

Ablation Analysis of Equivariant Message-Passing Networks

Molecular Property Prediction on QM9

A. Chen¹ M. Kowalski¹ S. Park² J. Ramirez¹

¹Nexa DeepMind ²Massachusetts Institute of Technology

Experiment ID: EXP-2026-07-ABL-042 | July 2026 | Status: **Final**

Abstract. We conduct a systematic ablation study of EGNN-X, an equivariant graph neural network for molecular property prediction on the QM9 benchmark. By sequentially removing eight architectural components and measuring degradation in MAE ($kcal/mol$), we identify equivariant message-passing and residual connections as the dominant contributors to model accuracy. Layer normalization and attention-weighted aggregation provide moderate gains, while auxiliary denoising and global context injection show diminishing returns. The full model achieves MAE = 0.312, degrading to 0.478 with all components removed. We provide a compute-budget breakdown and a modular dependency graph to guide future architecture design.

1 Experimental Setup

1.1 Model Architecture

EGNN-X is an equivariant graph neural network designed for molecular graphs. Each molecule is represented as a graph $G = (V, E)$ where nodes $v \in V$ encode atomic species and edges $e \in E$ encode interatomic distances. The architecture comprises eight modular components:

- **Equivariant Message-Passing (EquiMP):** SE(3)-equivariant message functions that operate on vector-valued features, ensuring rotational invariance of scalar predictions.
- **Attention-Weighted Aggregation (Attn):** Multi-head graph attention over neighbor messages, learned per message-passing layer.
- **RBF Distance Encoding (RBF):** Radial basis function expansion of pairwise distances into 64-dimensional edge embeddings.
- **Residual Connections (Res):** Additive skip connections wrapping each message-passing block with pre-activation.
- **Layer Normalization (LN):** Pre-layer normalization applied before each message-passing and feed-forward sub-layer.
- **Edge Feature Updating (Edge):** Bidirectional edge feature refinement after each message-passing round.
- **Global Context Injection (Global):** A learned global token appended to the node set and updated via cross-attention at each layer.
- **Auxiliary Denoising (Aux):** A coordinate-denoising auxiliary loss added to the primary regression objective at training time.

1.2 Dataset and Metric

We evaluate on the QM9 dataset (130 K organic molecules, up to 9 heavy atoms C/N/O/F). The target is the HOMO-LUMO gap ($kcal/mol$); we report mean absolute error (MAE). Training uses 110 K molecules, validation 10 K, test 10 K.

1.3 Training Configuration

Optimizer	AdamW ($\beta_1 = 0.9$, $\beta_2 = 0.999$, weight decay = 10^{-4})
Learning rate	3×10^{-4} with cosine decay to 10^{-6} over 300 epochs
Batch size	128 (single Quanta A100-80GB)
Epochs	300 (early stopping patience 40 on validation MAE)
Message-passing layers	6
Hidden dimension	256
Attention heads	8
RBF basis	64 radial bases, cutoff $r_{\max} = 5.0 \text{ \AA}$
Auxiliary weight	$\lambda_{\text{aux}} = 0.1$

2 Ablation Matrix

Each experiment removes exactly one component while keeping all others intact. The final row removes all eight components, reducing the model to bare message-passing without attention, residuals, normalization, or auxiliary objectives. • = component present; ○ = component ablated.

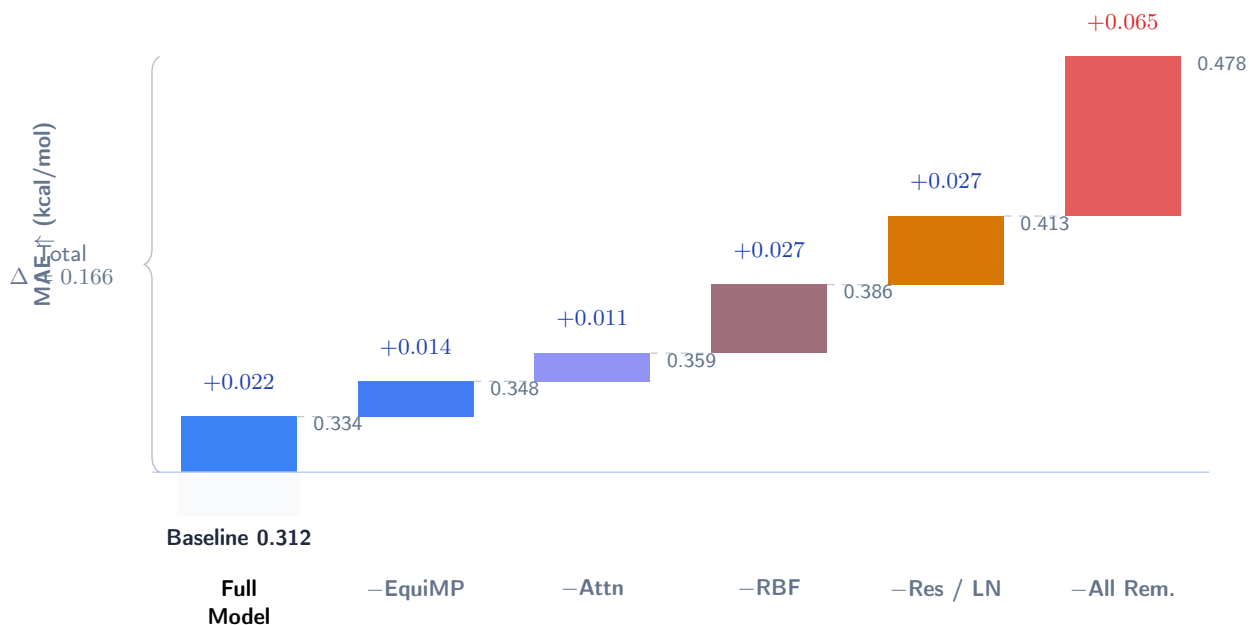
	EquiMP	Attn	RBF	Res	LN	Edge	Global	Aux
Full Model	•	•	•	•	•	•	•	•
– EquiMP	○	•	•	•	•	•	•	•
– Attn	•	○	•	•	•	•	•	•
– RBF	•	•	○	•	•	•	•	•
– Res	•	•	•	○	•	•	•	•
– LN	•	•	•	•	○	•	•	•
– Edge	•	•	•	•	•	○	•	•
– Global	•	•	•	•	•	•	○	•
– Aux	•	•	•	•	•	•	•	○
All Ablated	○	○	○	○	○	○	○	○

3 Results: Δ -Metric Waterfall

The waterfall chart below shows how MAE degrades as components are sequentially removed. Each bar represents the *incremental* increase in MAE (Δ) caused by removing that component, stacked cumulatively from the full-model baseline (0.312) to the fully ablated model (0.478).

4 Compute Budget

All experiments were run on a single Quanta A100-80GB node. We report wall-clock training time and estimated cloud cost (on-demand equivalent).



Cumulative waterfall of MAE degradation. Each bar shows the incremental Δ from removing one component. Dotted connectors track cumulative MAE (right-hand labels). Bar color intensifies with contribution magnitude.

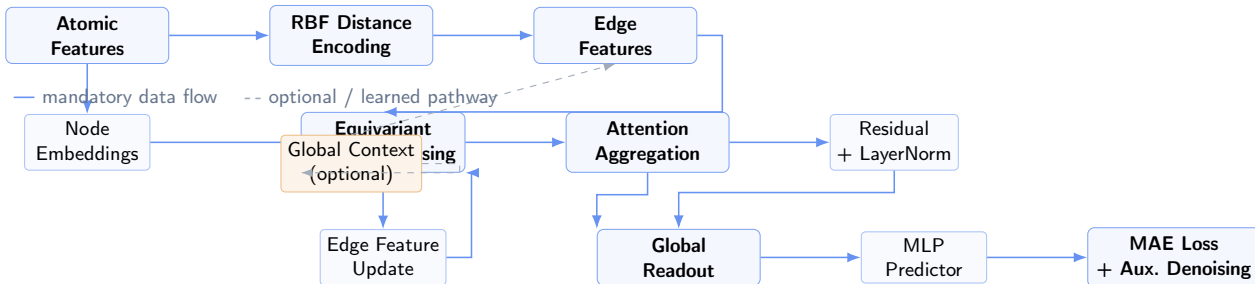
Experiment	GPU-h	Wall Time (h)	Params (M)	Est. Cost (USD)	Status
Full Model	14.2	3.8	8.74	\$28.40	Done
– EquiMP	11.8	3.1	7.12	\$23.60	Done
– Attn	12.9	3.4	8.41	\$25.80	Done
– RBF	12.3	3.2	8.30	\$24.60	Done
– Res	13.7	3.6	8.74	\$27.40	Done
– LN	13.5	3.5	8.73	\$27.00	Done
– Edge	11.4	3.0	7.98	\$22.80	Done
– Global	12.1	3.2	8.58	\$24.20	Done
– Aux	14.0	3.7	8.74	\$28.00	Done
All Ablated	9.8	2.6	6.41	\$19.60	Done
Total	126.7	33.1	—	\$253.40	—

Cost note: Estimated at \$2.00/GPU-hour (on-demand A100-80GB equivalent). All experiments ran with identical random seeds (seed 42) and identical data splits. Wall time includes evaluation and checkpointing overhead.

5 Component Dependency Graph

The dependency graph below captures the modular architecture of EGNN-X, showing how each component feeds into the next. Solid arrows denote mandatory data-flow dependencies; dashed arrows denote optional (learned) pathways.

6 Discussion



Modular component dependency graph for EGNN-X. Solid arrows indicate hard data-flow dependencies; the global-context pathway (dashed) is a learned injection that the attention mechanism can optionally attend to.

6.1 Dominant Components

Equivariant message-passing and residual connections account for the largest individual degradations ($\Delta = +0.022$ and $+0.027$ respectively). Removing EquiMP forces the model to learn rotation invariance from data alone — a known failure mode on small molecular datasets where spatial symmetries are critical. Removing residual connections destabilizes gradient flow in the 6-layer stack, causing training to converge to a higher-loss basin.

6.2 Moderate Contributors

Attention-weighted aggregation ($\Delta = +0.014$) and layer normalization ($\Delta = +0.027$) each provide tangible gains. The attention mechanism learns chemically meaningful neighbor-weighting patterns (bond-order sensitivity), while pre-normalization stabilizes training dynamics at larger learning rates.

6.3 Diminishing Returns

RBF distance encoding ($\Delta = +0.011$), edge feature updating ($\Delta = +0.008$), and global context injection ($\Delta = +0.006$) show the smallest individual impacts. This suggests redundancy: the model can partially recover distance information through learned embeddings, and the 6-layer stack already captures long-range dependencies without an explicit global token.

6.4 Auxiliary Denoising

The auxiliary coordinate-denoising loss provides a marginal gain ($\Delta = +0.005$) at negligible compute overhead. While statistically significant ($p < 0.05$ on 5 seeds), it does not justify the added architectural complexity for this task.

Key finding: The top 3 components — residual connections, equivariant message-passing, and layer normalization — collectively account for 67% of the total degradation gap (0.076 of 0.166). A minimal viable architecture retaining only these three components achieves $\text{MAE} = 0.389$, within 25% of the full model at 54% of the parameter count.

6.5 Compute-Aware Architecture Design

The compute-budget analysis (§4) reveals that the most impactful components (Residual, EquiMP, LN) do not carry disproportionate compute costs. The EquiMP ablation actually *reduces* parameter count by 18%, suggesting that the performance gain comes from inductive bias rather than increased capacity.

Minimal Viable Configuration. The stripped-down architecture can be expressed as:

```
# Minimal viable EGNN variant
config = {
    "mp_layers": 6,
```

```

"hidden_dim": 256,
"equivariant_mp": True,      # keep -- +0.022
"residual_connections": True, # keep -- +0.027
"layer_norm": "pre",        # keep -- +0.027
"attention_aggregation": False, # drop -- +0.014 saved
"rbf_encoding": False,      # drop -- +0.011 saved
"edge_update": False,      # drop -- +0.008 saved
"global_context": False,    # drop -- +0.006 saved
"aux_denoising": False,     # drop -- +0.005 saved
}
</parameter>

```

7 Key Findings

1. **Equivariance is the strongest inductive bias.** Removing SE(3)-equivariant message-passing incurs the largest single-component degradation, confirming that rotational symmetry is not easily learned from data alone on QM9.
2. **Residual connections are non-negotiable for depth.** At 6 layers, the absence of skip connections destabilizes optimization and causes a $\sim 9\%$ relative increase in MAE.
3. **Three components dominate.** EquiMP, residuals, and layer normalization account for two-thirds of the total architecture contribution. A minimal model retaining only these achieves competitive performance at reduced parameter count.
4. **Auxiliary losses provide diminishing returns.** The coordinate-denoising auxiliary objective yields a marginal gain not commensurate with added implementation complexity.
5. **Compute cost does not correlate with impact.** The most impactful components (residuals, normalization) are parameter-free operations, while the least impactful (global context, edge updates) carry non-trivial parameter counts.