

A GENERAL AND EFFICIENT SE(3)-EQUIVARIANT GRAPH FRAMEWORK: ENCODING SYMMETRIES WITH COMPLETE DIFFERENTIAL INVARIANTS AND FRAMES

Anonymous authors

Paper under double-blind review

ABSTRACT

Equivariant graph neural networks (Equiv-GNNs) have demonstrated effectiveness in modeling dynamics of multi-object systems by explicitly encoding symmetries. Among them, scalarization-based methods are widely adopted for their computational efficiency, particularly in comparison to high-steerable models. However, most existing scalarization-based approaches rely on empirical design of invariant functions, lacking rigorous theoretical guarantees. Moreover, these methods typically only consider directional information from object positions, neglecting that from higher-order differential components. To address these limitations, we propose a general and efficient SE(3)-equivariant graph framework with Complete Differential Invariants and Frames (CDIF). Specifically, we show how to construct a set of differential invariants to universally express any invariant functions through network layers. Additionally, we illustrate the complete recovery of directional information from the aforementioned invariants via frames that integrate both positional and differential components. Extensive experiments across diverse domains, including molecular dynamics, formation control, motion capture and particle simulation, validate that our method is simple, scalable, and outperforming state-of-the-art baselines.

1 INTRODUCTION

Multi-object systems are ubiquitous in diverse scientific domains, from multi-body systems (Shabana, 2020) and molecular dynamics (Karplus & McCammon, 2002) to motion planning of robots (Zheng & Li, 2024; Chen, 2024). Modeling the dynamics of these systems remains a challenging problem. Specifically, our goal is to predict future states accurately from the observed data. However, the absence of prior knowledge about system topologies and inter-object interactions leads to an ill-posed problem, i.e., the set of solutions consistent with observed data is typically non-unique and, in most cases, infinite. This fundamental challenge renders conventional approaches (Li et al., 2023; 2025; Prasse & Van Miegheem, 2022) incapable of reconstructing the system or forecasting trajectories precisely.

Equivariant graph neural networks (Equiv-GNNs) (Han et al., 2024) have emerged as a powerful paradigm to tackle this challenge by explicitly encoding symmetries for precise dynamics modeling. Compared with non-equivariant neural networks, Equiv-GNNs narrow the solution space by incorporating symmetries, such as permutation invariance and transformation equivariance. This facilitates models in learning to approximate ground-truth dynamics. Existing Equiv-GNNs can be divided into two categories, i.e., high-steerable models and scalarization models. High-steerable approaches employ group representation theory to lift geometric vectors into high-dimensional spaces, which enables symmetry representation. For example, TFN (Thomas et al., 2018) and SE(3)-Transformer (Fuchs et al., 2020) parameterize such vectors utilizing spherical harmonics, while CEGNN (Ruhe et al., 2023) employs Clifford algebra. Though these methods ensure strong expressiveness, they incur substantial computational costs, especially in large-scale scenarios. By contrast, scalarization-based methods like EGNN (Satorras et al., 2021) and GMN (Huang et al., 2022) construct invariant scalars and update states directly in the original space, which significantly reduce computational costs while achieving state-of-the-art performance across diverse tasks.

However, existing scalarization-based methods rarely explore the relationship between the design of invariant scalars and model expressiveness. Most approaches rely on experience to design invariant scalars, such as relative distances (Satorras et al., 2021) and angles (Gasteiger et al., 2021), which serve as model inputs. Though Villar et al. (2021) demonstrated that any invariant/equivariant function of input vectors can be represented via scalars, they inherently overlook structural constraints imposed by group actions. As a result, their estimated number of scalars does not capture the optimal constraints. Moreover, current methods (Jing et al., 2020) treat differential components of object states merely as input channels to compute invariant scalars, for instance, computing relative distance alongside the magnitude of relative velocities. Actually, the combinations of these components are also critical to object dynamics, e.g., object j ’s position and object i ’s velocity jointly influence the evolution of object i . Finally, some models (Du et al., 2023) only recover directional information from scalars using only object positions, failing to fully leverage all differential components. This leads to cumulative deviations of state updates from ground-truth data during training, ultimately degrading model performance.

In this paper, we propose a general and efficient framework dubbed Equivariant Graph Neural Network with **C**omplete **D**ifferential **I**nvariants and **F**rames (CDIF). Different from previous methods, (1) we demonstrate that the number of invariant scalars (hereafter referred to as invariants) sufficient to universally approximate all invariant/equivariant functions of the multi-object system can be determined. These invariants, which incorporate object differential components, are thus termed *complete differential invariants*. Compared with Villar et al. (2021), we provide an exact and tighter upper bound on the number of invariants, along with a systematic method to construct them. Notably, computing complete differential invariants incurs the same computational complexity as existing scalarization-based methods (Satorras et al., 2021), i.e., $O(m)$. (2) We construct a set of frames on each edge by fully leveraging object states, i.e., embedding both positional and differential components. These frames are theoretically proven to span the entire state space, covering all directions for object state updates, and are thus named *complete differential frames*. We conduct extensive experiments, benchmarking multiple domains including molecular dynamics, formation control, motion capture and physical particle simulation. The results demonstrate that CDIF outperforms state-of-the-art baselines. Moreover, both complete differential invariants and complete differential frames contribute to learning complex system dynamics, and they can be readily adapted to other models.

2 PRELIMINARIES

Multi-object System The multi-object system considered here is generalized to any system composed of interacting objects or nodes, such as multi-agent systems (Amirkhani & Barshooi, 2022) and N-body systems (Kipf et al., 2018). Given an N -object system evolving in the three-dimensional Euclidean space $\mathbb{R}^3$. At time t , the state of each object is formalized as a vector integrating positional and differential dynamics,

$$s(t) = [q(t)^\top, \dot{q}(t)^\top, \ddot{q}(t)^\top, \dots]^\top, \quad (1)$$

where $q(t) \in \mathbb{R}^3$ denotes the position vector, with $\dot{q}(t)$ and $\ddot{q}(t)$ representing its first and second time derivatives, respectively. The interaction topology is characterized by the graph $\mathcal{G} = (\mathcal{V}, \mathcal{E})$, where $\mathcal{V} = \{1, \dots, n\}$ and $\mathcal{E} \subseteq \mathcal{V} \times \mathcal{V}$ define the node and edge sets. In practical applications, direct observation of second-order and higher-order differential states, e.g., accelerations and jerks, is often infeasible. Thus we focus on the scenarios where the object state includes only position and velocity, i.e., $s(t) = [q(t)^\top, \dot{q}(t)^\top]^\top$. Notably, our framework naturally generalizes to incorporate higher-order dynamics, with detailed extensions deferred to the Appendix A.

Functional Dependence Consider a family of smooth $\mathbb{R}$ -valued functions $F = \{f_1, f_2, \dots, f_n\}$. Let $p_0 \in M$ be a point on manifold M with coordinates $(x^1, x^2, \dots, x^m)$. The *rank* of F at p_0 is defined by the rank of the Jacobian matrix $\partial f_i / \partial x^k$. The family F is *regular* if its rank is constant across M , ensuring a uniform structure of dependencies throughout the manifold. Building on the concept of rank, we characterize *functional dependence* among functions in F . The functions $\{f_i\} \subset F$ are *functionally dependent* if, for each $p_0 \in M$ and its neighborhood U , there exists a non-trivial smooth function H , which is not identically zero on any subset of $\mathbb{R}^n$, such that

$$H(f_1(x), f_2(x), \dots, f_n(x)) = 0, \quad \text{for all } p \in U. \quad (2)$$

Group Action Consider a Lie group G with identity element e . To describe the left action of a Lie group G on a manifold M , we introduce the permutation operator $T_g : M \rightarrow M$ that for each $g \in G$ and $p \in M$, we have $T_g(p) \in M$. The *orbit* of the point p under this action is the set

$$\mathcal{O}_p = \{T_g(p) : g \in G\}, \quad (3)$$

which captures all points reachable from p via left actions by elements of G . The action is termed *semi-regular* if every orbit is a smooth submanifold of M with uniform dimension.

Invariance and Equivariance Let T_g, S_g denote the left group actions of G on manifolds M, N , respectively, and consider a smooth map $\phi : M \rightarrow N$. Map ϕ is called *G-equivariant* if it commutes with group actions

$$\phi(T_g(p)) = S_g(\phi(p)), \quad \text{for all } p \in M \text{ and } g \in G. \quad (4)$$

Intuitively, this property ensures that applying a group action before or after mapping ϕ yields the same result. Then, we say ϕ is *G-invariant* if group actions leave its output unchanged.

$$\phi(T_g(p)) = \phi(p), \quad \text{for all } p \in M \text{ and } g \in G. \quad (5)$$

Scalarization-Based Equivariant GNNs Let $\mathcal{I}_i$ denote the set of invariants only relevant to the state s_i of object i , while $\mathcal{I}_{ij}$ represents the invariants dependent on the states of both object i and its neighbor j . Scalarization-based Equiv-GNNs construct messages by

$$m_{ij} = \phi_m(\{h_i, \mathcal{I}_i\}, \{h_j, \mathcal{I}_j\}, \{a_{ij}, \mathcal{I}_{ij}\}), \quad \mathbf{m}_{ij} = \phi_e(m_{ij})\mathbf{f}(s_i, s_j), \quad (6)$$

where h_i, h_j and a_{ij} denote initial node and edge features, and $\mathbf{f}(s_i, s_j)$ outputs a vector. These messages are then aggregated according to the topology, with object features and states updated as follows:

$$h'_i = \phi_h(h_i, \mathcal{I}_i, \sum_{j \in \mathcal{N}_i} m_{ij}), \quad s'_i = s_i + c \sum_{j \in \mathcal{N}_i} \mathbf{m}_{ij}. \quad (7)$$

Here, c serves as a normalization coefficient, ϕ_m, ϕ_e, ϕ_h denote arbitrary mappings, and $\mathcal{N}_i$ specifies the neighbor set of object i .

3 METHODOLOGY

In this section, we introduce the derivation of complete differential invariants and complete differential frames in the three-dimensional Euclidean space under the action of the special Euclidean group $SE(3)$. Specifically, we consider global group actions rather than gauge transformations (He et al., 2021), where the latter apply transformations independently to separated objects. Additional scenarios, such as those in two-dimensional space, are discussed in Appendix E.

3.1 COMPLETE DIFFERENTIAL INVARIANTS

In modeling dynamics of multi-object systems, messages m_{ij} characterize inter-object interactions and thus motivate the construction of edge-wise invariants rather than node-wise ones. These invariants depend on the states of two objects connected by an edge, formalized as $I(s_i, s_j)$. Translational effects can be readily eliminated via centralization, reducing the problem to constructing invariants under the $SO(3)$ group. Thus, our goal is to find a set of $SO(3)$ -invariants $\mathcal{I} = \{I_1(s_i, s_j), \dots, I_m(s_i, s_j)\}$ such that any other $SO(3)$ -invariant $\mathbb{R}$ -valued function is functionally dependent on $\mathcal{I}$. The following proposition specifies the maximal number of non-trivial functionally independent invariants in this setting.

Proposition 3.1. *Under $SO(3)$ group actions on object states s_i and s_j , the maximal number of non-trivial functionally independent invariants is nine. Specifically, any $SO(3)$ -invariant $\mathbb{R}$ -valued function $I \in \mathcal{F}(\mathbb{R}^3)^{SO(3)}$ can be expressed by a function of these invariants, i.e., $I = H(I_1, \dots, I_9)$.*

A detailed proof is provided in Appendix D.1. This implies that any additional invariant can be expressed through a function H of these nine invariants, which can be learned and approximated by

neural networks. Explicit constructions of these invariants can be directly derived via vector inner products and cross products.

$$\begin{aligned} I_{p_i} &= \mathbf{q}_i(t)^\top \mathbf{q}_i(t), & I_{v_i} &= \dot{\mathbf{q}}_i(t)^\top \dot{\mathbf{q}}_i(t), & I_{p_j} &= \mathbf{q}_j(t)^\top \mathbf{q}_j(t), & I_{v_j} &= \dot{\mathbf{q}}_j(t)^\top \dot{\mathbf{q}}_j(t), \\ I_{p_i v_i} &= \mathbf{q}_i(t)^\top \dot{\mathbf{q}}_i(t), & I_{p_i p_j} &= \mathbf{q}_i(t)^\top \mathbf{q}_j(t), & I_{p_j v_j} &= \mathbf{q}_j(t)^\top \dot{\mathbf{q}}_j(t), \\ I_{v_i v_j} &= \dot{\mathbf{q}}_i(t)^\top \dot{\mathbf{q}}_j(t), & I_{v_i p_i p_j} &= \dot{\mathbf{q}}_i(t)^\top (\mathbf{q}_i(t) \times \mathbf{q}_j(t)), \end{aligned} \quad (8)$$

Note that this set $\mathcal{I} = \{I_{p_i}, I_{v_i}, \dots, I_{v_i p_i p_j}\}$ is not unique. Any invariant can be replaced by another, provided that the invariants in the set remain functionally independent.

Lemma 3.2. *Vector cross products are $SO(3)$ -equivariant, and vector inner products are $SO(3)$ -invariant. Specifically,*

- 1) For any $\mathbf{u}, \mathbf{v} \in \mathbb{R}^3$ and $g \in SO(3)$, $(g \cdot \mathbf{u}) \times (g \cdot \mathbf{v}) = g \cdot (\mathbf{u} \times \mathbf{v})$.
- 2) For any $\mathbf{u}, \mathbf{v} \in \mathbb{R}^3$ and $g \in SO(3)$, $(g \cdot \mathbf{u})^\top (g \cdot \mathbf{v}) = \mathbf{u}^\top \mathbf{v}$

Thus, the invariants constructed in (8) are readily verified to be $SO(3)$ -invariant functions, i.e.,

$$\forall g \in SO(3), \quad I(g \cdot \mathbf{s}_i, g \cdot \mathbf{s}_j) = I(\mathbf{s}_i, \mathbf{s}_j). \quad (9)$$

The following corollary extends this property to arbitrary functions of the $SO(3)$ -invariant functions.

Corollary 3.3. *Any function f with $SO(3)$ -invariant functions as inputs is $SO(3)$ -invariant.*

Proof. Let $f = f(I_1, I_2, \dots, I_n)$ where each I_k is an $SO(3)$ -invariant function. For any $g \in SO(3)$, we have:

$$\begin{aligned} f(g \cdot \mathbf{s}_1, \dots, g \cdot \mathbf{s}_m) &= f(I_1(g \cdot \mathbf{s}_1, \dots), \dots, I_n(g \cdot \mathbf{s}_1, \dots)) \\ &= f(I_1(\mathbf{s}_1, \dots), \dots, I_n(\mathbf{s}_1, \dots)) \\ &= f(\mathbf{s}_1, \dots, \mathbf{s}_m) \end{aligned} \quad (10)$$

Thus, f is $SO(3)$ -invariant. $\square$

Building on the Corollary 3.3, we can design $SO(3)$ -invariant layers by taking invariants as inputs. As shown by Leshno et al. (1993), multi-layer neural networks with an adequate number of neurons enable universal approximation of any continuous function. Accordingly, by feeding the invariants $\mathcal{I} = \{I_{p_i}, I_{v_i}, \dots, I_{v_i p_i p_j}\}$ from (8) into the model, we achieve universal approximation of arbitrary $SO(3)$ -invariant functions associated with the state vectors $\mathbf{s}_i$ and $\mathbf{s}_j$.

3.2 COMPLETE DIFFERENTIAL FRAMES

Utilizing the invariants mentioned above, we can learn $SO(3)$ -invariant messages m_{ij} to approximate the actual interactions between objects. However, these messages can only represent the magnitude of interactions and lack directional information. Thus, this section focuses on how to endow messages m_{ij} with direction, thereby using message vectors $\mathbf{m}_{ij}$ to update object states. Specifically, our goal is to derive a suitable expression for $\mathbf{f}(\mathbf{s}_i, \mathbf{s}_j)$ in (6). Du et al. (2022) proposed complete local frames to recover directional information. Nevertheless, they overlooked differential components, rendering these frames unable to directly map to changes in object states. To fully recover directional information, we thus augment these local frames into complete differential frames. For object i and its neighbor j , the positional frame $\mathcal{F}_{ij}^p(t)$ and velocity frame $\mathcal{F}_{ij}^v(t)$ are constructed as

$$\mathcal{F}_{ij}^p(t) = \{e_{ij}^{p_x}(t), e_{ij}^{p_y}(t), e_{ij}^{p_z}(t)\}, \quad \mathcal{F}_{ij}^v(t) = \{e_{ij}^{v_x}(t), e_{ij}^{v_y}(t), e_{ij}^{v_z}(t)\}. \quad (11)$$

The explicit forms are as follows:

$$\begin{aligned} e_{ij}^{p_x}(t) &= \frac{\mathbf{q}_i(t) - \mathbf{q}_j(t)}{\|\mathbf{q}_i(t) - \mathbf{q}_j(t)\|}, & e_{ij}^{p_y}(t) &= \frac{\mathbf{q}_i(t) \times \mathbf{q}_j(t)}{\|\mathbf{q}_i(t) \times \mathbf{q}_j(t)\|}, & e_{ij}^{p_z}(t) &= e_{ij}^{p_x}(t) \times e_{ij}^{p_y}(t), \\ e_{ij}^{v_x}(t) &= \frac{\dot{\mathbf{q}}_i(t) - \dot{\mathbf{q}}_j(t)}{\|\dot{\mathbf{q}}_i(t) - \dot{\mathbf{q}}_j(t)\|}, & e_{ij}^{v_y}(t) &= \frac{\dot{\mathbf{q}}_i(t) \times \dot{\mathbf{q}}_j(t)}{\|\dot{\mathbf{q}}_i(t) \times \dot{\mathbf{q}}_j(t)\|}, & e_{ij}^{v_z}(t) &= e_{ij}^{v_x}(t) \times e_{ij}^{v_y}(t). \end{aligned} \quad (12)$$

It is evident that both the positional frame $\mathcal{F}_{ij}^p(t)$ and velocity frame $\mathcal{F}_{ij}^v(t)$ are orthonormal, where orthogonality arises from the properties of vector cross products. Moreover, by Lemma 3.2, cross products are $SO(3)$ -equivariant, leading to the following proposition with the detailed proof provided in the Appendix D.2.

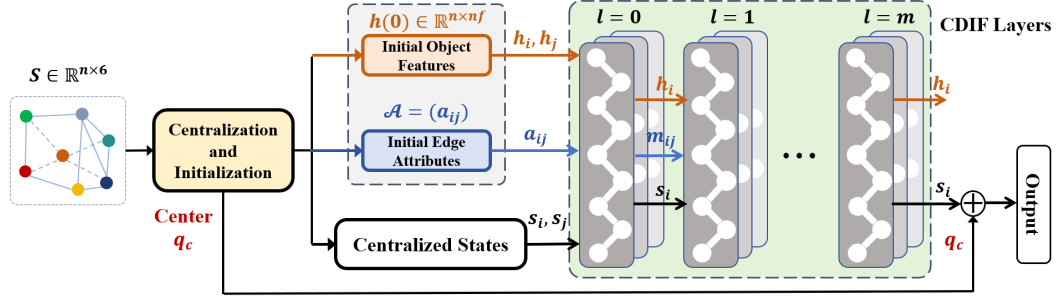


Figure 1: An overview of CDIF. The centralized object states are derived via the Centralization and Initialization module to get rid of translational effects. Then, we feed them along with the initial object features and initial edge attributes into CDIF layers to update the node features, edge messages, and object states.

Proposition 3.4. Both frames $\mathcal{F}_{ij}^p(t)$ and $\mathcal{F}_{ij}^v(t)$ are equivariant under actions of the $SO(3)$ group.

It is worth noting that the frames defined in (12) may suffer from degeneracy. First, degeneracy can occur within $\mathcal{F}_{ij}^p(t)$ and $\mathcal{F}_{ij}^v(t)$ when vectors like $q_i(t)$ and $q_j(t)$ are parallel, which causes the vector cross product to yield zero and reduces the frame’s dimensionality. Additionally, degeneracy may arise between $\mathcal{F}_{ij}^p(t)$ and $\mathcal{F}_{ij}^v(t)$ if their basis vectors, e.g., $e_{ij}^{p_x}(t)$ and $e_{ij}^{v_x}(t)$, are parallel. However, in most practical scenarios, the generic non-collinearity inherent in multi-object systems naturally mitigates such degenerate cases. Furthermore, in applications, a tiny artificial disturbance is typically introduced to proactively avoid the frame degeneracy. With the degeneracy carefully mitigated, the following proposition formalizes that the frames span the state space.

Proposition 3.5. Under non-degenerate conditions, $\mathcal{F}_{ij}^p(t)$ and $\mathcal{F}_{ij}^v(t)$ form a complete basis for the state space. Consequently, any state variation δs can be expressed as

$$\delta s = \lambda_{p_x} e_{ij}^{p_x} + \lambda_{p_y} e_{ij}^{p_y} + \lambda_{p_z} e_{ij}^{p_z} + \lambda_{v_x} e_{ij}^{v_x} + \lambda_{v_y} e_{ij}^{v_y} + \lambda_{v_z} e_{ij}^{v_z}, \quad (13)$$

where $\{\lambda_{p_k}, \lambda_{v_k}\}$ are learnable scalars.

Proof. Under non-degenerate conditions, the six basis vectors $\{e_{ij}^{p_x}, e_{ij}^{p_y}, e_{ij}^{p_z}, e_{ij}^{v_x}, e_{ij}^{v_y}, e_{ij}^{v_z}\}$ from $\mathcal{F}_{ij}^p(t)$ and $\mathcal{F}_{ij}^v(t)$ are linearly independent. The state space here, incorporating position and velocity components, is a six-dimensional vector space. By the fundamental theorem of linear algebra, a set of n linearly independent vectors in an n -dimensional space necessarily spans the space. Thus, $\mathcal{F}_{ij}^p(t)$ and $\mathcal{F}_{ij}^v(t)$ together span the entire state space, and any vector within this space, including the state variation δs , can be expressed as a linear combination of their basis vectors. $\square$

In summary, under the $SO(3)$ group action in 3D Euclidean space, we construct complete differential invariants to encode all symmetries in system dynamics. Models taking these invariants as inputs can learn to approximate any relevant invariant functions. Additionally, complete differential frames preserve all directional information in the state space. In the following sections, we demonstrate that the two modules jointly act to guarantee the expressiveness and generalization capability of the Equiv-GNNs.

4 CDIF FRAMEWORK

In this section, we introduce CDIF, a general and efficient equivariant graph framework. The overall architecture of CDIF is illustrated in Figure 1. To enforce translation invariance in object states, we first perform centralization that allows us to focus merely on rotational symmetry, i.e., $SO(3)$ -invariance/equivariance. Since velocities are inherently translation-invariant, centralization is applied only to object positions. Then the centralized states and initial features are fed into CDIF layers to

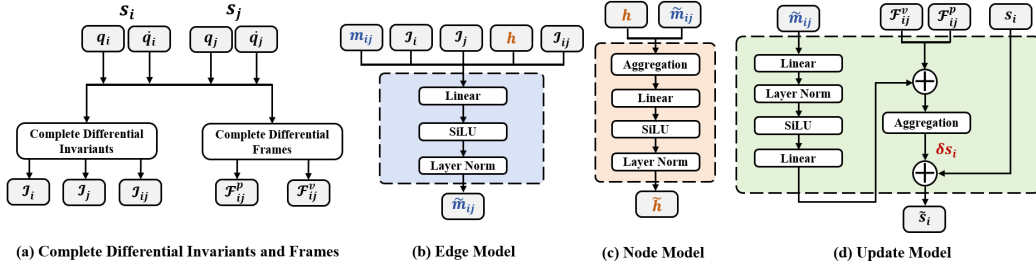


Figure 2: The detailed implementations of the CDIF Layer. The CDIF layer primarily consists of four modules. First, we construct complete differential invariants and frames, which then serve as inputs to the edge model for deriving messages. By aggregating these messages, the node model and update model subsequently update node features and object states, respectively. Variables marked with a tilde, such as $\tilde{h}$, denote their updated counterparts.

learn to predict system future states. Within each CDIF layer illustrated in Figure 2, we first construct complete differential invariants and frames using centralized states, as detailed in (8) and (12). The CDIF layer leverages the message-passing mechanism (Gilmer et al., 2020) to model inter-object interactions. The edge model and node model follow the design of conventional GNNs (Kipf et al., 2018), but with a key distinction that the edge model’s input incorporates the constructed complete differential invariants. Moreover, we replace the activation function with SiLU and introduce layer normalization to further enhance training stability. Next, with the constructed differential frames and messages $\tilde{m}_{ij}$ updated by the edge model, the CDIF layer harness update model to learn object state changes. To improve fitting accuracy, the model predicts the state variations δs rather than the direct future state $\tilde{s}$. Finally, since we initially centralized object positions, we reintroduce the centroid q_c at the output of CDIF to ensure translation equivariance. Details of the CDIF implementation are provided in the Appendix C. By Corollary 3.3 and Proposition 3.4, we present the following proposition regarding the CDIF layer, with its proof detailed in the Appendix D.3.

Proposition 4.1. *The CDIF layer is equivariant under actions of the $SO(3)$ group. Specifically, for any $g \in SO(3)$, the following holds:*

$$g \cdot \text{CDIFL}(s(l)) = \text{CDIFL}(g \cdot s(l)). \quad (14)$$

Compared with other scalarization-based methods (Satorras et al., 2021), CDIF incurs additional computational overhead in constructing complete differential invariants and frames. Let us consider a system topology with m edges. Note that the number of these invariants and frames depends merely on the dimension of space and the orbit regardless of the system topology. This ensures that the per-edge computational cost for the construction remains constant. Since the invariants and frames are constructed along each edge, the associated computational complexity scales as $O(m)$. In message-passing GNNs (Kipf et al., 2018), the computation of edge messages itself already exhibits a complexity of $O(m)$. Thus, the construction of complete differential invariants and frames does not introduce an increase in computational complexity.

5 RELATED WORK

Graph neural networks (GNNs) (Wu et al., 2020; Skarding et al., 2021; Kipf et al., 2018) offer promising solutions for modeling system dynamics. Label permutation invariance in GNNs enables the learning of interactions between objects regardless of their labeling order. To enhance model expressiveness and robustness, Equiv-GNNs (Han et al., 2022; 2024) further incorporate geometric symmetries into GNN frameworks. Recent approaches (Gilmer et al., 2017; Schütt et al., 2017; Finzi et al., 2020; Liu et al., 2022; Klicpera et al., 2020) have made substantial efforts to encode these symmetries, ensuring outputs remain invariant under group actions, and thus termed as invariant graph neural networks. For instance, SchNet (Schütt et al., 2018) constructs convolutional filters using relative distances, while GemNet (Gasteiger et al., 2021) encodes relative distances, angles, and dihedral angles to capture structural invariance. Equivariant graph neural networks go further by

explicitly designing equivariant update processes, making them widely applicable in scenarios such as modeling system dynamics (Xu et al., 2023) and force field prediction (Xu et al., 2024). Existing Equiv-GNNs primarily fall into two categories, i.e., scalarization-based models and high-degree steerable models (Han et al., 2024). High-degree steerable models (Anderson et al., 2019; Fuchs et al., 2020; Batatia et al., 2022; Yu et al., 2023) treat geometric vectors as steerable tensors, leveraging high-dimensional tensor operations to ensure equivariance. Among these methods, TFN (Thomas et al., 2018) and NequIP (Batzner et al., 2022) employs spherical harmonics, while LieTransformer (Hutchinson et al., 2021) combines Lie convolution with transformers and extends the framework to Lie groups. CEGNN (Ruhe et al., 2023) utilizes Clifford algebra to handle geometric symmetries. While these methods achieve strong expressiveness, their intricate tensor manipulations lead to high computational complexity, limiting scalability to large-scale scenarios (Zitnick et al., 2022; Musaelian et al., 2023). In contrast, scalarization-based models (Wang et al., 2023; Battiloro et al., 2024; Köhler et al., 2019; Schütt et al., 2021; Jing et al., 2020; Zhang et al., 2024) map vectors to invariant scalars directly in the original space and introduce equivariance during object state updating or message construction. For instance, EGNN (Satorras et al., 2021) and GMN (Huang et al., 2022) use relative states to derive invariant scalars and update object states along radial directions. Other methods (Du et al., 2022; Kofinas et al., 2021) construct local frames to represent directional information, with Du et al. (2023) providing guidelines for efficient frame-based model design. Some approaches additionally consider object orientations when encoding symmetries, for example, PONITA (Bekkers et al., 2024) achieves equivariance under group actions on both positions and orientations via group convolutions. Nevertheless, these methods lack theoretical guarantees in designing invariant scalars, especially when considering high-order differential components of the system.

In modeling system dynamics, there are also many recent researches devoted to integrating Equiv-GNNs with ODE-based approaches (Poli et al., 2019; Bishnoi et al., 2022; Jin et al., 2022) or trajectory prediction methods (Xu et al., 2023; 2024) to capture temporal correlations. For instance, SEGNO (Liu et al., 2024) improves prediction accuracy by incorporating physical bias, while NCGNN (Guo et al., 2023) leverages Newton-Cotes formulas to generate high-fidelity trajectories. However, these approaches do not focus on studying Equiv-GNNs themselves. Instead, they directly employ Equiv-GNNs like EGNN (Satorras et al., 2021) as the backbone. Since this line does not align with the core contribution of our work, more discussion about it is left for future work.

6 EXPERIMENTS

To validate the effectiveness of CDIF, we benchmark it on four datasets, MD17 (Chmiela et al., 2017), CMU Motion (CMU, 2003), Formation Consensus Control, and N-body simulation (Satorras et al., 2021). Furthermore, we evaluate CDIF against a series of state-of-the-art equivariant models, including EGNN (Satorras et al., 2021), ClofNet (Du et al., 2022), TFN (Thomas et al., 2018), Radial Field (Köhler et al., 2019), GMN (Huang et al., 2022), SE(3)-Transformer (Fuchs et al., 2020), SEGNN (Brandstetter et al., 2021) and PONITA (Bekkers et al., 2024). Ablation studies demonstrate that complete differential invariants and complete differential frames collectively enhance the model’s expressiveness. Additionally, experiments on the impact of invariant quantity reveal that model’s performance tends to be better as the number of invariants increases, but a excessive number of invariants may even degrade model performance. Detailed results of the ablation studies, experiments on the impact of invariant quantity, and other experiment details are provided in Appendix G.

6.1 MOLECULAR DYNAMIC

Since complete differential invariants already encode all features related to atomic states, we only use the atomic index as inputs to obtain initial node features. In this dataset, the system topology is determined by inter-atomic distances, where atoms with a distance less than 1.6 are treated as neighbors. Our task is to predict the positions of all atoms after 2000 data frames. As shown in Table 1, the CDIF framework achieves the best performance on 6 out of 10 molecular systems and ranks top 2 across all tasks. Note that the backbone of our method consists merely of MLPs. Nevertheless, our method remains competitive with SOTA approaches and even outperforms them on most tasks.

Table 1: Mean Square Error ($\times 10^{-2}$) for position prediction on MD17. The best results are **bolded** and the second best are underlined.

Molecule	EGNN	ClofNet	RF	GMN	TFN	SEGNN	PONITA	CDIF(ours)
Aspirin	6.71	6.29	12.18	6.29	7.19	5.89	2.92	<u>4.25</u> ± 0.04
Azobenzene	5.77	4.74	15.36	4.65	5.88	5.19	<u>3.97</u>	3.77 ± 0.06
Benzene	0.17	<u>0.16</u>	0.33	0.17	0.17	0.21	0.61	0.15 ± 0.01
Ethanol	3.95	3.95	4.66	3.96	3.95	3.87	<u>2.95</u>	2.70 ± 0.18
Malonaldehyde	12.23	12.29	21.29	12.13	12.30	12.31	9.18	<u>11.87</u> ± 0.75
Naphthalene	0.27	0.26	0.90	0.25	0.27	0.26	<u>0.24</u>	0.23 ± 0.00
Paracetamol	16.66	16.18	23.63	16.03	16.88	10.01	<u>6.60</u>	4.61 ± 0.47
Salicylic	0.64	0.63	1.62	0.98	0.68	0.63	<u>0.55</u>	0.53 ± 0.01
Toluene	5.15	4.84	6.20	4.36	4.70	5.31	1.87	<u>1.98</u> ± 0.10
Uracil	0.44	0.46	1.06	0.43	0.50	0.39	0.27	<u>0.34</u> ± 0.01

6.2 CMU MOTION CAPTURE

Unlike the other three datasets that require manual construction of system topology, the CMU Motion Capture dataset (CMU, 2003) provides a predefined topology that forms a human skeleton-like structure. Our objective is to predict the positions of all nodes after 10, 20, and 30 timesteps. As shown in Table 2, CDIF achieves the best performance on 2 out of 3 tasks, particularly in long-range prediction tasks. During experiments, we found that PONITA struggled to converge stably despite considerable efforts to stabilize the training process. Across these three tasks, the resulting prediction errors were 71.25, 80.05, and 89.13. This is likely due to misalignment between the orientations constructed by PONITA and the inherent orientations of human motion. Consequently, we omitted PONITA from the comparisons in the table.

Table 2: Mean Square Error ($\times 10^{-2}$) for position prediction across different motion steps on CMU Motion. The best results are **bolded** and the second best are underlined.

Motion	EGNN	ClofNet	RF	SEGNN	GMN	TFN	CDIF(ours)
10	1.495	1.731	8.296	11.017	1.275	1.177	<u>1.247</u> ± 0.020
20	5.314	5.222	63.267	19.158	<u>4.041</u>	4.681	3.584 ± 0.132
30	9.766	7.636	199.475	26.547	<u>6.523</u>	7.018	5.971 ± 0.216

6.3 FORMATION CONSENSUS CONTROL

Formation consensus control, a fundamental problem in control theory, aims to coordinate a group of agents to achieve unified states through inter-agent interactions. In this study, we revisit the conventional formation consensus (FC) framework by adding a velocity cross product term into the agent control inputs, and further partition the datasets into three levels, "Easy", "Medium", and "Difficult", based on the complexity of the agent interactions. The details for dataset generation are provided in Appendix F, and each level includes tasks with 5 agents (FC(5)) and 10 agents (FC(10)). As shown in Table 3, CDIF achieves the best performance on both FC(5) and FC(10) tasks across three levels. Regarding computational efficiency, CDIF introduces a marginal increase in forward time compared with other scalarization-based baselines like EGNN (Satorras et al., 2021) and ClofNet (Du et al., 2022), which is primarily due to the construction of complete differential invariants and frames. However, it significantly outpaces high-degree steerable models such as TFN (Thomas et al., 2018) and SE(3)-Transformer (Fuchs et al., 2020), particularly in FC(10) tasks. Moreover, a critical observation is that the forward time of CDIF remains remarkably close between FC(5) and FC(10) tasks, demonstrating its potential for applications in large-scale scenarios.

6.4 N-BODY SIMULATION

The n -body simulation Satorras et al. (2021) is widely adopted for evaluating the performance of models in position prediction tasks. In this experiment, we test model performance at 10, 20, 40

Table 3: Mean Square Error and forward time ($\times 10^{-2}$) for position prediction on FC(5) and FC(10) tasks across different levels. The best results are **bolded** and the second best are underlined.

Task	FC(5)				FC(10)			
Level	Easy	Medium	Difficult		Easy	Medium	Difficult	
	MSE	MSE	MSE	t	MSE	MSE	MSE	t
EGNN	0.202	0.458	0.696	0.10	0.178	0.293	4.138	0.11
ClofNet	0.111	0.329	0.496	0.17	0.050	0.180	3.850	0.16
TFN	0.129	<u>0.323</u>	0.514	0.64	0.047	0.163	3.807	1.39
SEGNN	0.134	0.395	0.582	0.48	0.098	0.176	3.742	1.03
Radial Field	0.194	0.460	0.696	0.07	0.187	0.321	5.145	0.07
SE(3)-Tr.	0.138	0.324	0.518	1.38	<u>0.047</u>	0.174	<u>3.689</u>	3.71
GMN	0.161	0.405	0.659	0.14	<u>0.139</u>	0.263	<u>4.376</u>	0.14
PONITA	<u>0.105</u>	0.326	<u>0.475</u>	0.33	0.054	<u>0.138</u>	3.594	0.48
CDIF(ours)	0.094	0.215	0.433	0.32	0.034	0.099	3.571	0.33

timesteps on 5-agent (Charged(5)) and 10-agent (Charged(10)) tasks. To demonstrate that complete differential invariants can be readily transferred to other models, we conduct additional experiments where these invariants are fed only as initial inputs to EGNN Satorras et al. (2021), ClofNet Du et al. (2022), and GMN Huang et al. (2022) (denoted as EGNN+CDIs, ClofNet+CDIs, GMN+CDIs), without modifying their other original settings. The results presented in Table 4 show that CDIF achieves top performance on 3 out of 6 tasks and ranks within the top 2 across all evaluations. Notably, CDIF demonstrates superior performance in longer-term predictions compared with short-range forecasts, owing to the embedded differential components. Additionally, integrating complete differential invariants into baseline models improves their prediction accuracy by an average of 6.35%, verifying the critical role of input representation completeness.

Table 4: Mean Square Error for position prediction over Charged(5) and Charged(10) tasks at different prediction time steps. The best results are **bolded** and the second best are underlined

Task	Charged(5)			Charged(10)		
Timestep	10	20	40	10	20	40
Model	MSE	MSE	MSE	MSE	MSE	MSE
EGNN	0.0859	0.5739	3.0840	0.1129	0.7893	4.6088
ClofNet	0.1307	0.6727	3.6656	0.1491	0.9180	5.0206
TFN	0.2259	1.2246	5.7761	0.2901	1.5932	7.8833
Radial Field	0.1573	0.9562	5.3055	0.2401	1.4727	7.6604
SE(3)-Tr.	0.2739	1.3820	6.6055	0.3394	1.7705	8.5021
GMN	0.1021	0.5809	3.1795	0.1178	0.7968	4.6437
PONITA	0.0459	0.4757	2.9013	0.0702	<u>0.6635</u>	<u>4.2281</u>
EGNN+CDIs	0.0845	0.5363	<u>2.8754</u>	0.1058	<u>0.7457</u>	<u>4.3436</u>
ClofNet+CDIs	0.1063	0.6561	3.2664	0.1250	0.8578	5.7007
GMN+CDIs	0.0902	0.5424	3.0426	0.1056	0.7712	4.4403
CDIF(ours)	<u>0.0788</u>	<u>0.4947</u>	2.7328	<u>0.0908</u>	0.6416	3.8583

7 CONCLUSION

In this paper, we tackle the limitations of existing scalarization-based Equiv-GNNs, which lack theoretical guarantees in designing invariant functions and overlook high-order differential components. By constructing complete differential invariants and complete differential frames, we propose a general and efficient framework called CDIF. Theoretically, we prove that complete differential invariants enable universal approximation of any invariant functions, while complete differential frames can represent any vectors in the state space. Extensive experiments across various domains validate the

effectiveness and scalability of the CDIF. Ablation studies further demonstrate that both the complete differential invariants modules and the complete differential frames modules enhance the model’s performance. For future work, we plan to extend the construction of complete differential invariants and frames beyond the $SE(3)$ group to more general Lie groups. Moreover, a promising direction is to explore the construction of functionally independent invariants with differential components for gauge transformations (He et al., 2021). Additionally, since CDIF currently employs MLPs as the backbone, we will investigate integrating the proposed invariants and frames with advanced architectures like Neural ODEs (Liu et al., 2024) for more precise modeling of system dynamics, as well as flow-based models (Hassan et al., 2024; Yim et al., 2023) for generation tasks.

REPRODUCIBILITY STATEMENT

The code of experiments is provided in the Supplementary Material. The main experimental results can be reproduced by adhering to the guidelines outlined in the README.

REFERENCES

- Abdollah Amirkhani and Amir Hossein Barshooi. Consensus in multi-agent systems: a review. *Artificial Intelligence Review*, 2022.
- Emma Andersdotter, Daniel Persson, and Fredrik Ohlsson. Equivariant manifold neural odes and differential invariants. *arXiv preprint arXiv:2009.01411*, 2024.
- Brandon Anderson, Truong Son Hy, and Risi Kondor. Cormorant: Covariant molecular neural networks. *Advances in neural information processing systems*, 2019.
- Ilyes Batatia, David P Kovacs, Gregor Simm, Christoph Ortner, and Gábor Csányi. Mace: Higher order equivariant message passing neural networks for fast and accurate force fields. *Advances in neural information processing systems*, 2022.
- Claudio Battiloro, Mauricio Tec, George Dasoulas, Michelle Audirac, Francesca Dominici, et al. E (n) equivariant topological neural networks. *arXiv preprint arXiv:2405.15429*, 2024.
- Simon Batzner, Albert Musaelian, Lixin Sun, Mario Geiger, Jonathan P Mailoa, Mordechai Kornbluth, Nicola Molinari, Tess E Smidt, and Boris Kozinsky. E(3)-equivariant graph neural networks for data-efficient and accurate interatomic potentials. *Nature communications*, 2022.
- Erik J Bekkers, Sharvaree P Vadgama, Rob Hesselink, Putri A van der Linden, and David W Romero. Fast, expressive se (n) equivariant networks through weight-sharing in position-orientation space. In *ICLR*, 2024.
- Suresh Bishnoi, Ravinder Bhattoo, Sayan Ranu, and NM Krishnan. Enhancing the inductive biases of graph neural ode for modeling dynamical systems. *arXiv preprint arXiv:2209.10740*, 2022.
- Johannes Brandstetter, Rob Hesselink, Elise van der Pol, Erik J Bekkers, and Max Welling. Geometric and physical quantities improve e (3) equivariant message passing. *arXiv preprint arXiv:2110.02905*, 2021.
- Hanxiao Chen. Graph neural networks with model-based reinforcement learning for multi-agent systems. *arXiv preprint arXiv:2407.09249*, 2024.
- Stefan Chmiela, Alexandre Tkatchenko, Huziel E Sauceda, Igor Poltavsky, Kristof T Schütt, and Klaus-Robert Müller. Machine learning of accurate energy-conserving molecular force fields. *Science advances*, 2017.
- CMU. Carnegie-mellon motion capture database, 2003. URL <http://mocap.cs.cmu.edu>.
- Weitao Du, He Zhang, Yuanqi Du, Qi Meng, Wei Chen, Nanning Zheng, Bin Shao, and Tie-Yan Liu. Se (3) equivariant graph neural networks with complete local frames. In *International Conference on Machine Learning*, 2022.

- Yuanqi Du, Limei Wang, Dieqiao Feng, Guifeng Wang, Shuiwang Ji, Carla P Gomes, Zhi-Ming Ma, et al. A new perspective on building efficient and expressive 3d equivariant graph neural networks. *Advances in neural information processing systems*, 2023.
- Marc Finzi, Samuel Stanton, Pavel Izmailov, and Andrew Gordon Wilson. Generalizing convolutional neural networks for equivariance to lie groups on arbitrary continuous data. In *International Conference on Machine Learning*, 2020.
- Fabian Fuchs, Daniel Worrall, Volker Fischer, and Max Welling. Se (3)-transformers: 3d roto-translation equivariant attention networks. *Advances in neural information processing systems*, 2020.
- Johannes Gasteiger, Florian Becker, and Stephan Günnemann. Gemnet: Universal directional graph neural networks for molecules. *Advances in Neural Information Processing Systems*, 2021.
- Justin Gilmer, Samuel S Schoenholz, Patrick F Riley, Oriol Vinyals, and George E Dahl. Neural message passing for quantum chemistry. In *International conference on machine learning*, 2017.
- Justin Gilmer, Samuel S Schoenholz, Patrick F Riley, Oriol Vinyals, and George E Dahl. Message passing neural networks. In *Machine learning meets quantum physics*. Springer, 2020.
- Lingbing Guo, Weiqing Wang, Zhuo Chen, Ningyu Zhang, Zequn Sun, Yixuan Lai, Qiang Zhang, and Huajun Chen. Newton–cotes graph neural networks: On the time evolution of dynamic systems. *Advances in Neural Information Processing Systems*, 2023.
- Jiaqi Han, Yu Rong, Tingyang Xu, and Wenbing Huang. Geometrically equivariant graph neural networks: A survey. *arXiv preprint arXiv:2202.07230*, 2022.
- Jiaqi Han, Jiacheng Cen, Liming Wu, Zongzhao Li, Xiangzhe Kong, Rui Jiao, Ziyang Yu, Tingyang Xu, Fandi Wu, Zihe Wang, et al. A survey of geometric graph neural networks: Data structures, models and applications. *arXiv preprint arXiv:2403.00485*, 2024.
- Majdi Hassan, Nikhil Shenoy, Jungyoon Lee, Hannes Stärk, Stephan Thaler, and Dominique Beaini. Et-flow: Equivariant flow-matching for molecular conformer generation. *Advances in Neural Information Processing Systems*, 2024.
- Lingshen He, Yiming Dong, Yisen Wang, Dacheng Tao, and Zhouchen Lin. Gauge equivariant transformer. *Advances in Neural Information Processing Systems*, 2021.
- Wenbing Huang, Jiaqi Han, Yu Rong, Tingyang Xu, Fuchun Sun, and Junzhou Huang. Equivariant graph mechanics networks with constraints. *arXiv preprint arXiv:2203.06442*, 2022.
- Michael J Hutchinson, Charline Le Lan, Sheheryar Zaidi, Emilien Dupont, Yee Whye Teh, and Hyunjik Kim. Lietransformer: Equivariant self-attention for lie groups. In *International Conference on Machine Learning*, 2021.
- Ming Jin, Yu Zheng, Yuan-Fang Li, Siheng Chen, Bin Yang, and Shirui Pan. Multivariate time series forecasting with dynamic graph neural odes. *IEEE Transactions on Knowledge and Data Engineering*, 2022.
- Bowen Jing, Stephan Eismann, Patricia Suriana, Raphael JL Townshend, and Ron Dror. Learning from protein structure with geometric vector perceptrons. *arXiv preprint arXiv:2009.01411*, 2020.
- Martin Karplus and J Andrew McCammon. Molecular dynamics simulations of biomolecules. *Nature structural biology*, 2002.
- Thomas Kipf, Ethan Fetaya, Kuan-Chieh Wang, Max Welling, and Richard Zemel. Neural relational inference for interacting systems. In *International conference on machine learning*, 2018.
- Johannes Klicpera, Janek Groß, Stephan Günnemann, et al. Directional message passing for molecular graphs. In *International Conference on Learning Representations*, 2020.
- Miltiadis Kofinas, Naveen Nagaraja, and Efstratios Gavves. Roto-translated local coordinate frames for interacting dynamical systems. *Advances in Neural Information Processing Systems*, 2021.

- Jonas Köhler, Leon Klein, and Frank Noé. Equivariant flows: sampling configurations for multi-body systems with symmetric energies. *arXiv preprint arXiv:1910.00753*, 2019.
- Moshe Leshno, Vladimir Ya Lin, Allan Pinkus, and Shimon Schocken. Multilayer feedforward networks with a nonpolynomial activation function can approximate any function. *Neural networks*, 1993.
- Yushan Li, Jianping He, Cailian Chen, and Xinpeng Guan. Topology inference for network systems: Causality perspective and nonasymptotic performance. *IEEE Transactions on Automatic Control*, 2023.
- Yushan Li, Jianping He, and Dimos V Dimarogonas. Inverse inference on cooperative control of networked dynamical systems. *arXiv preprint arXiv:2504.13701*, 2025.
- Yang Liu, Jiashun Cheng, Haihong Zhao, Tingyang Xu, Peilin Zhao, Fugee Tsung, Jia Li, and Yu Rong. Segno: Generalizing equivariant graph neural networks with physical inductive biases. In *ICLR*, 2024.
- Yi Liu, Limei Wang, Meng Liu, Yuchao Lin, Xuan Zhang, Bora Oztekin, and Shuiwang Ji. Spherical message passing for 3d molecular graphs. In *International Conference on Learning Representations*, 2022.
- Albert Musaelian, Simon Batzner, Anders Johansson, Lixin Sun, Cameron J Owen, Mordechai Kornbluth, and Boris Kozinsky. Learning local equivariant representations for large-scale atomistic dynamics. *Nature Communications*, 2023.
- Peter J Olver. *Equivalence, invariants and symmetry*. Cambridge University Press, 1995.
- Michael Poli, Stefano Massaroli, Junyoung Park, Atsushi Yamashita, Hajime Asama, and Jinkyoo Park. Graph neural ordinary differential equations. *arXiv preprint arXiv:1911.07532*, 2019.
- Bastian Prasse and Piet Van Mieghem. Predicting network dynamics without requiring the knowledge of the interaction graph. *Proceedings of the National Academy of Sciences*, 2022.
- David Ruhe, Johannes Brandstetter, and Patrick Forré. Clifford group equivariant neural networks. *Advances in Neural Information Processing Systems*, 2023.
- Victor Garcia Satorras, Emiel Hoogetboom, and Max Welling. E (n) equivariant graph neural networks. In *International conference on machine learning*, 2021.
- Kristof Schütt, Oliver Unke, and Michael Gastegger. Equivariant message passing for the prediction of tensorial properties and molecular spectra. In *International Conference on Machine Learning*, 2021.
- Kristof T Schütt, Farhad Arbabzadah, Stefan Chmiela, Klaus R Müller, and Alexandre Tkatchenko. Quantum-chemical insights from deep tensor neural networks. *Nature communications*, 2017.
- Kristof T Schütt, Huziel E Sauceda, P-J Kindermans, Alexandre Tkatchenko, and K-R Müller. Schnet—a deep learning architecture for molecules and materials. *The Journal of Chemical Physics*, 2018.
- Ahmed A Shabana. *Dynamics of multibody systems*. Cambridge university press, 2020.
- Joakim Skarding, Bogdan Gabrys, and Katarzyna Musial. Foundations and modeling of dynamic networks using dynamic graph neural networks: A survey. *IEEE Access*, 2021.
- Nathaniel Thomas, Tess Smidt, Steven Kearnes, Lusann Yang, Li Li, Kai Kohlhoff, and Patrick Riley. Tensor field networks: Rotation-and translation-equivariant neural networks for 3d point clouds. *arXiv preprint arXiv:1802.08219*, 2018.
- Soledad Villar, David W Hogg, Kate Storey-Fisher, Weichi Yao, and Ben Blum-Smith. Scalars are universal: Equivariant machine learning, structured like classical physics. *Advances in neural information processing systems*, 2021.

- Zun Wang, Guoqing Liu, Yichi Zhou, Tong Wang, and Bin Shao. Efficiently incorporating quintuple interactions into geometric deep learning force fields. *Advances in Neural Information Processing Systems*, 2023.
- Zonghan Wu, Shirui Pan, Fengwen Chen, Guodong Long, Chengqi Zhang, and S Yu Philip. A comprehensive survey on graph neural networks. *IEEE transactions on neural networks and learning systems*, 2020.
- Chenxin Xu, Robby T Tan, Yuhong Tan, Siheng Chen, Yu Guang Wang, Xinchao Wang, and Yanfeng Wang. Eqmotion: Equivariant multi-agent motion prediction with invariant interaction reasoning. In *Proceedings of the IEEE/CVF conference on computer vision and pattern recognition*, 2023.
- Zhao Xu, Haiyang Yu, Montgomery Bohde, and Shuiwang Ji. Equivariant graph network approximations of high-degree polynomials for force field prediction. *Transactions on machine learning research*, 2024.
- Jason Yim, Brian L Trippe, Valentin De Bortoli, Emile Mathieu, Arnaud Doucet, Regina Barzilay, and Tommi Jaakkola. Se (3) diffusion model with application to protein backbone generation. *arXiv preprint arXiv:2302.02277*, 2023.
- Haiyang Yu, Zhao Xu, Xiaofeng Qian, Xiaoning Qian, and Shuiwang Ji. Efficient and equivariant graph networks for predicting quantum hamiltonian. In *International Conference on Machine Learning*, 2023.
- Yuelin Zhang, Jiacheng Cen, Jiaqi Han, Zhiqiang Zhang, Jun Zhou, and Wenbing Huang. Improving equivariant graph neural networks on large geometric graphs via virtual nodes learning. In *Forty-first International Conference on Machine Learning*, 2024.
- Jiafeng Zheng and Jianzhen Li. Affine formation control of discrete-time second-order multi-agent systems with a fixed communication period. *Journal of Physics: Conference Series*, 2024.
- Larry Zitnick, Abhishek Das, Adeesh Kolluru, Janice Lan, Muhammed Shuaibi, Anuroop Sriram, Zachary Ulissi, and Brandon Wood. Spherical channels for modeling atomic interactions. *Advances in Neural Information Processing Systems*, 2022.

A MATHEMATICAL SUPPLEMENTS

In Section 2, we discuss the functional dependence. Now we further present a lemma that bridges the concepts of rank and functional dependence.

Lemma A.1 (Theorem 1.10 Olver (1995)). *If F is regular of rank k , then in a neighborhood of any $p_0 \in M$, there exist k functionally independent functions $f_1, \dots, f_k \in F$ such that any other function $f \in F$ can be expressed as $f = H(f_1, \dots, f_k)$ for some smooth function H .*

Lemma A.1 shows that regular families of rank k admit a local basis of k independent functions generating the entire family. On an m -dimensional manifold M , locally, at most m functions can be found functionally independent. This implies that for the action of the group G on the manifold M , and a point $p_0 \in M$, there must exist a finite number (not exceeding m) of functionally independent invariants. To formalize the group actions on system dynamics with differential components, we introduce the jet bundle framework. Let $u(t) : \mathbb{R} \rightarrow M$ be a smooth curve on the manifold M , and $u^{(k)}(t)$ be the k -th order derivative of $u(t)$ with respect to t . Define $E = \mathbb{R} \times M$ as the *total space*, where $\mathbb{R}$ is the *fiber* and M is the *base space*. The *graph* of $u(t)$ is $\gamma_u = \{t, u(t)\} \subset E$. To include the derivatives $u^{(1)}(t), \dots, u^{(k)}(t)$ of $u(t)$, we introduce the *jet bundle* $J^{(k)}E = \mathbb{R} \times M^{(k)}$. Here, $M^{(k)}$ contains all points in M along with their derivatives up to order k . Similar to the graph γ_u in the total space E , the *prolonged graph* of the jet bundle $J^{(k)}E$ is

$$\gamma_u^{(k)} = \{t, u(t), \dots, u^{(k)}(t)\} \subset J^{(k)}E. \quad (15)$$

Group actions typically leave the time variable t unchanged. Thus, the left-action of the group G on the total space E is restricted to the base space M . For $\gamma_u = \{t, u(t)\}$, the group action becomes $\{t, T_g u(t)\}$. Extending this, the *prolonged group action* $T_g^{(k)}$ on $\gamma_u^{(k)}$ is

$$T_g^{(k)} \gamma_u^{(k)} = \{t, T_g u(t), (T_g u(t))^{(1)}, \dots, (T_g u(t))^{(k)}\}. \quad (16)$$

For a Lie group G with Lie algebra $\mathfrak{g}$, induced vector fields X_i form a basis of $\mathfrak{g}$, with $T_g p = (\exp X_g)p$. The *prolonged vector field* $X^{(k)}$ describes actions on derivatives.

For a node i and its neighbors $\mathcal{N}_i$, the prolonged graph of the extended jet bundle $J^{(k)}E_\Lambda$ is

$$\gamma_\Lambda^{(k)} = \{t, u_i(t), \dots, u_i^{(k)}(t), u_{i_1}(t), \dots, u_{i_1}^{(k)}(t), \dots, u_{i_m}(t), \dots, u_{i_m}^{(k)}(t)\} \subset J^{(k)}E_\Lambda. \quad (17)$$

where $i_1, \dots, i_m \in \mathcal{N}_i$ and Λ denotes the set of $\{i, \mathcal{N}_i\}$. The prolonged group action $T_g^{(k)}$ on $\gamma_\Lambda^{(k)}$ acts simultaneously on node and neighbor states,

$$\begin{aligned} T_g^{(k)} \gamma_\Lambda^{(k)} = & \{t, T_g u_i(t), (T_g u_i(t))^{(1)}, \dots, (T_g u_i(t))^{(k)}, \\ & T_g u_{i_1}(t), (T_g u_{i_1}(t))^{(1)}, \dots, (T_g u_{i_1}(t))^{(k)}, \\ & T_g u_{i_2}(t), (T_g u_{i_2}(t))^{(1)}, \dots, (T_g u_{i_2}(t))^{(k)}, \\ & \dots, \\ & T_g u_{i_m}(t), (T_g u_{i_m}(t))^{(1)}, \dots, (T_g u_{i_m}(t))^{(k)}\}. \end{aligned} \quad (18)$$

B ROTATIONS IN EUCLIDEAN SPACE

In this section, we discuss the rotations in Euclidean spaces, specifically in the 2D and 3D Euclidean spaces $\mathbb{R}^2$ and $\mathbb{R}^3$. The special orthogonal group $\text{SO}(n)$ is defined as:

$$\text{SO}(n) \triangleq \{R \in \mathbb{R}^{n \times n} \mid R^\top R = I, \det(R) = 1\}, \quad (19)$$

where I denotes the n -dimensional identity matrix. Geometrically, $\text{SO}(n)$ represents the group of orientation-preserving isometries (i.e., rotations) on the Euclidean space $\mathbb{R}^n$. In 2D Euclidean space $\mathbb{R}^2$, rotations are fully characterized by the special orthogonal group $\text{SO}(2)$. Every rotation matrix $R \in \text{SO}(2)$ can be parameterized by a unique angle θ , and is explicitly given by

$$R(\theta) = \begin{bmatrix} \cos \theta & -\sin \theta \\ \sin \theta & \cos \theta \end{bmatrix}. \quad (20)$$

As illustrated in Figure 3, for any vector $\mathbf{p}_i \in \mathbb{R}^2$, the action of $\text{SO}(2)$ induces an orbit with cardinality $s_k = 1$, a circle in the 2D plane. Geometrically, this corresponds to rotating $\mathbf{p}_i$ counterclockwise around the origin by an angle θ . When considering multiple vectors, e.g., $\mathbf{p}_i$ and $\mathbf{p}_j$, the orbit cardinality remains $s_k = 1$, as all vectors are rotated by the same angle θ . Certain functions dependent on $\mathbf{p}_i$ and $\mathbf{p}_j$, such as $\|\mathbf{p}_i\|$, $\|\mathbf{p}_i - \mathbf{p}_j\|$, and relative angle α , remain invariant under rotations. Hence, these functions are so-called $\text{SO}(2)$ -invariant. In 3D Euclidean space $\mathbb{R}^3$, we first

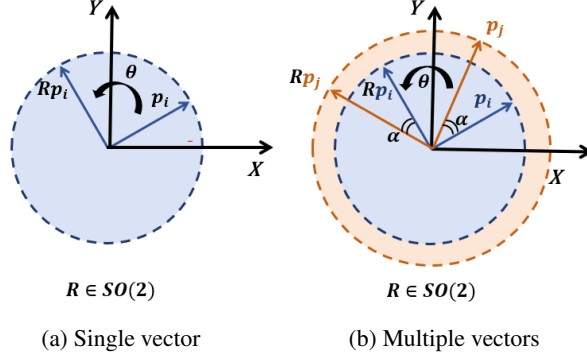


Figure 3: The dashed lines represent the orbits. Both $\mathbf{p}_i$ and $\mathbf{p}_j$ are rotated counterclockwise by the same angle θ , maintaining their relative angle α .

consider a single vector $\mathbf{p} \in \mathbb{R}^3$. As shown in Figure 4, the orbit induced by the action of the $\text{SO}(3)$ group on the single vector $\mathbf{p}$ is a sphere, whose cardinality $s_k = 2$. Geometrically, the actions of the $\text{SO}(3)$ group preserve a vector’s norm while altering its orientation to another direction on the sphere, which can be parameterized by two parameters θ and ϕ . However, unlike the 2D case, the orbit dimension $s_k = 3$ for multiple vectors, which aligns with the dimension of the $\text{SO}(3)$ group itself. This stems from the fundamental principle in Lie group actions; namely the orbit dimension is jointly determined by the geometric constraints of the acted-upon object and the intrinsic degrees of freedom of the group. For a single vector $\mathbf{p} \in \mathbb{R}^3$, its orbit under $\text{SO}(3)$ is a 2-dimensional submanifold, which is constrained by the norm and allows only directional changes parameterized by two angles (θ, ϕ). In contrast, when $\text{SO}(3)$ acts on multiple linearly independent vectors, the group governs their global orientation, which requires 3 independent parameters, corresponding to rotations about the three orthogonal axes (e.g., Euler angles α, β, γ). Since $\text{SO}(3)$ is a 3-dimensional Lie group, the orbit dimension here is fully determined by the group’s dimension, i.e., $s_k = 3$. This aligns with the intuition that a rigid body’s rotation in 3D space demands three independent angles to specify its orientation uniquely.

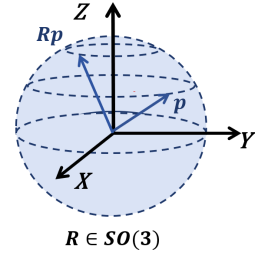


Figure 4: $\text{SO}(3)$ action on a single vector

C IMPLEMENTATION DETAILS OF CDIF

Centralization. Let t_0 be the initial time. For the positional configuration of the system

$$\mathbf{q}(t_0) = [\mathbf{q}_1^\top(t_0), \dots, \mathbf{q}_n^\top(t_0)], \quad (21)$$

the centroid $\mathbf{q}_c$ is calculated as

$$\mathbf{q}_c = \frac{1}{n} \sum_{i=1}^n \mathbf{q}_i(t_0). \quad (22)$$

By redefining the coordinate system origin to this centroid, each object’s position is centralized via

$$\tilde{\mathbf{q}}_i(t_0) = \mathbf{q}_i(t_0) - \mathbf{q}_c. \quad (23)$$

For any time t , the centralized position $\tilde{\mathbf{q}}_i(t)$ remains invariant under translations applied at t_0 . Specifically, for an arbitrary translation vector $\mathbf{T}$ acting on the centralized position $\mathbf{T}(\tilde{\mathbf{q}}_i(t))$, we have

$$\mathbf{T}(\tilde{\mathbf{q}}_i(t)) = (\mathbf{q}_i(t) + \mathbf{T}) - (\mathbf{q}_c + \mathbf{T}) = \mathbf{q}_i(t) - \mathbf{q}_c = \tilde{\mathbf{q}}_i(t), \quad (24)$$

which demonstrates that centralized object positions are invariant under global translations.

Initialization. To encode object-specific attributes unrelated to geometric information, such as object charges, we employ a linear layer to map these attributes into initial object feature embeddings $h(0) \in \mathbb{R}^{n \times n_f}$, where n_f denotes the dimension of hidden features. Meanwhile, the initial edge attributes are denoted as $\mathcal{A} = (a_{ij})$, which incorporate topological attributes such as connectivity and edge weights.

CDIF Layer. Building upon the centralized object states and initial features, we introduce the implementation details of the CDIF layer:

$$\mathcal{F}_{ij}^p(l), \mathcal{F}_{ij}^v(l) = \mathbf{CompFrame}(s_i(l), s_j(l)) \quad (25)$$

$$\mathcal{I}_i(l), \mathcal{I}_j(l), \mathcal{I}_{ij}(l) = \mathbf{CompInvar}(s_i(l), s_j(l)) \quad (26)$$

$$m_{ij}(l) = \phi_m^l(h_i(l), h_j(l), a_{ij}(l), \mathcal{I}_i(l), \mathcal{I}_j(l), \mathcal{I}_{ij}(l)) \quad (27)$$

$$s_i(l+1) = s_i(l) + C \sum_{j \in \mathcal{N}_i} \mathbf{Update}(m_{ij}(l), \mathcal{F}_{ij}^p(l), \mathcal{F}_{ij}^v(l)) \quad (28)$$

$$h_i(l+1) = \phi_h^l(h_i(l), \mathcal{I}_i(l), \sum_{j \in \mathcal{N}_i} m_{ij}(l)) \quad (29)$$

From the input-output perspective, at layer l , the CDIF layer takes as input node feature embeddings $h(l) = [h_1(l)^\top; \dots; h_n(l)^\top] \in \mathbb{R}^{n \times n_f}$, object states $s(l) \in \mathbb{R}^{n \times 6}$, and initial edge attributes $\mathcal{A}$, outputting updated features and states such that

$$(h(l+1), s(l+1)) = \text{CDIFL}(h(l), s(l), \mathcal{A}). \quad (30)$$

Specifically, for each object i and one of its neighbors j , we construct complete differential frames and invariants via the functions $\mathbf{CompFrame}(\cdot)$ and $\mathbf{CompInvar}(\cdot)$ in (25) and (26), whose explicit forms are provide in (8) and (12). Next, inter-object messages $m_{ij}(l)$ are generated by an MLP ϕ_m^l that integrates object features $h_i(l)$, $h_j(l)$, edge attributes $a_{ij}(l)$, and the extracted invariants in (27). Then, the state update model in (28) leverages the complete differential frames to fully recover directional information from invariants. The function $\mathbf{Update}(\cdot)$ employs MLPs $\phi_p^l(m_{ij}), \phi_v^l(m_{ij}) : \mathbb{R}^{n_f} \rightarrow \mathbb{R}^6$ to map messages $m_{ij}(l)$ to scalars $\{\lambda_{p1}^l, \dots, \lambda_{p6}^l\}$ and $\{\lambda_{v1}^l, \dots, \lambda_{v6}^l\}$. These scalars serve as weights for the frame basis vectors $\{e_{ij}^{p_x}, e_{ij}^{p_y}, e_{ij}^{p_z}, e_{ij}^{v_x}, e_{ij}^{v_y}, e_{ij}^{v_z}\}$. By linearly combining this basis with the learnable scalars, the interaction between object i and its neighbor i can be formulated as

$$\begin{aligned} \delta q_{ij} &= \lambda_{p1} e_{ij}^{p_x} + \lambda_{p2} e_{ij}^{p_y} + \lambda_{p3} e_{ij}^{p_z} + \lambda_{p4} e_{ij}^{v_x} + \lambda_{p5} e_{ij}^{v_y} + \lambda_{p6} e_{ij}^{v_z}, \\ \delta \dot{q}_{ij} &= \lambda_{v1} e_{ij}^{v_x} + \lambda_{v2} e_{ij}^{v_y} + \lambda_{v3} e_{ij}^{v_z} + \lambda_{v4} e_{ij}^{p_x} + \lambda_{v5} e_{ij}^{p_y} + \lambda_{v6} e_{ij}^{p_z}. \end{aligned} \quad (31)$$

The variations δq_{ij} and $\delta \dot{q}_{ij}$ are aggregated across all neighbors of object i , scaled by constants C_p and C_v to regulate magnitude, yielding the updated states

$$q_i(l+1) = q_i(l) + C_p \sum_{j \in \mathcal{N}_i} \delta q_{ij}, \quad \dot{q}_i(l+1) = \dot{q}_i(l) + C_v \sum_{j \in \mathcal{N}_i} \delta \dot{q}_{ij}. \quad (32)$$

At last, node feature embeddings are updated in (29) by combining the previous embeddings $h_i(l)$, object-specific invariants $\mathcal{I}_i(l)$, and aggregated messages from the neighbors.

D PROOFS OF THINGS

D.1 PROOF OF PROPOSITION 3.1

Proof. We first introduce an important lemma before proceeding with the proof.

Lemma D.1 (Theorem 2.4 Andersdotter et al. (2024)). *If the prolonged group $G^{(k)}$ of a group G acts semi-regularly on $J^{(k)}E$ with orbit dimension s_k , then there are $\dim J^{(k)}E - s_k$ functionally independent local differential invariants of order k .*

The notations in this lemma is detailed in Section 2 and Appendix A. Lemma D.1 indicates that the maximal number of functionally independent invariants is determined by the dimension of the jet bundle and the prolonged group action orbit. Now let us back to the the proof of the proposition.

When constructing invariants, we only consider the positions and velocities of objects i and j , implying that the order k equals one. The corresponding prolonged graph $\gamma_{ij}^{(1)}$ can be formulated as

$$\gamma_{ij}^{(1)} = \{t, x_i, y_i, z_i, x_j, y_j, z_j, \dot{x}_i, \dot{y}_i, \dot{z}_i, \dot{x}_j, \dot{y}_j, \dot{z}_j\}, \quad (33)$$

which means the dimension of $J^{(1)}E_{ij}$ is 13. Moreover, the orbit dimension of $\text{SO}(3)$ acting on $J^{(1)}E_{ij}$ is 3. Thus, by Lemma D.1, the number of functionally independent invariants is $13 - 3 = 10$. Note that t is a trivial invariant, reducing the number of non-trivial functionally independent differential invariants to $10 - 1 = 9$. Then, by Lemma A.1, any other function $I \in \mathcal{F}(\mathbb{R}^3)^{\text{SO}(3)}$ can be expressed by these invariants. $\square$

D.2 PROOF OF PROPOSITION 3.4

Proof. Let $R \in \text{SO}(3)$ be a rotation matrix, and the transformed positions are denoted as $\tilde{q}_i = Rq_i$ and $\tilde{q}_j = Rq_j$. Denote the transformed positional frame as $\tilde{\mathcal{F}}_{ij}^p(t) = \{\tilde{e}_{ij}^{p_x}, \tilde{e}_{ij}^{p_y}, \tilde{e}_{ij}^{p_z}\}$, by Lemma 3.2, we have

$$\begin{aligned} \tilde{e}_{ij}^{p_x} &= \frac{\tilde{q}_i - \tilde{q}_j}{\|\tilde{q}_i - \tilde{q}_j\|} = \frac{R(q_i - q_j)}{\|q_i - q_j\|} = Re_{ij}^{p_x}, \\ \tilde{e}_{ij}^{p_y} &= \frac{\tilde{q}_i \times \tilde{q}_j}{\|\tilde{q}_i \times \tilde{q}_j\|} = \frac{R(q_i \times q_j)}{\|q_i \times q_j\|} = Re_{ij}^{p_y}, \\ \tilde{e}_{ij}^{p_z} &= \tilde{e}_{ij}^{p_x} \times \tilde{e}_{ij}^{p_y} = (Re_{ij}^{p_x}) \times (Re_{ij}^{p_y}) = R(e_{ij}^{p_x} \times e_{ij}^{p_y}) = Re_{ij}^{p_z}, \end{aligned} \quad (34)$$

Thus, $\mathcal{F}_{ij}^p(t)$ is $\text{SO}(3)$ -equivariant, i.e.,

$$\mathcal{F}_{ij}^p(Rq_i, Rq_j) = \{Re_{ij}^{p_x}, Re_{ij}^{p_y}, Re_{ij}^{p_z}\} = R\mathcal{F}_{ij}^p(q_i, q_j) \quad (35)$$

The proof of the $\text{SO}(3)$ -equivariance of $\mathcal{F}_{ij}^v(t)$ is similar to that of $\mathcal{F}_{ij}^p(t)$, and it only requires replacing the superscript p with v . Therefore both $\mathcal{F}_{ij}^p(t)$ and $\mathcal{F}_{ij}^v(t)$ are equivariant under actions of the $\text{SO}(3)$ group. $\square$

D.3 PROOF OF PROPOSITION 4.1

Proof. Since $h_i(l), h_j(l), a_{ij}(l)$ are scalars independent of object states and the invariants $\mathcal{I}$ are $\text{SO}(3)$ -invariant functions, the construction of messages (27) remains invariant under $\text{SO}(3)$ group actions.

$$\begin{aligned} \forall g \in \text{SO}(3), \quad m_{ij}(g \cdot s(l)) &= \phi_m^l(\mathcal{I}_i(g \cdot s(l)), \mathcal{I}_j(g \cdot s(l)), \mathcal{I}_{ij}(g \cdot s(l))) \\ &= \phi_m^l(\mathcal{I}_i(s(l)), \mathcal{I}_j(s(l)), \mathcal{I}_{ij}(s(l))) = m_{ij}(s(l)). \end{aligned} \quad (36)$$

Moreover, by Proposition 3.4 the frames $\{\mathcal{F}_{ij}^p(t), \mathcal{F}_{ij}^v(t)\}$ are $\text{SO}(3)$ -equivariant functions. Thus, the state update process (28) is equivariant with respect to the $\text{SO}(3)$ group.

$$\begin{aligned} \forall g \in \text{SO}(3), \quad g \cdot s(l+1) &= g \cdot s(l) + C \sum_{j \in \mathcal{N}_i} \text{Update}(m_{ij}(s(l)), g \cdot \mathcal{F}_{ij}^p(s(l)), g \cdot \mathcal{F}_{ij}^v(s(l))) \\ &= g \cdot s(l) + C \sum_{j \in \mathcal{N}_i} \text{Update}(m_{ij}(g \cdot s(l)), \mathcal{F}_{ij}^p(g \cdot s(l)), \mathcal{F}_{ij}^v(g \cdot s(l))) \\ &= g \cdot s(l) + C \sum_{j \in \mathcal{N}_i} \text{Update}(g \cdot s(l)) \end{aligned} \quad (37)$$

Since the node model ϕ_h^l takes as inputs only $\text{SO}(3)$ -invariant functions, by Corollary 3.3, Object feature embeddings (29) is $\text{SO}(3)$ -invariant, i.e.,

$$\forall g \in \text{SO}(3), \quad g \cdot h_i(l+1) = \phi_h^l(g \cdot h_i(l), g \cdot \mathcal{I}_i(l), \sum_{j \in \mathcal{N}_i} g \cdot m_{ij}(l)) = h_i(l+1). \quad (38)$$

Thus, the CDIF layer is $\text{SO}(3)$ -equivariant, i.e.,

$$\forall g \in \text{SO}(3), \quad g \cdot \text{CDIFL}(s(l)) = \text{CDIFL}(g \cdot s(l)). \quad (39)$$

$\square$

E ANALYSIS ON MORE CASES

In this section, we demonstrate the construction of complete differential invariants and frames for additional cases. Moreover, we show that these invariants can be systematically derived leveraging characteristic equations.

Two-dimensional Euclidean space with object positions only First, we construct the complete invariants. The vector field induced by $SO(2)$ on $\mathbb{R}^2$ is as following:

$$X = -y \frac{\partial}{\partial x} + x \frac{\partial}{\partial y}, \quad (40)$$

and the graph of jet bundle JE_{ij} can be expressed as

$$\gamma_{ij} = \{t, x_i(t), y_i(t), x_j(t), y_j(t)\}. \quad (41)$$

By Lemma D.1, the number of functionally independent invariants is $\dim JE_{ij} - s_k = 4$, where $\dim JE_{ij} = 5$ and $s_k = 1$. A trivial invariant is time t . Thus, our goal is to find other three invariants. The vector field X can be written as

$$X = -y_i \frac{\partial}{\partial x_i} + x_i \frac{\partial}{\partial y_i} - y_j \frac{\partial}{\partial x_j} + x_j \frac{\partial}{\partial y_j}, \quad (42)$$

The invariants should satisfy $X(\mathcal{I}) = 0$. Then, the corresponding characteristic equation can be expressed as

$$\frac{dx_i}{-y_i} = \frac{dy_i}{x_i} = \frac{dx_j}{-y_j} = \frac{dy_j}{x_j} \quad (43)$$

From the characteristic equation (43), the object-specific invariants are

$$\begin{aligned} \frac{dx_i}{-y_i} = \frac{dy_i}{x_i} &\implies I_{p_i} = r_i^2 = x_i^2 + y_i^2 \\ \frac{dx_j}{-y_j} = \frac{dy_j}{x_j} &\implies I_{p_j} = r_j^2 = x_j^2 + y_j^2 \end{aligned} \quad (44)$$

and the edge differential invariants are as following:

$$\frac{dx_i}{-y_i} = \frac{dy_j}{x_j}, \frac{dy_i}{x_i} = \frac{dx_j}{-y_j} \implies I_{p_i p_j} = x_i x_j + y_i y_j \quad (45)$$

Therefore, we have found all 3 non-trivial differential invariants as

$$\mathcal{I} = \{I_{p_i}, I_{p_j}, I_{p_i p_j}\}. \quad (46)$$

It is clear to find that the rank of $\mathbf{J}(\mathcal{I})$ is 3, indicating that we have found all functionally independent non-trivial invariants. Next, we construct complete frames. To construct them, we first extend 2D object positions to 3D via zero-padding $p_i(t) = [x_i(t), y_i(t), 0]^\top$. Analogous to the construction in (12), the complete frames are formulated as:

$$e_{ij}^{p_x}(t) = \frac{\mathbf{q}_i(t) - \mathbf{q}_j(t)}{\|\mathbf{q}_i(t) - \mathbf{q}_j(t)\|}, \quad e_{ij}^{p_y}(t) = \frac{\mathbf{q}_i(t) \times \mathbf{q}_j(t)}{\|\mathbf{q}_i(t) \times \mathbf{q}_j(t)\|} \times e_{ij}^{p_x}(t). \quad (47)$$

After constructing these frames, we project them back to the 2D space by discarding the zero-padding component.

Two-dimensional Euclidean space with object positions and velocities First, we construct complete differential invariants. The vector field induced by $SO(2)$ on $\mathbb{R}^2$ is in (40), and the prolonged graph of jet bundle $J^{(1)}E_{ij}$ can be expressed as

$$\gamma_{ij}^{(1)} = \{t, x_i(t), y_i(t), \dot{x}_i(t), \dot{y}_i(t), x_j(t), y_j(t), \dot{x}_j(t), \dot{y}_j(t)\}. \quad (48)$$

According to Lemma D.1, the number of functionally independent local differential invariants is $\dim J^{(1)}E - s_k = 8$, where $\dim J^{(1)}E = 9$ and $s_k = 1$. Obviously, a trivial differential invariant is time t . Thus, our goal is to find other 7 invariants. The prolonged vector field $X^{(1)}$ can be written as

$$X^{(1)} = -y_i \frac{\partial}{\partial x_i} + x_i \frac{\partial}{\partial y_i} - y_j \frac{\partial}{\partial x_j} + x_j \frac{\partial}{\partial y_j} - \dot{y}_i \frac{\partial}{\partial \dot{x}_i} + \dot{x}_i \frac{\partial}{\partial \dot{y}_i} - \dot{y}_j \frac{\partial}{\partial \dot{x}_j} + \dot{x}_j \frac{\partial}{\partial \dot{y}_j}, \quad (49)$$

Then, the corresponding characteristic equation can be expressed as

$$\frac{dx_i}{-y_i} = \frac{dy_i}{x_i} = \frac{d\dot{x}_i}{-\dot{y}_i} = \frac{d\dot{y}_i}{\dot{x}_i} = \frac{dx_j}{-y_j} = \frac{dy_j}{x_j} = \frac{d\dot{x}_j}{-\dot{y}_j} = \frac{d\dot{y}_j}{\dot{x}_j} \quad (50)$$

From the characteristic equation (50), the object-specific differential invariants are found to be

$$\begin{aligned} \frac{dx_i}{-y_i} = \frac{dy_i}{x_i} &\implies I_{p_i} = r_i^2 = x_i^2 + y_i^2 \\ \frac{d\dot{x}_i}{-\dot{y}_i} = \frac{d\dot{y}_i}{\dot{x}_i} &\implies I_{v_i} = v_i^2 = \dot{x}_i^2 + \dot{y}_i^2 \\ \frac{dx_j}{-y_j} = \frac{dy_j}{x_j} &\implies I_{p_j} = r_j^2 = x_j^2 + y_j^2 \\ \frac{d\dot{x}_j}{-\dot{y}_j} = \frac{d\dot{y}_j}{\dot{x}_j} &\implies I_{v_j} = v_j^2 = \dot{x}_j^2 + \dot{y}_j^2 \\ \frac{dy_i}{x_i} = \frac{d\dot{x}_i}{-\dot{y}_i}, \frac{dx_i}{-y_i} = \frac{d\dot{y}_i}{\dot{x}_i} &\implies I_{p_i v_i} = x_i \dot{x}_i + y_i \dot{y}_i, \end{aligned} \quad (51)$$

and the edge differential invariants are as following:

$$\begin{aligned} \frac{dx_i}{-y_i} = \frac{dy_j}{x_j}, \frac{dy_i}{x_i} = \frac{dx_j}{-y_j} &\implies I_{p_i p_j} = x_i x_j + y_i y_j \\ \frac{d\dot{x}_i}{-\dot{y}_i} = \frac{d\dot{y}_j}{\dot{x}_j}, \frac{d\dot{y}_i}{\dot{x}_i} = \frac{d\dot{x}_j}{-\dot{y}_j} &\implies I_{v_i v_j} = \dot{x}_i \dot{x}_j + \dot{y}_i \dot{y}_j. \end{aligned} \quad (52)$$

Up to this point, we have found all 7 differential invariants as

$$\mathcal{I} = \{I_{p_i}, I_{v_i}, I_{p_j}, I_{v_j}, I_{p_i v_i}, I_{p_i p_j}, I_{v_i v_j}\}. \quad (53)$$

It is easy to check the Jacobian matrix $\mathbf{J}(\mathcal{I})$ is full rank. Next, we construct complete differential frames. Similarly, we first extend 2D states to 3D via zero-padding

$$p_i(t) = [x_i(t), y_i(t), 0]^\top, v_i(t) = [\dot{x}_i(t), \dot{y}_i(t), 0]^\top. \quad (54)$$

Analogous to the 3D frame construction in (12), the 2D differential frames are defined as:

$$\begin{aligned} e_{ij}^{p_x}(t) &= \frac{\mathbf{q}_i(t) - \mathbf{q}_j(t)}{\|\mathbf{q}_i(t) - \mathbf{q}_j(t)\|}, & e_{ij}^{p_y}(t) &= \frac{\mathbf{q}_i(t) \times \mathbf{q}_j(t)}{\|\mathbf{q}_i(t) \times \mathbf{q}_j(t)\|} \times e_{ij}^{p_x}(t), \\ e_{ij}^{v_x}(t) &= \frac{\dot{\mathbf{q}}_i(t) - \dot{\mathbf{q}}_j(t)}{\|\dot{\mathbf{q}}_i(t) - \dot{\mathbf{q}}_j(t)\|}, & e_{ij}^{v_z}(t) &= \frac{\dot{\mathbf{q}}_i(t) \times \dot{\mathbf{q}}_j(t)}{\|\dot{\mathbf{q}}_i(t) \times \dot{\mathbf{q}}_j(t)\|} \times e_{ij}^{v_x}(t). \end{aligned} \quad (55)$$

After constructing these frames in 3D, we project them back to the 2D space by discarding the zero-padding component. It is straightforward to verify that the frames defined in (55) are equivariant under $\text{SO}(2)$ group actions.

Three-dimensional Euclidean space with object positions only Compared with the 2D case, constructing complete invariants in 3D space is more challenging. The vector field induced by $\text{SO}(2)$ is single, while the vector fields induced by $\text{SO}(3)$ are multiple, which can be expressed as

$$X_1 = -y \frac{\partial}{\partial z} + z \frac{\partial}{\partial y}, \quad X_2 = -z \frac{\partial}{\partial x} + x \frac{\partial}{\partial z}, \quad X_3 = -x \frac{\partial}{\partial y} + y \frac{\partial}{\partial x}. \quad (56)$$

The graph of jet bundle JE_{ij} is

$$\gamma_{ij} = \{t, x_i(t), y_i(t), z_i(t), x_j(t), y_j(t), z_j(t)\}. \quad (57)$$

By Lemma D.1, the number of functionally independent invariants is $\dim J^{(k)}E - s_k = 4$ with $\dim J^{(k)}E = 7$ and $s_k = 3$. Similarly, t serves as a trivial invariant, leaving us to identify the remaining 3 functionally independent invariants. For the jet bundle JE_{ij} , the vector fields are defined as

$$\begin{aligned} X_1 &= -y_i \frac{\partial}{\partial z_i} + z_i \frac{\partial}{\partial y_i} - y_j \frac{\partial}{\partial z_j} + z_j \frac{\partial}{\partial y_j}, \\ X_2 &= -z_i \frac{\partial}{\partial x_i} + x_i \frac{\partial}{\partial z_i} - z_j \frac{\partial}{\partial x_j} + x_j \frac{\partial}{\partial z_j}, \\ X_3 &= -x_i \frac{\partial}{\partial y_i} + y_i \frac{\partial}{\partial x_i} - x_j \frac{\partial}{\partial y_j} + y_j \frac{\partial}{\partial x_j}. \end{aligned} \quad (58)$$

The invariants $\mathcal{I}$ must satisfy

$$X_i(\mathcal{I}) = 0, \quad i = 1, 2, 3. \quad (59)$$

Leveraging the properties of Lie groups, we can reduce this problem to satisfying just two of these three equations. For group $SO(3)$, the Lie bracket relations hold:

$$[X_1, X_2] = X_3, \quad [X_2, X_3] = X_1, \quad [X_3, X_1] = X_2. \quad (60)$$

With the property (60), we can reduce the problem (59) to satisfying just two of these three equations. Specifically, if the invariants satisfy $X_1(\mathcal{I}) = 0$ and $X_2(\mathcal{I}) = 0$, then the third condition $X_3(\mathcal{I}) = 0$ follows automatically

$$X_3(\mathcal{I}) = [X_1, X_2](\mathcal{I}) = X_1(X_2(\mathcal{I})) - X_2(X_1(\mathcal{I})) = 0. \quad (61)$$

Thus, we first compute $X_1(\mathcal{I}) = 0$. The corresponding characteristic equation is

$$\frac{dz_i}{-y_i} = \frac{dy_i}{z_i} = \frac{dz_j}{-y_j} = \frac{dy_j}{z_j} \quad (62)$$

The same as $SO(2)$ acting on 2D plane, there are 3 differential invariants at this stage, as shown in (63).

$$\begin{aligned} \frac{dz_i}{-y_i} = \frac{dy_i}{z_i} &\implies I_1 = y_i^2 + z_i^2 \\ \frac{dz_j}{-y_j} = \frac{dy_j}{z_j} &\implies I_2 = y_j^2 + z_j^2 \\ \frac{dz_i}{-y_i} = \frac{dy_j}{z_j}, \frac{dy_i}{z_i} = \frac{dz_j}{-y_j} &\implies I_3 = y_i y_j + z_i z_j. \end{aligned} \quad (63)$$

Then, the invariants $\mathcal{I}$ can be treated as functions of $(x_i, x_j, I_1, I_2, I_3)$, i.e., $\mathcal{I} = F(x_i, x_j, I_1, I_2, I_3)$. Substitute the invariants I to $X_2 \mathcal{I} = 0$:

$$\begin{aligned} X_2 \mathcal{I} &= -z_i \frac{\partial F}{\partial x_i} - z_j \frac{\partial F}{\partial x_j} \\ &+ 2x_i z_i \frac{\partial F}{\partial I_1} + 2x_j z_j \frac{\partial F}{\partial I_2} + (x_i z_j + x_j z_i) \frac{\partial F}{\partial I_3} = 0, \end{aligned} \quad (64)$$

whose corresponding characteristic equation can be expressed as

$$\frac{dx_i}{-z_i} = \frac{dx_j}{-z_j} = \frac{dI_1}{2x_i z_i} = \frac{dI_2}{2x_j z_j} = \frac{dI_3}{x_i z_j + x_j z_i} \quad (65)$$

From this, 3 functionally independent invariants are derived as:

$$\begin{aligned} \frac{dx_i}{-z_i} = \frac{dI_1}{2x_i z_i} &\implies I_{p_i} = x_i^2 + y_i^2 + z_i^2 \\ \frac{dx_j}{-z_j} = \frac{dI_2}{2x_j z_j} &\implies I_{p_j} = x_j^2 + y_j^2 + z_j^2 \\ \frac{dx_i}{-z_i} = \frac{dx_j}{-z_j} = \frac{dI_3}{x_i z_j + x_j z_i} &\implies I_{p_i p_j} = x_i x_j + y_i y_j + z_i z_j. \end{aligned} \quad (66)$$

Up to this point, we have found all 3 functionally nontrivial differential invariants as

$$\mathcal{I} = \{I_{p_i}, I_{p_j}, I_{p_i p_j}\}. \quad (67)$$

It is easy to check the Jacobian matrix $\mathbf{J}(\mathcal{I})$ is full rank. Next, we construct complete frames as follows

$$\mathbf{e}_{ij}^{p_x}(t) = \frac{\mathbf{q}_i(t) - \mathbf{q}_j(t)}{\|\mathbf{q}_i(t) - \mathbf{q}_j(t)\|}, \quad \mathbf{e}_{ij}^{p_y}(t) = \frac{\mathbf{q}_i(t) \times \mathbf{q}_j(t)}{\|\mathbf{q}_i(t) \times \mathbf{q}_j(t)\|}, \quad \mathbf{e}_{ij}^{p_z}(t) = \mathbf{e}_{ij}^{p_x}(t) \times \mathbf{e}_{ij}^{p_y}(t). \quad (68)$$

It is straightforward to verify that the frames defined in (68) are equivariant under $SO(3)$ group actions.

Complete differential invariants on more complex cases Complete differential invariants can be applied to more complex scenarios. For instance, when considering two-hop neighbors, the jet bundle can be denoted as $J^{(1)}E_{ijk}$, where $j \in \mathcal{N}_i$ and $k \in \mathcal{N}_j$. The corresponding prolonged graph $\gamma_{ijk}^{(1)}$ is expressed as

$$\gamma_{ijk}^{(1)} = \{t, x_i(t), y_i(t), z_i(t), \dot{x}_i(t), \dot{y}_i(t), \dot{z}_i(t), \\ x_j(t), y_j(t), z_j(t), \dot{x}_j(t), \dot{y}_j(t), \dot{z}_j(t), \\ x_k(t), y_k(t), z_k(t), \dot{x}_k(t), \dot{y}_k(t), \dot{z}_k(t)\}. \quad (69)$$

By Lemma D.1, the number of non-trivial functionally independent invariants is $\dim J^{(1)}E_{ijk} - s_k = 19 - 3 - 1 = 15$. Another scenario arises when higher-order derivatives are incorporated, such as observable object states including second-order derivatives (i.e., acceleration $\mathbf{a}(t)$). In this case, the prolonged graph $\gamma_{ij}^{(2)}$ of the jet bundle $J^{(2)}E_{ij}$ can be formulated as (70) and the number of non-trivial functionally independent invariants is 15.

$$\gamma_{ij}^{(2)} = \{t, x_i(t), y_i(t), z_i(t), x_j(t), y_j(t), z_j(t), \\ \dot{x}_i(t), \dot{y}_i(t), \dot{z}_i(t), \ddot{x}_i(t), \ddot{y}_i(t), \ddot{z}_i(t), \\ \ddot{x}_j(t), \ddot{y}_j(t), \ddot{z}_j(t)\}. \quad (70)$$

It is evident that the number of non-trivial functionally independent invariants can be readily determined by Lemma D.1 in such cases. While this count increases as scenarios become more complex, constructing explicit forms of complete differential invariants remains not difficult but rather tedious, as it entails a systematic, step-by-step derivation of each invariant to ensure functional independence.

F A NOVEL BENCHMARK: FORMATION CONSENSUS CONTROL

In this section, we introduce the implementation of a novel benchmark, formation consensus control, designed to evaluate equivariant graph neural networks across hierarchical interaction complexities. The core task of formation consensus control is coordinating agents to achieve unified states, such as synchronized positions, velocities, through inter-agent interactions. In this system, the agents follow double-integrator dynamics, where their accelerations depend on neighbors' positions and velocities

$$\ddot{\mathbf{v}}_i = - \sum_{j \in \mathcal{N}_i} \omega_{ij} (k_p(\mathbf{q}_i - \mathbf{q}_j) + k_v(\dot{\mathbf{q}}_i - \dot{\mathbf{q}}_j) + k_d(\dot{\mathbf{q}}_j \times \dot{\mathbf{q}}_i)), \quad (71)$$

with stability ensured by $k_p > 0$ and $k_v > 0$. The dataset is divided into three levels based on the complexity of interaction coefficients ω_{ij} . At the easy level, ω_{ij} involves only edge-wise position and velocity dot products, forming low-complexity interactions with linearity:

$$\omega_{ij} = 0.2 \times \|\mathbf{q}_i \cdot \mathbf{q}_j + \dot{\mathbf{q}}_i \cdot \dot{\mathbf{q}}_j\|. \quad (72)$$

The medium level introduces a velocity-position cross-term ($\dot{\mathbf{q}}_i \cdot (\mathbf{q}_i \times \mathbf{q}_j)$), introducing bilinear interactions and geometric dependencies:

$$\omega_{ij} = 0.2 \times \|\mathbf{q}_i \cdot \mathbf{q}_j + \dot{\mathbf{q}}_i \cdot \dot{\mathbf{q}}_j + \dot{\mathbf{q}}_i \cdot (\mathbf{q}_i \times \mathbf{q}_j)\| \quad (73)$$

The difficult level further incorporates higher-order cross interactions between positions and velocities, such as triple products and cross product norms, creating a highly non-linear, multi-modal interaction structure:

$$\omega_{ij} = 0.15 \times \|\mathbf{q}_i \cdot \mathbf{q}_j + \dot{\mathbf{q}}_i \cdot \dot{\mathbf{q}}_j + \dot{\mathbf{q}}_i \cdot (\mathbf{q}_i \times \mathbf{q}_j) \\ + \dot{\mathbf{q}}_j \cdot (\mathbf{q}_i \times \mathbf{q}_j) + \mathbf{q}_i \cdot (\dot{\mathbf{q}}_i \times \dot{\mathbf{q}}_j) \\ + \mathbf{q}_j \cdot (\dot{\mathbf{q}}_i \times \dot{\mathbf{q}}_j) + (\mathbf{q}_i \times \mathbf{q}_j) \cdot (\dot{\mathbf{q}}_i \times \dot{\mathbf{q}}_j)\|. \quad (74)$$

This hierarchy progressively challenges equivariant graph neural networks to model increasingly intricate spatial and velocity interactions. Moreover, the agent dynamics in (71) transcend radial

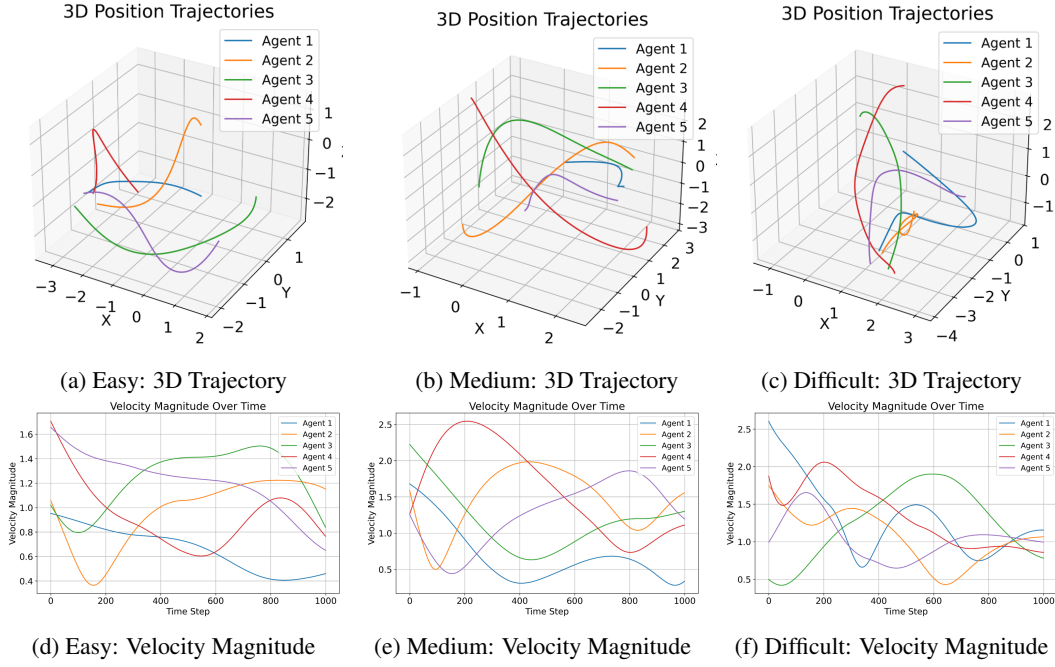


Figure 5: Formation dynamics across different difficulty levels, organized by column: First column (Easy), Second column (Medium), Third column (Difficult). Each column presents a 3D trajectory (top row) and corresponding velocity magnitude profile (bottom row). Trajectories across all levels are stable, with velocities reflecting the increasing complexity of inter-agent interactions encoded in ω_{ij} .

interactions $(q_i - q_j)$ by introducing a velocity cross-term $\dot{q}_j \times \dot{q}_i$, which drives state updates along tangential directions in the state space. This term requires equivariant graph neural networks to handle state update process in entire state space rather than merely relying on a single direction.

Figure 5 visualizes trajectories and velocity magnitudes across levels, illustrating system behaviors under varying interaction complexities. To ensure data validity, we truncate sequences at the first 1000 timesteps, excluding post-convergence data where agents reach unified states. Importantly, all trajectories exhibit stable and convergent behavior, confirming that the benchmark provides a rigorous testbed for evaluating network expressiveness.

G IMPLEMENTATION DETAILS, ABLATION STUDIES, AND ADDITIONAL EXPERIMENTS

G.1 IMPLEMENTATION DETAILS.

For the molecular dynamics experiment, the dataset is partitioned into 3000 training, 2000 validation, and 2000 test samples. In the CMU motion capture experiment, the corresponding splits are 2000 training, 600 validation, and 600 test samples. For the formation consensus control experiment, the dataset is divided into 6000 training samples, 2000 validation samples, and 1000 test samples, with the objective of predicting object positions after 600 timesteps. The n-body simulation experiments evaluate models on two scenarios, charged(5) and charged(10). Each scenario employs 3000 training trajectories, 2000 validation trajectories, and 2000 test trajectories, with each trajectory containing 10,000 timesteps. Across all experiments, the training procedure maintains a batch size of 100 and a maximum of 500 training epochs. Moreover, Models are optimized using the AdamW optimizer with mean squared error (MSE) loss. Additionally, models are optimized using the AdamW optimizer with mean squared error (MSE) loss. To ensure fair comparisons, hyperparameters are individually tuned for each model and task. All experiments are conducted on a single NVIDIA GeForce RTX

4090 24GB GPU. Detailed hyperparameters specific to CDIF across different experimental setups are listed in Table 5.

Table 5: Model and training hyperparameters for our method on different tasks.

Hyperparameter	Scenarios			
	FC	MD17	Charged	CMU Motion
Number of layers	4	4	4	4
Feature embeddings	64	64	32	64
Epochs	500	500	500	500
Batch size	100	100	100	100
Learning rate	1e-3	1e-3	1e-3	8e-4
Learning rate scheduler	steplr	steplr	steplr	steplr
Learning rate decay factor	0.9	0.9	0.9	0.9
Learning rate decay epochs	100	100	100	100

G.2 ANALYSIS OF INVARIANT COMPLEXITY.

This section investigates the impact of varying numbers of invariants on model performance for the FC(5) task. Frames are constructed using the complete differential frame in (12) to preserve directional information in the entire state space. For fair comparison, the initialization stage encodes no agent geometric features, including only intrinsic properties like mass and charge. We design six models with incremental invariant complexity:

Model 1: the inputs include only pairwise dot products $\mathbf{q}_i \cdot \mathbf{q}_j$.

Model 2: we replace dot products with squared Euclidean distances $\|\mathbf{q}_i - \mathbf{q}_j\|^2$ to compare positional encoding strategies.

Model 3: we introduce complete positional invariants to capture all geometric symmetries dependent on agent positions.

$$\{\|\mathbf{q}_i\|^2, \|\mathbf{q}_j\|^2, \|\mathbf{q}_i - \mathbf{q}_j\|^2\}, \quad (75)$$

Model 4: we extend the case 3 with velocity-related invariants

$$\{\|\mathbf{q}_i\|^2, \|\mathbf{q}_j\|^2, \|\mathbf{q}_i - \mathbf{q}_j\|^2\}, \quad \{\|\dot{\mathbf{q}}_i\|^2, \|\dot{\mathbf{q}}_j\|^2, \|\dot{\mathbf{q}}_i - \dot{\mathbf{q}}_j\|^2\}, \quad (76)$$

Model 5: we employ complete differential invariants in (8) as inputs of the layers.

Model 6: we further augment the complete differential invariants to investigate the impact of incorporating additional invariants on model performance.

$$\begin{aligned} & \{\|\mathbf{q}_i\|^2, \|\dot{\mathbf{q}}_i\|^2, \mathbf{q}_i \cdot \dot{\mathbf{q}}_i, \|\mathbf{q}_i \times \dot{\mathbf{q}}_i\|^2\}, \quad \{\|\mathbf{q}_j\|^2, \|\dot{\mathbf{q}}_j\|^2, \mathbf{q}_j \cdot \dot{\mathbf{q}}_j, \|\mathbf{q}_j \times \dot{\mathbf{q}}_j\|^2\}, \\ & \{\mathbf{q}_i \cdot \mathbf{q}_j, \|\mathbf{q}_i - \mathbf{q}_j\|^2, \dot{\mathbf{q}}_i \cdot \dot{\mathbf{q}}_j, \|\dot{\mathbf{q}}_i - \dot{\mathbf{q}}_j\|^2\}, \quad \{(\mathbf{q}_i \times \mathbf{q}_j) \cdot (\dot{\mathbf{q}}_i \times \dot{\mathbf{q}}_j)\} \\ & \{\dot{\mathbf{q}}_i \cdot (\mathbf{q}_i \times \mathbf{q}_j), \dot{\mathbf{q}}_j \cdot (\mathbf{q}_i \times \mathbf{q}_j), \mathbf{q}_i \cdot (\dot{\mathbf{q}}_i \times \dot{\mathbf{q}}_j), \mathbf{q}_j \cdot (\dot{\mathbf{q}}_i \times \dot{\mathbf{q}}_j)\}, \\ & \{\mathbf{q}_j \cdot (\mathbf{q}_i \times \dot{\mathbf{q}}_i), \dot{\mathbf{q}}_j \cdot (\mathbf{q}_i \times \dot{\mathbf{q}}_i), \mathbf{q}_i \cdot (\mathbf{q}_j \times \dot{\mathbf{q}}_j), \dot{\mathbf{q}}_i \cdot (\mathbf{q}_j \times \dot{\mathbf{q}}_j)\}, \end{aligned} \quad (77)$$

Then, we evaluate the six proposed models across the easy, medium, and difficult levels of the FC(5) task. The results are shown in Figure 6 and Table 6.

Models 1 and 2 exhibit nearly identical performance across all levels, indicating that $\mathbf{q}_i \cdot \mathbf{q}_j$ and $\|\mathbf{q}_i - \mathbf{q}_j\|^2$ serve similar roles in encoding the relative angular relationships between agents i and j . A significant performance jump occurs between model 2 and 3, with average improvements of 18.03%. This highlights the critical role of complete positional invariants, which enable smooth function approximation of all position-dependent invariants. Incorporating velocity-related invariants in model 4 yields average improvements of 12.4% over model 3, which demonstrates that invariants based on velocities enhance model’s expressiveness. Utilizing complete differential invariants in model 5 achieves an average improvement of 11.57% over model 4. This corroborates the importance

Table 6: Mean Square Error for position prediction on FC(5) tasks across three different levels. The best results are **bolded** and the second best are underlined.

Level	Easy	Medium	Difficult
	MSE	MSE	MSE
Model 1	0.1457	0.3359	0.5308
Model 2	0.1412	0.3604	0.5472
Model 3	0.1138	0.2969	0.5101
Model 4	0.0919	0.2622	0.4778
Model 5	0.0885	0.2085	<u>0.4277</u>
Model 6	<u>0.0937</u>	<u>0.2287</u>	0.4128

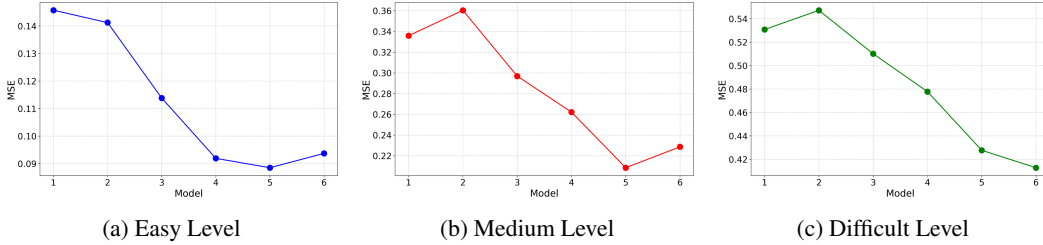


Figure 6: Mean squared error (MSE) is on the vertical axis against model index on the horizontal axis. Figures (a), (b), and (c) depict the predictive error as a function of invariant complexity for the easy, medium, and difficult levels, respectively.

of differential completeness, which significantly boosts modeling of complex dynamics. Despite including substantial invariants, model 6 performs comparably or sometimes even worse than model 5. This suggests that excessive invariants introduce redundant or irrelevant features, degrading model expressivity due to increased input dimensionality, especially in simple tasks.

G.3 ABLATION STUDY.

We introduce two core modules to ensure completeness, namely complete differential invariants and complete differential frames. To demonstrate the importance of each module, we conduct experiments on FC(5) and FC(10) tasks, with results summarized in Table 7. The results show that introducing complete differential invariants alone improves the performance of the model lacking both complete invariants and frames by an average of 35.26%. Furthermore, adding complete positional frames in (68) and complete differential frames in (12) yields additional average improvements of 14.78% and 32.99%, respectively. The results highlight the critical role of two modules in enhancing model performance.

Table 7: Ablation study on FC(5) and FC(10) tasks. Best results are **bolded**. CDIF(w/o CDIs and CDFs) corresponds to the EGNN model. CDIF(CDIs only) incorporates complete differential invariants as inputs at each layer while retaining radial direction via relative displacements. CDIF(CDIs + positional CFs) constructs positional complete frames (68) for update model.

Task	FC(5)			FC(10)		
	Easy	Medium	Difficult	Easy	Medium	Difficult
Model	MSE	MSE	MSE	MSE	MSE	MSE
CDIF(w/o CDIs and CDFs)	0.2024	0.4584	0.6956	0.1783	0.2931	4.1375
CDIF(CDIs only)	0.1295	0.2604	0.5135	0.0648	0.1778	4.0012
CDIF(CDIs + positional CFs)	0.1069	0.2557	0.4694	0.0436	0.1346	3.8470
CDIF	0.0944	0.2152	0.4331	0.0346	0.0991	3.5742