
Harnessing large language models for data-scarce learning of polymer properties

In the format provided by the
authors and unedited

1	Table of Contents
2	Algorithm 1: The key steps of the two-phase training strategy.

Supplementary Algorithm 1 The key steps of the two-phase training strategy.

- 1: **LLM encoder pretraining:**
 - 2: Use a large dataset of unlabeled SMILES representations of polymers to pretrain an LLM encoder $\tilde{\mathcal{M}}_{\text{encode}}$.
 - 3: **Phase-1 supervised pretraining:**
 - 4: Use physics-based hypothetical polymer generation methods, such as group contribution (GC), to generate a large dataset of physically meaningful synthetic polymer structures $\{\mathbf{X}_i\}_{i=1}^{S_{GC}}$ with the correlation of fundamental thermophysical properties;
 - 5: Build a physics-based model $\mathcal{M}_{\text{physics}}$ of the real-world physical process, by leveraging on the fundamental properties calculated from physically meaningful synthetic polymers;
 - 6: Construct a physically meaningful synthetic dataset $\mathcal{D}_{GC} := \{(\mathbf{X}_i, \mathcal{M}_{\text{physics}}(\mathbf{X}_i))\}_{i=1}^{S_{GC}}$;
 - 7: Apply supervised pretraining to the LLM decoder/predictor using the synthetic dataset $\mathcal{D}_{GC}$, to obtain an LLM decoder $\tilde{\mathcal{M}}_{\text{decode}}$ with physically consistent initial state;
 - 8: **Phase-2 finetuning:**
 - 9: Collect a (usually small) set of high-fidelity measurements from experiments, denoted as $\mathcal{D}_{HF} := \{(\mathbf{X}_i^{HF}, \mathbf{Y}_i^{HF})\}_{i=1}^{S_{HF}}$;
 - 10: Split the high-fidelity experimental dataset $\mathcal{D}_{HF}$ as a training set $\mathcal{D}_{HF}^{\text{train}}$ and a test set $\mathcal{D}_{HF}^{\text{test}}$, and finetune the phase-1 LLM $\tilde{\mathcal{M}}_{\text{decode}}$ using $\mathcal{D}_{HF}^{\text{train}}$;
 - 11: Obtain the final physics-guided LLM, and report the prediction accuracy on the test dataset $\mathcal{D}_{HF}^{\text{test}}$.
-