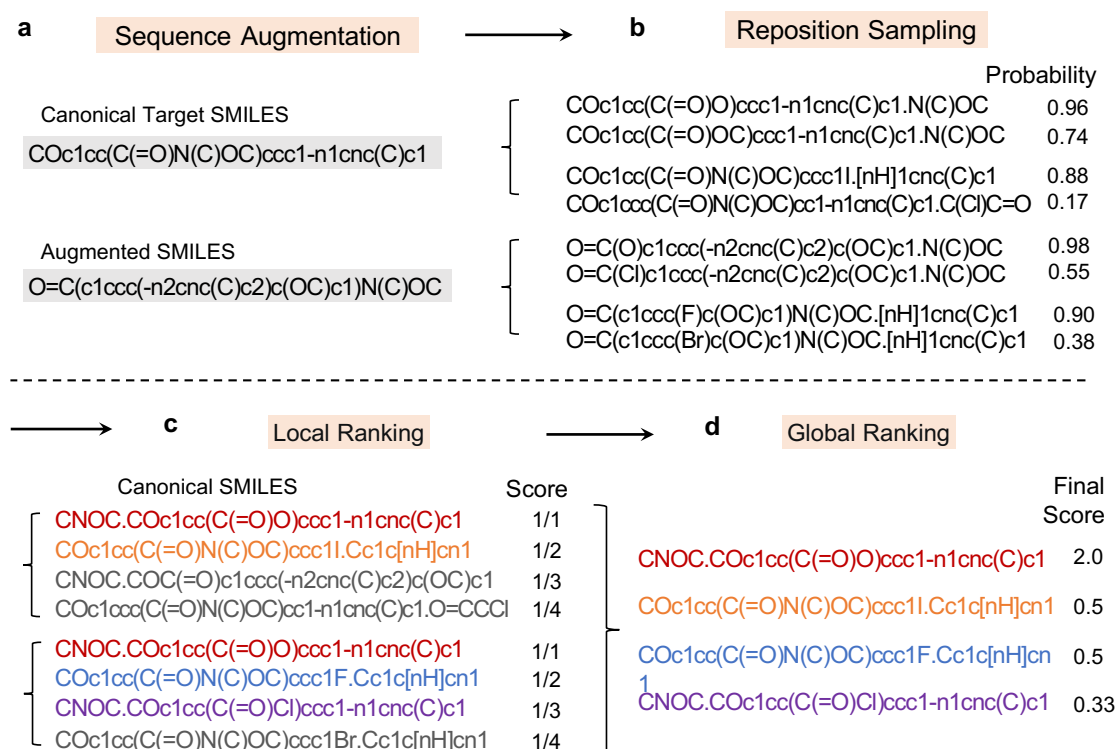
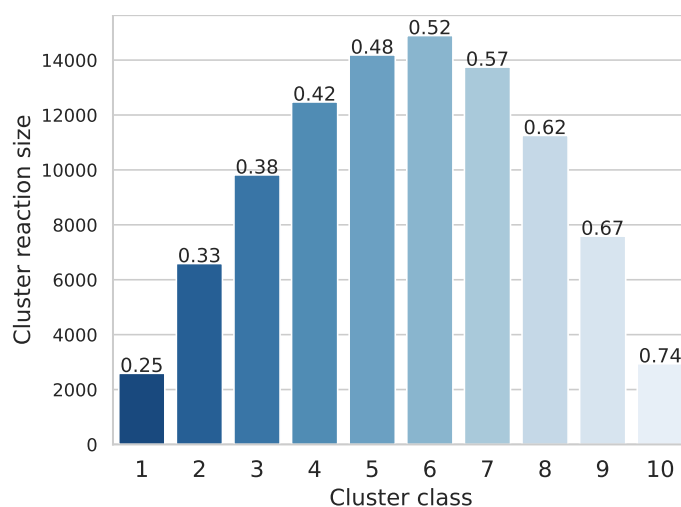


Supplementary Information for Retrosynthesis Prediction with an Iterative String Editing Model

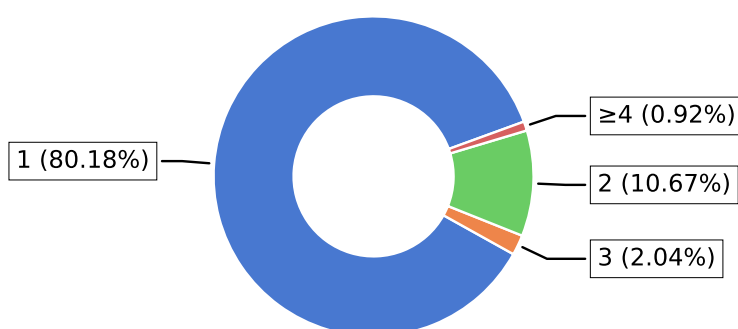
Supplementary Figures



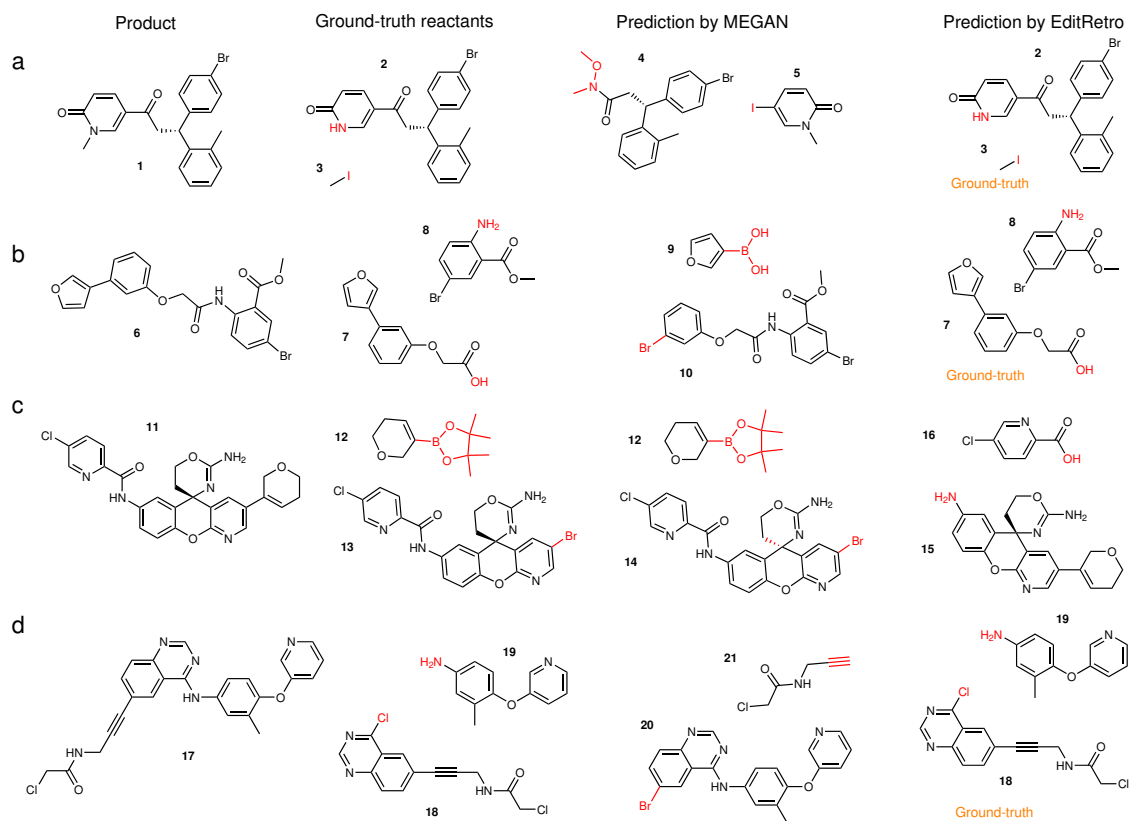
Supplementary Fig. 1: The workflow of the inference module with reposition sampling and data augmentation. In this example, we apply 2x SMILES augmentation, set $k_r = 2$ and $k_t = 2$ in reposition sampling to generate four predictions for each input sequence, and select the final top four predictions. **(a)** Sequence augmentation selects different starting atoms and direction of the molecular graph enumeration to generate variants of the canonical SMILES. We always set the canonical SMILES as the first input. **(b)** Reposition sampling generates multiple outputs for an input SMILES. **(c)** Local ranking ranks candidate sequences from the same input based on their probabilities. **(d)** In global ranking, each sequence is initially assigned a score using the reciprocal of its local rank. Afterwards, all canonical SMILES are ranked based on their final scores. The computation details of the final score are described in the Inference Module section. In the example, different colored strings represent different canonical SMILES. The top-4 predictions are selected as the final generations.



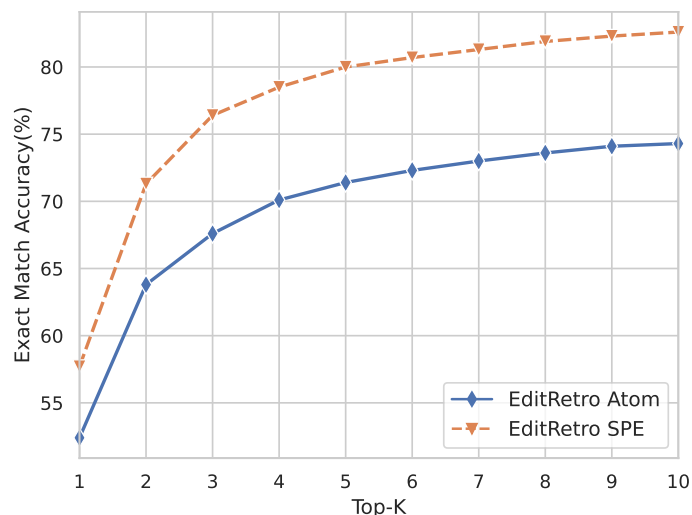
Supplementary Fig. 2: Cluster analysis of predicted reactants on the USPTO-FULL test set. The values displayed above the bars indicate the average similarity of the predicted reactants in the cluster, with lower values indicating higher diversity among the predicted reactants. The predictions in the first three clusters demonstrate high diversity, as they exhibit lower similarities (0.25, 0.33, 0.38, and 0.42) and account for approximately 33% of the test set. The predictions in the middle three clusters show medium diversity, with average similarities of 0.48, 0.52, and 0.57, accounting for nearly 45% of the test set. On the other hand, the predictions in the last three clusters are considered to have relatively low diversity, as they exhibit higher similarities (0.62, 0.67, and 0.74) and account for approximately 22% of the test set. These results demonstrate that EditRetro continues to exhibit promising diversity on larger and more diverse reaction datasets. Source data are provided as a Source Data file.



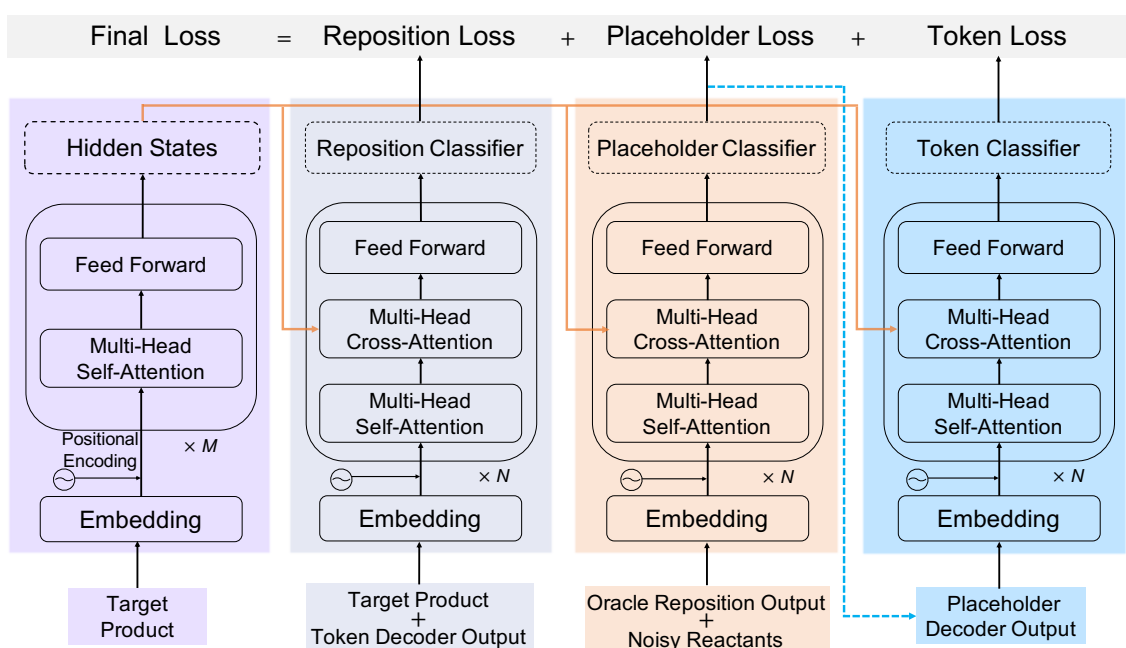
Supplementary Fig. 3: The distribution of refinement iterations for correctly predicted reactions in the test set of USPTO-50K. The majority of reactions achieve accurate predictions after just one refinement iteration. More refinement iterations are needed for other cases to achieve accurate predictions. Source data are provided as a Source Data file.



Supplementary Fig. 4: Comparison of Top-1 predictions between MEGAN and EditRetro Models across four error categories identified by MEGAN. **(a)** Only possible in multiple steps, **(b)** Low yield or side products, **(c)** Incorrect chirality, and **(d)** A reactive functional group ignored. EditRetro can accurately predict the ground-truth reactants as the top-ranked choice for the first, second, and fourth cases. In the third case, EditRetro predicts the ground truth as the second-ranked choice. However, it also identifies an alternative reaction center and predicts a feasible top-1 prediction.



Supplementary Fig. 5: An analysis of the effects of the SPE tokenizer. The SPE tokenizer can enhance the performance of our edit-based model compared to atom-wise tokenizer. The models are trained from scratch using different tokenizers on the USPTO-50K dataset with 10x data augmentation. Source data are provided as a Source Data file.



Supplementary Fig. 6: Training the model using imitation learning. The roll-in policy generates inputs for each module, which are used to train the model to mimic the expert's behavior. We share the parameters of the Transformer-Block across three decoders, with only the parameters of the classifiers being different.

Supplementary Tables

Supplementary Table 1: Top-1 exact match accuracy of different reaction classes on USPTO-50K with reaction class unknown.

Class	Reaction name	No.	Frac.(%)	Acc. (%)
1	Heteroatom alkylation and arylation	1,516	30.3	60.8
2	Acylation and related processes	1,190	23.8	70.9
3	C-C bond formation	567	11.3	44.3
4	Heterocycle formation	91	1.8	54.9
5	Protections	68	1.3	72.1
6	Deprotections	824	16.5	56.9
7	Reductions	462	9.2	63.9
8	Oxidations	82	1.6	74.4
9	Functional group interconversion(FGI)	184	3.7	43.5
10	Functional group addition(FGA)	23	0.5	78.3

Supplementary Table 2: Comparison of generated SMILES validity on USPTO-50K dataset.

Model	Top-k validity (%)			
	1	3	5	10
Liu's Seq2Seq [1]	87.8	84.7	81.6	78.0
Transformer	97.2	87.9	82.4	73.1
Graph2SMILES [2]	99.4	90.9	84.9	74.9
RetroPrime [3]	98.9	98.2	97.1	92.5
Retroformer [4]	99.2	98.5	97.4	96.7
SCROP [5]	99.3	98.6	98.2	97.7
Graph2Edits [6]	100	99.29	99.14	99.13
EditRetro	100	100	100	99.99

Supplementary Table 3: The inference latency of EditRetro and R-SMILES. It is calculated as the average time (in milliseconds (ms)) required to predict the output of one molecule in the test set using a batch size of one on one NVIDIA GeForce RTX 2080 ti GPU (excluding the model loading time). The results clearly indicate that EditRetro exhibits higher efficiency compared to R-SMILES, which is based on the standard Transformer architecture.

Method	Latency (ms)
R-SMILES	292.1
EditRetro	177.4

Supplementary Notes

EditRetro shows varied performance on different types of reactions

Table 1 provides a detailed breakdown of the top-1 exact match accuracy of our model across various reaction types. The results reveal that the model’s accuracy varies depending on the specific reaction types. This variation can be attributed to the structural diversity among potential reactants, particularly for substrates that can yield the same products. For example, we observe that EditRetro achieves relatively higher accuracy in predicting oxidation reactions, with a notable top-1 accuracy of 74.4% for reactions with an unknown type. This higher accuracy can be attributed to the limited reaction methods available for oxidations, where only a restricted set of substrates can produce ketones, primarily through the oxidation of alcohols. The relatively constrained reaction pathways for oxidations contribute to a higher prediction success rate in these cases. In contrast, carbon-carbon bond formation reactions offer multiple ways to construct a bond from different substrates, leading to a diverse set of reaction pathways. This diversity could lead to a lower prediction success rate for such reactions (e.g., 44.3% top-1 accuracy with reaction type unknown). The diversity in possible reaction pathways poses a challenge for accurate prediction, and this limitation is observed not only in our method but also in other approaches. The variation in accuracy across different reaction types highlights the inherent challenge in predicting reactions with diverse structural possibilities. Our method, like others, faces limitations in handling all possible chemical reaction types effectively. Addressing this challenge and developing more effective methods to improve prediction diversity would be of great interest for future research.

EditRetro generates more chemically valid molecules

Template-free approaches for generating chemical structures may produce predictions that are grammatically incorrect and not valid, highlighting the need for more robust methods. Our method exhibits robust generation capability due to three main contributions. Firstly, our approach benefits from edit-based generation, which is more reliable than generating a structure from scratch. Our analysis shows that the average Levenshtein edit distance at the fragment level between the product and reactants on training data is 5.4, indicating that we can obtain the reactants from the product with minimal modifications in most cases. This fragment-based refinement process reduces the risk of generating invalid molecules. Secondly, our model benefits from jointly masked sequence-to-sequence pre-training, which involves masking out certain tokens during pre-training to encourage the model to learn the underlying chemical rules and syntax of SMILES strings. The last one is data augmentation, which can help stabilize the model’s learning by introducing more randomness and flexibility into the network [7]. It functions similarly to ensemble learning by generating more valid candidates and filtering out invalid ones, resulting in a higher number of valid generations. We compare the validity of our model with several SMILES generative baselines. As shown in Table 2, EditRetro exhibits better molecule validity than the others. Especially, when $k = 1, 3, 5$, the predictions of our method are 100% valid chemical structures, indicating its generation robustness. Note that Graph2Edits is a recent semi-template-based method that uses an end-to-end graph generative architecture. Our results consistently outperform Graph2Edits, suggesting that sequence-based models, particularly with the assistance of large language models, may have some advantages over graph-based methods.

Analysis of the effects of SMILES Pair Encoding

A recent study [8] found that in retrosynthesis prediction on the USPTO-50K dataset, atom-wise tokenization outperformed SPE tokenization with top-1 accuracy rates of 42.0% and 19.82%, respectively. In this study, we evaluate the performance of EditRetro using both token-wise and SPE tokenization methods. For a fair comparison, we trained the model from scratch on the augmented training set of USPTO-50K, employing different tokenizers. As depicted in Supplementary Fig. 6, our results clearly demonstrate the significant impact of SPE tokenization on the model’s performance. Specifically, we observed an improvement in the top-1 accuracy from 52.4% to 57.3% and a notable enhancement in the top-10 accuracy from 74.3% to 82.2%. These findings strongly indicate that SPE tokenization is better suited for the edit-based model, effectively enhancing its overall performance.

Learning algorithm for EditRetro

Algorithm 1 Learning for EditRetro

```

1: Initialize: Reaction training set  $T$ , model policy  $\pi^\theta$ , expert policy  $\pi^*$ , random shuffle deletion policy  $\pi^{\text{rnd}}$ ,  $\alpha, \beta$ 
2: repeat
3:   Sample a training pair  $(s^0, s^*) \sim S$ 
4:   Sample  $u, v \sim \text{Uniform}[0, 1]$ 
5:   if  $v < \beta$  then
6:      $s_{\text{ins}} = E(s^0, \tilde{r})$ , where  $\tilde{r} = \arg \min_r D(s^*, E(s^0, r))$ 
7:   else
8:      $s_{\text{ins}} = E(s^*, \tilde{r})$ , where  $\tilde{r} \sim \pi^{\text{rnd}}(\cdot | s^*)$ 
9:   end if
10:   $s'_{\text{ins}} = E(s_{\text{ins}}, p^*)$ , where  $p^*, t^* = \arg \min_{p, t} D(s^*, E(s_{\text{ins}}, \{p, t\}))$ 
11:  if  $u < \alpha$  then
12:     $s_{\text{rps}} = s^0$ 
13:  else
14:     $s_{\text{rps}} = E(s'_{\text{ins}}, \hat{t})$ , where  $\hat{t} = \arg \max_t \sum_{s_i \in s'_{\text{ins}}} \log \pi_{\text{tok}}^\theta(t_i | i, s'_{\text{ins}})$ 
15:  end if
16:   $L_{\text{plh}}^\theta = - \sum_{s_i \in s_{\text{ins}}, p_i^* \in p^*} \log \pi_{\text{plh}}^\theta(p_i^* | i, s_{\text{ins}})$ 
17:   $L_{\text{tok}}^\theta = - \sum_{s'_i \in s'_{\text{ins}}, t_i^* \in t^*} \log \pi_{\text{tok}}^\theta(t_i^* | i, s'_{\text{ins}})$ 
18:   $L_{\text{ins}}^\theta = L_{\text{plh}}^\theta + L_{\text{tok}}^\theta$ 
19:   $L_{\text{rps}}^\theta = - \sum_{s_i \in s_{\text{rps}}, r_i^* \in r^*} \log \pi_{\text{rps}}^\theta(r_i^* | i, s_{\text{rps}})$ , where  $r^* = \arg \min_r D(s^*, E(s_{\text{rps}}, r))$ 
20:   $\theta = \theta - \lambda \cdot \nabla_\theta [L_{\text{ins}}^\theta + L_{\text{rps}}^\theta]$ 
21: until maximum training steps reached

```

Supplementary References

- [1] Liu, B. *et al.* Retrosynthetic reaction prediction using neural sequence-to-sequence models. *ACS central science* **3**, 1103–1113 (2017).
- [2] Tu, Z. & Coley, C. W. Permutation invariant graph-to-sequence model for template-free retrosynthesis and reaction prediction. *Journal of chemical information and modeling* **62**, 3503–3513 (2022).
- [3] Wang, X. *et al.* Retroprime: A diverse, plausible and transformer-based method for single-step retrosynthesis predictions. *Chemical Engineering Journal* **420**, 129845 (2021).
- [4] Wan, Y., Hsieh, C.-Y., Liao, B. & Zhang, S. *Retroformer: Pushing the limits of end-to-end retrosynthesis transformer*. *International Conference on Machine Learning, ICML 2022, 17-23 July 2022, Baltimore, Maryland, USA*, Vol. 162 of *Proceedings of Machine Learning Research*, 22475–22490 (PMLR, 2022).
- [5] Zheng, S., Rao, J., Zhang, Z., Xu, J. & Yang, Y. Predicting retrosynthetic reactions using self-corrected transformer neural networks. *Journal of chemical information and modeling* **60**, 47–55 (2019).
- [6] Zhong, W., Yang, Z. & Chen, C. Y.-C. Retrosynthesis prediction using an end-to-end graph generative architecture for molecular graph editing. *Nature Communications* **14**, 3009 (2023).
- [7] Tetko, I. V., Karpov, P., Van Deursen, R. & Godin, G. State-of-the-art augmented nlp transformer models for direct and single-step retrosynthesis. *Nature communications* **11**, 1–11 (2020).
- [8] Ucak, U. V., Ashyrmamatov, I. & Lee, J. Improving the quality of chemical language model outcomes with atom-in-smiles tokenization. *J. Cheminformatics* **15**, 55 (2023).