
Multiscale and Multi-Timestep Switching of Multiple Machine Learning Force Fields for Artificial Intelligence-Driven Materials Simulations

Ryuya Kanda

Graduate School of Science & Technology
Kobe University, Japan
kanda.ryuya@cs28.cs.kobe-u.ac.jp

Meguru Yamazaki

Fujitsu Limited
Kawasaki, Japan
yamazaki.meguru@fujitsu.com

Yuta Yoshimoto

Fujitsu Limited
Kawasaki, Japan
yoshimoto.yuta@fujitsu.com

Shintaro Izumi

Graduate School of Science & Technology
Kobe University, Japan
shin@cs28.cs.kobe-u.ac.jp

Atsuki Inoue

Graduate School of Science & Technology
Kobe University, Japan
ainoue@godzilla.kobe-u.ac.jp

Hiroshi Kawaguchi

Graduate School of Science & Technology
Kobe University, Japan
kawapy@godzilla.kobe-u.ac.jp

Yasufumi Sakai

Fujitsu Limited
Kawasaki, Japan
sakaiyasufumi@fujitsu.com

Abstract

Molecular dynamics (MD) is a crucial technique in materials science; though its application to large systems and long timescales remains constrained by the prohibitive computational cost of high-accuracy simulations. To address this issue, we propose a multiscale MD approach that switches between two deep potential (DP) models, a type of machine learning force field (MLFF), with different precisions and speeds to optimally balance efficiency and accuracy. A high-precision DP model with a 6 Å cutoff and a faster, lower-precision DP model with a 4 Å cutoff are applied in a 1:3 ratio during integration. Evaluated on a TiO_2 crystal system, the proposed method preserved structural accuracy with Pearson correlation coefficients ≥ 0.995 for radial distribution functions (RDFs), while delivering $1.32\times$ speedup higher than the high-precision baseline. Similarly, in a liquid polyethylene glycol (PEG) system, the method maintained RDF correlations ≥ 0.997 with a $1.27\times$ speedup. Furthermore, when combining the switching scheme with network-size reduction (model compression) and mixed-precision (fp16) inference, RDF correlations were maintained at ≥ 0.99 , while delivering a $3.95\times$ speedup. These results demonstrate that the proposed method can substantially accelerate MD simulations without compromising accuracy, thereby offering a robust approach for artificial intelligence (AI)-assisted material design and large-scale simulations.

1 Introduction

Molecular dynamics (MD) is a key computational method in materials science and chemistry for analyzing atomic-scale phenomena [1]. However, performing high-accuracy force calculations comparable to first-principles methods [2] incurs prohibitive computational costs, rendering simulations of large systems and long timescales impractical. Recently, machine learning force fields (MLFFs), exemplified by deep potential (DeePMD) [3], have emerged as potential approaches that provide near-first-principles accuracy at a significantly reduced computational cost. This has enabled simulations involving tens of thousands to millions of atoms that were previously infeasible [4, 5]. Nevertheless, a trade-off between accuracy and computational efficiency remains: high-accuracy models faithfully reproduce structures and properties but suffer from slower inference speeds, reducing the efficiency of long-timescale simulations [6].

In this context, researchers at MIT have proposed a multiscale framework based on Allegro MLIPs [7–9]. Their method separates short-range (fast) and long-range (slow) interactions and combines them with a multiple time-step (MTS) integrator [10], achieving up to a three-fold acceleration while maintaining accuracy in energy and structural properties. However, this approach requires the co-training of the two models and involves a relatively complex workflow.

To address this issue, we propose a simpler and practical acceleration strategy. Specifically, our method is based on DeePMD [11] and alternates between two independently trained models: a high-accuracy model with a 6 Å cutoff and a fast model with a 4 Å cutoff. Because DeePMD generally provides substantially faster inference than Allegro and is widely adopted in large-scale simulation environments [12], this approach achieves substantial computational savings while maintaining the overall accuracy, offering an efficient and practical acceleration scheme.

The effectiveness of the proposed method is validated on multiple material systems, including both solid and liquid phases. This paper presents a novel acceleration strategy for MLFF-based MD simulations with the potential to contribute to more efficient material design and molecular discovery.

2 Method

2.1 Overview

This study accelerated molecular dynamics (MD) simulations by switching between two Deep Potential (DP) models with different precisions and speeds. Short-range interactions change rapidly and are evaluated more frequently using a fast model, whereas long-range interactions evolve more slowly and are evaluated less frequently using a high-precision model. This approach reduces the computational cost while maintaining the overall precision (Fig. 1).

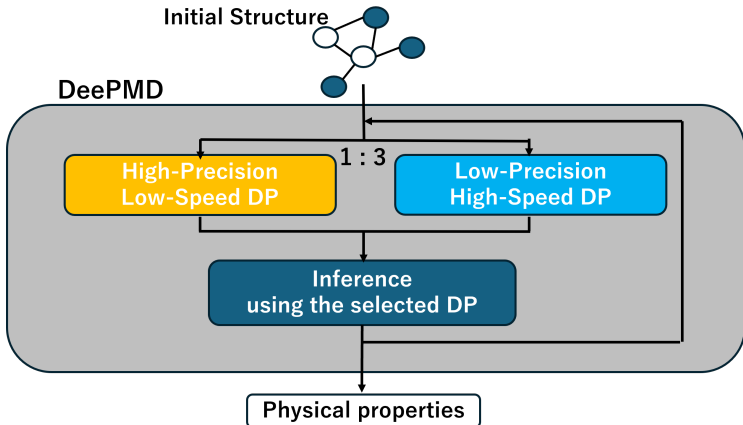


Figure 1: Overview of the proposed model-switching strategy in molecular dynamics simulations. A low-precision, high-speed DP model (cutoff = 4 Å) is applied for three consecutive steps, followed by one step using a high-precision, low-speed DP model (cutoff = 6 Å), yielding a 1:3 switching ratio. At each timestep, inference is performed using the selected DP, and the resulting energy, forces, and virial stress are used to update the system configuration and compute physical properties.

2.2 Models

DeePMD is a machine learning force field that predicts energy, forces, and virial stress by feeding the relative positions of neighboring atoms within a cutoff radius into a neural network [13–15]. A larger cutoff radius allows a more accurate reproduction of long-range interactions but increases the computational cost owing to the increased neighboring atoms. In this study, we selected 6 Å as the cutoff radius for the high-accuracy model, as it is considered a standard value in prior DeePMD studies that balance accuracy and computational costs. By contrast, the fast model employed a cutoff radius of 4 Å. Compared to 6 Å, this substantially reduces both the number of neighboring atoms and thus computational cost, while still preserving the main structural features (e.g., RDF peak positions and intensities). The accuracy degradation remained within an acceptable range. Both models were independently trained on the same dataset.

2.3 Switching strategy

During the time integration, the fast model was applied for several consecutive steps, followed by one step with the high-precision model. Furthermore, we tested multiple switching patterns and adopted a 1:3 ratio because it offered the optimal balance between accuracy and efficiency [16]. This ratio was chosen as a practical compromise between maintaining structural fidelity (in terms of RDF agreement) and improving computational efficiency. It also reflects the physical intuition that short-range interactions vary more rapidly than long-range ones, making it reasonable to employ the fast model more frequently. Switching was performed automatically at every simulation step and did not require manual intervention from the user.

2.4 Implementation

The proposed method was implemented in LAMMPS using the DeePMD-kit [17, 18]. We modified `pair_deepmd.cpp` to load multiple models and performed inferences according to a specified switching ratio. The neighbor list update frequency was shared between the two models to avoid redundant list construction during switching. All the simulations were performed on GPUs.

2.5 Additional optimization

We further evaluated the combination of model compression and mixed-precision (fp16) inference. Model compression was performed by reducing the number of network layers and nodes, whereas fp16 inference was used to accelerate GPU computation. In this configuration, we applied both model compression and mixed-precision in tandem: the network was lightweighted by reducing layers and nodes, and internal calculations (`compute_prec`) were performed in fp16 while outputs (`output_prec`) were kept in fp32, enabling Tensor Core usage and lowering memory consumption. Under these optimizations, MD simulations maintained stable energy and temperature, and no numerical instabilities or divergence were observed. These optimizations can be combined with the switching strategy to further reduce the computation time (details are provided in the Appendix).

3 Experiments

We conducted MD simulations on two representative systems to evaluate the effectiveness of the proposed method. The first system was a solid-phase anatase TiO_2 crystal with 12,000 atoms, which is a stable polymorph of titanium dioxide [19]. The second system was a liquid-phase polyethylene glycol (PEG) 150-mer consisting of 10,530 atoms. The PEG 150-mer corresponds to a polymer chain composed of 150 repeating ethylene oxide units ($-\text{CH}_2-\text{CH}_2-\text{O}-$) [20].

For each system, equilibration simulations were performed in an NPT ensemble at 300 K and 1 bar for 10 ps (10,000 time steps with a time step of 1 fs). All simulations were performed using the DeePMD-kit with LAMMPS on the GPU hardware.

Evaluation metrics: This study evaluated performance from two perspectives. For structural accuracy, the obtained structures were compared using radial distribution functions (RDFs) [21]. In addition to a direct comparison of RDF shapes, the Pearson correlation coefficient between the RDFs obtained by the proposed method and the reference data reported in the literature was calculated. For computational performance, we measured the inference time required for each simulation and reported the relative speedup compared to the high-precision 6 Å DP model. In addition, we evaluated

the time evolution of temperature, density, and potential energy during equilibration and confirmed that the simulations were thermodynamically stable.

4 Results

4.1 TiO_2

Figure 2 shows a comparison of the RDFs of anatase TiO_2 . The RDFs were compared with the reference data reported in the literature. When using the 4 Å model alone, discrepancies in peak positions were observed, with a particularly notable loss of accuracy for the Ti–Ti pair (Pearson correlation coefficient = 0.957). In contrast, the proposed method, which intermittently applies the 6 Å model, successfully corrected these discrepancies and reproduced RDFs nearly identical to those of the 6 Å model alone. The comparison of Pearson correlation coefficients (Table 1) also confirms that the proposed method achieves consistently high agreement, with all values above 0.995. This represents a substantial improvement over the 4 Å model alone.

Furthermore, Table 2 summarizes the computational performance. The proposed method achieved a $1.32\times$ speedup over the high-precision 6 Å model. In addition, combining the switching scheme with model compression and mixed-precision inference yielded a $3.95\times$ speedup, while the RDF correlation coefficients remained above 0.990 (details are provided in the Appendix). These results demonstrate that the proposed method quantitatively preserves accuracy while significantly improving computational efficiency.

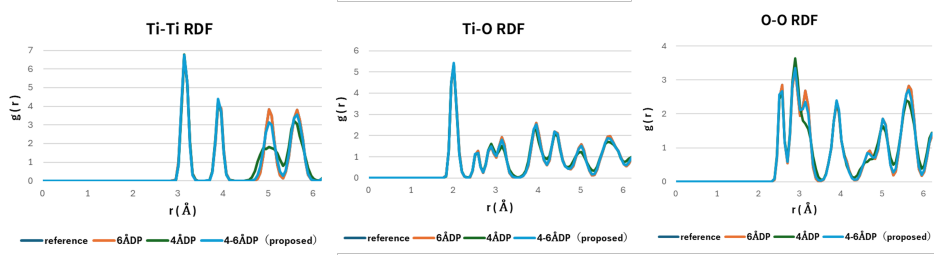


Figure 2: Comparison of RDFs for anatase TiO_2 under different models. A high-speed DP model (cutoff = 4 Å) was used alone, and in the proposed method combined with a high-precision DP model (cutoff = 6 Å) at a 1:3 ratio. The proposed method successfully reproduces RDFs consistent with the 6 Å DP model.

Table 1: Pearson correlation coefficients of RDFs for anatase TiO_2 .

Method	Ti–Ti	Ti–O	O–O
6 Å DP	1.000	1.000	1.000
4 Å DP	0.957	0.989	0.977
4–6 Å DP (proposed)	0.996	0.999	0.995

Table 2: Comparison of MD performance for anatase TiO_2 .

Method	Timesteps/s	Speedup
6 Å DP (baseline)	12.03	1.00
4 Å DP	17.70	1.47
4–6 Å DP (proposed)	15.85	1.32

To evaluate the thermodynamic stability, we analyzed the time evolution of temperature, density, and potential energy during the NPT simulations of the TiO_2 system, as shown in Fig. 3. As illustrated in the figure, all models (4 Å DP, 6 Å DP, and 4–6 Å DP) exhibited overall stable behavior throughout the simulations. The temperature quickly converged to the target value of 300 K after the initial equilibration and then fluctuated within a physically reasonable range. Similarly, the density stabilized after several thousand steps and remained nearly constant, with no significant differences observed among the models. Regarding the potential energy, no noticeable drift was observed over the long timescale, indicating stable energy conservation. On the shorter timescale, the 4–6 Å model showed small-amplitude periodic variations corresponding to the switching cycle with a 3:1 ratio. These oscillations reflect the periodic redistribution of energy but remained bounded, showing no cumulative increase or decrease over time. Overall, these results confirm that the proposed method exhibits thermodynamically stable behavior and does not introduce numerical instabilities or unphysical fluctuations during the switching process.

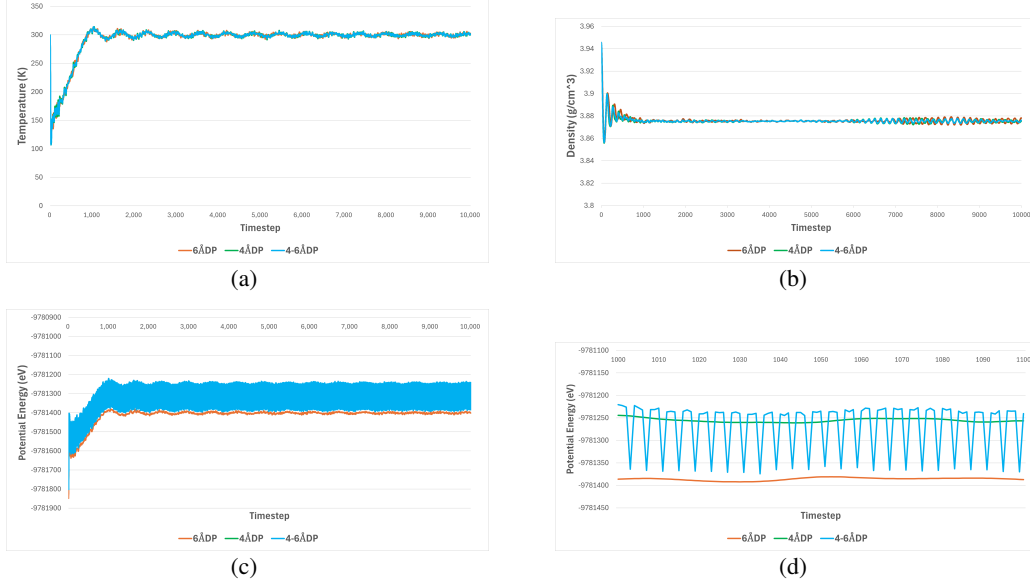


Figure 3: Time evolution of thermodynamic quantities for the anatase TiO_2 system under the NPT ensemble (300 K, 1 bar). All quantities show stable behavior after the initial equilibration, confirming that the proposed 4–6 Å DP switching scheme maintains thermodynamic stability. (a) Temperature, (b) Density, (c) Potential energy (long-term), and (d) Potential energy (short-term).

The proposed approach also maintained linear scaling with the system size, indicating its applicability to large-scale simulations.

4.2 Polyethylene Glycol (PEG)

For the PEG system, the comparison of RDFs showed that the 4 Å model alone already provided reasonably good accuracy overall. However, closer inspection reveals small discrepancies in the peak positions and intensities. The proposed method successfully resolved these discrepancies, reproducing RDFs that achieved near-perfect alignment with the reference 6 Å model. The comparison of Pearson correlation coefficients further confirmed that the proposed method consistently achieved higher agreement than the 4 Å model alone. For computational performance, the proposed method achieved a $1.27\times$ speedup relative to the high-precision 6 Å model. Moreover, during the simulations, thermodynamic parameters such as temperature, density, and potential energy remained stable, confirming that the method preserved a physically consistent behavior. These results demonstrated that even for liquid systems such as PEG, the proposed method effectively balances accuracy and efficiency (details are provided in the Appendix).

4.3 Discussion

The effectiveness of the proposed method depends on factors such as the proportion of high-precision steps and switching frequency, which may require optimization depending on the system and simulation conditions. Future work should explore adaptive switching algorithms that are dynamically adjusted based on the system behavior.

5 Conclusion

This study proposes a multiscale MD method that switches between two DP models with different cutoff radii. The effectiveness of the proposed method was validated on TiO_2 and PEG systems. It achieves a $1.32\times$ acceleration compared to the high-precision model alone, and a $3.95\times$ when combined with model compression and mixed-precision inference. While maintaining RDF accuracy and thermodynamic stability, the method also demonstrated linear scaling, suggesting its applicability to large-scale simulations. These results indicate that the proposed approach is a robust strategy for balancing accuracy and efficiency, contributing to the acceleration of AI-driven material design and molecular discovery.

References

- [1] B. J. Alder and T. E. Wainwright, "Studies in Molecular Dynamics. I. General Method," *Journal of Chemical Physics*, Vol. 31, pp. 459–466, 1959.
- [2] R. Car and M. Parrinello, "Unified Approach for Molecular Dynamics and Density Functional Theory," *Physical Review Letters*, Vol. 55, No. 22, pp. 2471–2474, 1985.
- [3] L. Zhang, J. Han, H. Wang, R. Car, and W. E, "Deep Potential Molecular Dynamics: A Scalable Model with the Accuracy of Quantum Mechanics," *Physical Review Letters*, Vol. 120, p. 143001, 2018.
- [4] W. Jia, H. Wang, M. Chen, D. Lu, L. Lin, R. Car, W. E, and L. Zhang, "Pushing the Limit of Molecular Dynamics with Ab Initio Accuracy to 100 Million Atoms with Machine Learning," *Proceedings of the International Conference for High Performance Computing (SC'20)*, pp. 1–14, 2020.
- [5] L. Klein, J. Smith, and K. R. Müller, "Timewarp: Transferable Acceleration of Molecular Dynamics by Learning Time-Coarsened Dynamics," *Advances in Neural Information Processing Systems (NeurIPS)*, 2023.
- [6] O. T. Unke, S. Chmiela, H. E. Sauceda, M. Gastegger, I. Poltavsky, K. T. Schütt, A. Tkatchenko, K.-R. Müller, and F. Noé, "Machine Learning Force Fields," *Chemical Reviews*, Vol. 121, No. 16, pp. 10142–10186, 2021.
- [7] X. Fu, A. Musaelian, A. Johansson, T. Jaakkola, and B. Kozinsky, "Learning Interatomic Potentials at Multiple Scales," *NeurIPS Workshop on AI for Science*, arXiv:2310.13756, 2023.
- [8] S. Batzner, A. Musaelian, L. Sun, M. Geiger, J. P. Mailoa, M. Kornbluth, N. Molinari, T. E. Smidt, and B. Kozinsky, "E(3)-equivariant graph neural networks for data-efficient and accurate interatomic potentials," *Nature Communications*, Vol. 13, p. 2453, 2022.
- [9] A. Musaelian, S. Batzner, A. Johansson, L. Sun, C. J. Owen, M. Kornbluth, and B. Kozinsky, "Learning local equivariant representations for large-scale atomistic dynamics," *Nature Communications*, Vol. 14, p. 579, 2023.
- [10] M. E. Tuckerman, B. J. Berne, and G. J. Martyna, "Reversible multiple time scale molecular dynamics," *Journal of Chemical Physics*, Vol. 97, No. 3, pp. 1990–2001, 1992.
- [11] H. Wang, L. Zhang, J. Han, and W. E, "DeePMD-kit: A deep learning package for many-body potential energy representation and molecular dynamics," *Computer Physics Communications*, Vol. 228, pp. 178–184, 2018.
- [12] Y. Zhang, H. Wang, W. Chen, J. Zeng, L. Zhang, H. Wang, and W. E, "DP-GEN: A concurrent learning platform for the generation of reliable deep learning based potential energy models," *Computer Physics Communications*, Vol. 253, p. 107206, 2020.
- [13] J. Behler and M. Parrinello, "Generalized Neural-Network Representation of High-Dimensional Potential-Energy Surfaces," *Physical Review Letters*, Vol. 98, No. 14, p. 146401, 2007.
- [14] A. P. Bartók, M. C. Payne, R. Kondor, and G. Csányi, "Gaussian Approximation Potentials: The Accuracy of Quantum Mechanics, without the Electrons," *Physical Review Letters*, Vol. 104, No. 13, p. 136403, 2010.
- [15] K. T. Schütt, H. E. Sauceda, P.-J. Kindermans, A. Tkatchenko, and K.-R. Müller, "SchNet: A deep learning architecture for molecules and materials," *Journal of Chemical Physics*, Vol. 148, No. 24, p. 241722, 2018.
- [16] R. Galvelis, A. Varela-Rial, S. Doerr, R. Fino, P. Eastman, T. E. Markland, J. D. Chodera, and G. De Fabritiis, "NNP/MM: Accelerating Molecular Dynamics Simulations with Machine Learning Potentials and Molecular Mechanics," *Journal of Chemical Information and Modeling*, Vol. 63, No. 18, pp. 5701–5708, 2023.

- [17] J. Zeng, D. Zhang, D. Lu, P. Mo, Z. Li, Y. Chen, M. Rynik, L. Huang, Z. Li, S. Shi, Y. Wang, H. Ye, P. Tuo, J. Yang, Y. Ding, Y. Li, D. Tisi, Q. Zeng, H. Bao, and H. Wang, "DeePMD-kit v2: A software package for deep potential models," *The Journal of Chemical Physics*, Vol. 159, No. 5, p. 054801, 2023.
- [18] A. P. Thompson, H. M. Aktulga, R. Berger, D. S. Bolintineanu, W. M. Brown, P. S. Crozier, P. J. in 't Veld, A. Kohlmeyer, S. G. Moore, T. D. Nguyen, R. Shan, M. J. Stevens, J. Tranchida, C. Trott, and S. J. Plimpton, "LAMMPS - a flexible simulation tool for particle-based materials modeling at the atomic, meso, and continuum scales," *Computer Physics Communications*, Vol. 271, p. 108171, 2022.
- [19] M. F. Calegari Andrade and A. Selloni, "Structure of disordered TiO_2 phases from ab initio based deep neural network simulations," *Physical Review Materials*, Vol. 4, No. 11, p. 113803, 2020.
- [20] N. Matsumura, T. Ito, and H. Tanaka, "A Generator of Neural Network Potential for Molecular Dynamics: Constructing Robust and Accurate Potentials with Active Learning for Nanosecond-Scale Simulations," *arXiv preprint arXiv:2411.17191*, 2024.
- [21] L. Zhang, H. Wang, R. Car, and W. E, "Phase Diagram of a Deep Potential Water Model," *Physical Review Letters*, Vol. 126, p. 236001, 2021.

A Appendix

A.1 Supplementary Results for PEG System

Due to space limitations, we present the supplementary results for the liquid polyethylene glycol (PEG, 150-mer) system in this Appendix. Figure 4 compares the RDFs under different models, Table 3 summarizes the Pearson correlation coefficients for each atomic pair, and Table 4 lists the MD performance metrics.

These results confirm that the proposed method maintains high accuracy (Pearson correlation coefficients ≥ 0.997) while achieving approximately a $1.27\times$ speedup for the PEG system.

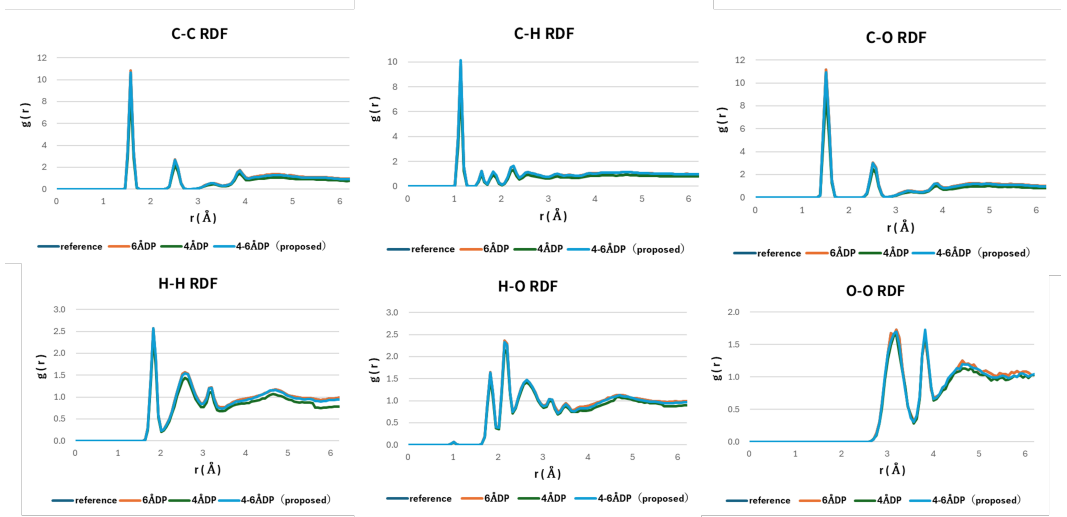


Figure 4: Comparison of RDFs for liquid polyethylene glycol (PEG, 150-mer) under different models. A high-speed DP model (cutoff = 4 Å) was evaluated on its own, and in the proposed method, this model was combined with a high-precision DP model (cutoff = 6 Å) at a 1:3 switching ratio. The proposed method successfully reproduces RDFs consistent with those obtained using the 6 Å DP model.

Table 3: Pearson correlation coefficients of RDFs for PEG 150-mer.

Method	C-C	C-O	O-O	C-H	O-H	H-H
6 Å DP (reference)	1.000	1.000	1.000	1.000	1.000	1.000
4 Å DP	0.998	0.997	0.993	0.996	0.995	0.993
4-6 Å DP (proposed)	0.999	0.999	0.997	0.998	0.997	0.997

Table 4: Comparison of MD performance for PEG 150-mer.

Method	Timesteps/s	Speedup
6 Å DP (baseline)	7.54	1.00
4 Å DP	10.86	1.44
4-6 Å DP (proposed)	9.61	1.27

A.2 Supplement: Details of Model Compression and Mixed-Precision Inference

In addition to the model-switching strategy, we incorporated model compression and mixed-precision inference to further accelerate the simulations.

A.2.1 Model Compression

The computational cost of a Deep Potential model depends heavily on the architecture of its networks: the descriptor network (which encodes local atomic environments) and the fitting network (which predicts energies and forces). In this study, we reduced the number of layers and neurons in both networks to create a lightweight model. This simplification improves the inference speed without reducing the number of neighboring atoms considered, enabling more efficient simulations.

A.2.2 Mixed-Precision Inference

To better utilize GPU resources, we employed mixed-precision computation. Specifically, internal calculations (`compute_prec`) were performed in 16-bit floating point (fp16), while outputs (`output_prec`) were kept in 32-bit (fp32). This configuration enabled the use of Tensor Cores for acceleration and reduced memory usage.

A.2.3 Integrated Effect

By combining model compression and mixed-precision inference with the proposed model-switching strategy, we observed additional performance improvements. In general, these additional optimization techniques further improved the efficiency of MD simulations while maintaining accuracy at a practically acceptable level. Figure 5, Table 5, and Table 6 present the RDF comparisons and performance results achieved with this integrated optimization scheme.

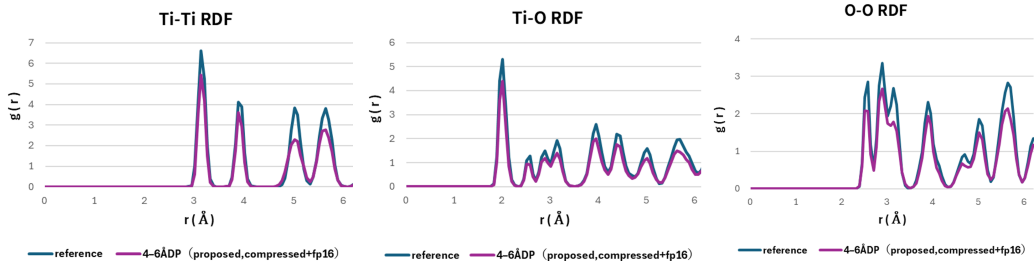


Figure 5: RDF comparison for anatase TiO_2 under the integrated optimization scheme. In this scheme, the proposed 4-6 Å DP model is combined with model compression and mixed-precision (fp16) inference.

Table 5: Pearson correlation coefficients of RDFs with the integrated optimization scheme.

Method	Ti-Ti	Ti-O	O-O
6 Å DP (reference)	1.000	1.000	1.000
4-6 Å DP (proposed, compressed + fp16)	0.991	0.997	0.994

Table 6: MD performance for anatase TiO_2 with the integrated optimization scheme.

Method	Timesteps/s	Speedup
6 Å DP (baseline)	12.03	1.00
4-6 Å DP (proposed, compressed + fp16)	47.5	3.95