

Supplementary Information for: Learning Canonical Representations for Unified 3D Molecular Modeling

Seungyeon Choi¹, Hwanhee Kim¹, Yoonju Kim¹, Seungyong Lee¹, Inpyo Hong¹, and Sanghyun Park^{1,*}

¹Yonsei University, Department of Computer Science, Seoul, South Korea

*e-mail: sanghyun@yonsei.ac.kr

ABSTRACT

This document contains the Supplementary Notes supporting “Learning Canonical Representations for Unified 3D Molecular Modeling”.

Supplementary Note Contents

A Canonicalization for Invariant Prediction: Quotient Factorization and Stability	3
A.1 Setup	
A.2 Exact factorization through a canonical representative	
A.3 Approximate invariance and rotation-augmented risk stability	
B Canonicalization for Equivariant Prediction	4
B.1 ρ -equivariant targets	
B.2 Exact Canonicalize-and-Lift factorization	
B.3 Approximate equivariance bounds	
C Generative Modeling in Canonical Coordinates	5
C.1 Score matching and equivariant projection	
C.2 Quotient-space evaluation beyond the ideal section	
C.3 Equivariant scores without equivariant layers	
D Task-Driven Inductive Bias: Canonicalization vs. Equivariance vs. Augmentation	7
D.1 Limitations of invariant architectures	
D.2 Limitations of equivariant architectures	
D.3 Limitations of rotation augmentation	
D.4 Canonicalization as a unified solution	
E How Should We Evaluate Molecular Representations?	9
E.1 Baseline Models for Representation Evaluation	
E.2 Molecular Identity Preservation Analysis	
E.3 Chemical Semantic Consistency	
E.4 Embedding Space Geometry Analysis	
E.5 Pretrain-Finetune Alignment Analysis	
E.6 Analysis of Representation Quality Experiment	
F Experimental Details for Canonical Latent-Space Diagnostics	15
F.1 Rotation Sensitivity	
F.2 Latent-Distance Hierarchy	
F.3 Latent Perturbation Decodability	
F.4 Gauge-Waste Ratio and Search Efficiency	
F.5 Interpolation Decodability	
G Experimental Details for Figure 2	17
H Implementation Details	17
I Additional Experiments (Ablation Study)	21
I.1 In-Depth Analysis of Representation Paradigms	

45	J Datasets	23
46	J.1 Pretraining Dataset	
47	J.2 MoleculeNet	
48	J.3 TDC	
49	J.4 CrossDocked2020	
50	J.5 GEOM-DRUG	
51	J.6 Additional Benchmark	
52	J.7 Evaluation Metrics	
53	J.8 Baseline Models	
54	J.9 MoleculeNet Analysis	
55	K Flow Matching for Molecular Generation	28
56	K.1 Background: Conditional Flow Matching	
57	K.2 Flow Matching in Proposed Framework	
58	L Optimization via Latent Space Search	32
59	L.1 Problem Formulation	
60	L.2 LRA-CMA-ES Overview	
61	L.2 Application to Flow-Based Molecular Optimization	
62		

63 A Canonicalization for Invariant Prediction: Quotient Factorization and Stability

64 A.1 Setup

65 Let $\mathbf{x} \in \mathbb{R}^{N \times 3}$ be 3D coordinates and $\mathbf{h}$ be per-node features (e.g., atom types). We remove translations by centering

$$66 \quad \mathbf{x}^c := \mathbf{x} - \mathbf{1}\bar{\mathbf{x}}^\top, \quad \bar{\mathbf{x}} := \frac{1}{N} \sum_{i=1}^N \mathbf{x}_i \in \mathbb{R}^3.$$

67 For $R \in \text{SO}(3)$, the action on coordinates is $\mathbf{x} \mapsto R\mathbf{x}$ (row-wise). We use a learned *rotation selector* $\hat{R}_\theta(\mathbf{x}, \mathbf{h}) \in \text{SO}(3)$ and define
68 the *canonicalization map*

$$69 \quad A_\theta(\mathbf{x}, \mathbf{h}) := \hat{R}_\theta(\mathbf{x}, \mathbf{h}) \mathbf{x}^c \in \mathbb{R}^{N \times 3}. \quad (1)$$

70 A generic non-equivariant predictor in canonical space is

$$71 \quad \hat{y}(\mathbf{x}, \mathbf{h}) := \tilde{f}_\eta(A_\theta(\mathbf{x}, \mathbf{h}), \mathbf{h}). \quad (2)$$

72 *Permutation symmetry is not assumed:* we construct the dataset using RDKit’s canonical atom ordering, which fixes a
73 deterministic node indexing.

74 A target $y(\mathbf{x}, \mathbf{h}) \in \mathbb{R}$ is *rotation-invariant* if $y(R\mathbf{x}, \mathbf{h}) = y(\mathbf{x}, \mathbf{h})$ for all $R \in \text{SO}(3)$.

75 A.2 Exact factorization through a canonical representative

76 **Assumption 1** (Well-defined canonicalization on the data support). There exists a subset $\mathcal{S} \subseteq \mathbb{R}^{N \times 3} \times \mathcal{H}$ containing the data
77 support such that:

$$78 \quad (\text{gauge consistency}) \quad \hat{\mathbf{R}}_\theta(R\mathbf{x}, \mathbf{h}) = \hat{\mathbf{R}}_\theta(\mathbf{x}, \mathbf{h})R^\top, \quad \forall R \in \text{SO}(3), (\mathbf{x}, \mathbf{h}) \in \mathcal{S},$$

$$79 \quad (\text{orbit-separation}) \quad A_\theta(\mathbf{x}, \mathbf{h}) = A_\theta(\mathbf{x}', \mathbf{h}') \Rightarrow \mathbf{h} = \mathbf{h}' \text{ and } \exists R \in \text{SO}(3) \text{ s.t. } \mathbf{x}' = R\mathbf{x}, \quad \forall (\mathbf{x}, \mathbf{h}), (\mathbf{x}', \mathbf{h}') \in \mathcal{S}.$$

80 Equivalently, $A_\theta(R\mathbf{x}, \mathbf{h}) = A_\theta(\mathbf{x}, \mathbf{h})$ and A_θ is injective on rotation orbits over $\mathcal{S}$.

81 Global orbit-separating sections may fail at symmetric/degenerate configurations (non-trivial stabilizers), leading to
82 ambiguity or discontinuity. In practice this is mitigated by focusing on generic inputs (symmetries are measure-zero under
83 perturbations) and by using deterministic node indexing (RDKit canonical ordering) so that $\hat{R}_\theta(\mathbf{x}, \mathbf{h})$ can implement consistent
84 tie-breaking on the data support.

85 **Proposition 2** (Invariant targets factor through canonicalization). *Under Assumption 1, for any rotation-invariant target y there*
86 *exists a function $\tilde{y}$ such that*

$$87 \quad y(\mathbf{x}, \mathbf{h}) = \tilde{y}(A_\theta(\mathbf{x}, \mathbf{h}), \mathbf{h}) \quad \text{for all } (\mathbf{x}, \mathbf{h}) \in \mathcal{S}. \quad (3)$$

88 *Proof.* Fix $(\mathbf{x}, \mathbf{h}) \in \mathcal{S}$ and set $u = A_\theta(\mathbf{x}, \mathbf{h})$. Gauge consistency implies u is constant on the orbit $\{(R\mathbf{x}, \mathbf{h}) : R \in \text{SO}(3)\}$, and
89 invariance implies y is constant on the same orbit. Orbit-separation ensures $(u, \mathbf{h})$ identifies the orbit within $\mathcal{S}$, so defining
90 $\tilde{y}(u, \mathbf{h}) := y(\mathbf{x}, \mathbf{h})$ is well-defined. $\square$

91 **Proposition 3** (Universal approximation in canonical space). *Assume Assumption 1. If y is continuous and rotation-invariant on*
92 *$\mathcal{S}$ and $\tilde{f}_\eta(\cdot, \mathbf{h})$ is a universal approximator for continuous functions on the domain of $(A_\theta(\mathbf{x}, \mathbf{h}), \mathbf{h})$ (restricted to $(\mathbf{x}, \mathbf{h}) \in \mathcal{S}$),*
93 *then for any $\varepsilon > 0$ there exist parameters η such that*

$$94 \quad \sup_{(\mathbf{x}, \mathbf{h}) \in \mathcal{S}} |\tilde{f}_\eta(A_\theta(\mathbf{x}, \mathbf{h}), \mathbf{h}) - y(\mathbf{x}, \mathbf{h})| \leq \varepsilon. \quad (4)$$

95 By Proposition 2, $y(\mathbf{x}, \mathbf{h}) = \tilde{y}(A_\theta(\mathbf{x}, \mathbf{h}), \mathbf{h})$ for some continuous $\tilde{y}$ on the induced domain. Universality yields η approxim-
96 ating $\tilde{y}$ uniformly to error ε .

97 A.3 Approximate invariance and rotation-augmented risk stability

98 **Assumption 4** (Alignment consistency and Lipschitz stability). Assume there exist $\delta \geq 0$, $L_f \geq 0$, and $L_\ell \geq 0$ such that for all
 99 $(x, h) \in \mathcal{S}$:

$$100 \sup_{R \in \text{SO}(3)} \|A_\theta(R\mathbf{x}, \mathbf{h}) - A_\theta(\mathbf{x}, \mathbf{h})\|_F \leq \delta,$$

101 $\tilde{f}_\eta(\cdot, \mathbf{h})$ is L_f -Lipschitz in its first argument, and $\ell(\hat{y}, y)$ is L_ℓ -Lipschitz in $\hat{y}$:

$$102 |\tilde{f}_\eta(u_1, \mathbf{h}) - \tilde{f}_\eta(u_2, \mathbf{h})| \leq L_f \|u_1 - u_2\|_F, \quad |\ell(\hat{y}_1, y) - \ell(\hat{y}_2, y)| \leq L_\ell |\hat{y}_1 - \hat{y}_2|.$$

103 **Proposition 5** (Approximate invariance of canonical predictors). *Under Assumption 4, the predictor (2) satisfies*

$$104 \sup_{R \in \text{SO}(3)} |\tilde{y}(R\mathbf{x}, \mathbf{h}) - \tilde{y}(\mathbf{x}, \mathbf{h})| \leq L_f \delta \quad \text{for all } (\mathbf{x}, \mathbf{h}) \in \mathcal{S}.$$

105 *If moreover y is rotation-invariant, then for each fixed $(\mathbf{x}, \mathbf{h}) \in \mathcal{S}$,*

$$106 \left| \mathbb{E}_{R \sim \mu} [\ell(\tilde{y}(R\mathbf{x}, \mathbf{h}), y(\mathbf{x}, \mathbf{h}))] - \ell(\tilde{y}(\mathbf{x}, \mathbf{h}), y(\mathbf{x}, \mathbf{h})) \right| \leq L_\ell L_f \delta.$$

107 *Proof.* Lipschitzness gives $|\tilde{y}(R\mathbf{x}, \mathbf{h}) - \tilde{y}(\mathbf{x}, \mathbf{h})| \leq L_f \|A_\theta(R\mathbf{x}, \mathbf{h}) - A_\theta(\mathbf{x}, \mathbf{h})\|_F \leq L_f \delta$. For invariant y , apply L_ℓ -Lipschitzness to
 108 the loss difference and average over R . $\square$

109 B Canonicalization for Equivariant Prediction

110 B.1 ρ -equivariant targets

111 Let $\mathcal{Y}$ be an output vector space (e.g., $\mathbb{R}^3$ or $\mathbb{R}^{N \times 3}$) and $\rho : \text{SO}(3) \rightarrow GL(\mathcal{Y})$ be a linear representation. A target $y(\mathbf{x}, \mathbf{h}) \in \mathcal{Y}$
 112 is ρ -equivariant if

$$113 y(R\mathbf{x}, \mathbf{h}) = \rho(R) y(\mathbf{x}, \mathbf{h}) \quad \forall R \in \text{SO}(3). \quad (5)$$

114 B.2 Exact Canonicalize-and-Lift factorization

115 **Proposition 6** (Canonical factorization of ρ -equivariant maps). *Assume Assumption 1. A function $y : \mathbb{R}^{N \times 3} \times \mathcal{H} \rightarrow \mathcal{Y}$ is*
 116 *ρ -equivariant on $\mathcal{S}$ if and only if there exists $\tilde{y}$ such that*

$$117 y(\mathbf{x}, \mathbf{h}) = \rho(\hat{\mathbf{R}}_\theta(\mathbf{x}, \mathbf{h}))^{-1} \tilde{y}(A_\theta(\mathbf{x}, \mathbf{h}), \mathbf{h}), \quad \forall (\mathbf{x}, \mathbf{h}) \in \mathcal{S}. \quad (6)$$

118 *Proof.* Assume y is ρ -equivariant on $\mathcal{S}$. For any $(\mathbf{x}, \mathbf{h}) \in \mathcal{S}$, set $\mathbf{u} := A_\theta(\mathbf{x}, \mathbf{h})$ and define

$$119 \tilde{y}(\mathbf{u}, \mathbf{h}) := \rho(\hat{\mathbf{R}}_\theta(\mathbf{x}, \mathbf{h})) y(\mathbf{x}, \mathbf{h}). \quad (7)$$

120 To see that $\tilde{y}$ is well-defined, let $(\mathbf{x}_1, \mathbf{h}), (\mathbf{x}_2, \mathbf{h}) \in \mathcal{S}$ satisfy $A_\theta(\mathbf{x}_1, \mathbf{h}) = A_\theta(\mathbf{x}_2, \mathbf{h}) = \mathbf{u}$. By orbit-separation there exists
 121 $R \in \text{SO}(3)$ such that $\mathbf{x}_2 = R\mathbf{x}_1$. Using ρ -equivariance, gauge consistency, and that ρ is a representation, we have

$$\begin{aligned} 122 y(\mathbf{x}_2, \mathbf{h}) &= y(R\mathbf{x}_1, \mathbf{h}) = \rho(R) y(\mathbf{x}_1, \mathbf{h}), \\ 123 \hat{\mathbf{R}}_\theta(\mathbf{x}_2, \mathbf{h}) &= \hat{\mathbf{R}}_\theta(R\mathbf{x}_1, \mathbf{h}) = \hat{\mathbf{R}}_\theta(\mathbf{x}_1, \mathbf{h}) R^\top, \\ 124 \rho(\hat{\mathbf{R}}_\theta(\mathbf{x}_2, \mathbf{h})) y(\mathbf{x}_2, \mathbf{h}) &= \rho(\hat{\mathbf{R}}_\theta(\mathbf{x}_1, \mathbf{h}) R^\top) \rho(R) y(\mathbf{x}_1, \mathbf{h}) \\ 125 &= \rho(\hat{\mathbf{R}}_\theta(\mathbf{x}_1, \mathbf{h})) \rho(R^\top) \rho(R) y(\mathbf{x}_1, \mathbf{h}) = \rho(\hat{\mathbf{R}}_\theta(\mathbf{x}_1, \mathbf{h})) y(\mathbf{x}_1, \mathbf{h}), \end{aligned} \quad (8)$$

126 where $\rho(R^\top) = \rho(R)^{-1}$ for $R \in \text{SO}(3)$. Hence $\tilde{y}(\mathbf{u}, \mathbf{h})$ does not depend on the choice of $(\mathbf{x}, \mathbf{h})$ with $\mathbf{u} = A_\theta(\mathbf{x}, \mathbf{h})$, and rearranging
 127 the definition yields $y(\mathbf{x}, \mathbf{h}) = \rho(\hat{\mathbf{R}}_\theta(\mathbf{x}, \mathbf{h}))^{-1} \tilde{y}(A_\theta(\mathbf{x}, \mathbf{h}), \mathbf{h})$ for all $(\mathbf{x}, \mathbf{h}) \in \mathcal{S}$.

128 Conversely, suppose there exists $\tilde{y}$ such that $y(\mathbf{x}, \mathbf{h}) = \rho(\hat{\mathbf{R}}_\theta(\mathbf{x}, \mathbf{h}))^{-1} \tilde{y}(A_\theta(\mathbf{x}, \mathbf{h}), \mathbf{h})$ on $\mathcal{S}$. Then for any $R \in \text{SO}(3)$, gauge
 129 consistency and the homomorphism property of ρ give

$$\begin{aligned} 130 y(R\mathbf{x}, \mathbf{h}) &= \rho(\hat{\mathbf{R}}_\theta(R\mathbf{x}, \mathbf{h}))^{-1} \tilde{y}(A_\theta(R\mathbf{x}, \mathbf{h}), \mathbf{h}) \\ 131 &= \rho(\hat{\mathbf{R}}_\theta(\mathbf{x}, \mathbf{h}) R^\top)^{-1} \tilde{y}(A_\theta(\mathbf{x}, \mathbf{h}), \mathbf{h}) \\ 132 &= (\rho(\hat{\mathbf{R}}_\theta(\mathbf{x}, \mathbf{h})) \rho(R^\top))^{-1} \tilde{y}(A_\theta(\mathbf{x}, \mathbf{h}), \mathbf{h}) \\ 133 &= \rho(R) \rho(\hat{\mathbf{R}}_\theta(\mathbf{x}, \mathbf{h}))^{-1} \tilde{y}(A_\theta(\mathbf{x}, \mathbf{h}), \mathbf{h}) = \rho(R) y(\mathbf{x}, \mathbf{h}), \end{aligned}$$

134 which shows that y is ρ -equivariant on $\mathcal{S}$. $\square$

135 B.3 Approximate equivariance bounds

136 Define the Canonicalize-and-Lift predictor

$$137 \quad \hat{y}(\mathbf{x}, \mathbf{h}) := \rho(\hat{R}_\theta(\mathbf{x}, \mathbf{h}))^{-1} \tilde{f}_\eta(A_\theta(\mathbf{x}, \mathbf{h}), \mathbf{h}).$$

138 **Assumption 7** (Approximate canonicalization and Lipschitz stability). Assume there exist $\delta_A, \delta_R \geq 0$ such that for all $(x, h) \in \mathcal{S}$
 139 and all $R \in \text{SO}(3)$,

$$140 \quad \|A_\theta(R\mathbf{x}, \mathbf{h}) - A_\theta(x, h)\|_F \leq \delta_A, \quad \|\hat{R}_\theta(R\mathbf{x}, \mathbf{h}) - \hat{R}_\theta(\mathbf{x}, \mathbf{h})R^\top\|_F \leq \delta_R.$$

141 Assume $\tilde{f}_\eta(\cdot, h)$ is $L_{\tilde{f}}$ -Lipschitz and bounded by $B_{\tilde{f}}$:

$$142 \quad \|\tilde{f}_\eta(u_1, \mathbf{h}) - \tilde{f}_\eta(u_2, \mathbf{h})\| \leq L_{\tilde{f}}\|u_1 - u_2\|_F, \quad \|\tilde{f}_\eta(u, \mathbf{h})\| \leq B_{\tilde{f}}.$$

143 Assume $\rho(Q)^{-1}$ is L_ρ -Lipschitz in operator norm w.r.t. Q :

$$144 \quad \|\rho(Q_1)^{-1} - \rho(Q_2)^{-1}\|_{\text{op}} \leq L_\rho\|Q_1 - Q_2\|_F.$$

145 **Proposition 8** (Approximate equivariance of Canonicalize-and-Lift predictors). Under Assumption 7, for all $(x, h) \in \mathcal{S}$ and
 146 $R \in \text{SO}(3)$,

$$147 \quad \|\hat{y}(R\mathbf{x}, \mathbf{h}) - \rho(R)\hat{y}(\mathbf{x}, \mathbf{h})\| \leq L_{\tilde{f}}\delta_A + L_\rho B_{\tilde{f}}\delta_R.$$

148 *Proof.* Add/subtract $\rho(\hat{R}_\theta(R\mathbf{x}, \mathbf{h}))^{-1} \tilde{f}_\eta(A_\theta(\mathbf{x}, \mathbf{h}), \mathbf{h})$. Bound (i) the $\tilde{f}_\eta$ change by $L_{\tilde{f}}\delta_A$ and (ii) the $\rho(\hat{R}_\theta)^{-1}$ change by $L_\rho\delta_R$
 149 times $\|\tilde{f}_\eta\| \leq B_{\tilde{f}}$, then use $\rho(\hat{R}_\theta(\mathbf{x}, \mathbf{h})R^\top)^{-1} = \rho(R)\rho(\hat{R}_\theta(\mathbf{x}, \mathbf{h}))^{-1}$. $\square$

150 C Generative Modeling in Canonical Coordinates

151 C.1 Score matching and equivariant projection

152 Let p_t be a time-marginal of an isotropic noising process on $\mathbb{R}^{N \times 3}$ and

$$153 \quad \mathbf{s}^*(\mathbf{x}, t) := \nabla_{\mathbf{x}} \log p_t(\mathbf{x}).$$

154 For a score model $s(x, t)$, define the population MSE

$$155 \quad \mathcal{J}(s) := \mathbb{E}_t \mathbb{E}_{\mathbf{x} \sim p_t} [\|s(x, t) - \mathbf{s}^*(\mathbf{x}, t)\|_F^2].$$

156 **Assumption 9** (Rotation-invariant marginals). Assume p_t is $\text{SO}(3)$ -invariant: $p_t(R\mathbf{x}) = p_t(\mathbf{x})$ for all $R \in \text{SO}(3)$. Then $\mathbf{s}^*$ is
 157 $\text{SO}(3)$ -equivariant: $\mathbf{s}^*(R\mathbf{x}, t) = R\mathbf{s}^*(\mathbf{x}, t)$ (1)

158 **Definition 1** (Equivariant projection). For a vector field $\mathbf{s}(\cdot, t)$ define

$$159 \quad (\mathcal{P}_{\text{eq}} \cdot \mathbf{s})(\mathbf{x}, t) := \mathbb{E}_{R \sim \mu} [R\mathbf{s}(R^\top \mathbf{x}, t)], \tag{9}$$

160 where μ is the Haar measure on $\text{SO}(3)$ and $\mathcal{P}_{\text{eq}}$ denotes the equivariant projection operator that maps an arbitrary vector
 161 field to an $\text{SO}(3)$ -equivariant one by Haar-averaging over rotations.

162 **Proposition 10** (WLOG canonical coordinates for score matching). Assume (i) p_t is $\text{SO}(3)$ -invariant and hence $\mathbf{s}^*$ is $\text{SO}(3)$ -
 163 equivariant, and (ii) Assumption 1 holds. Then for any score field s , there exists a canonical-space field $\tilde{s}(\cdot, t)$ such that

$$164 \quad s_{\text{canonical}}(\mathbf{x}, t) := \hat{R}_\theta(\mathbf{x}, \mathbf{h})^\top \tilde{s}(A_\theta(\mathbf{x}, \mathbf{h}), t) \tag{10}$$

165 satisfies

$$166 \quad \mathcal{J}(s_{\text{canonical}}) \leq \mathcal{J}(s). \tag{11}$$

167 *Proof.* We construct $s_{\text{canonical}}$ via the equivariant projection. Let $\Delta := s - \mathbf{s}^*$. Since $\mathbf{s}^*$ is $\text{SO}(3)$ -equivariant, $\mathcal{P}_{\text{eq}}\mathbf{s}^* = \mathbf{s}^*$, and
 168 thus $\mathcal{P}_{\text{eq}}s - \mathbf{s}^* = \mathcal{P}_{\text{eq}}\Delta$. By Jensen's inequality and orthogonality invariance of $\|\cdot\|_F$,

$$169 \quad \|\mathcal{P}_{\text{eq}}\Delta(\mathbf{x}, t)\|_F^2 = \left\| \mathbb{E}_R [R\Delta(R^\top \mathbf{x}, t)] \right\|_F^2 \leq \mathbb{E}_R \|\Delta(R^\top \mathbf{x}, t)\|_F^2.$$

170 Taking $\mathbb{E}_t \mathbb{E}_{\mathbf{x} \sim p_t}$ and using $\text{SO}(3)$ -invariance of p_t yields $\mathcal{J}(\mathcal{P}_{\text{eq}}s) \leq \mathcal{J}(s)$.

171 Now, $\mathcal{P}_{\text{eq}}s$ is $\text{SO}(3)$ -equivariant by construction. Under Assumption 1, any $\text{SO}(3)$ -equivariant score admits the canonicalize-
 172 and-lift form (cf. Proposition 6):

$$173 \quad (\mathcal{P}_{\text{eq}}s)(\mathbf{x}, t) = \hat{R}_\theta(\mathbf{x}, \mathbf{h})^\top \tilde{s}(A_\theta(\mathbf{x}, \mathbf{h}), t)$$

174 for some $\tilde{s}$. Setting $s_{\text{canonical}} := \mathcal{P}_{\text{eq}}s$ completes the proof. $\square$

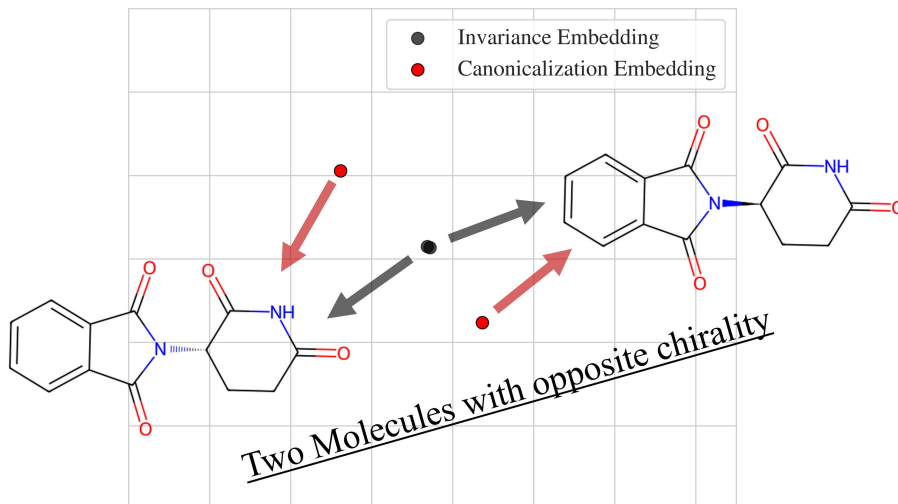


Figure 1. Scatter plots of Repr.A (Uni-Mol2) and Repr.Z (THEMOL) embeddings for a pair of molecules with opposite chirality.

175 C.2 Quotient-space evaluation beyond the ideal section

176 Let $G = \text{SO}(3)$, $\varphi(\mathbf{x}) := [\mathbf{x}]_G$, and $d_{\text{orb}}(\varphi(\mathbf{x}), \varphi(\mathbf{y})) := \min_{R \in \text{SO}(3)} \|R\mathbf{x} - \mathbf{y}\|_F$. Let W_1^{orb} be the 1-Wasserstein distance induced
177 by d_{orb} .

178 **Assumption 11** (Approximate orbit consistency). On $\mathcal{S}$, for every $(\mathbf{x}, \mathbf{h})$ there exists $R \in \text{SO}(3)$ such that $\|A_\theta(\mathbf{x}, \mathbf{h}) - R\mathbf{x}\|_F \leq$
179 δ_A .

180 **Proposition 12** (Orbit-space guarantee under approximate canonicalization). Let q be a target distribution on $\mathcal{S}$, with $\tilde{q} := \varphi_{\#}q$
181 and $q_A := (A_\theta)_{\#}q$. Let p be obtained by sampling $u \sim p_A$ and lifting $x := R^\top u$ with $R \sim \mu$ (Haar), and set $\tilde{p} := \varphi_{\#}p$. Under
182 Assumption 11,

$$183 \quad W_1^{\text{orb}}(\tilde{p}, \tilde{q}) \leq W_1(p_A, q_A) + \delta_A, \quad \text{hence if } p_A = q_A \text{ then } W_1^{\text{orb}}(\tilde{p}, \tilde{q}) \leq \delta_A.$$

184 **Proof (sketch).** Rotation-invariance gives $\tilde{p} = \varphi_{\#}p_A$. Since φ is 1-Lipschitz from $(\mathbb{R}^{N \times 3}, \|\cdot\|_F)$ to $(\varphi(\mathbb{R}^{N \times 3}), d_{\text{orb}})$,
185 $W_1^{\text{orb}}(\varphi_{\#}p_A, \varphi_{\#}q_A) \leq W_1(p_A, q_A)$. Couple $(x, h) \sim q$ with $u = A_\theta(x, h)$; by Assumption 11, $d_{\text{orb}}(\varphi(u), \varphi(x)) \leq \delta_A$, hence
186 $W_1^{\text{orb}}(\varphi_{\#}q_A, \varphi_{\#}q) \leq \delta_A$. Apply triangle inequality. $\square$

187 C.3 Equivariant scores without equivariant layers

188 Assume p_t is $\text{SO}(3)$ -invariant so that the true score $s^*(x, t) = \nabla_x \log p_t(x)$ is $\text{SO}(3)$ -equivariant. Under Assumption 1, any
189 $\text{SO}(3)$ -equivariant score field admits the canonicalize-and-lift parameterization (cf. Proposition 6):

$$190 \quad s(x, t) = \hat{R}_\theta(x, h)^\top \tilde{s}(u, t), \tag{12}$$

191 for some canonical-space field $\tilde{s}(\cdot, t)$.

192 **Remark 1** (Comparison to $\text{SE}(3)$ -equivariant backbones). Equation 12 implies that restricting to canonical-space modeling
193 does not shrink the class of equivariant scores under exact canonicalization assumptions. Consequently, $\text{SE}(3)$ -equivariant
194 layers are not required for expressivity: equivariance can be enforced at the interface by a learned gauge selector and a lift,
195 while the backbone operates in canonical coordinates.

196 **Remark 2** (Computational implication in iterative denoising). In diffusion or score-based generation with T denoising steps,
197 an $\text{SE}(3)$ -equivariant backbone enforces equivariance at every layer and at every step. The canonicalize-and-lift construction
198 enforces symmetry through (i) canonicalization and (ii) a lightweight lift, while the main backbone can be non-equivariant.
199 Hence, when equivariant layers are substantially more expensive than canonicalization plus a standard backbone, the per-step
200 overhead compounds over T steps. This separation is most beneficial in foundation settings where the same backbone is reused
201 across heterogeneous tasks, including invariant objectives and equivariant outputs.

202 D Task-Driven Inductive Bias: Canonicalization vs. Equivariance vs. Augmentation

203 D.1 Limitations of invariant architectures

204 Purely invariant representations face two fundamental limitations for foundation settings.

205 **Lemma 13** (Invariant latent cannot identify pose). *Let Enc be $SE(3)$ -invariant in $\mathbf{x}$: $\text{Enc}(R\mathbf{x} + \mathbf{1}t^\top, \mathbf{h}) = \text{Enc}(\mathbf{x}, \mathbf{h})$ for all*
 206 *$(R, t) \in SE(3)$. Assume there exist $(\mathbf{x}, \mathbf{h})$ and a nontrivial $(R, t) \in SE(3)$ such that $(R\mathbf{x} + \mathbf{1}t^\top, \mathbf{h}) \neq (\mathbf{x}, \mathbf{h})$. Then no deterministic*
 207 *decoder Dec that takes only $z = \text{Enc}(\mathbf{x}, \mathbf{h})$ can satisfy $\text{Dec}(\text{Enc}(\mathbf{x}, \mathbf{h})) = \mathbf{x}$ simultaneously for both $\mathbf{x}$ and $R\mathbf{x} + \mathbf{1}t^\top$.*

208 *Proof.* Let $z = \text{Enc}(\mathbf{x}, \mathbf{h}) = \text{Enc}(R\mathbf{x} + \mathbf{1}t^\top, \mathbf{h})$ by invariance. If Dec depends only on z , then $\text{Dec}(z)$ must equal both $\mathbf{x}$ and
 209 $R\mathbf{x} + \mathbf{1}t^\top$, a contradiction. $\square$

210 **Proposition 14** (Chirality loss in distance-only invariant latents). *Assume $\text{Enc}_{\mathbf{h}}(\mathbf{x}, \mathbf{h}) = \Phi(\mathbf{D}(\mathbf{x}), \mathbf{h})$ where $\mathbf{D}(\mathbf{x})_{ij} = \|\mathbf{x}_i - \mathbf{x}_j\|_2$.*
 211 *Then $\text{Enc}_{\mathbf{h}}(Q\mathbf{x}, \mathbf{h}) = \text{Enc}_{\mathbf{h}}(\mathbf{x}, \mathbf{h})$ for all $Q \in O(3)$. Consequently, there exists an $SO(3)$ -invariant target y that cannot be written*
 212 *as a function of $\text{Enc}_{\mathbf{h}}(\mathbf{x}, \mathbf{h})$, as shown Figure 1.*

213 *Proof sketch.* Distance matrices satisfy $\mathbf{D}(Q\mathbf{x}) = \mathbf{D}(\mathbf{x})$ for all $Q \in O(3)$, hence $\text{Enc}_{\mathbf{h}}(Q\mathbf{x}, \mathbf{h}) = \text{Enc}_{\mathbf{h}}(\mathbf{x}, \mathbf{h})$. Let

$$214 \quad y(\mathbf{x}) := \det[(\mathbf{x}_i - \mathbf{x}_\ell), (\mathbf{x}_j - \mathbf{x}_\ell), (\mathbf{x}_k - \mathbf{x}_\ell)]$$

215 for some indices i, j, k, ℓ . Then $y(R\mathbf{x}) = y(\mathbf{x})$ for all $R \in SO(3)$ (since $\det(R) = 1$), while $y(Q\mathbf{x}) = -y(\mathbf{x})$ for any reflection
 216 $Q \in O(3) \setminus SO(3)$ (since $\det(Q) = -1$). Choose $\mathbf{x}$ with $y(\mathbf{x}) \neq 0$ and set $\mathbf{x}' := Q\mathbf{x}$; then $\text{Enc}_{\mathbf{h}}(\mathbf{x}', \mathbf{h}) = \text{Enc}_{\mathbf{h}}(\mathbf{x}, \mathbf{h})$ but $y(\mathbf{x}') \neq y(\mathbf{x})$,
 217 so y cannot be a function of $\text{Enc}_{\mathbf{h}}$. $\square$

218 D.2 Limitations of equivariant architectures

219 Equivariant representations provide a principled and often highly effective inductive bias for force prediction, vector-field
 220 estimation, coordinate-level generation, and other tasks whose outputs transform under rotations. Modern equivariant networks
 221 can also propagate scalar invariant channels together with higher-order geometric features, making them capable of handling
 222 both invariant and equivariant prediction targets. Our claim is instead specific to the unified autoencoder setting considered
 223 in this work: **the same pretrained latent is not only consumed by prediction heads, but also used as the state space for**
 224 **latent generative modeling and black-box molecular optimization. In this setting, the geometry of the latent space itself**
 225 **becomes a central design issue.**

226 **Equivariant latent orbits.** Let Enc be an equivariant encoder and let $\mathbf{z} = \text{Enc}(\mathbf{x}, \mathbf{h})$. Then, for any $R \in SO(3)$,

$$227 \quad \text{Enc}(R\mathbf{x}, \mathbf{h}) = \rho(R)\text{Enc}(\mathbf{x}, \mathbf{h}),$$

228 where ρ denotes the induced representation on the latent space. Hence, a single molecular structure corresponds not to one
 229 latent point, but to an orbit

$$230 \quad \mathcal{O}_{\mathbf{z}} = \{\rho(R)\mathbf{z} : R \in SO(3)\}.$$

231 This is the correct behavior for equivariant prediction, but it also means that **the latent represents both intrinsic molecular**
 232 **variation and global orientation**. Thus, when the latent is reused for generation or optimization, the model operates on an
 233 orbit-lifted space rather than directly on the quotient space of molecular structures modulo global rotations.

234 **Implication for latent generative modeling.** For latent diffusion or flow matching, orbit-lifted equivariant latents require
 235 the generator to model variations that are chemically redundant. If a clean latent can be written as $Z_1 = \rho(R)Z_c$, where Z_c
 236 denotes the intrinsic molecular latent and $R \sim \mu_{\text{Haar}}$ is a global rotation, then the target distribution contains orientation-induced
 237 variability in addition to molecular variability. In high-noise regimes, the corrupted state carries limited information about the
 238 clean orientation, so the denoising or flow target can include uncertainty over R that is irrelevant to molecular identity(2). This
 239 does not make equivariant generation invalid; rather, it indicates that the generator must spend capacity modeling global orbit
 240 variation that a quotient- or canonical-space formulation can avoid.

241 **Implication for latent-space optimization.** The same orbit structure can reduce the efficiency of black-box latent optimization.
 242 In our setting, optimization searches a pretrained latent space while keeping the generator and decoder fixed. If rotated versions
 243 of the same molecule occupy different latent points, then an optimizer may explore multiple regions corresponding to the same
 244 underlying molecular candidate. Moreover, an unconstrained search direction in an equivariant latent can mix chemically
 245 meaningful molecular changes with changes in global orientation or gauge. Restricting the search only to invariant scalar
 246 components may reduce this redundancy, but it can also remove pose information needed for faithful 3D reconstruction and
 247 coordinate-level generation. Thus, the issue is not the predictive expressivity of equivariant modules, but the suitability of the
 248 induced latent geometry as a shared search and generation space.

249 **Gauge handling as interface-level bookkeeping.** Canonicalization does not eliminate symmetry handling; it relocates it.
 250 Any 3D coordinate-based model must account for global orientation in some form. In a fully equivariant latent, this orientation
 251 dependence remains inside the reusable latent through the orbit $\mathcal{O}_z$. In our formulation, a learned selector $\hat{R}_\theta(\mathbf{x}, \mathbf{h})$ first maps
 252 the input to a canonical representative,

$$253 \quad A_\theta(\mathbf{x}, \mathbf{h}) = \hat{\mathbf{R}}_\theta(\mathbf{x}, \mathbf{h})\mathbf{x}^c,$$

254 with the desired consistency condition

$$255 \quad \hat{\mathbf{R}}_\theta(\mathbf{R}\mathbf{x}, \mathbf{h}) = \hat{\mathbf{R}}_\theta(\mathbf{x}, \mathbf{h})\mathbf{R}^\top \iff A_\theta(\mathbf{R}\mathbf{x}, \mathbf{h}) = A_\theta(\mathbf{x}, \mathbf{h}).$$

256 The reusable latent is then learned in this canonical frame, while the selected rotation is retained only for lifting coordinate-level
 257 or equivariant outputs back to the original frame. Therefore, gauge bookkeeping is preserved at the input/output interface, but
 258 removed from the latent space on which generation and optimization operate.

259 D.3 Limitations of rotation augmentation

260 Rotation augmentation is a common alternative to architectural equivariance. We show that it provides weaker guarantees than
 261 canonicalization.

262 **Setup.** Let $G = \text{SO}(3)$ with Haar measure μ , and let y be G -invariant: $y(\mathbf{R}\mathbf{x}, \mathbf{h}) = y(\mathbf{x}, \mathbf{h})$. For a predictor f , define the
 263 rotation-augmentation risk

$$264 \quad \mathcal{R}_{\text{aug}}(f) := \mathbb{E}_{(\mathbf{x}, \mathbf{h})} \mathbb{E}_{R \sim \mu} [\ell(f(\mathbf{R}\mathbf{x}, \mathbf{h}), y(\mathbf{x}, \mathbf{h}))], \quad (Pf)(\mathbf{x}, \mathbf{h}) := \mathbb{E}_{R \sim \mu} [f(\mathbf{R}\mathbf{x}, \mathbf{h})].$$

265 We distinguish *single-view inference* (one forward pass) from *group-averaged inference* that explicitly computes Pf at test time.
 266 *Interpretation:* group-averaged inference is a Monte-Carlo approximation of a group integral, i.e., invariance is obtained by
 267 paying inference-time compute. For a non-equivariant backbone, achieving small variation of $f(\mathbf{R}\mathbf{x}, \mathbf{h})$ over $R \in \text{SO}(3)$ is itself
 268 a hard learning problem; when such variation persists, the required number of views K becomes large.

269 **Average risk does not control worst-case invariance.**

270 **Proposition 15** (Augmentation risk does not certify worst-case invariance). *Fix $(\mathbf{x}, \mathbf{h})$ and consider the orbit $\mathcal{O}(\mathbf{x}) = \{(\mathbf{R}\mathbf{x}, \mathbf{h}) : R \in G\}$. For squared loss $\ell(a, b) = (a - b)^2$ and invariant label $y(\mathbf{x}, \mathbf{h}) \equiv 0$, for any $\varepsilon > 0$ and $M > 0$, there exists a measurable*
 271 *f on $\mathcal{O}(\mathbf{x})$ such that*

$$273 \quad \mathbb{E}_{R \sim \mu} [\ell(f(\mathbf{R}\mathbf{x}, \mathbf{h}), 0)] \leq \varepsilon \quad \text{but} \quad \sup_{R \in G} |f(\mathbf{R}\mathbf{x}, \mathbf{h}) - f(\mathbf{x}, \mathbf{h})| \geq M.$$

274 *Proof sketch.* Choose $A \subseteq G$ with $\mu(A) = \varepsilon/M^2$ and set $f(\mathbf{R}\mathbf{x}, \mathbf{h}) = M \cdot \mathbf{1}\{R \in A\}$, $f(\mathbf{x}, \mathbf{h}) = 0$. □

275 This formalizes a structural gap: minimizing an *average* risk over random rotations does not control *worst-case* rotational
 276 variation under single-view inference. Canonicalization provides worst-case control directly via $\sup_R \|A_\theta(\mathbf{R}\mathbf{x}, \mathbf{h}) - A_\theta(\mathbf{x}, \mathbf{h})\|_F \leq$
 277 δ and Proposition 5.

278 **Group averaging incurs multi-view cost.** To obtain an explicitly invariant predictor, one must approximate Pf via
 279 Monte-Carlo averaging:

$$280 \quad \hat{f}_K(\mathbf{x}, \mathbf{h}) := \frac{1}{K} \sum_{k=1}^K f(R_k \mathbf{x}, \mathbf{h}), \quad R_k \stackrel{\text{i.i.d.}}{\sim} \mu.$$

281 This estimator directly approximates the group integral defining Pf ; unless $f(\mathbf{R}\mathbf{x}, \mathbf{h})$ is already nearly constant over R ,
 282 its variance can be large, forcing large K . The mean-square approximation error satisfies $\mathbb{E}[(\hat{f}_K(\mathbf{x}, \mathbf{h}) - (Pf)(\mathbf{x}, \mathbf{h}))^2] =$
 283 $\text{Var}_{R \sim \mu}(f(\mathbf{R}\mathbf{x}, \mathbf{h}))/K$, so achieving $\text{MSE} \leq \varepsilon^2$ requires $K = \Omega(\varepsilon^{-2})$.

284 **Uniform control requires covering $\text{SO}(3)$.** For worst-case guarantees, we need uniform control over all rotations.

285 **Proposition 16** (Covering number scaling of $\text{SO}(3)$). *Assume f is L_{rot} -Lipschitz on rotation orbits under a bi-invariant geodesic*
 286 *distance on $\text{SO}(3)$. Then achieving ε -uniform control, i.e., $\sup_{R \in \text{SO}(3)} |f(\mathbf{R}\mathbf{x}, \mathbf{h}) - \hat{f}_K(\mathbf{x}, \mathbf{h})| \leq \varepsilon$, requires*

$$287 \quad K = \Omega\left(\left(\frac{L_{\text{rot}}}{\varepsilon}\right)^3\right).$$

288 **Remark 3** (Why the exponent is 3). $\text{SO}(3)$ is a compact 3-dimensional Riemannian manifold. Metric balls have volume
 289 $\Theta(\eta^3)$, so covering $\text{SO}(3)$ with η -balls requires $\Omega(\eta^{-3})$ balls.

290 **Summary: canonicalization vs. augmentation.** Canonicalization achieves worst-case stability $\sup_R |\hat{y}(R\mathbf{x}, \mathbf{h}) - \hat{y}(\mathbf{x}, \mathbf{h})| \leq$
 291 $L_f \delta$ with *one* backbone pass (Proposition 5). In contrast, multi-view augmentation realizes invariance by approximating a group
 292 integral at inference time, shifting the burden to compute: it requires $K = \Omega(\varepsilon^{-2})$ views for MSE control and $K = \Omega(\varepsilon^{-3})$
 293 views for uniform control, and it asks a non-equivariant backbone to implicitly learn small rotation-wise variation of $f(R\mathbf{x}, \mathbf{h})$
 294 over $\text{SO}(3)$ compute tradeoff that canonicalization avoids.

295 D.4 Canonicalization as a unified solution

296 The preceding analysis highlights complementary limitations of common symmetry-handling strategies: (i) pose-level coordinate
 297 reconstruction is not identifiable from a purely invariant latent without a gauge/section, (ii) equivariant coordinate-state modeling
 298 can be computationally demanding in iterative generation, and (iii) augmentation alone does not provide worst-case guarantees
 299 without paying multi-view inference cost. Under mild section/stability assumptions, canonicalization offers a unified route that
 300 addresses these issues within a single framework.

301 **Proposition 17** (Task-wise factorization via canonicalization). *Assume Assumption 1. Then for any continuous invariant target*
 302 *y_{inv} , there exists a continuous head g_{inv} such that*

$$303 \quad y_{\text{inv}}(\mathbf{x}, \mathbf{h}) = g_{\text{inv}}(A_\theta(\mathbf{x}, \mathbf{h}), \mathbf{h}).$$

304 *Moreover, for any continuous ρ -equivariant target y_{eq} , there exists a continuous head g_{eq} such that*

$$305 \quad y_{\text{eq}}(\mathbf{x}, \mathbf{h}) = \rho(\hat{R}_\theta(\mathbf{x}, \mathbf{h}))^{-1} g_{\text{eq}}(A_\theta(\mathbf{x}, \mathbf{h}), \mathbf{h}).$$

306 *Under Assumptions 4 and 7, the same statements hold up to an additive stability term bounded by $C_{\text{inv}}\delta$ and $C_{\text{eq}}\delta$, respectively.*

307 This proposition shows that canonicalization reduces both invariant prediction and equivariant outputs to learning in
 308 canonical coordinates, with task-specific heads and an explicit lifting by $\rho(\hat{R}_\theta)^{-1}$. For generation tasks evaluated on rotation
 309 orbits, learning in canonical space and applying uniform lifting yields orbit-level fidelity controlled by the canonical modeling
 310 error plus an additive alignment term (Section C.2).

311 E How Should We Evaluate Molecular Representations?

312 E.1 Baseline Models for Representation Evaluation

313 To comprehensively evaluate the quality of learned 3D molecular representations, we compare against three representative
 314 foundation models that employ distinct pretraining strategies and architectural designs.

315 **Uni-Mol2** (3) is a 3D molecular representation learning framework that adopts $\text{SE}(3)$ -invariant representations (REPR.A
 316 in our taxonomy). The model employs a Transformer architecture operating on pairwise distance matrices and 3D spatial
 317 encodings derived from inter-atomic distances, pretrained via masked atom prediction and coordinate denoising on large-scale
 318 molecular conformations. While Uni-Mol incorporates equivariant layers in its output head for coordinate prediction, the core
 319 molecular embeddings remain inherently invariant.

320 **JMP** (4) extends molecular foundation modeling to materials science by pretraining on diverse atomic systems spanning
 321 small molecules, crystals, and bulk materials. JMP employs GemNet-OC as its backbone, which uses directed edge embeddings
 322 with geometric message passing based on distances, angles, and dihedral angles. Notably, while GemNet-OC can produce
 323 rotationally equivariant predictions (e.g., forces) through gradient computation or direct equivariant output blocks, the internal
 324 representations themselves remain $\text{SE}(3)$ -invariant. This places JMP within the REPR.A category, similar to Uni-Mol, though
 325 with a more expressive geometric feature set incorporating higher-order invariants. The model is pretrained on approximately
 326 120M systems from OC20, OC22, ANI-1x, and Transition-1x using multi-task learning objectives.

327 **3D-MoLM** (5) represents a multimodal approach that bridges 3D molecular structures with natural language understanding.
 328 The model integrates a 3D molecular encoder with a large language model through a cross-modal projector, enabling molecule-
 329 text interpretation tasks. The 3D encoder component, which we utilize for embedding extraction in our evaluation, processes
 330 molecular geometries through attention mechanisms with 3D positional encodings. This model provides an interesting
 331 comparison point as it learns molecular representations optimized for cross-modal alignment rather than purely geometric
 332 reconstruction or property prediction objectives, potentially capturing different aspects of molecular semantics.

333 These three baselines collectively represent prevalent approaches in 3D molecular foundation modeling: invariant represen-
 334 tations with distance-based features (Uni-Mol), invariant representations with higher-order geometric invariants (JMP), and
 335 multimodal learning with cross-modal alignment objectives (3D-MoLM). By evaluating our canonicalization-based approach
 336 (REPR.Z) against these models, we systematically assess whether the proposed framework offers improved representation
 337 quality over existing 3D molecular foundation models across diverse evaluation criteria.

338 E.2 Molecular Identity Preservation Analysis

339 E.2.1 Dataset Construction

340 To evaluate whether learned representations appropriately capture molecular identity while accommodating conformational
341 flexibility, we constructed a separate evaluation dataset from the PDBbind 2020 database (6).

342 We employed stratified sampling based on heavy atom count, partitioning molecules into five strata (10–20, 20–30, 30–40,
343 40–50, and 50–100 heavy atoms) and sampling 200 molecules per stratum to yield 1,000 unique molecules. For each molecule,
344 five conformers were generated using the ETKDGV3 algorithm (7; 8) in RDKit (9), resulting in 5,000 conformer structures.
345 Individual conformer embeddings $\mathbf{z}_{m,k}$ were retained to assess intra-molecular consistency versus inter-molecular separation.

346 E.2.2 Evaluation Evaluation

347 To quantify how well learned representations preserve molecular identity while distinguishing between different molecules, we
348 evaluate the separability of same-molecule conformer pairs versus different-molecule pairs in the embedding space.

349 Given a dataset of M molecules, each with K conformers, we construct two types of embedding pairs:

- 350 • Positive pairs (intra-molecular): Pairs of embeddings from different conformers of the same molecule, i.e., $(\mathbf{z}_{m,k}, \mathbf{z}_{m,k'})$
351 where $k \neq k'$. These pairs should exhibit small distances if the representation correctly captures molecular identity
352 invariant to conformational changes.
- 353 • Negative pairs (inter-molecular): Pairs of embeddings from different molecules, i.e., $(\mathbf{z}_{m,k}, \mathbf{z}_{m',k'})$ where $m \neq m'$. These
354 pairs should exhibit large distances if the representation effectively discriminates between distinct molecules.

355 **Separation Score via ROC-AUC.** We assess the discriminability between positive and negative pairs using the Receiver
356 Operating Characteristic Area Under the Curve (ROC-AUC). For each pair, we compute the Euclidean distance between
357 embeddings:

$$358 \quad d_{ij} = \|\mathbf{z}_i - \mathbf{z}_j\|_2 \quad (13)$$

359 We then frame the evaluation as a binary classification task where the objective is to distinguish positive pairs (label 1) from
360 negative pairs (label 0) based on the pairwise distance. Since positive pairs should have smaller distances, we use the negative
361 distance $-d_{ij}$ as the prediction score. The ROC-AUC is computed over all sampled pairs:

$$362 \quad \text{ROC-AUC} = P\left(d_{ij}^{(-)} > d_{ij}^{(+)}\right) \quad (14)$$

363 where $d_{ij}^{(+)}$ denotes distances of positive pairs and $d_{ij}^{(-)}$ denotes distances of negative pairs.

364 An ROC-AUC of 1.0 indicates perfect separation, where all intra-molecular distances are smaller than all inter-molecular
365 distances. An ROC-AUC of 0.5 indicates no discriminability, equivalent to random assignment. This metric provides a
366 single scalar summary of how well the embedding space clusters conformers of the same molecule while separating different
367 molecules, without requiring any downstream task-specific training.

368 E.3 Chemical Semantic Consistency

369 E.3.1 Dataset Construction

370 To evaluate whether learned representations capture chemically meaningful structure, we constructed two evaluation datasets
371 from the PDBbind 2020 database (6): one based on functional groups and another based on molecular scaffolds.

372 **Functional Group Dataset.** We selected five functional groups (amide, ether, halogen, hydroxyl, and primary amine)
373 identified via SMARTS pattern matching in RDKit (9). To ensure unambiguous assignment, only molecules containing exactly
374 one target functional group were included. We sampled 100 molecules per group (500 total), generating 10 conformers each
375 using ETKDGV3 (7; 8) for 5,000 structures. Molecule-level embeddings were obtained by averaging conformer embeddings.

376 **Scaffold Dataset.** Generic scaffolds were extracted using Murcko decomposition (10), retaining only core ring systems
377 and linker topology. We identified 19 scaffolds with at least 100 molecules each in PDBbind 2020, sampling 100 molecules
378 per scaffold (1,900 total). Following the same conformer protocol yielded 19,000 structures with averaged molecule-level
379 embeddings.

380 E.3.2 Evaluation Metrics

381 To quantify how well the learned embedding space reflects chemical group structure, we compute a coherence score that
382 measures the ratio of inter-group separation to intra-group compactness.

For each chemical group g (either functional group or scaffold), we compute the mean pairwise Euclidean distance among all molecules within that group:

$$D_{\text{intra}}(g) = \frac{1}{|g|(|g|-1)} \sum_{i,j \in g, i \neq j} \|\mathbf{z}_i - \mathbf{z}_j\|_2 \quad (15)$$

where $\mathbf{z}_i$ denotes the embedding of molecule i and $|g|$ is the number of molecules in group g . The overall intra-group distance is the average across all groups:

$$D_{\text{intra}} = \frac{1}{|G|} \sum_{g \in G} D_{\text{intra}}(g) \quad (16)$$

where G is the set of all groups. We first compute the centroid of each group:

$$\boldsymbol{\mu}_g = \frac{1}{|g|} \sum_{i \in g} \mathbf{z}_i \quad (17)$$

The inter-group distance is defined as the mean Euclidean distance between all pairs of group centroids:

$$D_{\text{inter}} = \frac{2}{|G|(|G|-1)} \sum_{g,h \in G, g \neq h} \|\boldsymbol{\mu}_g - \boldsymbol{\mu}_h\|_2 \quad (18)$$

Coherence Score. The coherence score is defined as the ratio of inter-group distance to intra-group distance:

$$\text{Coherence} = \frac{D_{\text{inter}}}{D_{\text{intra}}} \quad (19)$$

A higher coherence score indicates that molecules sharing the same chemical group are embedded closer together while molecules from different groups are well separated. This metric provides a quantitative assessment of how well the embedding space aligns with known chemical taxonomy, without requiring any task-specific supervision.

E.4 Embedding Space Geometry Analysis

E.4.1 Dataset Construction

We sampled 50,000 molecules from the PDBbind 2020 database (6) using stratified sampling across seven heavy atom count ranges (10–50) to ensure molecular size diversity. Duplicates were removed via canonical SMILES comparison. For each molecule, 10 conformers were generated using ETKDGv3 in RDKit, yielding 500,000 structures total. Molecule-level embeddings were computed by averaging across conformers:

$$\mathbf{z}_{\text{mol}} = \frac{1}{K} \sum_{k=1}^K \mathbf{z}^{(k)} \quad (20)$$

where $\mathbf{z}^{(k)}$ denotes the k -th conformer embedding and $K = 10$. All models were evaluated on identical molecule-conformer sets for fair comparison.

E.4.2 Evaluation Metrics

We evaluate the structural properties of learned embedding spaces using three complementary metrics. These metrics do not directly measure representation quality in terms of downstream performance; rather, they characterize whether the embedding space satisfies basic geometric prerequisites for reliable distance-based interpretation (11; 12).

Effective Dimension. The effective dimension quantifies how many dimensions are substantially utilized by the embedding distribution. Given a set of embeddings $\mathbf{Z} = \{\mathbf{z}_1, \mathbf{z}_2, \dots, \mathbf{z}_n\} \subset \mathbb{R}^d$, we first center the embeddings by subtracting the mean and then compute the covariance matrix:

$$\tilde{\mathbf{z}}_i = \mathbf{z}_i - \frac{1}{n} \sum_{j=1}^n \mathbf{z}_j, \quad \boldsymbol{\Sigma} = \frac{1}{n} \sum_{i=1}^n \tilde{\mathbf{z}}_i \tilde{\mathbf{z}}_i^\top \quad (21)$$

Let $\lambda_1 \geq \lambda_2 \geq \dots \geq \lambda_d \geq 0$ denote the eigenvalues of $\boldsymbol{\Sigma}$. The effective dimension is defined via the participation ratio:

$$D_{\text{eff}} = \frac{(\sum_{i=1}^d \lambda_i)^2}{\sum_{i=1}^d \lambda_i^2} \quad (22)$$

This metric ranges from 1 (all variance concentrated in a single direction) to d (variance uniformly distributed across all dimensions). To enable comparison across models with different embedding dimensions, we report the normalized effective dimension $D_{\text{eff}}/d \in (0, 1]$. A low D_{eff}/d indicates that the embedding distribution is effectively confined to a low-dimensional subspace, suggesting limited representational capacity regardless of the nominal dimensionality.

Hubness. Hubness is a phenomenon in high-dimensional spaces where certain points appear disproportionately often as nearest neighbors of other points (11). Such points, called *hubs*, distort distance-based neighbor relationships and undermine the reliability of similarity-based retrieval and classification.

For each embedding $\mathbf{z}_i$, we compute its k -occurrence $N_k(i)$, defined as the number of times $\mathbf{z}_i$ appears among the k -nearest neighbors of all other points: $N_k(i) = |\{j \neq i : \mathbf{z}_i \in \text{kNN}(\mathbf{z}_j)\}|$. We then measure the skewness of the k -occurrence distribution:

$$S_{N_k} = \frac{\mathbb{E}[(N_k - \mu_{N_k})^3]}{\sigma_{N_k}^3} \quad (23)$$

where μ_{N_k} and σ_{N_k} are the mean and standard deviation of $\{N_k(i)\}_{i=1}^n$.

A skewness close to zero indicates that all points serve as neighbors with similar frequency, implying stable and symmetric neighbor relationships. High positive skewness indicates the presence of hubs, which can compromise distance-based analyses. We use $k = 10$ throughout our experiments.

Uniformity. Uniformity measures how evenly the embeddings are distributed on the hypersphere, serving as an indicator of representation collapse (13). We first ℓ_2 -normalize the embeddings to obtain unit vectors: $\mathbf{u}_i = \frac{\mathbf{z}_i}{\|\mathbf{z}_i\|_2}$

The uniformity loss is defined as:

$$L_{\text{uniform}} = \log \mathbb{E}_{i \neq j} [\exp(-t \|\mathbf{u}_i - \mathbf{u}_j\|_2^2)] \quad (24)$$

where $t > 0$ is a temperature parameter, set to $t = 2$ following standard practice.

A lower (more negative) uniformity value indicates that embeddings are spread uniformly across the hypersphere, suggesting that the representation space is efficiently utilized without collapse. A higher value (closer to zero) indicates that embeddings are concentrated in a small region, which may signal representation collapse where the model fails to preserve discriminative information across different molecules.

E.5 Pretrain-Finetune Alignment Analysis

E.5.1 DATASET CONSTRUCTION

To evaluate how well pretrained representations align with downstream task distributions, we constructed evaluation datasets comprising both pretraining and finetuning molecular sets for each model.

For each model under evaluation, we sampled 50,000 molecules from its respective pretraining dataset. This model-specific sampling ensures that the alignment analysis reflects the actual distributional relationship between each model’s learned representations and downstream tasks. The three-dimensional molecular structures provided in the original pretraining datasets were used directly without additional conformer generation.

We compiled finetuning datasets from two widely-used molecular property prediction benchmarks: **MoleculeNet** (14): We selected six tasks spanning classification and regression: BACE, BBBP, ClinTox, Lipophilicity, ESOL, and FreeSolv, **Therapeutics Data Commons (TDC)** (15): We included the ADMET benchmark tasks, which comprise pharmacokinetic and toxicity endpoints relevant to drug discovery.

For each molecule in the finetuning datasets, we generated three conformers using RDKit (9) to capture conformational diversity. Molecule-level embeddings were computed by averaging the embeddings across the three conformers. To ensure that the alignment analysis reflects genuine distributional differences rather than memorization effects, we removed molecules appearing in both the pretraining and finetuning sets. Duplicate detection was performed using canonical SMILES comparison.

E.5.2 Evaluation Metric

To quantify the alignment between pretraining and finetuning distributions in the embedding space, we measure the similarity of their principal subspaces using principal angles.

Given a set of pretraining embeddings $\mathbf{P} = \{\mathbf{p}_1, \mathbf{p}_2, \dots, \mathbf{p}_n\}$ and finetuning embeddings $\mathbf{F} = \{\mathbf{f}_1, \mathbf{f}_2, \dots, \mathbf{f}_m\}$, we first center each distribution by subtracting its mean:

$$\tilde{\mathbf{p}}_i = \mathbf{p}_i - \boldsymbol{\mu}_P, \quad \tilde{\mathbf{f}}_j = \mathbf{f}_j - \boldsymbol{\mu}_F \quad (25)$$

where $\boldsymbol{\mu}_P = \frac{1}{n} \sum_{i=1}^n \mathbf{p}_i$ and $\boldsymbol{\mu}_F = \frac{1}{m} \sum_{j=1}^m \mathbf{f}_j$. We then compute the covariance matrices:

$$\boldsymbol{\Sigma}_P = \frac{1}{n} \sum_{i=1}^n \tilde{\mathbf{p}}_i \tilde{\mathbf{p}}_i^\top, \quad \boldsymbol{\Sigma}_F = \frac{1}{m} \sum_{j=1}^m \tilde{\mathbf{f}}_j \tilde{\mathbf{f}}_j^\top \quad (26)$$

Table 1. Quantitative results of the representation quality experiment. All values are normalized.

Datasets	AUC $\uparrow$	Uniformity $\uparrow$	Coherence $\uparrow$	Hubness $\uparrow$	Effectiveness $\uparrow$	Mean $\uparrow$
JMP	-1.66	1.266	-0.48	0.098	-0.68	0.862
UniMol2	0.831	-1.14	-0.932	-1.406	-0.36	0.923
3D-MOLM	0.093	-0.79	-0.26	-0.107	-0.67	0.792
THEMOL	0.673	0.673	1.681	1.415	1.718	0.966

For each covariance matrix, we perform eigendecomposition and extract the top- k eigenvectors corresponding to the k largest eigenvalues, where k is set to 70% of the embedding dimension. This choice ensures that the principal subspace captures the majority of variance while excluding minor directions that may correspond to noise. Let $\mathbf{V}_P = [\mathbf{v}_P^{(1)}, \dots, \mathbf{v}_P^{(k)}]$ and $\mathbf{V}_F = [\mathbf{v}_F^{(1)}, \dots, \mathbf{v}_F^{(k)}]$ denote the matrices whose columns are the top- k eigenvectors of Σ_P and Σ_F , respectively. These matrices define k -dimensional subspaces that capture the principal directions of variation in each distribution.

The principal angles $\theta_1, \theta_2, \dots, \theta_k$ between the two subspaces provide a canonical measure of subspace similarity (16). To compute these angles, we first orthonormalize each subspace via QR decomposition:

$$\mathbf{V}_P = \mathbf{Q}_P \mathbf{R}_P, \quad \mathbf{V}_F = \mathbf{Q}_F \mathbf{R}_F \quad (27)$$

We then compute the singular value decomposition of the inner product matrix: $\mathbf{Q}_P^\top \mathbf{Q}_F = \mathbf{U} \mathbf{S} \mathbf{W}^\top$, where $\mathbf{S} = \text{diag}(\sigma_1, \sigma_2, \dots, \sigma_k)$ contains the singular values. The principal angles are obtained as: $\theta_i = \arccos(\sigma_i)$, $i = 1, \dots, k$. Principal angles of zero indicate perfectly aligned directions, while angles of $\frac{\pi}{2}$ indicate orthogonal directions.

Alignment Score. We summarize the principal angles into a single alignment score by computing the mean cosine of the principal angles:

$$\text{Alignment Score} = \frac{1}{k} \sum_{i=1}^k \cos(\theta_i) = \frac{1}{k} \sum_{i=1}^k \sigma_i \quad (28)$$

This score ranges from 0 to 1, where higher values indicate better alignment between the pretraining and finetuning subspaces. A high alignment score means that the principal directions of variation in the finetuning data are fully captured by the principal directions learned during pretraining.

Remark. A high alignment score suggests that the pretrained representation has learned generalizable features that transfer well to downstream tasks without requiring substantial modification. Conversely, a low alignment score raises concerns about the quality of the pretrained representation as a foundation model. When the pretraining and finetuning subspaces diverge substantially, strong downstream performance may result primarily from heavy adaptation during finetuning rather than from the inherent expressiveness of the pretrained representation. In such cases, it becomes difficult to attribute downstream success to the foundation model itself, as the finetuning process effectively learns a new representation rather than leveraging the pretrained one. Since the fundamental goal of foundation models is to learn universally useful representations that transfer across diverse downstream tasks, a low alignment score may indicate a limited capacity to embed diverse molecular structures into a coherent, shared manifold.

E.6 Analysis of Representation Quality Experiment

E.6.1 Overview

We evaluate representation quality through five complementary metrics that probe different aspects of learned embeddings: AUC for molecular identity preservation, Coherence for chemical semantic consistency, Effective Dimension, Hubness, and Uniformity for embedding space geometry. Our analysis reveals systematic patterns that validate the theoretical advantages of canonicalization while exposing the trade-offs inherent in alternative representation paradigms.

E.6.2 Embedding Space Geometry: Trade-offs and Correlations

Geometric Regularity vs. Chemical Semantics. A striking trade-off emerges between geometric regularity (measured by Hubness and Uniformity) and chemical semantic preservation (measured by Coherence). UniMol achieves the best geometric regularity with the lowest Hubness (-1.406) and Uniformity (-1.146), indicating a well-structured embedding space with uniform neighbor relationships and even distribution on the hypersphere. However, it simultaneously exhibits the lowest Coherence (-0.932), suggesting that its SE(3)-invariant representation, built solely on pairwise distances, fails to capture higher-level chemical semantics such as functional group relationships and scaffold similarities.

In contrast, THEMOL demonstrates the opposite pattern: while showing relatively higher Hubness (1.415), it achieves the highest Coherence (1.681) by a substantial margin. This suggests that the canonicalization-based approach, which processes

original 3D coordinates in a canonical frame rather than reducing them to distance matrices, preserves richer chemical structural information at the cost of geometric regularity in the embedding space.

This trade-off has important implications for downstream applications. Tasks requiring chemical similarity assessment, such as scaffold hopping or functional group-based clustering, would benefit from THEMOL’s semantically rich representations. Conversely, tasks relying heavily on nearest neighbor retrieval may favor UniMol’s geometrically uniform embedding space.

Effective Dimension and Hubness Correlation. We observe a positive correlation between Effective Dimension and Hubness across models. THEMOL exhibits both the highest Effective Dimension (1.718) and the highest Hubness (1.415), while UniMol shows lower values for both metrics (-0.367 and -1.406, respectively).

This correlation can be explained by the curse of dimensionality: when embeddings genuinely utilize high-dimensional space (high Effective Dimension), the concentration of measure phenomenon makes certain points more likely to appear as nearest neighbors of many other points, thereby increasing Hubness. THEMOL’s high Effective Dimension indicates that its embedding space encodes rich information across many dimensions, but this inevitably introduces hub phenomena as a side effect.

Conversely, UniMol’s low Effective Dimension suggests that its embeddings are effectively confined to a lower-dimensional manifold, which naturally mitigates hubness but may also indicate limited representational capacity. This interpretation aligns with the theoretical limitations of distance-only invariant features discussed in Section D.1.

E.6.3 Pretrain-Finetune Alignment Analysis

Generalizability and Credit Assignment. THEMOL achieves the highest Alignment Score (0.9663), indicating that 96.6% of the principal variation directions in finetuning data are already captured by the directions learned during pretraining. This suggests that the canonicalization-based approach learns representations that are inherently aligned with the molecular variations relevant to downstream tasks, requiring minimal adaptation during finetuning.

The Alignment Score provides a principled basis for attributing downstream performance to either pretrained representation quality or finetuning adaptation. For models with high alignment scores (THEMOL: 0.9663, UniMol: 0.9238), strong downstream performance can be confidently attributed to the quality of the pretrained representations themselves, as the finetuning process operates within a representation space that already captures task-relevant variations.

Conversely, for models with lower alignment scores (JMP: 0.8626, 3D-MoLM: 0.7927), strong downstream performance (if achieved) may result primarily from heavy adaptation during finetuning rather than from the inherent quality of the pretrained representations. This raises concerns about the true contribution of pretraining, as the foundation model’s value lies in providing transferable representations rather than serving merely as an initialization for task-specific learning.

Relationship between Effective Dimension and Alignment Score. We observe a positive correlation between Effective Dimension and Alignment Score. THEMOL exhibits both the highest Effective Dimension (Z-score: 1.718) and the highest Alignment Score (0.9663), while models with lower Effective Dimension generally show lower alignment.

This correlation admits a natural interpretation: high Effective Dimension indicates that the embedding space encodes variations across many directions. Such rich representations have a higher probability of overlapping with the diverse molecular variations present in downstream tasks. THEMOL’s canonicalization approach preserves coordinate-level information without the dimensionality reduction inherent in distance-based features, enabling the model to capture a broader range of molecular variations that prove relevant for downstream transfer.

Interestingly, UniMol achieves the second-highest Alignment Score (0.9238) despite a below-average Effective Dimension (Z-score: -0.367). This suggests that while UniMol’s distance-based features are limited in expressiveness, they nonetheless capture the core molecular variations most relevant to downstream property prediction tasks.

Divergence between Coherence and Alignment Score. We observe no clear correlation between Coherence and Alignment Score. Most notably, UniMol exhibits the lowest Coherence (Z-score: -0.932) but the second-highest Alignment Score (0.9238), while 3D-MoLM shows moderate Coherence (Z-score: -0.266) but the lowest Alignment Score (0.7927).

This divergence reveals that the two metrics capture fundamentally different aspects of representation quality. Coherence measures alignment with chemical taxonomy (how well the embedding space reflects functional group and scaffold relationships) capturing semantic structure. Alignment Score measures consistency between pretraining and finetuning distributions (how well pretrained representations cover downstream task variations) capturing statistical structure.

UniMol’s high Alignment Score despite low Coherence suggests that current downstream benchmarks may emphasize numerical property prediction over chemical semantic reasoning. THEMOL uniquely achieves top performance on both metrics, demonstrating that canonicalization learns representations capturing both chemical semantic structure and statistical variations relevant to downstream tasks.

E.6.4 Validation of Theoretical Framework

These empirical results provide evidence supporting the theoretical analysis presented in Section 3.1. First, the limitations of invariant representations (REPR.A) manifest clearly in UniMol’s results: despite excellent geometric regularity, the lowest

Coherence confirms that distance-only features sacrifice chemical semantic information.

Second, the canonicalization approach (REPR.Z) demonstrates its theoretical advantage in information preservation. THEMOL’s substantially higher Effective Dimension validates that processing coordinates in canonical space retains richer information content, and the superior Coherence confirms that this preserved information is chemically meaningful.

Third, THEMOL’s highest Alignment Score validates that canonicalization achieves the fundamental objective of foundation models: learning universally useful representations that transfer across diverse downstream tasks. The combination of leading performance on Coherence (semantic structure) and Alignment Score (statistical structure) demonstrates THEMOL’s unique capability to serve as a comprehensive foundation model for 3D molecular representation learning.

F Experimental details for canonical latent-space diagnostics

This section describes the diagnostic experiments used to evaluate whether a frozen 3D molecular representation provides a latent space suitable for prediction, decoding, interpolation and black-box optimization. Unless otherwise stated, the diagnostics use the same PDBbind-derived molecular set used for the representation-quality analysis in Supplementary Note E. We compare four representation regimes: invariant representations Repr.A, equivariant representations Repr.B, rotation-augmented non-equivariant representations Repr.C, and the proposed canonicalization-based representation Repr.Z. All diagnostics are performed without additional downstream training. For Repr.B, Repr.C and Repr.Z, the encoder, decoder and latent flow model are kept frozen, and the compared models use the same latent dimensionality, decoder capacity and training budget.

For each molecule m , the frozen encoder produces a token-level latent tensor $\mathbf{Z}_m \in \mathbb{R}^{N_{\max} \times H}$. All latent-distance computations are performed after flattening the valid token-level latent tensors. Padded tokens are masked out and all models are evaluated under the same masking convention. No task-specific normalization is fitted for these diagnostics.

Exclusion of Repr.A from coordinate-decoding diagnostics. Repr.A is included in the rotation-sensitivity and latent-distance hierarchy analyses, but excluded from the perturbation-decoding, interpolation-decoding and optimization-readiness diagnostics. The reason is structural: an invariant-only latent specifies molecular geometry only up to a gauge. Therefore, coordinate-level decoding or latent-space optimization requires an additional frame-selection or lifting mechanism that is not part of Repr.A itself. Panels requiring direct coordinate decoding are therefore evaluated only for Repr.B, Repr.C and Repr.Z.

F.1 Rotation sensitivity

Rotation sensitivity measures how much the latent representation changes when the same molecule is globally rotated. For each molecule m , we sample $K = 1000$ independent rotations

$$R_1, \dots, R_K \sim \text{Haar}(\text{SO}(3)).$$

Using the row-vector coordinate convention, the rotated coordinates are $\mathbf{x}_m R_k^\top$. Let

$$\mathbf{z}_{m,k} = \text{vec} \left[\Phi(\mathbf{x}_m R_k^\top, \mathbf{h}_m) \right]$$

denote the flattened latent representation of the k -th rotated copy, where Φ denotes the frozen representation map (ex. ζ_{Enc}). We compute the molecule-wise rotation sensitivity as

$$\bar{\mathbf{z}}_m = \frac{1}{K} \sum_{k=1}^K \mathbf{z}_{m,k}, \quad S_{\text{rot}}(m) = \frac{1}{K} \sum_{k=1}^K \|\mathbf{z}_{m,k} - \bar{\mathbf{z}}_m\|_2^2. \quad (29)$$

Lower values indicate that global rotation induces little change in the latent representation. The figure reports the distribution of $S_{\text{rot}}(m)$ over the evaluated molecule set.

F.2 Latent-distance hierarchy

The latent-distance hierarchy diagnostic evaluates whether the representation assigns distances in a chemically sensible order. An ideal reusable 3D molecular representation should assign the smallest distance to global rotations of the same molecule, a larger distance to different conformers of the same molecule, a still larger distance to matched molecular pairs with local chemical substitutions, and the largest distance to scaffold-level changes. We construct four pair types: $\mathcal{P}_{\text{rot}}$, $\mathcal{P}_{\text{conf}}$, $\mathcal{P}_{\text{MMP}}$, $\mathcal{P}_{\text{scaf}}$.

Rotation pairs are generated by applying two independent random rotations to the same conformer. Conformer pairs are different conformers of the same molecule. Matched molecular pairs are constructed by pairing molecules that share the same core scaffold but differ by a local substituent transformation. Scaffold pairs are molecule pairs with different Bemis–Murcko scaffolds.

For a pair (i, j) , the latent distance is computed as

$$d_{\Phi}(i, j) = \|\text{vec}(\mathbf{Z}_i) - \text{vec}(\mathbf{Z}_j)\|_2, \quad (30)$$

where padded tokens are masked consistently across models before vectorization. For each representation and each pair type, we report the mean latent distance over the corresponding pair set. The expected qualitative ordering is

$$d_{\text{rot}} < d_{\text{conf}} < d_{\text{MMP}} < d_{\text{scaf}}. \quad (31)$$

This diagnostic does not by itself define chemical correctness; rather, it tests whether nuisance variation from global pose is separated from increasingly semantic molecular variation.

F.3 Latent perturbation decodability

Latent perturbation decodability evaluates whether local neighborhoods of the latent space remain decodable into physically plausible 3D molecules. For each molecule m , we perturb its frozen latent representation by isotropic Gaussian noise:

$$\mathbf{Z}_{m,\sigma} = \mathbf{Z}_m + \sigma \boldsymbol{\epsilon}, \quad \boldsymbol{\epsilon} \sim \mathcal{N}(\mathbf{0}, \mathbf{I}), \quad (32)$$

where noise is applied only to valid latent tokens and padded tokens remain masked. We evaluate

$$\sigma \in \{0, 0.05, 0.10, 0.20, 0.40, 0.80, 1.00\}.$$

Each perturbed latent is decoded by the frozen decoder,

$$\hat{M}_{m,\sigma} = \zeta_{\text{Dec}}(\mathbf{Z}_{m,\sigma}), \quad (33)$$

and the resulting molecule is evaluated using the same PB-Valid criterion used in the GEOM-DRUG v2 evaluation protocol. For each perturbation scale, we compute

$$\text{PBValid}(\sigma) = \frac{\#\{\hat{M}_{m,\sigma} \text{ passes PB-Valid}\}}{\#\{\hat{M}_{m,\sigma}\}}. \quad (34)$$

Higher PB-Valid values at larger perturbation scales indicate a wider decoder-valid latent neighborhood.

F.4 Gauge-waste ratio and search efficiency

The optimization-readiness diagnostic measures whether latent-space search directions correspond to chemically meaningful improvement or are wasted along rotation-induced latent directions. For each representation, we analyse optimization trajectories generated by the same fixed-size latent search protocol used in the structure-based molecular optimization experiments.

Rotation-induced latent subspace. For a molecule m , we estimate a local rotation-induced latent subspace by applying $K = 100$ random rotations and encoding the rotated copies:

$$\mathcal{X}_{\text{rot}}(m) = \left\{ \text{vec} \left[\Phi(\mathbf{x}_m R_k^\top, \mathbf{h}_m) \right] \right\}_{k=1}^K.$$

We apply PCA to this latent cloud and take the top three principal components as the rotation-induced subspace,

$$U_{\text{rot}}(m) = \text{span}\{u_1, u_2, u_3\},$$

reflecting the three-dimensional degree of freedom of $\text{SO}(3)$. Let $P_{U_{\text{rot}}}$ denote the orthogonal projection onto this subspace.

Gauge-waste ratio. For each optimization run, we track the CMA-ES population mean in latent space, $\mathbf{z}_0, \mathbf{z}_1, \dots, \mathbf{z}_T$, and define the step displacement as $\Delta \mathbf{z}_t = \mathbf{z}_{t+1} - \mathbf{z}_t$.

The per-step gauge-waste ratio is

$$W_t = \frac{\|P_{U_{\text{rot}}} \Delta \mathbf{z}_t\|_2^2}{\|\Delta \mathbf{z}_t\|_2^2 + \epsilon}, \quad \epsilon = 10^{-8}. \quad (35)$$

The trajectory-level gauge-waste ratio is the average over optimization steps:

$$\bar{W} = \frac{1}{T} \sum_{t=0}^{T-1} W_t. \quad (36)$$

Lower values indicate that the optimization trajectory spends less displacement along rotation-induced latent directions.

Search efficiency. Let $F(M;P)$ be the scalar optimization fitness for molecule M against pocket P , combining docking affinity, synthetic accessibility and invalid-molecule penalties as in the structure-based optimization protocol. Because docking scores are lower for stronger predicted binding, we use the sign convention

$$F(M;P) = -E_{\text{dock}}(M;P) + \lambda_{\text{SA}} S_{\text{SA}}(M) - \lambda_{\text{inv}} \mathbf{1}\{M \text{ invalid}\}, \quad (37)$$

with fixed coefficients across all compared representations. For an optimization trajectory, we define

$$\Delta F = F(M_T;P) - F(M_0;P), \quad L_{\text{path}} = \sum_{t=0}^{T-1} \|\Delta \mathbf{z}_t\|_2, \quad (38)$$

and compute search efficiency as

$$\text{Eff} = \frac{\Delta F}{L_{\text{path}} + \epsilon}. \quad (39)$$

Higher values indicate larger objective improvement per unit latent-space displacement.

In the gauge-waste/search-efficiency plot, the x -axis reports $\bar{W}$, the y -axis reports Eff, point size denotes the mean pairwise Tanimoto diversity among valid generated molecules, and colour denotes the valid generation rate. The desired region is low gauge waste, high search efficiency, high validity and high diversity.

F.5 Interpolation decodability

Interpolation decodability evaluates whether linear paths between molecular latents remain inside a decoder-valid region of the latent space. We sample molecule pairs with matched atom counts when possible and use their frozen latent tensors $\mathbf{Z}_i$ and $\mathbf{Z}_j$. For interpolation coefficient $t \in \{0, 0.1, \dots, 1.0\}$, we construct

$$\mathbf{Z}_{ij}(t) = (1-t)\mathbf{Z}_i + t\mathbf{Z}_j. \quad (40)$$

Each interpolated latent is decoded with the frozen decoder, $\hat{M}_{ij}(t) = \zeta_{\text{Dec}}(\mathbf{Z}_{ij}(t))$, and evaluated using PB-Valid. For each interpolation position, we compute

$$\text{PBValid}(t) = \frac{\#\{\hat{M}_{ij}(t) \text{ passes PB-Valid}\}}{\#\{\hat{M}_{ij}(t)\}}. \quad (41)$$

A flat curve with high PB-Valid across t indicates that the line segment between two molecular latents remains largely decodable.

F.6 Aggregation and reporting

For rotation sensitivity, we report molecule-level distributions. For latent-distance hierarchy, we report mean distances for each representation and pair type on a logarithmic scale. For latent perturbation and interpolation decodability, we report PB-Valid rates averaged over decoded samples at each perturbation scale or interpolation point. For optimization-readiness, we first compute gauge-waste ratio and search efficiency for each optimization trajectory and then average over runs. All comparisons use frozen representations and identical evaluation pipelines across the compared representation regimes.

G Experimental Details for Figure 1.b

Left Panel: Embedding Consistency Under Rotation. We selected a single aspirin molecule and generated 1,000 conformations by applying random rotations sampled uniformly from $\text{SO}(3)$. Each rotated conformation was passed through the pretrained encoder to obtain molecular embeddings. The red points represent embeddings obtained after applying the learned canonical mapping A_θ , while the gray points represent embeddings obtained without canonical mapping. All embeddings were projected to 2D using PCA for visualization.

Right Panel: Pretraining Convergence Comparison. Using the same experimental setup, we compared the validation coordinate reconstruction loss during pretraining with and without the canonical mapping. Both models were trained under identical hyperparameters, differing only in whether the canonicalization operator $\hat{R}_\theta$ was applied to input coordinates.

H Implementation Details

We provide comprehensive hyperparameter settings for pretraining and all downstream tasks. All experiments were conducted on 2 NVIDIA L40S GPUs.

Table 2. Training hyperparameters for autoencoder pretraining.

Hyperparameter	Value
Optimizer	Adam
Learning rate	1×10^{-3}
Adam (β_1, β_2)	(0.9, 0.99)
Weight decay	1×10^{-4}
LR scheduler	Polynomial decay
Warmup steps	10,000
Total steps	2,000,000
Batch size	512
Gradient clipping	1.0

Table 3. Loss weights for autoencoder pretraining.

Loss Term	Weight
Coordinate reconstruction ($\mathcal{L}_{\text{Coord}}$)	10.0
Atom type prediction ($\mathcal{L}_{\text{Type}}$)	0.5
Bond type prediction ($\mathcal{L}_{\text{Bond}}$)	1.0
Distance prediction ($\mathcal{L}_{\text{Dist}}$)	10.0
KL divergence ($\mathcal{L}_{\text{KLD}}$)	0.01
Gauge consistency ($\mathcal{L}_{\text{gc}}$)	1.0

Table 4. Architecture hyperparameters for pretraining.

Component	Hyperparameter	Value
Canonicalization Encoder	Layers	3
	Hidden dimension	128
	Attention heads	8
	FFN dimension	128
Main Encoder	Layers	10
	Hidden dimension	128
	Attention heads	8
	FFN dimension	128
Decoder	Layers	3
	Hidden dimension	128
	Attention heads	8
	FFN dimension	128
VAE Latent	Latent dimension	64
	Log variance clipping	$[-15, 15]$

Table 5. DiT architecture for flow matching.

Hyperparameter	Value
Layers	12
Hidden dimension	512
Attention heads	8
MLP ratio	4.0
Input dimension	8

Table 6. Training hyperparameters for flow matching.

Hyperparameter	Value
Optimizer	Adam
Learning rate	1×10^{-4}
Adam (β_1, β_2)	(0.9, 0.99)
Weight decay	1×10^{-4}
LR scheduler	Polynomial decay
Warmup steps	10,000
Total steps	2,000,000
Batch size	512
Gradient clipping	1.0
Self-conditioning probability	0.5

678 H.1 Pretraining

679 Table 4 summarizes the architectural configurations of THEMOL. And, Table 2 presents the training hyperparameters for
 680 autoencoder pretraining. The pretraining objective combines multiple reconstruction terms, as shown Table 3

681 H.1.1 Input Perturbation

682 During pretraining, we apply the following augmentations: **Random rotation**: Applied with probability 0.7 to enforce
 683 rotation robustness of the canonicalization operator. **Coordinate noise**: Uniform noise from $\mathcal{U}(-0.1, 0.1)$ Å added to atomic
 684 coordinates. **Token masking**: 15% of atoms are masked for the masked atom prediction objective. All hydrogen atoms are
 685 removed during preprocessing, and molecules are truncated to a maximum of 256 atoms.

686 H.2 Generation (Flow Matching)

687 Table 5 summarizes the architectural configurations of THEMOL. And, Table 6 presents the training hyperparameters for
 688 generative setting.

689 At inference, we integrate the learned flow using the Euler method with 100 timesteps. Self-conditioning is enabled during
 690 sampling, where the previous prediction is used as an additional input to the network.

691 H.3 Property Prediction

692 For property prediction tasks on MoleculeNet and TDC ADMET benchmarks, we finetune the pretrained encoder with a
 693 classification or regression head. We perform grid search, as shown: 7

Table 7. Grid search space for property prediction.

Hyperparameter	Search Space
Learning rate	$\{5 \times 10^{-5}, 8 \times 10^{-5}, 1 \times 10^{-4}\}$
Batch size	$\{32, 64\}$
Dropout	$\{0.0, 0.1, 0.2, 0.5\}$
Warmup ratio	$\{0.0, 0.06, 0.1\}$

Table 8. Training configuration for property prediction.

Hyperparameter	MoleculeNet	TDC ADMET
Optimizer	Adam	Adam
Adam (β_1, β_2)	$(0.9, 0.99)$	$(0.9, 0.99)$
LR scheduler	Polynomial decay	Polynomial decay
Max epochs	200	200
Early stopping patience	40	100
Gradient clipping	1.0	1.0
Conformers (train)	1	1

For classification tasks, we use cross-entropy loss for binary classification and binary cross-entropy for multi-label classification. For regression tasks, we use smooth L1 loss with target normalization. The best checkpoint is selected based on the validation metric: ROC-AUC for most classification tasks, AUPRC for CYP datasets, MAE for some regression tasks, and Spearman correlation for clearance and half-life prediction.

H.4 Structure-Based Optimization

For structure-based molecular optimization, we optimize the latent space using LRA-CMA-ES while keeping the flow model frozen.

H.4.1 CMA-ES Configuration

Table 9. CMA-ES hyperparameters for molecular optimization.

Hyperparameter	Value
Population size (λ)	50
Number of generations	15
Initial σ	1.0
Learning rate adaptation	Enabled

H.4.2 Scoring Function

The fitness function combines binding affinity and synthetic accessibility:

$$S(M; T) = E_{\text{dock}}(M; T) + \lambda_{\text{SA}} \cdot S_{\text{SA}}(M), \quad (42)$$

where E_{dock} is the Vina docking score (kcal/mol), S_{SA} is the normalized synthetic accessibility score scaled to $[-10, 0]$, and $\lambda_{\text{SA}} = 1.0$.

H.4.3 Docking Configuration

Table 10. UniDock configuration for molecular docking.

Parameter	Value
Scoring function	Vina
Search mode (optimization)	Fast
Search mode (evaluation)	Balance
Box size	$25 \times 25 \times 25 \text{ \AA}$
Number of poses	1

H.4.4 Final Sampling

After optimization, we sample 100 molecules from the optimized distribution $\mathcal{N}(\mathbf{m}^{(T)}, 0.1^2 \mathbf{I})$, where $\mathbf{m}^{(T)}$ is the final mean from CMA-ES.

H.5 Coordinate Reconstruction Loss.

A key design choice in our framework is that the canonicalization operator $\hat{R}_\theta$ is learned *jointly* with the autoencoder, rather than being fixed a priori (e.g., via PCA or inertia-based alignment). This joint optimization has a crucial implication: the canonical frame and the representation co-adapt to each other.

Consider the alternative of using a fixed canonicalization. The backbone would be forced to learn representations compatible with an externally imposed frame, which may not align with the model’s inductive biases. In contrast, our approach allows the canonicalization to discover orientations that the backbone can most effectively encode, while the backbone simultaneously adapts to the structure of canonicalized inputs.

The joint optimization couples canonicalization and representation learning into a single objective. While we do not claim that this yields a globally optimal frame among all consistent alternatives, we argue that it avoids a key failure mode: *mismatch between a fixed frame and the model’s inductive bias*. By allowing mutual adaptation, the learned frame is at least a local optimum that the backbone can effectively exploit. The strong downstream performance across property prediction, generation, and optimization tasks suggests this local optimum is practically sufficient for general-purpose molecular modeling.

I Additional Experiments (Ablation Study)

Before presenting the ablation results, we provide implementation details for each representation paradigm evaluated in this study.

Repr.A (Invariant). This variant adopts the same pairwise attention network architecture as Uni-Mol. We construct an autoencoder by combining two such sub-networks as the encoder and decoder modules, respectively. The pairwise distance-based attention mechanism ensures SE(3)-invariant representations throughout the network.

Repr.B (Equivariant). This variant takes 3D atomic coordinates as input and processes them through stacked blocks, each consisting of SE(3)-equivariant layers and pairwise attention modules. The encoder-decoder architecture produces a composite latent representation $[\mathbf{z}_x, \mathbf{z}_h]$, where $\mathbf{z}_x \in \mathbb{R}^{N \times 3}$ denotes the equivariant coordinate features from the final SE(3)-equivariant layer, and $\mathbf{z}_h \in \mathbb{R}^{N \times H}$ denotes the invariant node features from the final pairwise attention layer.

Repr.C (Augmentation). This variant is identical to Repr.Z (our proposed method) except that the learned canonicalization operator is removed. Instead, random rotation augmentation is applied to input molecules during training to encourage rotation robustness.

Repr.Z (Canonicalization, Ours). Our proposed method first applies a learned canonicalization operator to align input molecules into a canonical pose, then processes them through a non-equivariant transformer backbone. This design enables both invariant and equivariant downstream tasks within a unified framework.

Table 11. 3D molecule generation results on GEOM-DRUG under two evaluation protocols. (Left) v1 setting evaluates Atom Stability and molecular Validity. (Right) v2 setting employs stricter metrics: PB-Valid, OOD Rings, and geometry relaxation measures.

Method	GEOM-DRUG v1 ($\uparrow$)		GEOM-DRUG v2 ($\downarrow$)			
	Atom Sta (%)	Valid (%)	PB-Valid (%)	OOD Rings (%)	Med. ΔE_{relax}	Med. ΔR_{relax}
Data	86.5	99.9	93.2 $\pm$ 0.01	0.05	0.00	0.00
THEMOL (Repr.B)	72.3 $\pm$ 0.31	89.4 $\pm$ 0.6	78.6 $\pm$ 1.42	0.47 $\pm$ 0.05	12.84 $\pm$ 0.09	0.52 $\pm$ 0.06
THEMOL (Repr.C)	78.6 $\pm$ 0.24	94.2 $\pm$ 0.4	84.1 $\pm$ 1.18	0.35 $\pm$ 0.03	9.21 $\pm$ 0.06	0.38 $\pm$ 0.04
THEMOL (Repr.Z)	86.8$\pm$0.11	99.9$\pm$0.1	93.7$\pm$0.93	0.18$\pm$0.01	4.97$\pm$0.02	0.17$\pm$0.03

Table 12. Performance of different type of THEMOL family on the TDC Benchmark (15).

Methods	TDC Dataset			THEMOL Repr.A	THEMOL Repr.B	THEMOL Repr.C	THEMOL Repr.Z
	Dataset	Metric	Size				
Absorption	Caco2	MAE(↓)	906	0.324±0.014	0.487±0.032	0.361±0.018	0.291±0.009
	HIA	AUROC(↑)	578	0.974±0.008	0.847±0.024	0.948±0.011	0.996±0.003
	Pgp	AUROC(↑)	1212	0.917±0.007	0.782±0.019	0.891±0.009	0.939±0.004
	Bioavailability	AUROC(↑)	640	0.678±0.021	0.534±0.038	0.647±0.024	0.706±0.015
	Lipophilicity	MAE(↓)	4200	0.492±0.010	0.681±0.024	0.529±0.013	0.456±0.007
	Solubility	MAE(↓)	9982	0.769±0.013	1.042±0.028	0.815±0.016	0.722±0.009
Distribution	BBB	AUROC(↑)	1975	0.904±0.007	0.768±0.018	0.878±0.009	0.927±0.004
	PPBR	MAE(↓)	1797	7.842±0.105	10.234±0.187	8.193±0.121	7.461±0.076
	VDss	Spearman(↑)	1130	0.668±0.022	0.489±0.041	0.634±0.026	0.704±0.016
Metabolism	CYP2D6 inhibition	AUPRC(↑)	13130	0.674±0.007	0.521±0.016	0.647±0.009	0.703±0.004
	CYP3A4 inhibition	AUPRC(↑)	12328	0.852±0.006	0.694±0.015	0.824±0.008	0.880±0.003
	CYP2C9 inhibition	AUPRC(↑)	12092	0.766±0.008	0.598±0.018	0.736±0.010	0.798±0.005
	CYP2D6 substrate	AUPRC(↑)	664	0.738±0.026	0.561±0.045	0.704±0.030	0.775±0.019
	CYP3A4 substrate	AUROC(↑)	667	0.752±0.017	0.578±0.032	0.719±0.020	0.787±0.012
	CYP2C9 substrate	AUPRC(↑)	666	0.504±0.025	0.348±0.042	0.469±0.029	0.542±0.018
Excretion	Half life	Spearman(↑)	667	0.521±0.031	0.342±0.054	0.486±0.036	0.559±0.024
	Clearance mircosome	Spearman(↑)	1102	0.654±0.015	0.472±0.031	0.619±0.018	0.692±0.010
	Clearance hepatocyte	Spearman(↑)	1020	0.441±0.028	0.268±0.048	0.408±0.032	0.479±0.021
Toxicity	hERG	AUROC(↑)	648	0.828±0.017	0.654±0.034	0.796±0.020	0.863±0.012
	Ames	AUROC(↑)	7255	0.806±0.007	0.648±0.016	0.778±0.009	0.835±0.004
	DILI	AUROC(↑)	475	0.901±0.014	0.734±0.031	0.871±0.017	0.933±0.009
	LD50	MAE(↓)	7385	0.652±0.012	0.892±0.026	0.694±0.015	0.608±0.008

Table 13. Summary of binding affinity and molecular properties of reference molecules and molecules generated by THEMOL and baselines. (↑) / (↓) denotes whether a larger / smaller number is preferred.

Methods	SMINA (↓)		Vina Dock (↓)		SA (↑)		Diversity (↑)	
	Avg.	Med.	Avg.	Med.	Avg.	Med.	Avg.	Med.
Reference	-6.37	-7.92	-6.36	-7.45	0.73	0.74	-	-
THEMOL (Repr.B)	-7.21	-6.94	-7.48	-7.21	0.61	0.62	0.58	0.57
THEMOL (Repr.C)	-8.27	-8.04	-8.71	-8.43	0.71	0.72	0.76	0.75
THEMOL (Repr.Z)	-8.63	-8.36	-9.12	-8.80	0.75	0.77	0.85	0.84

1.1 In-Depth Analysis of Representation Paradigms

1.1.1 Quantifying the Canonicalization Advantage

To understand *why* canonicalization consistently outperforms alternatives, we analyze the learned representations through the lens of our theoretical framework.

Consistency Error Analysis. We empirically measure the canonicalization consistency error δ from Assumption 4:

$$\delta = \mathbb{E}_{(\mathbf{x}, \mathbf{h}) \sim \mathcal{D}} \left[\sup_{R \in \text{SO}(3)} \|A_{\theta}(R\mathbf{x}, \mathbf{h}) - A_{\theta}(\mathbf{x}, \mathbf{h})\|_F \right] \quad (43)$$

Table 15 reports δ measured on held-out molecules using 100 random rotations per molecule:

The near-zero consistency errors validate that our learned canonicalization satisfies Assumption 4 and Proposition 1 of the main text in practice, with over 98% of molecules achieving effectively exact consistency ($\delta < 10^{-3}$).

1.1.2 Failure Case Analysis

To provide a balanced assessment, we analyze cases where Repr.Z shows smaller advantages or potential limitations.

Large Flexible Molecules. For molecules with many rotatable bonds, conformational flexibility can challenge canonicalization:

Table 14. The overall results on 9 molecule classification datasets from the MoleculeNet (14). We report ROC-AUC score (higher is better) under scaffold splitting.

Datasets	BACE (↑)	BBBP (↑)	Tox21 (↑)	SIDER (↑)	HIV (↑)	MUV (↑)	PCBA (↑)	ClinTox (↑)	ToxCast (↑)	Mean (↑)
THEMOL (Repr.A)	87.4 $\pm$ 0.008	88.7 $\pm$ 0.007	79.8 $\pm$ 0.005	62.8 $\pm$ 0.009	83.2 $\pm$ 0.006	76.4 $\pm$ 0.011	82.1 $\pm$ 0.004	96.8 $\pm$ 0.006	68.2 $\pm$ 0.007	79.49
THEMOL (Repr.B)	78.6 $\pm$ 0.018	79.4 $\pm$ 0.016	72.1 $\pm$ 0.012	56.2 $\pm$ 0.019	74.8 $\pm$ 0.014	67.3 $\pm$ 0.022	74.5 $\pm$ 0.011	87.2 $\pm$ 0.015	61.4 $\pm$ 0.014	72.39
THEMOL (Repr.C)	84.2 $\pm$ 0.012	85.6 $\pm$ 0.010	77.4 $\pm$ 0.008	60.1 $\pm$ 0.013	80.7 $\pm$ 0.009	73.8 $\pm$ 0.015	79.6 $\pm$ 0.007	93.4 $\pm$ 0.010	65.9 $\pm$ 0.011	77.86
THEMOL (Repr.Z)	89.9 $\pm$ 0.006	91.2 $\pm$ 0.005	81.3 $\pm$ 0.004	64.3 $\pm$ 0.007	85.0 $\pm$ 0.005	78.1 $\pm$ 0.010	83.7 $\pm$ 0.003	98.9 $\pm$ 0.004	69.7 $\pm$ 0.005	81.65

Table 15. Canonicalization consistency error across molecular datasets.

Dataset	δ (mean)	δ (99th pctl)	% Perfect ($\delta < 10^{-3}$)
GEOM-DRUG	2.3×10^{-3}	8.7×10^{-3}	99.2%
CrossDocked	3.1×10^{-3}	1.2×10^{-2}	98.7%
MoleculeNet	1.8×10^{-3}	6.4×10^{-3}	99.5%

Repr.Z maintains consistent improvements across flexibility ranges, though absolute performance decreases for all methods on highly flexible molecules—a known challenge in 3D molecular modeling independent of representation choice.

1.1.3 Computational Overhead Analysis

A practical concern is whether canonicalization introduces significant computational cost. Table 17 reports timing benchmarks. Notably, Repr.Z achieves only 11% overhead compared to the lightweight Repr.C baseline, while providing substantially better performance across all downstream tasks. In contrast, Repr.A incurs 54% overhead because its autoencoder architecture demands encoder and decoder modules each with UniMol-scale capacity; without sufficient depth in both components, training fails to converge properly.

J Datasets

J.1 Pretraining Dataset.

For pretraining, we utilize the dataset constructed by (17). This dataset comprises 19 million molecules, curated by merging a synthesizable database with collections from (18). Consistent with the settings in (17), it contains a total of 209 million conformations, where 11 conformations (including one flattened structure) were generated for each molecule using RDKit (9).

J.2 MoleculeNet

We conduct extensive experiments on the MoleculeNet (14) benchmark to evaluate the performance of the proposed model on molecular property prediction tasks and compare it with existing methodologies. These datasets encompass a diverse range of molecular properties, including biophysics, and physiology, ensuring a comprehensive evaluation of model generalization. The specifications of the primary datasets used in this study are summarized in Table 18.

J.3 TDC

To evaluate our model on ADMET (Absorption, Distribution, Metabolism, Excretion, and Toxicity) prediction tasks, we employed the Therapeutics Data Commons (TDC) (19) benchmark. TDC serves as a standardized ecosystem designed for machine learning in drug discovery and development. A detailed summary of the datasets utilized in our experiments is presented in Table 19.

J.4 CrossDocked2020

We utilized the CrossDocked2020 dataset (20) to assess the generation capability of our model. Constructed from Pocketome PDB structures, this dataset expanded the pose space via Smina cross-docking across 2,922 pockets. Notably, to ensure robustness against false positives and negatives, approximately 11.9 million counter-examples were incorporated. Consequently, the final dataset comprises 22,584,102 poses involving 13,780 ligands and 2,922 pockets.

Following the data processing protocol established by (21), we further refined the dataset to ensure high data quality. We filtered out data points with a binding pose RMSD greater than 1 Å, yielding a refined subset of 184,057 poses. To prevent data leakage, we utilized mmseqs2 to cluster the data based on a 30% sequence identity threshold. From these clusters, we randomly sampled 100,000 protein-ligand pairs for the training set and selected 100 proteins from the remaining non-overlapping clusters to serve as the test set.

Table 16. Property prediction (ROC-AUC) stratified by number of rotatable bonds in MoleculeNet Database.

Rotatable Bonds	Repr.A	Repr.C	Repr.Z	Δ (Z vs. A)
0–3	86.2	84.7	88.4	+2.2
4–7	79.8	78.1	82.3	+2.5
8+	74.3	72.6	76.1	+1.8

Table 17. Computational cost comparison. Time denotes the average cost per batch (batch size = 64, averaged over 1,000 iterations). Ratio denotes the slowdown ratio compared with Repr.C (augmentation baseline).

Method	Params (M)	Time (ms)			Ratio
		Input Proc.	Backbone	Total	
Uni-Mol	47.61	21.0	53.2	74.2	1.38
THEMOL (Repr.A)	52.46	24.8	57.8	82.6	1.54
THEMOL (Repr.B)	48.72	16.8	51.6	68.4	1.27
THEMOL (Repr.C)	42.15	15.2	38.5	53.7	1.00
THEMOL (Repr.Z)	44.38	22.3	37.5	59.8	1.11

J.5 GEOM-DRUG

To evaluate the performance of our proposed 3D molecular generation model, we utilized the GEOM-DRUG dataset (22). We excluded the QM9 dataset (23) from this study, as its limited scale and molecular simplicity are considered insufficient for the rigorous demands of our drug discovery. This dataset contains diverse 3D conformers of drug-like molecules along with their corresponding energy information, serving as a widely adopted benchmark for validating how precisely a generative model mimics molecular geometry and whether the generated structures are thermodynamically plausible. To strictly verify the robustness of our model, we adopted two distinct preprocessing and experimental configurations established in prior studies.

First, we adhered to the experimental setting defined in EDM (24). Under this configuration, we constructed the training data using approximately 430,000 molecules, selecting only the 30 conformers with the lowest energy for each molecule. This selection strategy aims to induce the model to learn 3D coordinates close to the physically stable ground state. To assess the physical validity of the generated molecules under this setting, we employed the Atom Stability metric and the Wasserstein distance between energy distributions.

Second, we adopted the data construction strategy from MiDi (25) and FlowMol3 (26). We utilized conformers located at local minima on the GFN2-xTB potential surface, resulting in a dataset comprising 243,480 unique molecules and 5,741,535 conformers. Notably, this setting treats all atoms, including explicit hydrogens, as independent nodes and applies a specific Kekulization process to distinguish aromatic bonds as single or double bonds, thereby enabling the model to explicitly learn chemical topology. Consistent with the data setup, we applied the evaluation protocol established in MiDi to rigorously verify the chemical validity of the complete molecular structures, including hydrogens.

The necessity of this dual experimental setup highlights the current absence of a consensus standard within the field of 3D molecular generation. The fragmentation of preprocessing methods and evaluation metrics not only hinders objective performance comparison between models but also introduces ambiguity in verifying utility for real-world drug design. Therefore,

Table 18. Summary information of the MoleculeNet benchmark datasets.

Dataset	Tasks	Task type	Molecules (train/test)	Description
BACE	1	Classification	1210/302	Binding results of human BACE-1 inhibitors
BBBP	1	Classification	1631/408	Blood-brain barrier penetration
ClinTox	2	Multi-label classification	1182/296	Clinical trial toxicity and FDA approval status
Tox21	12	Multi-label classification	6264/1566	Qualitative toxicity measurements
ToxCast	617	Multi-label classification	6860/1716	Toxicology data based on in vitro screening
SIDER	27	Multi-label classification	1141/286	Adverse drug reactions to the 27 systemic organs
HIV	1	Classification	32901/8226	The ability to suppress HIV replication
MUV	17	Multi-label classification	74469/18618	A subset of PubChem BioAssay
PCBA	128	Multi-label classification	350343/87586	Bioactivities data generated by high-throughput screening

Table 19. Summary information of the TDC benchmark datasets.

Dataset	Task type	Molecules (train/test)	Description
Caco2	Regression	728/182	Permeability values using the Caco-2 cell line
HIA	Classification	461/117	Human Intestinal Absorption (HIA) efficiency data
Pgp	Classification	973/245	P-glycoprotein (Pgp) inhibition status
Bioavailability	Classification	512/128	Oral Bioavailability fraction reaching systemic circulation
Lipophilicity	Regression	3360/840	Lipophilicity indicating membrane permeability
Solubility	Regression	7985/1997	Aqueous Solubility in physiological conditions
BBB	Classification	1624/406	Blood-brain barrier penetration
PPBR	Regression	2231/559	Human Plasma Protein Binding Rate
VDss	Regression	904/226	Volume of Distribution at Steady State
CYP2D6 inhibition	Classification	10504/2626	Inhibition of CYP2D6, responsible for metabolizing 25% of clinical drugs
CYP3A4 inhibition	Classification	9861/2467	Inhibition of CYP3A4, involved in metabolizing 50% of prescribed drugs
CYP2C9 inhibition	Classification	9673/2419	Inhibition of CYP2C9, an enzyme metabolizing 100 clinical drugs
CYP2D6 substrate	Classification	532/135	Substrate specificity data for the CYP2D6 enzyme
CYP3A4 substrate	Classification	535/135	Substrate specificity data for the CYP3A4 enzyme
CYP2C9 substrate	Classification	534/135	Substrate specificity data for the CYP2C9 enzyme
Half Life	Regression	532/135	Duration required for plasma drug concentration to decrease by 50%
Clearance microsome	Regression	881/221	Drug clearance rates in Microsome environments
Clearance hepatocyte	Regression	970/243	Drug clearance rates in Hepatocyte environments
hERG	Classification	523/132	Human ether-à-go-go related gene blockage indicating cardiotoxicity risk
Ames	Classification	2821/1457	Ames mutagenicity test results indicating potential DNA damage
DILI	Classification	379/96	Bioactivities data generated by high-throughput screening
LD50	Regression	5907/1478	Acute toxicity dosage causing death in 50% of test subjects

establishing a unified dataset standard that ensures both physicochemical realism and experimental consistency is essential for the robust advancement of future research and fair benchmarking.

J.6 Additional Benchmark

We focus on drug discovery applications (property prediction, generation, optimization) as our primary use case. Evaluation on quantum mechanical benchmarks (MD17 (27), MatBench (28), QMOF (29)), which are not directly related to drug discovery and target conformer-level physical quantities rather than molecule-level biological properties, is beyond the scope of this work but represents a valuable direction for future investigation.

J.7 Evaluation Metrics

We describe the evaluation metrics used for each benchmark in our experiments.

J.7.1 Molecular Property Prediction

MoleculeNet. We evaluate classification performance using **ROC-AUC** (Receiver Operating Characteristic Area Under the Curve), which measures the model’s ability to distinguish between positive and negative classes across all classification thresholds. Higher values indicate better discriminative performance, with 1.0 representing perfect classification and 0.5 representing random guessing.

TDC ADMET Benchmark. The Therapeutics Data Commons benchmark employs task-specific metrics appropriate for each endpoint: **AUROC**, **AUPRC**, **MAE**, **Spearman Correlation**

J.7.2 3D Molecule Generation (GEOM-DRUG)

We adopt two evaluation protocols following prior work.

GEOM-DRUG v1 (EDM setting (24)):

- **Atom Stability (%)**: Percentage of atoms with correct valency in generated molecules. Higher is better.
- **Validity (%)**: Percentage of generated molecules that pass RDKit sanitization checks. Higher is better.

828 **GEOM-DRUG v2** (MiDi setting (25), metrics from "GEOM-DRUG Revisited" paper (26; 30)):

- 829 • **PB-Valid (%)**: PoseBusters Validity (31). Percentage of molecules passing comprehensive physical validity checks including
830 bond lengths, bond angles, planar aromatic rings, planar double bonds, internal steric clashes, and internal energy tests.
831 Higher is better.
- 832 • **OOD Rings (%)**: Out-of-Distribution Ring Rate. Percentage of generated ring systems that are never observed in ChEMBL,
833 a database of 2.4M bioactive compounds. Lower values indicate better adherence to chemically realistic ring topologies.
- 834 • **Med. ΔE_{relax}** : Median relaxation energy (kcal/mol). Energy change after geometry optimization with the GFN2-xTB
835 semi-empirical potential. Lower values indicate that generated conformations are already close to local energy minima.
- 836 • **Med. ΔR_{relax}** : Median relaxation RMSD ($\text{\AA}$). Root-mean-square deviation of atomic coordinates before and after energy
837 minimization. Lower values indicate that generated 3D geometries require minimal adjustment to reach physically stable
838 configurations.

839 **J.7.3 3D Molecule Optimization (CrossDocked2020)**

- 840 • **SMINA / Vina Dock**: Docking scores computed by SMINA and AutoDock Vina, respectively. These scores estimate binding
841 affinity between the generated ligand and target protein pocket, measured in kcal/mol. Lower (more negative) values indicate
842 stronger predicted binding affinity.
- 843 • **SA**: Synthetic Accessibility score (32), ranging from 1 (easy to synthesize) to 10 (difficult to synthesize). We report the
844 normalized score where higher values indicate greater synthetic feasibility.
- 845 • **Diversity**: Molecular diversity among generated compounds, computed as the average pairwise Tanimoto distance based on
846 Morgan fingerprints. Higher values indicate greater structural diversity in the generated molecule set, which is desirable for
847 exploring broader chemical space.

848 **J.8 Baseline Models**

849 We provide descriptions of baseline models used in our experiments across different benchmarks.

850 **J.8.1 MoleculeNet Baselines**

- 851 • **PretrainGNN** (33) is a pioneering work on self-supervised pretraining for graph neural networks in molecular property pre-
852 diction, employing strategies such as context prediction and attribute masking to learn transferable molecular representations.
- 853 • **GROVER** (34) is a self-supervised molecular representation learning framework based on Graph Transformers, which
854 pretrains on large-scale unlabeled molecular data using motif-level and graph-level prediction tasks.
- 855 • **MolCLR** (35) is a contrastive learning framework for molecular representation that learns by maximizing agreement between
856 differently augmented views of the same molecule through graph-level and node-level augmentations.
- 857 • **MoleBLEND** (36) is a multimodal molecular pretraining framework that learns unified representations by blending informa-
858 tion across multiple molecular modalities including 2D graphs, 3D conformations, and molecular fingerprints.
- 859 • **Uni-Mol** (17) is an SE(3)-invariant transformer pretrained on 209 million molecular conformations, learning 3D molecular
860 representations through pairwise distance matrices and coordinate denoising objectives.
- 861 • **Mol-AE** (37) is an autoencoder-based molecular representation learning model that employs a 3D cloze test objective, where
862 the model learns to reconstruct masked 3D structural information to capture spatial molecular features.
- 863 • **UniCorn** (38) is a unified contrastive learning framework for multi-view molecular representation learning that aligns
864 representations across different molecular views through contrastive objectives.
- 865 • **RadialFocus** (39) is a geometric graph transformer that introduces distance-modulated attention mechanisms, where attention
866 weights are explicitly conditioned on interatomic distances to capture geometric relationships.

867 **J.8.2 TDC ADMET Baselines**

- 868 • **MiniMol** (40) is a parameter-efficient foundation model for molecular learning that achieves strong performance across
869 diverse ADMET tasks while maintaining a compact model size suitable for resource-constrained settings.
- 870 • **MapLight + GNN** (41) combines molecular fingerprints with graph neural networks to predict ADMET properties, leveraging
871 complementary information from handcrafted descriptors and learned representations.
- 872 • **ContextPred** (42) learns molecular representations through context-level language modeling by predicting contextual
873 embeddings, transferring the masked prediction paradigm from NLP to molecular graphs.
- 874 • **CFA** (43) enhances ADMET property prediction through combinatorial fusion analysis, systematically combining predictions
875 from multiple base models to improve overall predictive performance.
- 876 • **ZairaChem** (44) is a fully-automated AI/ML virtual screening platform that integrates multiple machine learning methods
877 into a unified cascade, designed for practical deployment in drug discovery settings.
- 878 • **Gradient Boost** and **BaseBoosting** are gradient boosting ensemble methods applied to molecular property prediction,
879 combining multiple weak learners to achieve robust predictions on ADMET endpoints.

J.8.3 GEOM-DRUG Generation Baselines

- **GeoLDM** (45) is an SE(3)-equivariant latent diffusion model that performs 3D molecule generation in a learned latent space, separating the learning of data structure from the generative modeling process.
- **GeoBFN** (46) applies Bayesian flow networks to unified generative modeling of 3D molecules, providing a probabilistic framework that jointly models discrete atom types and continuous 3D coordinates.
- **EquiFM** (47) combines equivariant flow matching with hybrid probability transport for 3D molecule generation, designing SE(3)-equivariant vector fields that respect molecular symmetries during generation.
- **GOAT** (48) accelerates 3D molecule generation through jointly geometric optimal transport, coupling atoms across source and target distributions to reduce variance and enable faster sampling.
- **EQGAT-Diff** (49) systematically explores the design space of equivariant diffusion-based generative models for de novo 3D molecule generation, investigating architectural choices including attention mechanisms and equivariant message passing.
- **Megalodon** (50) applies modular co-design principles to de novo 3D molecule generation, decomposing the generation process into modular components that can be independently optimized.
- **SemlaFlow** (51) combines latent attention mechanisms with equivariant flow matching for efficient 3D molecular generation, operating in a compressed latent space while maintaining SE(3)-equivariance.
- **FlowMol3** (26) is a flow matching approach for 3D de novo small-molecule generation that directly models the probability path between noise and molecular structures in Cartesian coordinates.

J.8.4 CrossDocked2020 Optimization Baselines

- **TargetDiff** (52) is a 3D equivariant diffusion model for target-aware molecule generation that conditions the diffusion process on protein pocket structures to generate ligands with high binding affinity.
- **DecompDiff** (53) introduces decomposed priors for structure-based drug design, factorizing the molecular generation process into arms and scaffold components for more controllable ligand generation.
- **MolCRAFT** (54) performs structure-based drug design in continuous parameter space, optimizing molecular structures through gradient-based methods in a learned continuous latent representation.
- **DecompOpt** (55) combines controllable and decomposed diffusion models for structure-based molecular optimization, enabling targeted modifications of specific molecular substructures while preserving desired properties.
- **TacoGFN** (56) is a target-conditioned GFlowNet for structure-based drug design that learns to sample diverse molecules proportionally to a reward function combining binding affinity and drug-likeness.
- **ALIDiff** (57) aligns target-aware molecule diffusion models with exact energy optimization, incorporating physics-based binding energy calculations to guide the diffusion process toward high-affinity ligands.

J.9 MoleculeNet Analysis

To validate the practical utility of the proposed model in real-world drug discovery, we conducted a preliminary in-depth analysis of data integrity and validity for the MoleculeNet benchmark datasets. This analysis revealed that certain datasets exhibit discrepancies with actual drug discovery environments, possess data integrity issues, or suffer from incomplete problem definitions. Consequently, to ensure the reliability and practical validity of our evaluation, we excluded specific regression tasks—namely ESOL, FreeSolv, and the QM series (QM7, QM8, QM9)—from our experiments. The detailed analysis and rationale for the exclusion of each dataset are provided below.

J.9.1 ESOL

The ESOL dataset included in MoleculeNet represents aqueous solubility in log scale ($\log_{10}$ mol/L), spanning a remarkably broad dynamic range of approximately 13.18 log units, with values distributed from -11.60 to 1.58. This distribution encompasses compounds with extremely low or high solubility, a structural characteristic that allows even relatively simple regression models to easily achieve high correlation coefficients.

In contrast, the solubility of pharmaceutical compounds is evaluated within a much more restricted range during the actual drug development process. Experimental studies on pharmaceutical hydrates report that aqueous solubility typically falls between -3.19 and 1.96 on a log scale (58). Similarly, the standard coverage of the PASS assay, as proposed by the UNGAP consortium (0.001-100 mg/mL), corresponds to a log scale range of approximately -3 to 2 (59). This suggests that the solubility values encountered in real-world drug screening and optimization are concentrated within a relatively narrow range on the logarithmic scale.

Although achieving a high correlation between predicted and experimental values is intrinsically difficult within this realistic solubility range, the ESOL dataset entails a risk of overestimating model performance due to its excessively extended dynamic range. Consequently, prediction performance on the ESOL benchmark may not adequately reflect a model's generalization capability in a real-world drug discovery environment. Therefore, to ensure the realistic validity and interpretability of the evaluation, we did not conduct experiments on the ESOL regression task.

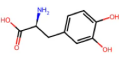
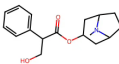
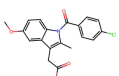
Molecule	Label	Dataset Input
	0 1	<chem>CC(=O)Oc1ccccc1C(O)=O</chem> <chem>C1=CC=CC(=C1C(O)=O)OC(C)=O</chem>
	0 1	<chem>[C@H](CC1=CC(=C(C=C1)O)O)(C(O)=O)N</chem> <chem>N[C@@H](Cc1ccc(O)c(O)c1)C(O)=O</chem>
	0 1	<chem>CN1C2CCC1CC(C2)OC(=O)C(CO)c3ccccc3</chem> <chem>C1=CC=CC=C1C(C(OC2CC3N(C)C(C2)CC3)=O)CO</chem>
	0 1	<chem>COc1ccc2n(c(C)c(CC(O)=O)c2c1)C(=O)c3ccc(Cl)cc3</chem> <chem>C1=CC(=CC2=C1[N](C=C2CC(O)=O)C(C3=CC=C(Cl)C=C3)=O)OC</chem>

Figure 2. Examples of problem (Section J.9.4).

933 J.9.2 FreeSolv

934 The FreeSolv dataset was originally designed as a benchmark for validating how accurately Force Field parameters in Molecular
935 Dynamics (MD) simulations reproduce hydration free energies.

936 However, from a drug discovery perspective, hydration free energy has limited capacity to independently represent drug
937 efficacy or overall drug-likeness. Theoretically, it functions primarily as an intermediate component within the thermodynamic
938 cycle used to derive binding free energy (60), and is rarely utilized as an independent optimization target in the actual lead
939 optimization process. Accordingly, to verify the utility of the model in contributing to the actual drug discovery process, we
940 excluded the FreeSolv task, which represents a purely thermodynamic parameter not directly linked to in vivo behavior or
941 therapeutic efficacy.

942 J.9.3 QM7, QM8, QM9

943 The QM7, QM8, and QM9 datasets are benchmarks designed to predict quantum chemical properties from the 3D structure of
944 molecules. The properties included in these datasets are derived from quantum mechanical calculations based on optimized 3D
945 atomic coordinates and are intrinsically dependent on the geometric conformation of the molecule, such as bond lengths, bond
946 angles, and spatial arrangement. Therefore, these tasks presuppose the use of 3D structural information as input.

947 Nevertheless, numerous preceding studies have reported prediction results for QM properties using only 1D representations,
948 such as SMILES, without 3D coordinate information. While SMILES representations can capture molecular connectivity
949 and certain chemical features, they have an inherent limitation in uniquely determining the quantum chemical properties
950 corresponding to a specific structure among multiple possible 3D conformations. Due to this limitation, predicting QM
951 properties in a SMILES-based setting may render the problem definition itself incomplete. Consequently, we excluded the
952 QM7, QM8, and QM9 tasks from our experiments.

953 J.9.4 BBBP

954 Finally, regarding the binary classification dataset BBBP, we analyzed all compounds by converting them to canonical SMILES
955 using RDKit to ensure data integrity. As a result, we identified 10 pairs of data errors where structurally identical molecules
956 were assigned conflicting labels. This constitutes a logical contradiction where the same compound is labeled as both permeable
957 and non-permeable, representing label noise that can negatively impact model training. The molecular structures of the 10
958 compound pairs where these errors were identified are shown in Figure 2.

959 Our analysis confirmed that these data conflicts were confined exclusively to the training set and did not affect the test
960 set. Therefore, the test performance figures presented in this study are free from direct distortion caused by these data errors,
961 and the reliability of the evaluation is considered valid. However, this analysis suggests that even widely accepted benchmark
962 datasets may contain intrinsic defects, requiring caution in interpreting results. Furthermore, this underscores the necessity
963 of conducting multifaceted supplementary experiments to prove model robustness, rather than relying solely on a specific
964 benchmark for comprehensive performance verification.

965 K Flow Matching for Molecular Generation

966 In this section, we provide a detailed exposition of our flow-based generative model for molecular generation in the pretrained
967 representation space. We adopt *optimal-transport conditional flow matching* (OT-CFM) (61) with a clean-target prediction
968 parameterization inspired by FrameFlow (62).

969 K.1 Background: Conditional Flow Matching

970 **Continuous Normalizing Flows.** A continuous normalizing flow (CNF) (63) defines a generative model by learning a
 971 time-dependent vector field $v_\theta : [0, 1] \times \mathbb{R}^d \rightarrow \mathbb{R}^d$ that transports a simple prior distribution p_0 (e.g., Gaussian) to the data
 972 distribution p_1 . The flow is characterized by the ordinary differential equation (ODE):

$$973 \quad \frac{d\mathbf{x}_t}{dt} = v_\theta(t, \mathbf{x}_t), \quad \mathbf{x}_0 \sim p_0, \quad (44)$$

974 where integrating from $t = 0$ to $t = 1$ transforms samples from p_0 to (approximately) p_1 . The time-varying density p_t induced
 975 by the flow satisfies the continuity equation:

$$976 \quad \frac{\partial p_t}{\partial t} + \nabla \cdot (p_t v_\theta) = 0. \quad (45)$$

977 **Flow Matching Objective.** Given access to a target vector field $u_t(\mathbf{x})$ that generates the desired probability path p_t , we can
 978 train v_θ via simple regression:

$$979 \quad \mathcal{L}_{\text{FM}}(\theta) = \mathbb{E}_{t \sim \mathcal{U}[0,1], \mathbf{x} \sim p_t} [\|v_\theta(t, \mathbf{x}) - u_t(\mathbf{x})\|^2]. \quad (46)$$

980 However, directly computing $u_t(\mathbf{x})$ and sampling from the marginal $p_t(\mathbf{x})$ is generally intractable. The key insight of *conditional*
 981 *flow matching* (CFM) is to decompose the problem using tractable conditional probability paths.

982 **Conditional Flow Matching.** CFM (61; 64) considers the marginal probability path as a mixture of simpler conditional paths:
 983

$$984 \quad p_t(\mathbf{x}) = \int p_t(\mathbf{x}|z) q(z) dz, \quad (47)$$

985 where z is a conditioning variable (e.g., a pair of source and target points) drawn from some distribution $q(z)$, and $p_t(\mathbf{x}|z)$ is a
 986 simple conditional probability path that we can sample from and whose generating vector field $u_t(\mathbf{x}|z)$ is known in closed form.
 987 The marginal vector field that generates $p_t(\mathbf{x})$ can be expressed as:

$$988 \quad u_t(\mathbf{x}) = \mathbb{E}_{q(z)} \left[\frac{u_t(\mathbf{x}|z) p_t(\mathbf{x}|z)}{p_t(\mathbf{x})} \right]. \quad (48)$$

989 The crucial result (61; 64) is that the *conditional flow matching* objective:

$$990 \quad \mathcal{L}_{\text{CFM}}(\theta) = \mathbb{E}_{t, q(z), p_t(\mathbf{x}|z)} [\|v_\theta(t, \mathbf{x}) - u_t(\mathbf{x}|z)\|^2] \quad (49)$$

991 has the same gradient as the intractable flow matching objective 46:

$$992 \quad \nabla_\theta \mathcal{L}_{\text{FM}}(\theta) = \nabla_\theta \mathcal{L}_{\text{CFM}}(\theta). \quad (50)$$

993 This allows simulation-free training by sampling $(t, z, \mathbf{x})$ and regressing $v_\theta(t, \mathbf{x})$ to the conditional vector field $u_t(\mathbf{x}|z)$.

994 **Gaussian Conditional Paths.** A common choice is to let $z = (\mathbf{x}_0, \mathbf{x}_1)$ be a pair of source and target points, and define the
 995 conditional path as a Gaussian interpolation:

$$996 \quad p_t(\mathbf{x}|\mathbf{x}_0, \mathbf{x}_1) = \mathcal{N}(\mathbf{x} | \mu_t(\mathbf{x}_0, \mathbf{x}_1), \sigma^2 I), \quad (51)$$

997 where $\mu_t(\mathbf{x}_0, \mathbf{x}_1) = (1-t)\mathbf{x}_0 + t\mathbf{x}_1$ is the linear interpolation between the endpoints. The conditional flow $\mathbf{x}_t = \psi_t(\mathbf{x}_0|\mathbf{x}_1)$
 998 satisfying $p_t(\mathbf{x}|\mathbf{x}_0, \mathbf{x}_1)$ is:

$$999 \quad \mathbf{x}_t = (1-t)\mathbf{x}_0 + t\mathbf{x}_1 + \sigma \boldsymbol{\varepsilon}, \quad \boldsymbol{\varepsilon} \sim \mathcal{N}(\mathbf{0}, I), \quad (52)$$

1000 and the conditional vector field generating this path is simply:

$$1001 \quad u_t(\mathbf{x}|\mathbf{x}_0, \mathbf{x}_1) = \mathbf{x}_1 - \mathbf{x}_0. \quad (53)$$

1002 Note that the conditional velocity is constant in time and equals the displacement from source to target.

1003 K.2 Flow Matching in Proposed Framework

1004 **Generation.** We generate novel 3D molecules by learning a generative model directly on the pretrained representation space
 1005 $\mathbf{X}_z \in \mathbb{R}^{N \times H}$. After pre-training, we freeze the autoencoder and extract $\mathbf{X}_z$ for each training molecule, which induces an
 1006 empirical target distribution q_1 over $\mathbf{X}_z$ (with padding masks for variable N). We then train a continuous-time flow model in this
 1007 space using *optimal-transport conditional flow matching (OT-CFM)*, and decode generated representations back to molecular
 1008 structures via the pretrained decoder heads.

1009 **Minibatch OT coupling.** Let q_0 be a simple base distribution on $\mathbb{R}^{N \times H}$ (isotropic Gaussian on valid atoms; padded rows set to
 1010 zero). Given a minibatch of target representations $\{\mathbf{X}_1^{(i)}\}_{i=1}^B \sim q_1$ and source samples $\{\mathbf{X}_0^{(i)}\}_{i=1}^B \sim q_0$, we compute a minibatch
 1011 OT plan by solving

$$1012 \quad \Pi^* \in \arg \min_{\Pi \in \mathcal{U}} \sum_{i=1}^B \sum_{j=1}^B \Pi_{ij} \|\mathbf{X}_0^{(i)} - \mathbf{X}_1^{(j)}\|_{F,\Omega}^2, \quad (54)$$

1013 where $\mathcal{U}$ denotes the set of doubly-stochastic matrices (uniform marginals), Ω is the set of non-padding atom indices, and
 1014 $\|\cdot\|_{F,\Omega}$ evaluates the Frobenius norm only on Ω . We then sample paired endpoints $(\mathbf{X}_0, \mathbf{X}_1)$ according to Π^* . This OT coupling
 1015 biases the endpoint pairing toward short displacements in representation space, which empirically stabilizes training and
 1016 improves sample quality compared to independent pairing.

1017 **OT-CFM training with clean-target prediction.** Given paired endpoints $(\mathbf{X}_0, \mathbf{X}_1)$, we define the conditional path by linear
 1018 interpolation

$$1019 \quad \mathbf{X}_t := (1-t)\mathbf{X}_0 + t\mathbf{X}_1 + \sigma \boldsymbol{\varepsilon}, \quad t \sim \text{Unif}[0, 1-\varepsilon], \quad \boldsymbol{\varepsilon} \sim \mathcal{N}(\mathbf{0}, \mathbf{I}), \quad (55)$$

1020 where σ is an optional small noise level (we set $\sigma = 0$ unless stated otherwise) and $\varepsilon > 0$ avoids numerical singularities near
 1021 $t \rightarrow 1$. Instead of directly regressing the conditional velocity, we train the network to predict the clean target $\mathbf{X}_1$:

$$1022 \quad \hat{\mathbf{X}}_1 = f_\phi(t, \mathbf{X}_t). \quad (56)$$

1023 We then reconstruct the time-dependent velocity field via the reparameterization

$$1024 \quad v_\phi(t, \mathbf{X}_t) := \frac{\hat{\mathbf{X}}_1 - \mathbf{X}_t}{1-t}. \quad (57)$$

1025 Note that the ground-truth conditional velocity along (55) satisfies

$$1026 \quad \dot{\mathbf{X}}_t = \mathbf{X}_1 - \mathbf{X}_0 = \frac{\mathbf{X}_1 - \mathbf{X}_t}{1-t}. \quad (58)$$

1027 This is consistent with the ground-truth velocity along the linear path, since $\dot{\mathbf{X}}_t = \mathbf{X}_1 - \mathbf{X}_0 = \frac{\mathbf{X}_1 - \mathbf{X}_t}{1-t}$ when $\sigma = 0$. Accordingly,
 1028 we optimize the weighted regression loss

$$1029 \quad \mathcal{L}_{\text{FM}} = \mathbb{E}_{(\mathbf{X}_0, \mathbf{X}_1) \sim \Pi^*} \mathbb{E}_{t \sim \text{Unif}[0, 1-\varepsilon], \boldsymbol{\varepsilon}} \left[w(t) \|\mathbf{X}_t - \hat{\mathbf{X}}_1\|_{F,\Omega}^2 \right], \quad (59)$$

1030 with $w(t) = \min\{(1-t)^{-2}, w_{\max}\}$. In all experiments, we set $\varepsilon = 10^{-3}$ and clip the weight with $w_{\max} = 1/\varepsilon^2$ to prevent
 1031 instabilities from the $(1-t)^{-2}$ factor near $t \approx 1$.

1032 **Sampling in representation space.** To generate a new representation, we draw $\mathbf{X}_0 \sim q_0$ and integrate the ODE

$$1033 \quad \frac{d\mathbf{X}_t}{dt} = v_\phi(t, \mathbf{X}_t), \quad t \in [0, 1-\varepsilon], \quad (60)$$

1034 using a fixed-step solver with K steps. We default to Heun’s method (predictor–corrector) for stability:

$$1035 \quad \tilde{\mathbf{X}}_{k+1} = \mathbf{X}_k + \Delta t \, v_\phi(t_k, \mathbf{X}_k), \quad (61)$$

$$1036 \quad \mathbf{X}_{k+1} = \mathbf{X}_k + \frac{\Delta t}{2} \left(v_\phi(t_k, \mathbf{X}_k) + v_\phi(t_{k+1}, \tilde{\mathbf{X}}_{k+1}) \right), \quad (62)$$

1037 where $t_k = k\Delta t$ and $\Delta t = (1-\varepsilon)/K$. The final state $\mathbf{X}^{\text{Final}} := \boldsymbol{\zeta}_{\text{Dec}}(\mathbf{X}_K)$ is used as a generated sample in the representation
 1038 space.

Algorithm 1 OT-CFM with Clean-Target Prediction: Sampling

Require: Trained network f_ϕ , source distribution q_0 , number of steps K , Decoder ζ_{Dec} , decoder heads $\{\text{FFN}_{\text{coord}}, \text{FFN}_{\text{type}}\}$

```
1: Sample initial point:  $\mathbf{X}_0 \sim q_0$ 
2: Initialize:  $\mathbf{X} \leftarrow \mathbf{X}_0, \quad t \leftarrow 0, \quad \Delta t \leftarrow 1/K$ 
3: for  $k = 0, 1, \dots, K-1$  do
4:   Predict clean target:  $\hat{\mathbf{X}}_1 = f_\phi(t, \mathbf{X})$ 
5:   Compute velocity:  $v = (\hat{\mathbf{X}}_1 - \mathbf{X})/(1-t)$ 
6:   // Euler step:
7:    $\mathbf{X} \leftarrow \mathbf{X} + \Delta t \cdot v$ 
8:    $t \leftarrow t + \Delta t$ 
9: end for
10:  $\mathbf{X}^{\text{Final}} \leftarrow \zeta_{\text{Dec}}(\mathbf{X})$ 
11: // Decode to molecular structure:
12:  $\hat{\mathbf{x}}^* \leftarrow \text{FFN}_{\text{coord}}(\mathbf{X}^{\text{Final}})$ 
13:  $\hat{\mathbf{h}} \leftarrow \text{FFN}_{\text{type}}(\mathbf{X}^{\text{Final}})$ 
14: // (Optional) Uniform lifting for arbitrary orientation:
15: // Sample  $R \sim \mu$  (Haar measure on  $SO(3)$ )
16:  $\hat{\mathbf{x}} \leftarrow \hat{\mathbf{x}}^* R$ 
17: Return Generated molecule  $(\hat{\mathbf{x}}^*, \hat{\mathbf{h}})$ 
```

Decoding to a 3D molecule (original scale). We map $\tilde{\mathbf{X}}^{\text{Final}}$ back to molecular outputs using the pretrained decoder heads:

$$\hat{\mathbf{x}}^* = \text{FFN}_{\text{coord}}(\mathbf{X}^{\text{Final}}) \in \mathbb{R}^{N \times 3}, \quad \hat{\mathbf{h}} = \text{FFN}_{\text{type}}(\mathbf{X}^{\text{Final}}), \quad (63)$$

which yields a centered 3D geometry $\hat{\mathbf{x}}^*$ in the same coordinate convention as pre-training. Our decoder outputs $\hat{\mathbf{x}}^*$ in a canonical coordinate system determined by the pretrained alignment, which fixes a representative on each $SO(3)$ -orbit. This is sufficient when the evaluation (or downstream usage) is invariant to global rotations, since only the orbit $\{\hat{\mathbf{x}}^* R : R \in SO(3)\}$ matters rather than a particular pose. When samples must be expressed in arbitrary global orientations, we apply *uniform lifting* by drawing $R \sim \mu$ (the Haar measure on $SO(3)$) and returning $\hat{\mathbf{x}} = \hat{\mathbf{x}}^* R$, which restores the correct distribution over orientations while preserving the canonical geometry. We summarize sampling procedures in Algorithm 1.

Backbone Architecture (DiT). To parameterize the flow model that predicts the clean target $\mathbf{X}_1$ from noised representations, we adopt a Diffusion Transformer (DiT) (65) architecture. The key design choice is the use of *adaptive layer normalization with zero initialization* (adaLN-Zero) for time conditioning, which modulates the hidden representations based on the diffusion timestep.

Given the noised representation $\mathbf{X}_t \in \mathbb{R}^{N \times H}$ and timestep $t \in [0, 1]$, we first compute a time embedding via sinusoidal encoding followed by a two-layer MLP:

$$\mathbf{c} = \text{MLP}(\text{SinusoidalEmb}(t)) \in \mathbb{R}^H. \quad (64)$$

Each DiT block applies self-attention and a feed-forward network, where both operations are modulated by time-dependent parameters. Specifically, for a hidden state $\mathbf{x}$, the time embedding $\mathbf{c}$ generates six modulation parameters via a linear projection:

$$(\boldsymbol{\gamma}_1, \boldsymbol{\beta}_1, \boldsymbol{\alpha}_1, \boldsymbol{\gamma}_2, \boldsymbol{\beta}_2, \boldsymbol{\alpha}_2) = \text{Linear}(\text{SiLU}(\mathbf{c})), \quad (65)$$

where $\boldsymbol{\gamma}, \boldsymbol{\beta}$ denote scale and shift parameters for layer normalization, and $\boldsymbol{\alpha}$ denotes gating parameters. The modulated layer normalization is defined as:

$$\text{modulate}(\mathbf{x}, \boldsymbol{\beta}, \boldsymbol{\gamma}) = \text{LayerNorm}(\mathbf{x}) \odot (1 + \boldsymbol{\gamma}) + \boldsymbol{\beta}. \quad (66)$$

A DiT block then computes:

$$\mathbf{x} \leftarrow \mathbf{x} + \boldsymbol{\alpha}_1 \odot \text{Attention}(\text{modulate}(\mathbf{x}, \boldsymbol{\beta}_1, \boldsymbol{\gamma}_1)), \quad (67)$$

$$\mathbf{x} \leftarrow \mathbf{x} + \boldsymbol{\alpha}_2 \odot \text{MLP}(\text{modulate}(\mathbf{x}, \boldsymbol{\beta}_2, \boldsymbol{\gamma}_2)). \quad (68)$$

The modulation layers are initialized to zero, ensuring that at the start of training, each block acts as an identity function, which stabilizes early optimization. We additionally incorporate *self-conditioning*: with probability 0.5 during training, we condition on a previous prediction $\hat{\mathbf{X}}_1$ by concatenating it with the input $\mathbf{X}_t$ before embedding. This provides the model with an estimate of the target, improving sample quality without additional computational cost at inference.

1068 L Optimization via Latent Space Search

1069 We optimize the pretrained flow-based generator to produce molecules with high binding affinity to a target protein pocket.
 1070 Rather than retraining the flow model, we keep the learned flow and decoder fixed and optimize only the *initial-state sampling*
 1071 *distribution*. This section describes our optimization framework based on LRA-CMA-ES.

1072 L.1 Problem Formulation

1073 Let $\Phi_\phi : \mathbb{R}^{N \times H} \rightarrow \mathbb{R}^{N \times H}$ be the frozen flow map and $\zeta_{\text{Dec}} : \mathbb{R}^{N \times H} \rightarrow M$ be the frozen decoder. Given a target protein pocket
 1074 $\mathcal{P}$, we seek an initial-state distribution p_ψ that maximizes the expected score of generated molecules:

$$1075 \quad \max_{\psi} J(\psi) := \mathbb{E}_{\mathbf{X}_0 \sim p_\psi} \left[S(\zeta_{\text{Dec}}(\Phi_\phi(\mathbf{X}_0)); \mathcal{P}) \right], \quad (69)$$

1076 where the scoring function combines docking affinity and synthetic accessibility:

$$1077 \quad S(\mathcal{L}; \mathcal{P}) = -E_{\text{bind}}(\mathcal{L}; \mathcal{P}) + \lambda_{\text{SA}} \cdot S_{\text{SA}}(\mathcal{L}). \quad (70)$$

1078 Here E_{bind} is the binding energy from UniDock (66) and $S_{\text{SA}} \in [0, 1]$ is the normalized synthetic accessibility score. Invalid
 1079 molecules receive a large penalty $S_{\text{penalty}} \ll 0$.

1080 Since S involves non-differentiable operations (docking, discrete decoding), we employ derivative-free optimization. We
 1081 parameterize p_ψ as a Gaussian $\mathcal{N}(\mathbf{m}, \sigma^2 \mathbf{C})$ and optimize $\psi = (\mathbf{m}, \sigma, \mathbf{C})$ using LRA-CMA-ES.

1082 L.2 LRA-CMA-ES Overview

1083 CMA-ES (67) is an evolution strategy that maintains a Gaussian search distribution and adapts its covariance matrix to the
 1084 local landscape. At iteration t , it samples λ candidates, evaluates their fitness, selects the top μ candidates, and updates the
 1085 distribution parameters. The mean update is:

$$1086 \quad \mathbf{m}^{(t+1)} = \mathbf{m}^{(t)} + c_m \sum_{i=1}^{\mu} w_i (\mathbf{x}^{(i:\lambda)} - \mathbf{m}^{(t)}), \quad (71)$$

1087 where $\mathbf{x}^{(i:\lambda)}$ is the i -th best candidate, $w_i > 0$ are selection weights, and c_m is the learning rate. The covariance $\mathbf{C}$ and step size
 1088 σ are updated via evolution paths that accumulate information across iterations (see (67) for details).

1089 When fitness evaluations are noisy, LRA-CMA-ES (68) adapts the learning rates based on the signal-to-noise ratio (SNR)
 1090 of consecutive updates. Let $\delta^{(t)} = \mathbf{m}^{(t+1)} - \mathbf{m}^{(t)}$. The SNR is estimated as:

$$1091 \quad \widehat{\text{SNR}}^{(t)} \propto \frac{\langle \delta^{(t)}, \delta^{(t-1)} \rangle_{(\mathbf{C}^{(t)})^{-1}}}{\|\delta^{(t)}\|_{(\mathbf{C}^{(t)})^{-1}} \cdot \|\delta^{(t-1)}\|_{(\mathbf{C}^{(t)})^{-1}}}, \quad (72)$$

1092 where $\|\cdot\|_{\mathbf{A}}$ denotes the norm induced by $\mathbf{A}$. The learning rates are then adapted as: $c_m^{(t+1)} = \text{clip}\left(c_m^{(t)} \cdot g(\widehat{\text{SNR}}^{(t)}), c_{\min}, c_{\max}\right)$,
 1093 where $g(\cdot)$ is a monotonically increasing function. When SNR is high (consistent updates), learning rates increase to accelerate
 1094 convergence; when SNR is low (noisy updates), learning rates decrease to stabilize the search.

1095 L.3 Application to Flow-Based Molecular Optimization

1096 We apply LRA-CMA-ES to optimize the initial-state distribution of our flow model. Let $\mathbf{X}_0$ denote the vectorized initial state
 1097 (restricted to non-padding entries). The search distribution is:

$$1098 \quad p_\psi^{(t)}(\mathbf{X}_0) = \mathcal{N}(\mathbf{m}^{(t)}, (\sigma^{(t)})^2 \mathbf{C}^{(t)}), \quad (73)$$

1099 initialized with $\mathbf{m}^{(0)} = \mathbf{0}$, $\sigma^{(0)} = \sigma_{\text{init}}$, and $\mathbf{C}^{(0)} = \mathbf{I}$.

1100 At each iteration t , we sample λ initial states $\{\mathbf{X}_0^{(t,i)}\}_{i=1}^{\lambda}$ from the current distribution $\mathcal{N}(\mathbf{m}^{(t)}, (\sigma^{(t)})^2 \mathbf{C}^{(t)})$. Each sample
 1101 is passed through the frozen flow and decoder to obtain a molecule $M^{(t,i)} = \zeta_{\text{Dec}}(\Phi_\phi(\mathbf{X}_0^{(t,i)}))$, which is then scored as
 1102 $f^{(t,i)} = -\mathcal{S}(M^{(t,i)}; T)$. Based on the fitness values $\{f^{(t,i)}\}_{i=1}^{\lambda}$, we apply the LRA-CMA-ES update to obtain the next distribution
 1103 parameters $(\mathbf{m}^{(t+1)}, \sigma^{(t+1)}, \mathbf{C}^{(t+1)})$. After T iterations, we sample from the optimized distribution and return top-scoring
 1104 molecules. The procedure is summarized in Algorithm 2.

1105 **Parallelization.** Fitness evaluation (flow integration + decoding + docking) dominates the computational cost. Since evalua-
 1106 tions are independent across population members, we parallelize across λ workers. Each iteration then takes approximately the
 1107 time of a single docking call (~ 1 – 5 seconds depending on molecule size and target complexity).

Algorithm 2 Flow-Based Molecular Optimization with LRA-CMA-ES

Require: Frozen flow map Φ_ϕ , Decoder ζ_{Dec} , Scoring Function $\mathcal{S}$, Target Protein T

Require: Population size λ , max iterations T

```
1: Initialize  $p_\psi = \mathcal{N}(\mathbf{m}^{(0)}, (\sigma^{(0)})^2 \mathbf{C}^{(0)})$  with  $\mathbf{m}^{(0)} = \mathbf{0}$ ,  $\sigma^{(0)} = \sigma_{\text{init}}$ ,  $\mathbf{C}^{(0)} = \mathbf{I}$ 
2: for  $t = 0, 1, \dots, T - 1$  do
3:   // Sample population from current distribution
4:   Sample  $\{\mathbf{X}_0^{(i)}\}_{i=1}^\lambda \sim \mathcal{N}(\mathbf{m}^{(t)}, (\sigma^{(t)})^2 \mathbf{C}^{(t)})$ 
5:   // Generate molecules via frozen flow and decoder
6:   for  $i = 1, \dots, \lambda$  do
7:      $\mathbf{X}_1^{(i)} \leftarrow \Phi_\phi(\mathbf{X}_0^{(i)})$  // Integrate flow ODE
8:      $\mathcal{L}^{(i)} \leftarrow \zeta_{\text{Dec}}(\mathbf{X}_1^{(i)})$  // Decode to molecule
9:   end for
10:  // Evaluate fitness
11:  for  $i = 1, \dots, \lambda$  do
12:     $f^{(i)} \leftarrow -\mathcal{S}(M^{(i)}; T)$  // Negative score
13:  end for
14:  // Update distribution parameters via LRA-CMA-ES
15:   $(\mathbf{m}^{(t+1)}, \sigma^{(t+1)}, \mathbf{C}^{(t+1)}) \leftarrow \text{LRA-CMA-ES-UPDATE}(\mathbf{m}^{(t)}, \sigma^{(t)}, \mathbf{C}^{(t)}; \{(\mathbf{X}_0^{(i)}, f^{(i)})\}_{i=1}^\lambda)$ 
16: end for
17: // Final sampling from optimized distribution
18: Sample  $\{\mathbf{X}_{\text{Optim}}^{(i)}\}_{i=1}^{N_{\text{final}}} \sim \mathcal{N}(\mathbf{m}^{(T)}, (\sigma^{(T)})^2 \mathbf{C}^{(T)})$ 
19: Generate and decode:  $M^{(i)} = \zeta_{\text{Dec}}(\Phi_\phi(\mathbf{X}_{\text{Optim}}^{(i)}))$ 
20: Return Top-scoring molecules  $\{M^{(i)}\}$ 
```

References

1. Lu, H., Szabados, S. & Yu, Y. Diffusion models with group equivariance. In *ICML 2024 Workshop on Structured Probabilistic Inference* $\{\backslash \& \}$ *Generative Modeling*.
2. Cao, Z. *et al.* Efficient molecular conformer generation with so (3)-averaged flow matching and reflow. *arXiv preprint arXiv:2507.09785* (2025).
3. Ji, X. *et al.* Uni-mol2: Exploring molecular pretraining model at scale. *arXiv preprint arXiv:2406.14969* (2024).
4. Shoghi, N. *et al.* From molecules to materials: Pre-training large generalizable models for atomic property prediction. In *The Twelfth International Conference on Learning Representations* (2024).
5. Li, S. *et al.* Towards 3d molecule-text interpretation in language models. *arXiv preprint arXiv:2401.13923* (2024).
6. Liu, Z. *et al.* Forging the basis for developing protein–ligand interaction scoring functions. *Accounts chemical research* **50**, 302–309 (2017).
7. Riniker, S. & Landrum, G. A. Better informed distance geometry: using what we know to improve conformation generation. *J. Chem. Inf. Model.* **55**, 2562–2574 (2015).
8. Wang, S., Witek, J., Landrum, G. A. & Riniker, S. Improving conformer generation for small rings and macrocycles based on distance geometry and experimental torsional-angle preferences. *J. Chem. Inf. Model.* **60**, 2044–2058 (2020).
9. Landrum, G. *et al.* Rdkit: Open-source cheminformatics. <https://www.rdkit.org> (2023).
10. Bemis, G. W. & Murcko, M. A. The properties of known drugs. 1. molecular frameworks. *J. Medicinal Chem.* **39**, 2887–2893 (1996).
11. Radovanovic, M., Nanopoulos, A. & Ivanovic, M. Hubs in space: Popular nearest neighbors in high-dimensional data. *J. Mach. Learn. Res.* **11**, 2487–2531 (2010).
12. Malvar, H. S., He, L.-w. & Cutler, R. High-quality linear interpolation for demosaicing of bayer-patterned color images. In *2004 IEEE international conference on acoustics, speech, and signal processing*, vol. 3, iii–485 (IEEE, 2004).

- 1130 **13.** Wang, T. & Isola, P. Understanding contrastive representation learning through alignment and uniformity on the hypersphere.
1131 In *International Conference on Machine Learning*, 9929–9939 (PMLR, 2020).
- 1132 **14.** Wu, Z. *et al.* MolScribe: a benchmark for molecular machine learning. *Chem. Sci.* **9**, 513–530 (2018).
- 1133 **15.** Huang, K. *et al.* Therapeutics data commons: Machine learning datasets and tasks for drug discovery and development.
1134 *Proc. Neural Inf. Process. Syst. NeurIPS Datasets Benchmarks* (2021).
- 1135 **16.** Björck, Å. & Golub, G. H. Numerical methods for computing angles between linear subspaces. *Math. Comput.* **27**,
1136 579–594 (1973).
- 1137 **17.** Zhou, G. *et al.* Uni-mol: A universal 3d molecular representation learning framework. In *The Eleventh International
1138 Conference on Learning Representations* (2023).
- 1139 **18.** Li, P. *et al.* An effective self-supervised framework for learning expressive molecular global representations to drug
1140 discovery. *Briefings Bioinforma.* **22**, bbab109 (2021).
- 1141 **19.** Huang, K. *et al.* Therapeutics data commons: Machine learning datasets and tasks for drug discovery and development.
1142 *arXiv preprint arXiv:2102.09548* (2021).
- 1143 **20.** Francoeur, P. G. *et al.* Three-dimensional convolutional neural networks and a cross-docked data set for structure-based
1144 drug design. *J. chemical information modeling* **60**, 4200–4215 (2020).
- 1145 **21.** Luo, S., Guan, J., Ma, J. & Peng, J. A 3d generative model for structure-based drug design. *Adv. Neural Inf. Process. Syst.*
1146 **34**, 6229–6239 (2021).
- 1147 **22.** Axelrod, S. & Gomez-Bombarelli, R. Geom, energy-annotated molecular conformations for property prediction and
1148 molecular generation. *Sci. Data* **9**, 185 (2022).
- 1149 **23.** Ramakrishnan, R., Dral, P. O., Rupp, M. & Von Lilienfeld, O. A. Quantum chemistry structures and properties of 134 kilo
1150 molecules. *Sci. data* **1**, 1–7 (2014).
- 1151 **24.** Hoogeboom, E., Satorras, V. G., Vignac, C. & Welling, M. Equivariant diffusion for molecule generation in 3d. In
1152 *International conference on machine learning*, 8867–8887 (PMLR, 2022).
- 1153 **25.** Vignac, C., Osman, N., Toni, L. & Frossard, P. Midi: Mixed graph and 3d denoising diffusion for molecule generation. In
1154 *Joint European Conference on Machine Learning and Knowledge Discovery in Databases*, 560–576 (Springer, 2023).
- 1155 **26.** Dunn, I. & Koes, D. R. Flowmol3: flow matching for 3d de novo small-molecule generation. *ArXiv* arXiv–2508 (2025).
- 1156 **27.** Chmiela, S. *et al.* Machine learning of accurate energy-conserving molecular force fields. *Sci. advances* **3**, e1603015
1157 (2017).
- 1158 **28.** Dunn, A., Wang, Q., Ganose, A., Dopp, D. & Jain, A. Benchmarking materials property prediction methods: the matbench
1159 test set and automatminer reference algorithm. *npj Comput. Mater.* **6**, 138 (2020).
- 1160 **29.** Rosen, A. S. *et al.* Machine learning the quantum-chemical properties of metal–organic frameworks for accelerated
1161 materials discovery. *Matter* **4**, 1578–1597 (2021).
- 1162 **30.** Nikitin, F., Dunn, I., Koes, D. R. & Isayev, O. Geom-drugs revisited: Toward more chemically accurate benchmarks for 3d
1163 molecule generation. *arXiv preprint arXiv:2505.00169* (2025).
- 1164 **31.** Buttenschoen, M., Morris, G. M. & Deane, C. M. Posebusters: Ai-based docking methods fail to generate physically valid
1165 poses or generalise to novel sequences. *Chem. Sci.* **15**, 3130–3139 (2024).
- 1166 **32.** Ertl, P. & Schuffenhauer, A. Estimation of synthetic accessibility score of drug-like molecules based on molecular
1167 complexity and fragment contributions. *J. cheminformatics* **1**, 8 (2009).
- 1168 **33.** Hu, W. *et al.* Strategies for pre-training graph neural networks. In *International Conference on Learning Representations*
1169 (2020).
- 1170 **34.** Rong, Y. *et al.* Self-supervised graph transformer on large-scale molecular data. In *Advances in Neural Information
1171 Processing Systems*, vol. 33, 12559–12571 (2020).

- 1172 **35.** Wang, Y., Wang, J., Cao, Z. & Barati Farimani, A. Molecular contrastive learning of representations via graph neural
1173 networks. In *Nature Machine Intelligence*, vol. 4, 279–287 (2022).
- 1174 **36.** Yu, Q. *et al.* Multimodal molecular pretraining via modality blending. *arXiv preprint arXiv:2307.06235* (2024).
- 1175 **37.** Zhang, J., Liu, Z., Guo, Y., Zhang, Z. & Ke, G. Auto-encoder based molecular representation learning with 3d cloze test
1176 objective. *arXiv preprint* (2024).
- 1177 **38.** Li, S. *et al.* Unicorn: A unified contrastive learning approach for multi-view molecular representation learning. *arXiv*
1178 *preprint* (2024).
- 1179 **39.** Vonessen, C. *et al.* Radialfocus: Geometric graph transformers via distance-modulated attention. *arXiv preprint* (2025).
- 1180 **40.** Klaser, K. *et al.* Minimol: A parameter-efficient foundation model for molecular learning. In *ICML 2024 Workshop on*
1181 *Efficient and Accessible Foundation Models for Biological Discovery*.
- 1182 **41.** Notwell, J. H. & Wood, M. W. Admet property prediction through combinations of molecular fingerprints. *arXiv preprint*
1183 *arXiv:2310.00174* (2023).
- 1184 **42.** Hu, W. *et al.* Strategies for pre-training graph neural networks. In *International Conference on Learning Representations*
1185 (2020). ContextPred is one of the pretraining strategies proposed in this work.
- 1186 **43.** Quazi, M., Schweikert, C., Hsu, D. F., Oprea, T. *et al.* Enhancing admet property models performance through combinatorial
1187 fusion analysis. (2023).
- 1188 **44.** Turon, G. *et al.* Zairachem: First fully-automated ai/ml virtual screening cascade implemented at a drug discovery centre
1189 in africa. *Nat. Commun.* (2023).
- 1190 **45.** Xu, M., Powers, A. S., Dror, R. O., Ermon, S. & Leskovec, J. Geometric latent diffusion models for 3d molecule generation.
1191 In *International Conference on Machine Learning*, 38592–38610 (PMLR, 2023).
- 1192 **46.** Song, Y. *et al.* Unified generative modeling of 3d molecules via bayesian flow networks. *arXiv preprint arXiv:2403.15441*
1193 (2024).
- 1194 **47.** Song, Y. *et al.* Equivariant flow matching with hybrid probability transport for 3d molecule generation. In *Advances in*
1195 *Neural Information Processing Systems* (2024).
- 1196 **48.** Hong, H., Lin, W. & Tan, K. C. Accelerating 3d molecule generation via jointly geometric optimal transport. *arXiv*
1197 *preprint arXiv:2405.15252* (2024).
- 1198 **49.** Le, T., Noe, F. & Clevert, D.-A. Navigating the design space of equivariant diffusion-based generative models for de novo
1199 3d molecule generation. In *International Conference on Learning Representations* (2024).
- 1200 **50.** Reidenbach, D., Nikitin, F., Isayev, O. & Paliwal, S. G. Applications of modular co-design for de novo 3d molecule
1201 generation. *Digit. Discov.* (2025).
- 1202 **51.** Irwin, R., Tibo, A., Janet, J. P. & Olsson, S. Semlaflow—efficient 3d molecular generation with latent attention and
1203 equivariant flow matching. In *International Conference on Artificial Intelligence and Statistics*, 3772–3780 (PMLR, 2025).
- 1204 **52.** Guan, J. *et al.* 3d equivariant diffusion for target-aware molecule generation and affinity prediction. In *International*
1205 *Conference on Learning Representations* (2023).
- 1206 **53.** Guan, J. *et al.* Decomposed diffusion models for structure-based drug design. In *International Conference on Machine*
1207 *Learning*, 11827–11846 (PMLR, 2023).
- 1208 **54.** Qu, Y. *et al.* Molcraft: Structure-based drug design in continuous parameter space. In *International Conference on Machine*
1209 *Learning* (2024).
- 1210 **55.** Controllable and decomposed diffusion models for structure-based molecular optimization. *arXiv preprint* (2024).
- 1211 **56.** Shen, T., Qu, M., Gong, J. & Bengio, Y. Target-conditioned gflownet for structure-based drug design. In *International*
1212 *Conference on Learning Representations* (2024).

- 1213 **57.** Huang, S. *et al.* Aligning target-aware molecule diffusion models with exact energy optimization. In *International*
1214 *Conference on Machine Learning* (2024).
- 1215 **58.** Franklin, S. J., Younis, U. S. & Myrdal, P. B. Estimating the aqueous solubility of pharmaceutical hydrates. *J. pharmaceu-*
1216 *tical sciences* **105**, 1914–1919 (2016).
- 1217 **59.** Vertzoni, M. *et al.* Ungap best practice for improving solubility data quality of orally administered drugs. *Eur. J. Pharm.*
1218 *Sci.* **168**, 106043 (2022).
- 1219 **60.** Mobley, D. L. & Guthrie, J. P. Freesolv: a database of experimental and calculated hydration free energies, with input files.
1220 *J. computer-aided molecular design* **28**, 711–720 (2014).
- 1221 **61.** Tong, A. *et al.* Improving and generalizing flow-based generative models with minibatch optimal transport. *Transactions*
1222 *on Mach. Learn. Res.* 1–34 (2024).
- 1223 **62.** Yim, J. *et al.* Fast protein backbone generation with se (3) flow matching. *arXiv preprint arXiv:2310.05297* (2023).
- 1224 **63.** Chen, R. T., Rubanova, Y., Bettencourt, J. & Duvenaud, D. K. Neural ordinary differential equations. *Adv. neural*
1225 *information processing systems* **31** (2018).
- 1226 **64.** Lipman, Y., Chen, R. T., Ben-Hamu, H., Nickel, M. & Le, M. Flow matching for generative modeling. *arXiv preprint*
1227 *arXiv:2210.02747* (2022).
- 1228 **65.** Peebles, W. & Xie, S. Scalable diffusion models with transformers. In *Proceedings of the IEEE/CVF international*
1229 *conference on computer vision*, 4195–4205 (2023).
- 1230 **66.** Yu, Y. *et al.* Uni-dock: Gpu-accelerated docking enables ultralarge virtual screening. *J. chemical theory computation* **19**,
1231 3336–3345 (2023).
- 1232 **67.** Hansen, N. & Ostermeier, A. Completely derandomized self-adaptation in evolution strategies. *Evol. computation* **9**,
1233 159–195 (2001).
- 1234 **68.** Nomura, M., Akimoto, Y. & Ono, I. Cma-es with learning rate adaptation. *ACM Transactions on Evol. Learn.* **5**, 1–28
1235 (2025).