-042018-052331>.
- [4] Simon Batzner, Albert Musaelian, Lixin Sun, et al. E(3)-equivariant graph neural networks for data-efficient and accurate interatomic potentials. *Nature Communications*, 13:2453, 2022. doi: 10.1038/s41467-022-29939-5. URL <https://doi.org/10.1038/s41467-022-29939-5>.
- [5] Ilyes Batatia, David Peter Kovacs, Gregor N. C. Simm, Christoph Ortner, and Gabor Csanyi. MACE: Higher order equivariant message passing neural networks for fast and accurate force fields. In Alice H. Oh, Alekh Agarwal, Danielle Belgrave, and Kyunghyun Cho, editors, *Advances in Neural Information Processing Systems*, 2022. URL <https://openreview.net/forum?id=YPpSngE-ZU>.
- [6] Ryan Jacobs, Dane Morgan, Siamak Attarian, Jun Meng, Chen Shen, Zhenghao Wu, Clare Yijia Xie, Julia H. Yang, Nongnuch Artrith, Ben Blaiszik, Gerbrand Ceder, Kamal Choudhary, Gabor Csanyi, Ekin Dogus Cubuk, Bowen Deng, Ralf Drautz, Xiang Fu, Jonathan Godwin, Vasant Honavar, Olexandr Isayev, Anders Johansson, Boris Kozinsky, Stefano Martiniani, Shyue Ping Ong, Igor Poltavsky, KJ Schmidt, So Takamoto, Aidan P. Thompson, Julia Westermayr, and Brandon M. Wood. A practical guide to machine learning interatomic potentials – status and future. *Current Opinion in Solid State and Materials Science*, 35:101214, March 2025. ISSN 1359-0286. doi: 10.1016/j.cossms.2025.101214. URL <http://dx.doi.org/10.1016/j.cossms.2025.101214>.
- [7] Niklas Leimeroth, Linus C. Erhard, Karsten Albe, and Jochen Rohrer. Machine-learning interatomic potentials from a users perspective: A comparison of accuracy, speed and data efficiency, 2025. URL <https://arxiv.org/abs/2505.02503>.
- [8] Brandon Anderson, Truong Son Hy, and Risi Kondor. Cormorant: Covariant molecular neural networks. In *Advances in Neural Information Processing Systems*, volume 32. Curran Associates, Inc., 2019. URL https://proceedings.neurips.cc/paper_files/paper/2019/file/03573b32b2746e6e8ca98b9123f2249b-Paper.pdf.
- [9] Philipp Thölke and Gianni De Fabritiis. Equivariant transformers for neural network based molecular potentials. In *International Conference on Learning Representations*, 2022. URL <https://openreview.net/forum?id=zNHqZ9wrRB>.

- [10] Yi-Lun Liao, Brandon Wood, Abhishek Das*, and Tess Smidt*. EquiformerV2: Improved Equivariant Transformer for Scaling to Higher-Degree Representations. In *International Conference on Learning Representations (ICLR)*, 2024. URL <https://openreview.net/forum?id=mCOBKZmrzD>.
- [11] Xiang Fu, Brandon M Wood, Luis Barroso-Luque, Daniel S. Levine, Meng Gao, Misko Dzamba, and C. Lawrence Zitnick. Learning smooth and expressive interatomic potentials for physical property prediction. In *Forty-second International Conference on Machine Learning*, 2025. URL <https://openreview.net/forum?id=R0PBjxIbgm>.
- [12] Fabian B. Fuchs, Daniel E. Worrall, Volker Fischer, and Max Welling. Se(3)-transformers: 3d roto-translation equivariant attention networks. In *Advances in Neural Information Processing Systems 34 (NeurIPS)*, 2020.
- [13] Saro Passaro and C. Lawrence Zitnick. Reducing so(3) convolutions to so(2) for efficient equivariant gnns. In *Proceedings of the 40th International Conference on Machine Learning*, ICML’23. JMLR.org, 2023.
- [14] Yi-Lun Liao and Tess Smidt. Equiformer: Equivariant graph attention transformer for 3d atomistic graphs. In *International Conference on Learning Representations*, 2023. URL <https://openreview.net/forum?id=KwmPfARgOTD>.
- [15] Moin Uddin Maruf, Sungmin Kim, and Zeeshan Ahmad. Learning long-range interactions in equivariant machine learning interatomic potentials via electronic degrees of freedom. *The Journal of Physical Chemistry Letters*, 16(35):9078–9087, August 2025. ISSN 1948-7185. doi: 10.1021/acs.jpclett.5c02352. URL <http://dx.doi.org/10.1021/acs.jpclett.5c02352>.
- [16] Chaitanya K. Joshi, Cristian Bodnar, Simon V Mathis, Taco Cohen, and Pietro Lio. On the expressive power of geometric graph neural networks. In *Proceedings of the 40th International Conference on Machine Learning*, 2023. URL <https://proceedings.mlr.press/v202/joshi23a.html>.
- [17] Jiacheng Cen, Anyi Li, Ning Lin, Yuxiang Ren, Zihe Wang, and Wenbing Huang. Are high-degree representations really unnecessary in equivariant graph neural networks? In *The Thirty-eighth Annual Conference on Neural Information Processing Systems*, 2024. URL <https://openreview.net/forum?id=M0ncNVuGYN>.
- [18] Yuyang Wang, Ahmed A. Elhag, Navdeep Jaitly, Joshua M. Susskind, and Miguel Ángel Bautista. Swallowing the bitter pill: Simplified scalable conformer generation. In *Forty-first International Conference on Machine Learning*, 2024.
- [19] Jacob Abramson, Jane Adler, John Dunger, et al. Accurate structure prediction of biomolecular interactions with alphafold 3. *Nature*, 630:493–500, 2024. doi: 10.1038/s41586-024-07487-w. URL <https://doi.org/10.1038/s41586-024-07487-w>.
- [20] Leo Zhang, Kianoosh Ashouritaklimi, Yee Whye Teh, and Rob Cornish. Symdiff: Equivariant diffusion via stochastic symmetrisation. In *The Thirteenth International Conference on Learning Representations*, 2025. URL <https://openreview.net/forum?id=i1NNCrRxdM>.
- [21] Chaitanya K. Joshi, Xiang Fu, Yi-Lun Liao, Vahe Gharakhanyan, Benjamin Kurt Miller, Anuroop Sriram, and Zachary W. Ulissi. All-atom diffusion transformers: Unified generative modelling of molecules and materials. In *International Conference on Machine Learning*, 2025.
- [22] Daniel S. Levine, Muhammed Shuaibi, Evan Walter Clark Spotte-Smith, Michael G. Taylor, Muhammad R. Hasyim, Kyle Michel, Ilyes Batatia, Gábor Csányi, Misko Dzamba, Peter Eastman, Nathan C. Frey, Xiang Fu, Vahe Gharakhanyan, Aditi S. Krishnapriyan, Joshua A. Rackers, Sanjeev Raja, Ammar Rizvi, Andrew S. Rosen, Zachary Ulissi, Santiago Vargas, C. Lawrence Zitnick, Samuel M. Blau, and Brandon M. Wood. The open molecules 2025 (omol25) dataset, evaluations, and models, 2025. URL <https://arxiv.org/abs/2505.08762>.
- [23] Jianlin Su, Yu Lu, Shengfeng Pan, Ahmed Murtadha, Bo Wen, and Yunfeng Liu. Roformer: Enhanced transformer with rotary position embedding, 2023. URL <https://arxiv.org/abs/2104.09864>.

- [24] Kristof T. Schütt, Huziel E. Sauceda, P.-J. Kindermans, Alexandre Tkatchenko, and Klaus-Robert Müller. Schnet: A continuous-filter convolutional neural network for modeling quantum interactions. In *NeurIPS*, 2017.
- [25] Stefan Chmiela, Valentin Vassilev-Galindo, Oliver T. Unke, Adil Kabylda, Huziel E. Sauceda, Alexandre Tkatchenko, and Klaus-Robert Müller. Accurate global machine learning force fields for molecules with hundreds of atoms, 2022. URL <https://arxiv.org/abs/2209.14865>.
- [26] Albert Musaelian, Sebastian Batzner, Anders Johansson, Simon Kozinsky, and Boris Kozinsky. Learning local 3d energetics with graph neural networks. *Nature Communications*, 2023.
- [27] Yi-Lun Liao et al. Equiformerv2: Improved equivariant transformer for scaling 3d molecular learning. *arXiv preprint arXiv:2402.xxxxx*, 2024.
- [28] Ziduo Yang, Xian Wang, Yifan Li, Qiuji Lv, Calvin Yu-Chian Chen, and Lei Shen. Efficient equivariant model for machine learning interatomic potentials. *npj Computational Materials*, 11(1):49, 2025. doi: 10.1038/s41524-025-01535-3. URL <https://doi.org/10.1038/s41524-025-01535-3>.
- [29] Eric C-Y Yuan, Yunsheng Liu, Junmin Chen, Peichen Zhong, Sanjeev Raja, Tobias Kreiman, Santiago Vargas, Wenbin Xu, Martin Head-Gordon, Chao Yang, et al. Foundation models for atomistic simulation of chemistry and materials. *arXiv preprint arXiv:2503.10538*, 2025.
- [30] Johannes Gasteiger, Janek Groß, and Stephan Günnemann. Directional message passing for molecular graphs. In *International Conference on Learning Representations*, 2020. URL <https://openreview.net/forum?id=B1eWbxStPH>.
- [31] Johannes Klicpera, Florian Becker, and Stephan Günnemann. Gemnet: Universal directional graph neural networks for molecules. In A. Beygelzimer, Y. Dauphin, P. Liang, and J. Wortman Vaughan, editors, *Advances in Neural Information Processing Systems*, 2021. URL https://openreview.net/forum?id=HS_s0axS9K-.
- [32] Victor Garcia Satorras, Emiel Hooeboom, and Max Welling. E(n) equivariant graph neural networks. In *Proceedings of the 38th International Conference on Machine Learning*, 2021.
- [33] Saro Passaro and C. Lawrence Zitnick. Reducing SO(3) convolutions to SO(2) for efficient equivariant GNNs. In *Proceedings of the 40th International Conference on Machine Learning*, Proceedings of Machine Learning Research. PMLR, 2023.
- [34] Benjamin Rhodes, Sander Vandenhaute, Vaidotas Šimkus, James Gin, Jonathan Godwin, Tim Duignan, and Mark Neumann. Orb-v3: atomistic simulation at scale, 2025. URL <https://arxiv.org/abs/2504.06231>.
- [35] Alvaro Sanchez-Gonzalez, Jonathan Godwin, Tobias Pfaff, Rex Ying, Jure Leskovec, and Peter W. Battaglia. Learning to simulate complex physics with graph networks. In *Proceedings of the 37th International Conference on Machine Learning*. JMLR.org, 2020.
- [36] Hiroshi Inoue. Data augmentation by pairing samples for images classification, 2018. URL <https://arxiv.org/abs/1801.02929>.
- [37] Alexey Dosovitskiy, Lucas Beyer, Alexander Kolesnikov, Dirk Weissenborn, Xiaohua Zhai, Thomas Unterthiner, Mostafa Dehghani, Matthias Minderer, Georg Heigold, Sylvain Gelly, Jakob Uszkoreit, and Neil Houlsby. An image is worth 16x16 words: Transformers for image recognition at scale. In *International Conference on Learning Representations*, 2021. URL <https://openreview.net/forum?id=YicbFdNTTy>.
- [38] Fazle Rahat, M Shifat Hossain, Md Rubel Ahmed, Sumit Kumar Jha, and Rickard Ewetz. Data augmentation for image classification using generative ai, 2024. URL <https://arxiv.org/abs/2409.00547>.
- [39] Misgana Negassi, Diane Wagner, and Alexander Reiterer. Smart(sampling)augment: Optimal and efficient data augmentation for semantic segmentation. *Algorithms*, 15(5), 2022. ISSN 1999-4893. doi: 10.3390/a15050165. URL <https://www.mdpi.com/1999-4893/15/5/165>.

- [40] Xinyi Yu, Guanbin Li, Wei Lou, Siqi Liu, Xiang Wan, Yan Chen, and Haofeng Li. Diffusion-based data augmentation for nuclei image segmentation. In *Medical Image Computing and Computer Assisted Intervention – MICCAI 2023*, 2023.
- [41] Omri Puny, Matan Atzmon, Edward J. Smith, Ishan Misra, Aditya Grover, Heli Ben-Hamu, and Yaron Lipman. Frame averaging for invariant and equivariant network design. In *International Conference on Learning Representations*, 2022. URL <https://openreview.net/forum?id=zIUyj55nXR>.
- [42] Alexandre Agm Duval, Victor Schmidt, Alex Hernández-García, Santiago Miret, Fragkiskos D. Malliaros, Yoshua Bengio, and David Rolnick. FAENet: Frame averaging equivariant GNN for materials modeling. In *Proceedings of the 40th International Conference on Machine Learning*, 2023. URL <https://proceedings.mlr.press/v202/duval23a.html>.
- [43] Yuchao Lin, Jacob Helwig, Shurui Gui, and Shuiwang Ji. Equivariance via minimal frame averaging for more symmetries and efficiency. In *Forty-first International Conference on Machine Learning*, 2024. URL <https://openreview.net/forum?id=guFsTBXsov>.
- [44] Tinglin Huang, Zhenqiao Song, Rex Ying, and Wengong Jin. Protein-nucleic acid complex modeling with frame averaging transformer. In *The Thirty-eighth Annual Conference on Neural Information Processing Systems*, 2024. URL <https://openreview.net/forum?id=Xngi3Z3wkN>.
- [45] Nadav Dym, Hannah Lawrence, and Jonathan W. Siegel. Equivariant frames and the impossibility of continuous canonicalization. In *ICML*, 2024. URL <https://openreview.net/forum?id=4iy0q0carb>.
- [46] Sékou-Oumar Kaba, Arnab Kumar Mondal, Yan Zhang, Yoshua Bengio, and Siamak Ravanbakhsh. Equivariance with learned canonicalization functions. In *NeurIPS 2022 Workshop on Symmetry and Geometry in Neural Representations*, 2022. URL <https://openreview.net/forum?id=pVD1k8ge25a>.
- [47] Justin Baker, Shih-Hsin Wang, Tommaso de Fernex, and Bao Wang. An explicit frame construction for normalizing 3d point clouds. In *Forty-first International Conference on Machine Learning*, 2024. URL <https://openreview.net/forum?id=SZ0JnRxi0x>.
- [48] George Ma, Yifei Wang, Derek Lim, Stefanie Jegelka, and Yisen Wang. A canonicalization perspective on invariant and equivariant learning. In *The Thirty-eighth Annual Conference on Neural Information Processing Systems*, 2024. URL <https://openreview.net/forum?id=jjcY92FX4R>.
- [49] Peter Lippmann, Gerrit Gerhartz, Roman Remme, and Fred A Hamprecht. Beyond canonicalization: How tensorial messages improve equivariant message passing. In *International Conference on Representation Learning*, volume 2025, 2025. URL https://proceedings.iclr.cc/paper_files/paper/2025/file/db7534a06ace69f4ec95bc89e91d5dbb-Paper-Conference.pdf.
- [50] Eric Qu and Aditi Krishnapriyan. The importance of being scalable: Improving the speed and accuracy of neural network interatomic potentials across chemical domains. *Advances in Neural Information Processing Systems*, 37:139030–139053, 2024.
- [51] Arslan Mazitov, Filippo Bigi, Matthias Kellner, Paolo Pegolo, Davide Tisi, Guillaume Fraux, Sergey Pozdnyakov, Philip Loche, and Michele Ceriotti. Pet-mad, a lightweight universal interatomic potential for advanced materials modeling, 2025. URL <https://arxiv.org/abs/2503.14118>.
- [52] Muhammed Shuaibi, Abhishek Das, Anuroop Sriram, Misko, Luis Barroso-Luque, Ray Gao, Siddharth Goyal, Zachary Ulissi, Brandon Wood, Tian Xie, Junwoong Yoon, Brook Wander, Adeesh Kolluru, Richard Barnes, Ethan Sunshine, Kevin Tran, Xiang, Daniel Levine, Nima Shoghi, Ilias Chair, , Janice Lan, Kaylee Tian, Joseph Musielewicz, clz55, Weihua Hu, , Kyle Michel, willis, and vbtchr. FAIRChem, 2025. URL <https://github.com/facebookresearch/fairchem>.

- [53] Johannes Gasteiger, Muhammed Shuaibi, Anuroop Sriram, Stephan Günnemann, Zachary Ward Ulissi, C. Lawrence Zitnick, and Abhishek Das. Gemnet-OC: Developing graph neural networks for large and diverse molecular simulation datasets. *Transactions on Machine Learning Research*, 2022. ISSN 2835-8856. URL <https://openreview.net/forum?id=u8tvSxm4Bs>.
- [54] Alethea Power, Yuri Burda, Harri Edwards, Igor Babuschkin, and Vedant Misra. Grokking: Generalization beyond overfitting on small algorithmic datasets. *arXiv preprint arXiv:2201.02177*, 2022.

A What TransIP Learns

To understand the structure of learned equivariance, we ask whether the effect of rotating different inputs can be explained by a *single* group action in the latent space; i.e., whether there exists a representation $\rho(g) : \mathbb{R}^d \rightarrow \mathbb{R}^d$ such that

$$f(\phi(g)(m)) \approx \rho(g) f(m),$$

where $f_\theta : \mathcal{M} \rightarrow \mathbb{R}^d$ denotes the embedding network, and $g \in \text{SO}(3)$ acts on a molecule m via the input representation $\phi(g)$ (rotation of atomic coordinates). Because $\rho(g)$ is unknown, we compute an approximate group action $\hat{\rho}(g) \in \text{O}(d)$ by solving an orthogonal Procrustes problem on embeddings from 100 validation samples (obtained from a trained TransIP model). Writing

$$Z = \begin{bmatrix} f(m_1)^\top \\ \vdots \\ f(m_n)^\top \end{bmatrix}, \quad Z_g = \begin{bmatrix} f(\phi(g)(m_1))^\top \\ \vdots \\ f(\phi(g)(m_n))^\top \end{bmatrix},$$

we first pool-whiten the two views (shared mean and standard deviation per channel) and then solve

$$\hat{\rho}(g) = \arg \min_{Q \in \text{O}(d)} \|\tilde{Z}Q - \tilde{Z}_g\|_F^2,$$

which has the closed form $\hat{\rho}(g) = UV^\top$ for the SVD of $\tilde{Z}^\top \tilde{Z}_g = U\Sigma V^\top$.

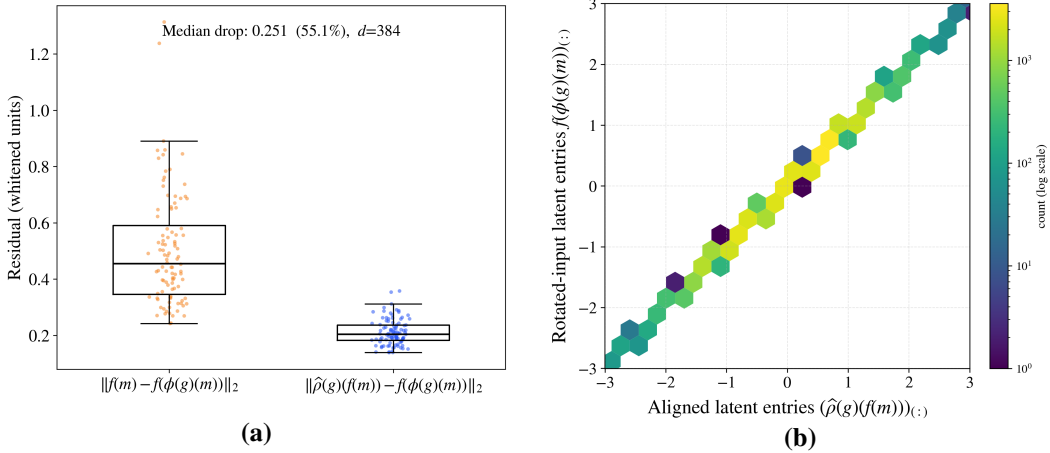


Figure 5: Group action in the embedding space. (a) Per-molecule residuals before alignment, $\|f(m) - f(\phi(g)(m))\|_2$, and after applying a single global orthogonal map $\hat{\rho}(g)$ estimated by an orthogonal Procrustes problem on pool-whitened latents, $\|\hat{\rho}(g)f(m) - f(\phi(g)(m))\|_2$. (b) Entrywise comparison: hexbin density of $(\hat{\rho}(g)f(m))_{(\cdot)}$ vs. $f(\phi(g)(m))_{(\cdot)}$, pooled over molecules’ embeddings.

In Figure 5a, we report per-molecule residuals before alignment, $\|f(m) - f(\phi(g)(m))\|_2$, and after applying the global orthogonal map, $\|\hat{\rho}(g)f(m) - f(\phi(g)(m))\|_2$. A left→right drop in the distribution indicates that a single orthogonal transform explains most of the rotation-induced change in the embedding. In Figure 5b, we compare the channel-level relation by plotting a hexbin density of all pairs

$$(\hat{\rho}(g)f(m))_k, \quad (f(\phi(g)(m)))_k, \quad k = 1, \dots, d, \quad m \in \text{val}.$$

where color encodes the log count of points in each hexagonal bin. A tight diagonal concentration after the single global alignment $\hat{\rho}(g)$ might suggest that the two views are almost identical at entrywise-level and the group action in latent space is *approximately orthogonal* and shared across different molecules.

Takeaways. Figure 5a shows that the magnitude of the rotation-induced discrepancy of different molecules drops after a single orthogonal alignment, and Figure 5b shows that the aligned channels match entrywise, concentrating along the identity. These results indicate that TransIP learns an embedding where input rotations act approximately as a shared orthogonal transformation, even though explicit equivariance was not enforced in the architecture.

B Implementation Details

B.1 Model Architecture

Table 2 provides the complete architectural specifications for TransIP’s model versions.

Table 2: TransIP model configurations. All versions share the same embedding method and activation functions.

Configuration	Small (S)	Medium (M)	Large (L)
Hidden dimension (d)	384	768	1024
Number of layers (L)	8	12	24
Number of heads	6	12	16
Total parameters	14M	85M	302M
<i>Shared configurations:</i>			
Coordinate embedding		MLP	
Activation function		GELU	
Context length		1024	
Projection dropout		0.01	
Attention dropout		0.0	
<i>Transformation network $\mathcal{T}_\tau$:</i>			
Number of layers		2	
Hidden dimension		$2 \times d$	
Activation		GELU	

B.2 Training Hyperparameters

Table 3 provides TransIP’s optimal hyperparameters.

Table 3: Training hyperparameters used for all TransIP experiments.

Hyperparameter	Value
<i>Optimization:</i>	
Optimizer	AdamW
Learning rate	5×10^{-4}
Weight decay	1×10^{-3}
Gradient clip norm	200
<i>Learning rate schedule:</i>	
Scheduler type	Cosine
Warmup fraction	0.01
Min LR factor	0.01
<i>Loss weights:</i>	
Energy (λ_E)	5
Forces (λ_F)	15
Equivariance (λ_{eq})	5 (selected from {1, 5, 10, 100})

B.3 Data Processing and Augmentation

TransIP processes molecular data with the following pipeline:

- **Coordinate centering:** Atomic coordinates are centered by subtracting the center of mass:

$$\mathbf{r}_i \leftarrow \mathbf{r}_i - \frac{1}{|m|} \sum_j \mathbf{r}_j$$
- **Equivariance pairs:** For training with learned equivariance, we create pairs $(m, \phi(g)(m))$ where g is sampled uniformly from $\text{SO}(3)$ per molecule.

B.4 Evaluation Metrics

We evaluate model performance using the following metrics:

Force Mean Absolute Error (MAE):

$$\text{Force MAE} = \frac{1}{3|m|} \sum_{i=1}^N \sum_{\alpha \in \{x,y,z\}} |\mathbf{F}_{i,\alpha} - \mathbf{F}_{i,\alpha}^*| \quad (\text{eV/\AA}) \quad (8)$$

Force Cosine Similarity:

$$\text{Force CosSim} = \frac{1}{|m|} \sum_{i=1}^{|m|} \frac{\mathbf{F}_i \cdot \mathbf{F}_i^*}{\|\mathbf{F}_i\| \|\mathbf{F}_i^*\|} \quad (9)$$

Energy per Atom MAE:

$$\text{Energy/atom MAE} = \frac{1}{|m|} |E - E^*| \quad (\text{eV/atom}) \quad (10)$$

Total Energy MAE:

$$\text{Total Energy MAE} = |E - E^*| \quad (\text{eV}) \quad (11)$$

where $\mathbf{F}$ and E denote predicted forces and energies, $\mathbf{F}^*$ and E^* are ground truth values, and $|m|$ is the total number of atoms. For energies, we use referenced targets following [22].

B.5 Computational Resources

- 5-epoch experiments: 8 NVIDIA 80GB GPUs.
- 80-epoch experiments: 64 NVIDIA 80GB GPUs.

B.6 Validation Splits

For 5-epoch runs, we evaluate on domain-specific validation subsets sampled from the OMol25 validation (Val-Comp) dataset:

- Metal complexes, Electrolytes, biomolecules, reactivity, and neutral organics (including ANI2x, OrbNet-Denali, GEOM, Trans1x, RGD): 20,000 samples from each subset.
- SPICE: 9,630 samples (complete subset).
- Full validation set: 20,000 samples.

We use the full (2M) Val Comp dataset to evaluate TransIP and TransAug in Table 1.

C Additional Results

In this section, we include additional scaling results for TransIP and TransAug.

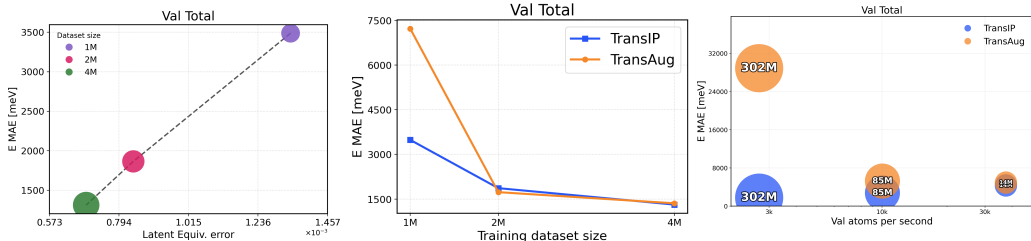


Figure 6: Total energy results. **Left:** Latent equiv. error vs. validation performance. **Middle:** Validation performance across dataset sizes (1M/2M/4M). **Right:** Speed/accuracy trade-off (atoms/s vs. performance).

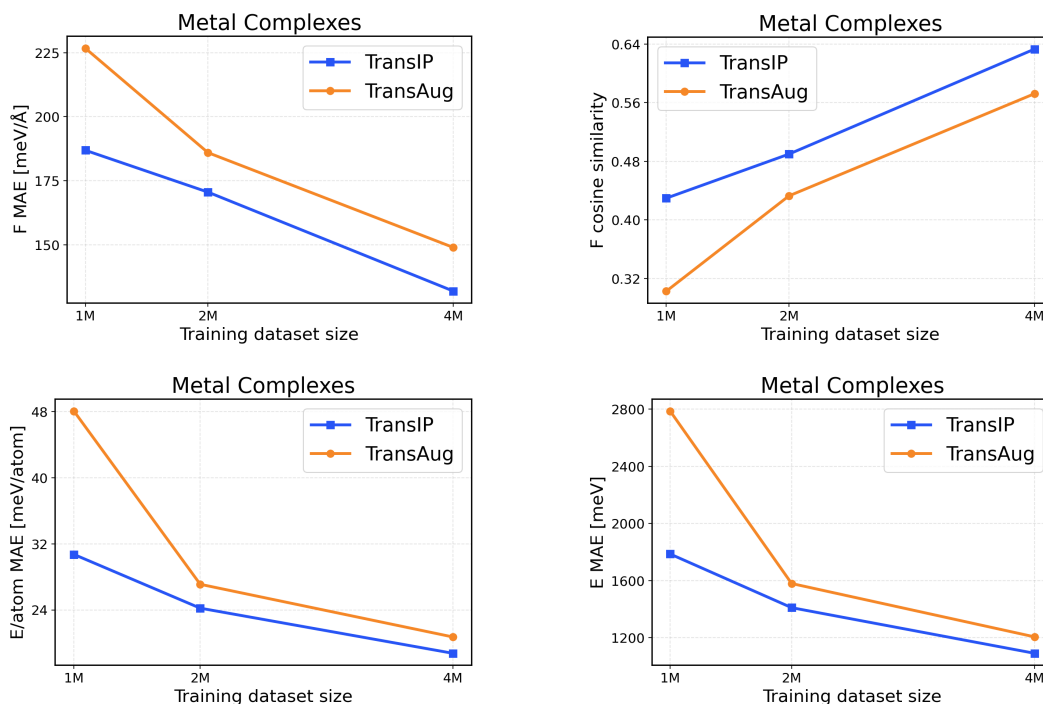


Figure 7: Metal Complexes scaling across training dataset sizes (1M / 2M / 4M). The top row presents force metrics, while the bottom row displays energy metrics.

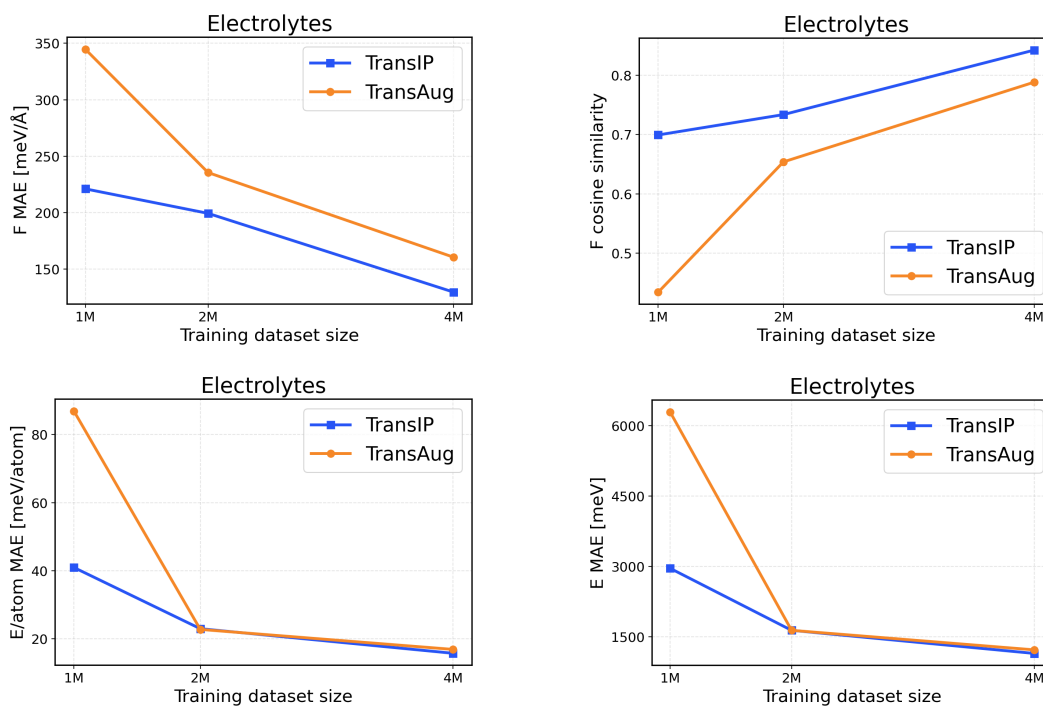


Figure 8: Electrolytes scaling across training dataset sizes (1M / 2M / 4M). The top row presents force metrics, while the bottom row displays energy metrics.

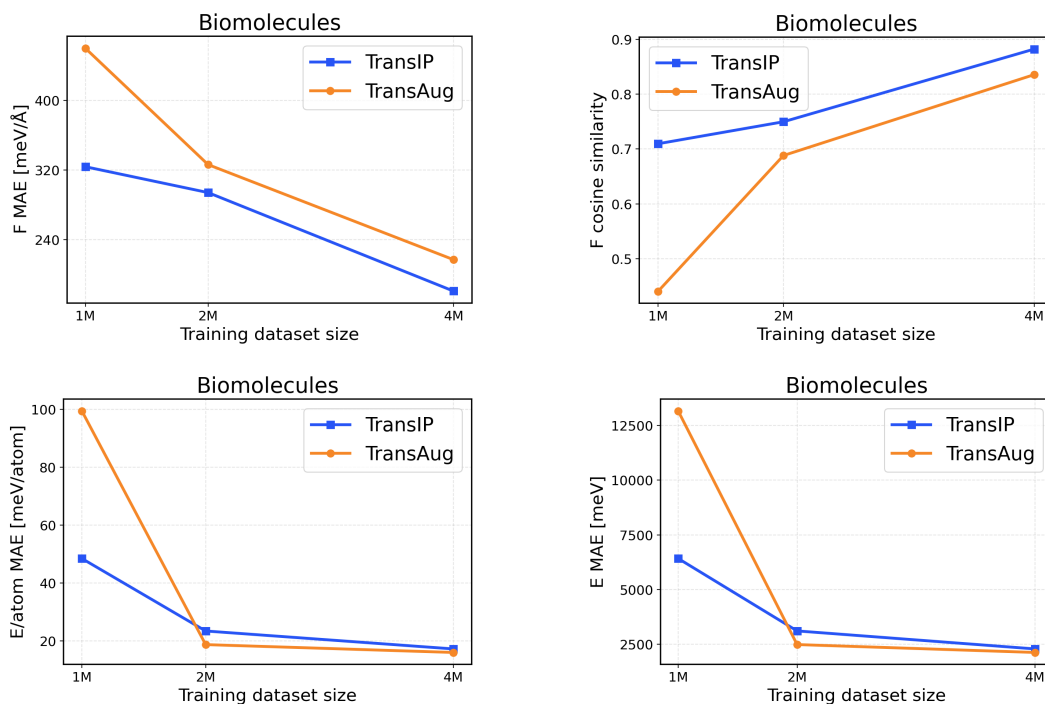


Figure 9: Biomolecules scaling across training dataset sizes (1M / 2M / 4M). The top row presents force metrics, while the bottom row displays energy metrics.

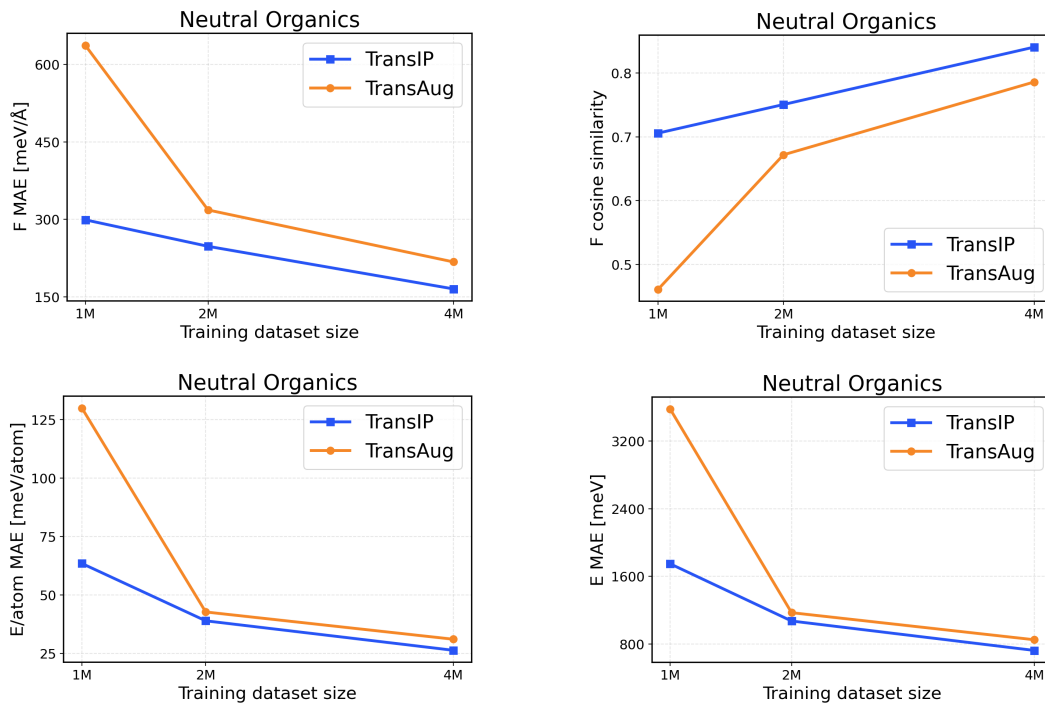


Figure 10: Neutral Organics scaling across training dataset sizes (1M / 2M / 4M). The top row presents force metrics, while the bottom row displays energy metrics.

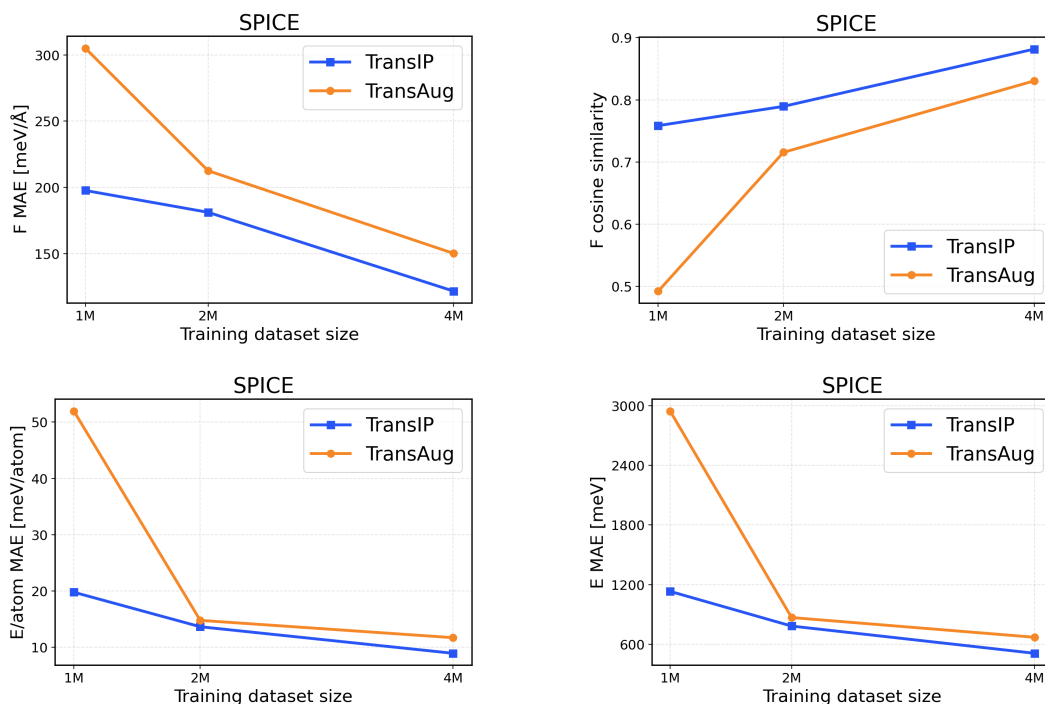


Figure 11: SPICE scaling across training dataset sizes (1M / 2M / 4M). The top row presents force metrics, while the bottom row displays energy metrics.

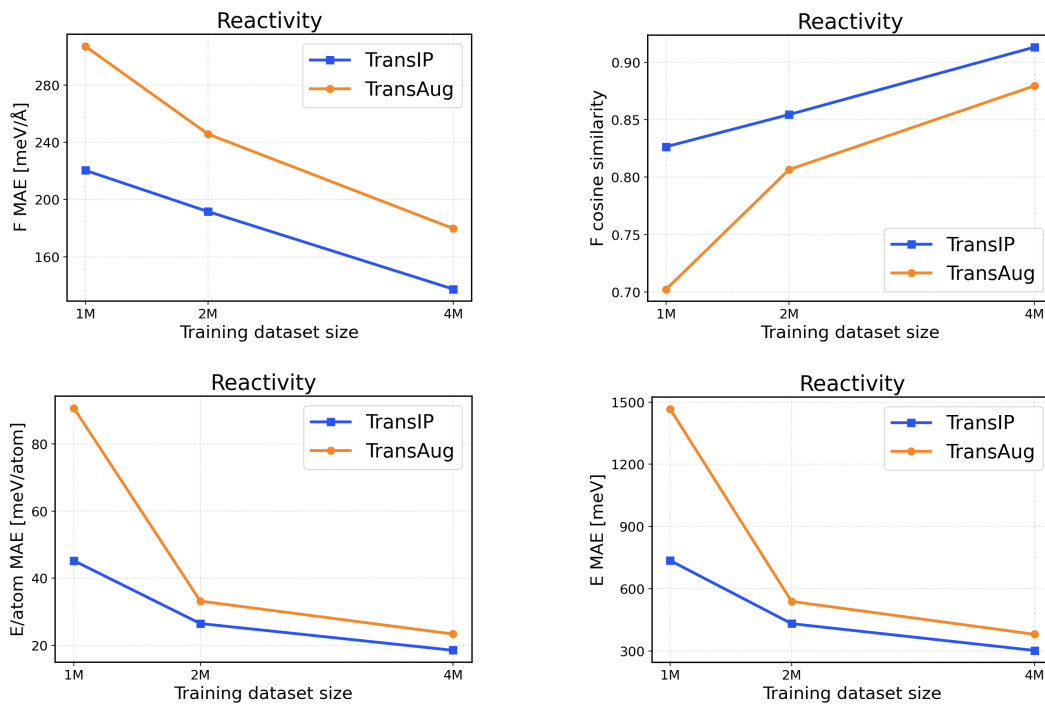


Figure 12: Reactivity scaling across training dataset sizes (1M / 2M / 4M). The top row presents force metrics, while the bottom row displays energy metrics.