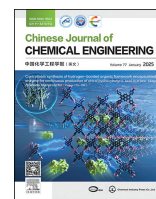




Contents lists available at ScienceDirect

Chinese Journal of Chemical Engineering

journal homepage: www.elsevier.com/locate/CJChE

Review

Machine learning-assisted retrosynthesis planning: Current status and future prospects

Yixin Wei^{1,2}, Leyu Shan¹, Tong Qiu^{1,2,*}, Diannan Lu¹, Zheng Liu¹¹ Department of Chemical Engineering, Tsinghua University, Beijing 100084, China² Beijing Key Laboratory of Industrial Big Data System and Application, Beijing 100084, China

ARTICLE INFO

Article history:

Received 4 September 2024

Received in revised form

30 October 2024

Accepted 31 October 2024

Available online 30 November 2024

Keywords:

Retrosynthesis planning

Machine learning

Artificial intelligence

Synthetic pathway

Chemoinformatics

ABSTRACT

Machine learning-assisted retrosynthesis planning aims to utilize machine learning (ML) algorithms to find synthetic pathways for target compounds. In recent years, with the development of artificial intelligence (AI), especially ML, researchers' interest in ML-assisted retrosynthesis planning has rapidly increased, bringing development and opportunities to the field. In this review, we aim to provide a comprehensive understanding of ML-assisted retrosynthesis planning. We first discuss the formal definition and the objective of retrosynthesis planning, and organize a modular framework which includes four modules: data preparation, data preprocessing, pathway generation and evaluation, and pathway verification. Then, we sequentially review the current status of the first three modules (except pathway verification) in the ML-assisted retrosynthesis planning framework, including ideas, methods, and latest progress. Following that, we specifically discuss large language models in retrosynthesis planning. Finally, we summarize the extant challenges that are faced by current ML-assisted retrosynthesis planning research and offer a perspective on future research directions and development.

© 2024 The Chemical Industry and Engineering Society of China, and Chemical Industry Press Co., Ltd.

This is an open access article under the CC BY-NC-ND license (<http://creativecommons.org/licenses/by-nc-nd/4.0/>).

1. Introduction

Over the past few decades, scientists have been working on exploring ways to artificially synthesize organics with specific functions (e.g., drugs), which is a high-investment, high-risk, and long-cycle process [1–3]. In order to reduce the cost of trial and error and improve efficiency, machine learning (ML) has been introduced in various parts of the organic synthesis planning process [1,3,4], including *de novo* design of compounds [5–7], retrosynthesis planning [8,9], and prediction of ligand–protein interactions [10], etc. Despite the significant advances, retrosynthesis planning remains an important focus in the field of chemistry and functional chemical discovery [11].

Retrosynthesis planning (also referred to as retrosynthesis analysis), which was first proposed by Corey in the 1960s [12,13], refers to the generation of a series of candidate synthetic pathways that reach the predetermined target compound. The original objective of the research on retrosynthesis planning is to devise optimal synthetic pathways from the target compound back to

readily available starting materials both efficiently and effectively. More specifically, retrosynthesis planning first defines the target compound as the end of the pathway, iteratively predicts the reactants of the current product against the direction of real reactions, and progressively converts the target compound into more easily synthesized precursors until the reactants meet the termination conditions, e.g., chemically stable, easily available, etc [14,15]. Fig. 1 is an example of retrosynthesis planning, which finds synthetic pathways for 2-amino-5-nitroterephthalic acid.

As a fundamental component of synthetic pathways, retrosynthesis planning involves two types of reactions, which can be categorized as enzymatic *versus* non-enzymatic (synthetic) according to the use of enzymes as catalysts or not [17]. Enzymatic and synthetic reactions have their own advantages, their synergistic combination can be seen in some reports of actual compound synthesis [18–20]. Compared with synthetic reactions, enzymatic reactions have good selectivity and mild reaction conditions, but require more careful reaction evaluation, including activity prediction [21] and safety assessment [22]. Co-use of enzymatic and synthetic reactions can extend the search space of retrosynthesis planning for more efficient and creative synthetic pathways. It is noteworthy that, from the perspectives of catalytic science and data science, enzymatic reactions and synthetic reactions are essentially

* Corresponding author. Department of Chemical Engineering, Tsinghua University, Beijing 100084, China.

E-mail address: qiutong@mail.tsinghua.edu.cn (T. Qiu).

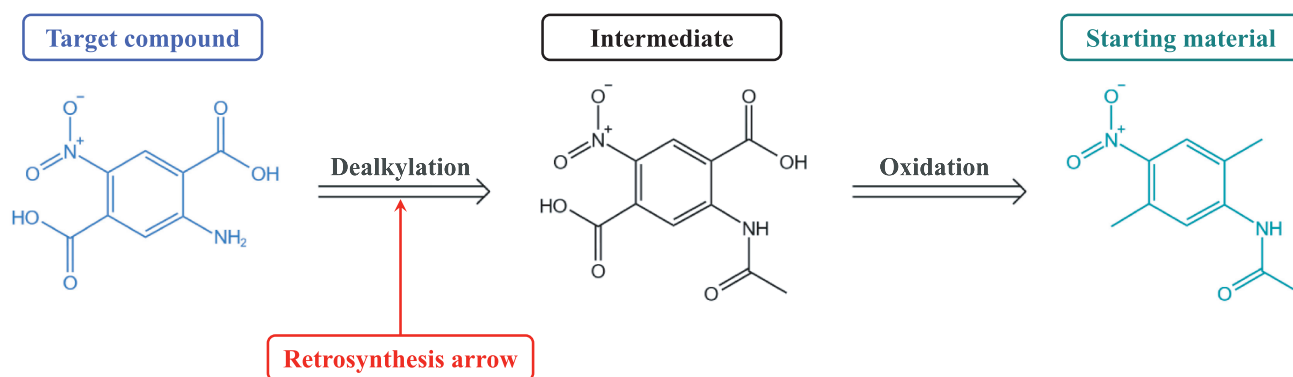


Fig. 1. An example of retrosynthesis planning. The target compound 2-amino-5-nitroterephthalic acid is in blue, the intermediate that cannot be purchased directly [16] are in black, and the reactant (starting material) that can be purchased directly and used as pathway starters are in green. The direction of the retrosynthesis arrow is opposite to the direction of the real reaction.

similar in logic, differing only in their catalysts, and are generally not distinguished separately during pathway generation. The most significant difference lies in the vastly dissimilar dataset sizes, which poses a challenge to the integration and equilibrium of enzymatic and synthetic reactions within the pathway.

Chemists have been exploring ways to use computers to help them with retrosynthesis planning for the past 60 years. The introduction of computers and ML has made retrosynthesis planning no longer an entirely manual and intuitive process [3,17,23]. With the introduction of advanced ML algorithms (including deep learning algorithms) in 2010s, chemists have begun to deeply integrate ML into retrosynthesis planning research. By leveraging ML algorithms, chemists have been able to construct data-driven models that assist in predicting chemical reactions and generating synthetic pathways [24–27], a process now referred to as ML-assisted retrosynthesis planning. Fig. 2 shows the development of ML and retrosynthesis planning during the past 60 years. The year 2017 is a pivotal moment in retrosynthesis planning, as it witnessed the beginning of the rise of data-driven models with the proposal of the first deep-learning based method [28].

Existing sophisticated ML algorithms (e.g. artificial neural networks (ANNs) [29], deep neural networks (DNNs) [30,31], and graph neural networks (GNNs) [32]) not only learn from reactions with enhanced efficiency but also adeptly navigate the substantial datasets. To imitate the process by which chemists acquire chemical knowledge and conduct retrosynthesis planning, the current ML-assisted retrosynthesis planning divides this task into four modules, jointly forming its mature, modular framework, as shown in Fig. 3. The data preparation module is tasked with provisioning an ample reservoir of data (chemical knowledge) for the ML-assisted retrosynthesis planning, laying the foundation for subsequent learning and training. To enable ML models to understand the reaction and compound data, the data preprocessing module represents the information supplied by the data preparation module in a format that is amenable to the model's input, with potential annotations. The pathway generation and evaluation module is the most pivotal component within the framework. Pathway generation part imitates the manner in which humans acquire chemical knowledge, learning the transformation of functional groups in actual reactions through the training of ML models. Grounded in this acquired chemical knowledge, the computer is capable of executing reaction prediction (single-step pathway extension) and pathway provision (multi-step pathway planning). Furthermore, in order to enhance the efficiency and feasibility of ML-assisted retrosynthesis planning and to circumvent issues of combinatorial explosion, pathway evaluation part imitates a human's assessment of synthetic pathways to refine the decision-

making processes within the pathway generation module. The efficacy of this module determines the extent to which an ML-assisted retrosynthesis planning tool (a computer expert) can substitute for a chemist. Finally, all candidate pathways are subject to the scrutiny of a pathway verification module to affirm their feasibility, thereby corroborating the aptitude of ML-assisted retrosynthesis planning tools in designing viable synthetic pathways. We believe, a sophisticated ML-assisted retrosynthesis planning tool should serve as a chemist's invaluable aide, capable of swiftly and comprehensively utilizing extant reactions to compose viable synthetic pathways, while also possessing the ingenuity to unearth novel reactions and pathways that have yet to be reported.

It is noteworthy that we describe the framework of ML-assisted retrosynthesis planning as a modular structure because the methods and models to realize each module can be freely combined, imparting flexibility to this field. Regrettably, extant research on ML-assisted retrosynthesis planning is predominantly focused on individual modules rather than the entirety of retrosynthesis planning. This leads to a significant discrepancy between the objectives of ML models and the primary goal of retrosynthesis planning. In this review, we aim to provide a comprehensive overview of the current state of ML-assisted retrosynthesis planning research, thereby assisting scholars in chemical engineering, chemistry, and computer science to better understand this domain. We sequentially and thoroughly delineate the prevailing ideas, methods, algorithms utilized, and the current state of progress for each module within ML-assisted retrosynthesis planning framework. Concurrently, we place an emphasis on discussing the challenges and limitations faced by current ML-assisted retrosynthesis planning research and identify potential directions for future developments, with the aspiration of stimulating subsequent research.

2. Data Preparation

Over the past 60 years, chemists have explored and accumulated a large amount of reaction data (including reaction condition) as well as experimental or computational data on compounds that provide the data base for ML-assisted retrosynthesis planning [33]. Reaction data refers to the information of a series of existing and verified reactions, which is considered the ground truth of the synthetic pathway. Reaction data generally includes a reaction index, the indices of compounds in the reaction, and the catalyst and condition of the reaction if available. In recent years, many mature and complete data resources have emerged, and the huge amount of rich chemical information and the application of big data technology have provided great convenience for chemists. Table 1

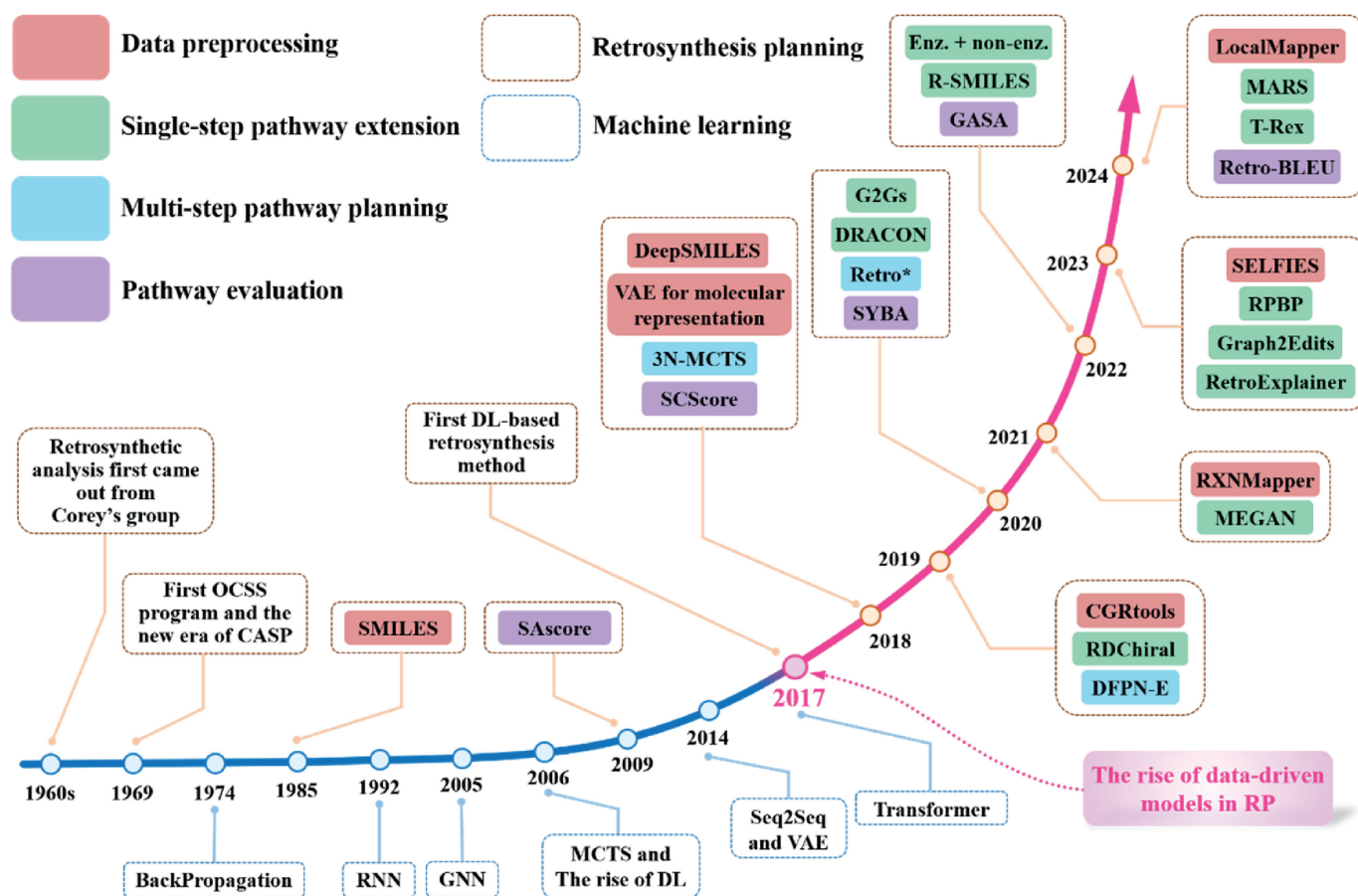


Fig. 2. Development of machine learning and retrosynthesis planning.

shows the databases and datasets that are popular in retrosynthesis planning. Compared to reaction and compound databases, utilizing enzyme and non-enzyme catalyst data in the generation of pathways is more challenging. On one hand, the necessity for catalyst data to be tied to specific reactions in order for it to be useful in reaction prediction, which limited its availability to some extent in the field of retrosynthesis planning. On the other hand, predicting the catalytic activity of a catalyst for any arbitrary reaction remains a formidable task. These two factors limit the current usage of catalyst data in retrosynthesis planning, resulting in a predominant focus on the transformation of compounds before and after the reaction in pathway generation, while neglecting the attention towards reaction catalysts and conditions.

Among all the database in Table 1, USPTO is the most widely used dataset in current retrosynthesis planning research, and the original dataset (USPTO-FULL) covers more than 1 million reactions extracted from US patents published between 1976 and September 2016 with low data quality [34]. On this basis, USPTO-50K [47] and USPTO-480K [35] (also called USPTO-MIT) were created. USPTO-50K was obtained by Nadine Schneider *et al.* [47] after a random extraction with data preprocessing on the full dataset, containing about 50000 reactions of 10 types with atom-atom mapping (AAM). USPTO-480K, on the other hand, was compiled by Jin *et al.* [35] at MIT after removing duplicates and incorrect reactions from the full dataset, which contains about 480000 reactions. As researchers have widely used the USPTO datasets in retrosynthesis planning, later researchers have also chosen from them for easier comparisons with previous studies.

3. Data Preprocessing

Data preprocessing employs cheminformatics tools to convert the data provided by the data preparation into input form for the ML models [48]. Data preprocessing consists of two levels: (1) representation of reaction and compound data into a computer-recognizable uniform format, which generally includes strings, molecular graphs, and descriptors (Section 3.1, 3.2); (2) identification of structural changes embedded in the origin reaction data, which typically refers to AAM (Section 3.3). The efficiency and accuracy of models in ML-assisted retrosynthesis planning are generally closely related to compound and reaction data preprocessing [49]. Fig. 4 shows a schematic diagram of the compound and reaction representation, with examples of AAM.

3.1. Compound representation

Different compound representations can extract features from different perspectives and levels of the molecule. In this section, we use the concept of molecular descriptors [50] to unify different types of compound representations commonly used in ML-assisted retrosynthesis planning. The calculation of compound representations and molecular descriptors is generally derived from some cheminformatics tools or pre-trained models, such as RDKit [51] and some quantitative structure-activity relationship (QSAR) models [52]. Specifically, molecular representations commonly used for retrosynthesis planning include: (1) linear molecular descriptors, (2) graphic molecular descriptors, and (3) molecular fingerprints [53]. Linear descriptors are the most straightforward,

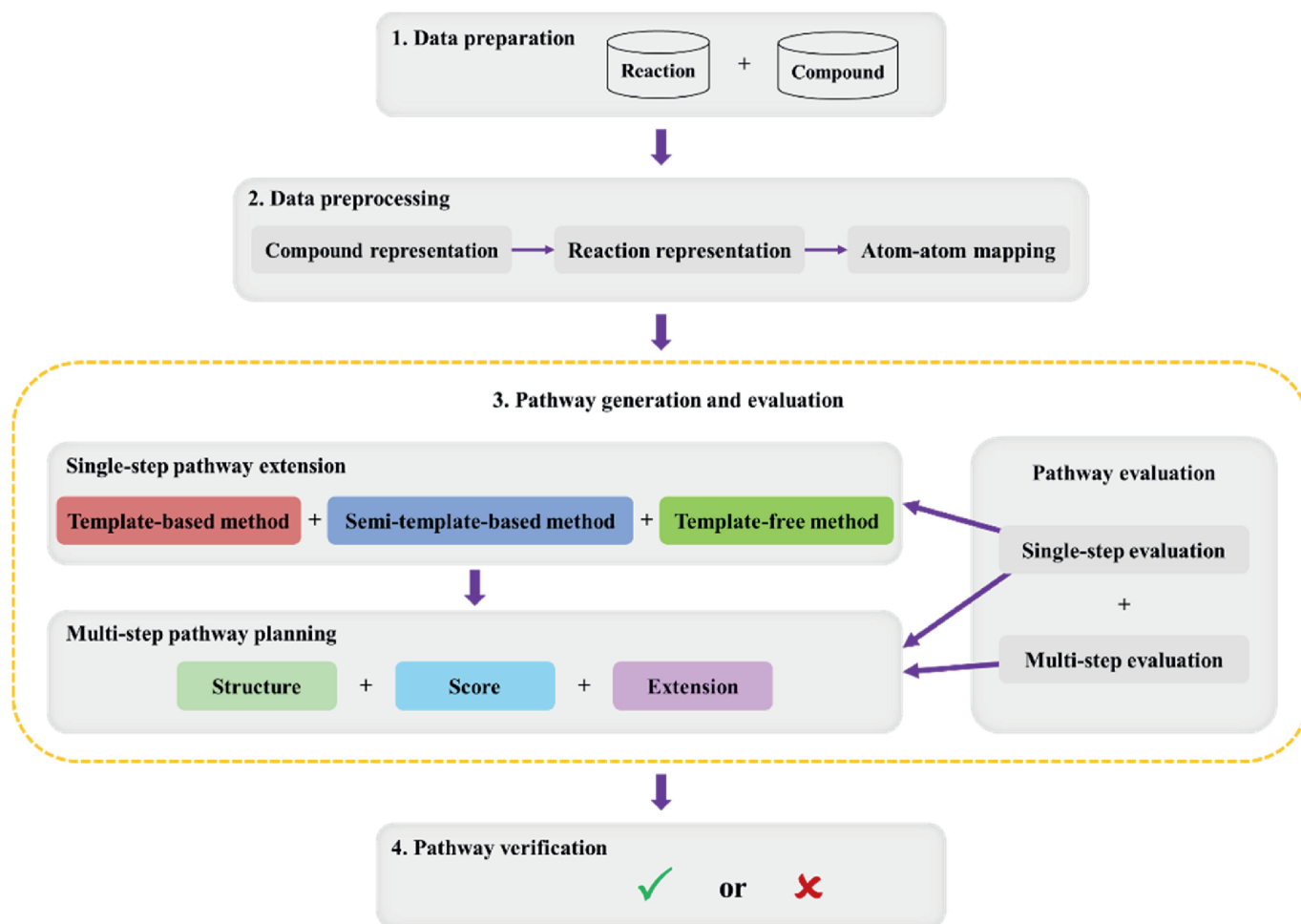


Fig. 3. Framework of ML-assisted retrosynthesis planning.

whereas graphical descriptors retain a greater abundance of molecular structural information.

Linear molecular descriptors, *i.e.*, the use of a one-dimensional string to represent a molecular structure. Prior to the 1970s, the Wiswesser Line Notation (WLN) [54] was the most commonly used symbolic format to accurately represent complex molecules. However, with the proposal of simplified molecular input line entry system (SMILES) [55] in the 1980s, the flexibility, readability, non-uniqueness (potential for data augmentation [49]), and interpretability of SMILES made it quickly become the most commonly used molecular representation up to now. SMILES specifies a rule for encoding the connection relationship of atoms within a molecule, making it possible to express two-dimensional molecular structures with one-dimensional strings, which greatly reduces the storage cost of complex molecular structures. To further represent the structural patterns of molecules, including atom environments and stereoisomerism, Daylight developed SMILES Arbitrary Target Specification (SMARTS) [56] based on SMILES, which uses more complex symbols to specify these molecular features. Based on SMILES, chemists have also developed DeepSMILES [57] and SELFIES [58] to address the errors that tend to occur when deals with rings and branched chains, and to improve the robustness.

Graphic molecular descriptors describe molecular structures using molecular graphs or matrices generated based on atomic connectivity, thereby offering better representation of topological

information and capturing the atomic environment within a molecule [59]. Using mathematical language, a molecular graph can be represented as $G = (V, E)$, where $V = \{v_i\}$ is the set of nodes of the graph and $E = \{e_{ij}\}$ is the set of edges of the graph. Each node represents an atom within the molecule, and each edge represents a bond between two atoms in the molecule. Compared to linear descriptors, graphic descriptors (1) offer a more lucid presentation of topological information, facilitating a better understanding and manipulation of chemical reactions; (2) enhance the concept of “bonds”, enabling GNNs to transmit information between nodes and bonds; (3) exhibit superior stability, ensuring chemically stable moieties within molecules remain undisturbed.

A molecular fingerprint is a fixed-length (*e.g.* 1024, 2048) vector that represents the structure of a molecule. Molecular fingerprints can be categorized into two types based on whether they depend on predefined substructures, represented by molecular ACCESS system (MACCS) fingerprints and extended-connectivity fingerprints [60] (ECFP), respectively. The former uses each element in the vector to indicate whether a certain predefined substructure exists within the molecule, while the latter considers the molecule localized at different radii, uses a hash algorithm to obtains a binary vector of fixed length. MACCS fingerprint is highly interpretable and flexible, which are often used in compound structure analysis and physical property prediction [61]. ECFP is often used in the calculation of compound similarity and pathway evaluation [29,62] because of its ease of computation.

Table 1
Popular databases and datasets in retrosynthesis planning.

Name	Source	Availability	Introduction
Reaction			
USPTO [34,35]	USPTO	Free	The USPTO family are open-source organic reaction datasets commonly used in ML, derived from patents granted by the US patent and trademark office (USPTO).
Reaxys [16]	Elsevier	Commercial	Reaxys contains data about the structure, physical and chemical properties, and biological activity of the compounds. Three sets of key information are provided for the compounds: (1) synthetic routes that can reach the target molecule; (2) information of reaction conditions in the synthetic pathways; and (3) literature or patents from which the reactions originated.
Pistachio	NextMove	Commercial	Pistachio is a document-centered, continuously updated subscription database with a larger volume of data and higher quality of information than USPTO.
ChEMU [36]	ChEMU lab	Free	ChEMU is a high-quality dataset containing 1500 manually annotated organic reaction texts from patents.
Rhea [37]	SIB	Free	Rhea records information in a tabular form containing more than 13000 individual enzymatic reactions, 12900 compounds with over 8000 enzymes. For each reaction, Rhea provides information of the enzyme and source of the reaction, as well as the compounds involved in the reaction.
ORD [38,39]	ORD group	Free	ORD is an initiative aimed at addressing the issues of dispersed, inaccessible, and non-standardized chemical reaction data in synthetic organic chemistry by supporting machine learning and related work through standardized data formats and open-access data repositories. Currently, ORD has grown to include more than 2 M reactions.
Compound			
PubChem [40]	NCBI	Free	PubChem is the largest open-source chemical database containing structure and bioactivity data for compounds. PubChem provides access to data in standardized text, as well as supply and procurement information for compounds.
KEGG [41]	ICR	Free	KEGG is a database generated by genome sequencing and other high-throughput experimental technologies. KEGG also provides the well-established network of biometabolic reactions.
ChEBI [42]	EMBL-EBI	Free	ChEBI is a database and ontology of chemical entities of biological interest containing a wide range of manually curated data items. ChEBI is widely used as a source of stable unique identifiers for chemicals in annotations in a wide range of bioinformatics databases, including Rhea and UniProt.
NIST Chemistry WebBook	NIST	Free	The NIST chemistry WebBook focuses on thermodynamics provides data compiled and distributed by NIST, including thermodynamic data for more than 7000 organic small molecules, thermodynamic data for over 8000 organic reactions, and a series of compound spectroscopy data.
Enzyme			
BRENDA [43]	Leibniz Institute DSMZ	Free	BRENDA is a database which focuses on enzyme and enzyme-ligand information.
UniProt [44]	EMBL-EBI, SIB and PIR	Free	UniProt is a knowledgebase which provides users with a comprehensive, high-quality and freely accessible set of protein sequences annotated with functional information.
Catalyst			
Materials Project [45]	LBNL and MIT	Free	The materials project represents a collaborative, international initiative aimed at calculating the properties of all inorganic substances. It strives to furnish comprehensive data and the corresponding analytical algorithms to the global materials research community.
Catalysis-Hub.org [46]	SUNCAT center, Stanford University	Free	Catalysis-hub.org is a web-platform for sharing data and software for computational catalysis research. It offers extensive data on catalysts and catalytic reactions, aiming to facilitate collaboration and information sharing in catalyst research.

3.2. Reaction representation

The preprocessing of reaction data contains two levels: reaction representation (Fig. 4(b)), focusing on the changes in compounds before and after the reaction, and the calculation of AAM (Fig. 4(c)), which delves deeper by examining the correspondence of atoms before and after the reaction, thereby providing insight into the reaction's nature. In this section, we will briefly introduce reaction representation and delve into a more detailed discussion on AAM, which is more commonly used in retrosynthetic planning, in the following section. Considering that reaction representation focuses only on conversions between compounds, the compound representation methods introduced in Section 3.1 can be effectively transferred. Both SMILES and molecular graphs are common methods of reaction representation in ML-assisted retrosynthetic planning. SMILES utilizes “>>” to space reactants and products and indicate the direction of a reaction, and uses “.” to space multiple reactants or multiple products in the reaction. Molecular graphs achieve reaction representation through directly combining the graphical depictions of reactants and products. To facilitate the message passing between compounds in GNNs, a commonly employed strategy involves augmenting additional hypernodes or pseudo-nodes in disconnected graphs [63]. The hypernode is linked to all substances relevant to the reaction to ensure the overall connectivity of the reaction molecular graph. Compared to employing SMILES, the use of molecular graph to represent reaction is more intuitive and expressive. Both methods for reaction

representation are currently being utilized in retrosynthesis planning research, leading to the development of sequence-based and graph-based approaches, respectively.

3.3. Atom-atom mapping

For chemical reactions, the understanding of the group changes in a reaction facilitates the understanding of the essence of the reaction. AAM is the task of identifying the position of each atom in molecules before and after a chemical reaction based on reaction representation. AAM contains more information than simple reaction representation and is more commonly used as input to ML models [64–66]. Unfortunately, except for a few reaction datasets (e.g., USPTO-50K), most reaction data do not contain AAM [34,35,37], and thus require separate calculation for AAM based on reaction representation. Furthermore, many reaction data record only the major reactants and products, which leads to difficulties to AAM [67].

Current methods for calculating AAM of chemical reactions generally include two categories: (1) methods based on structural changes [68–70] and (2) methods based on data-driven models [71,72].

Methods based on structural changes are essentially the identification of atoms and bonds that have changed in a reaction, which includes the approach based on maximum common substructure (MCS) and the approach based on optimization. MCS-based approaches associate the atoms from reactants and products by determining the MCS from their subgraphs and processing

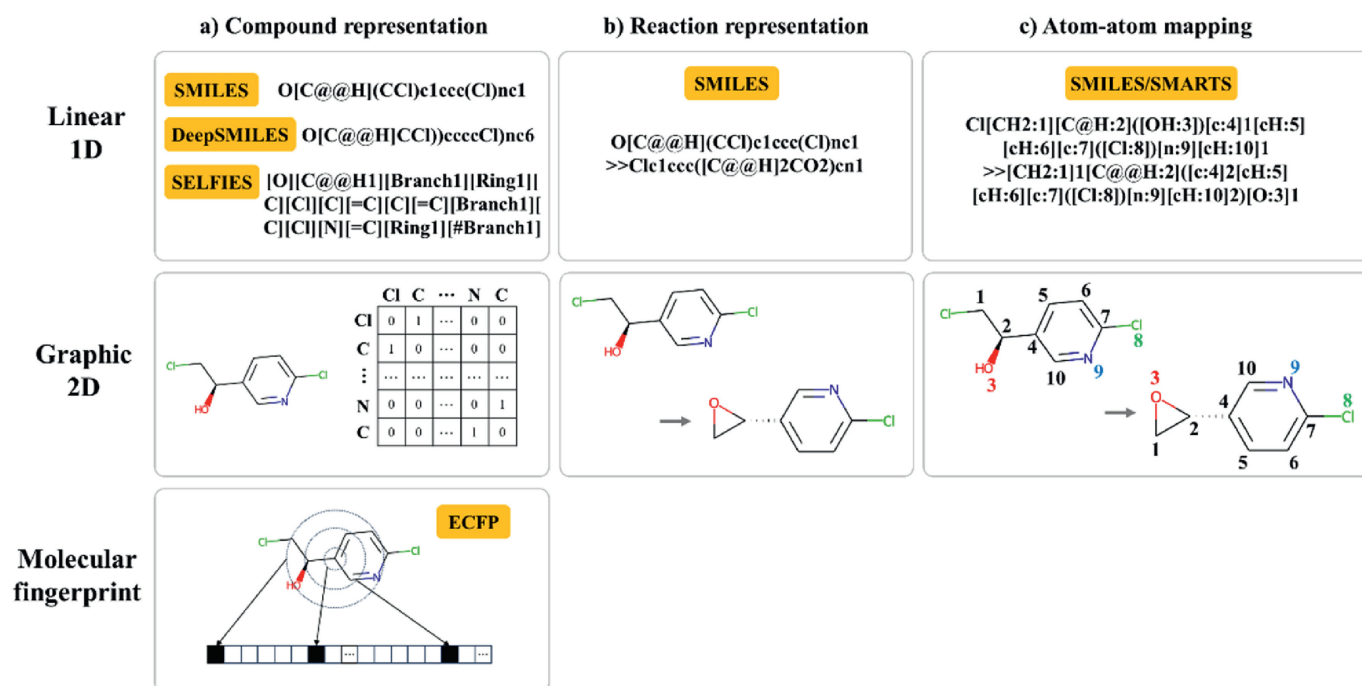


Fig. 4. An illustration of compound representation, reaction representation, and atom-atom mapping.

the remaining portions. Optimization-based approaches, on the other hand, capture structural changes in reactions by minimizing predefined functions (used to measure the number of dynamical changes in a reaction) through methods such as mixed-integer linear programming (MILP). The typical MCS-based approaches include RDT [73] by Rahman and ICMAP [74] by Kraut, while optimization-based approach is represented by DREAM [75] and MWED [69]. Fooshee *et al.* [76] combined the two approaches and proposed ReactionMap, which first finds the maximal map based on the substructure and then deals with the remaining atom through an optimization strategy, obtaining an AAM result that outperform all the previous models, while taking only 2 s per reaction on average.

Methods based on data-driven models expects computers to directly learn the AAM through ML models, especially deep learning models. The Transformer-based RXNMapper developed by Schwaller *et al.* [71] is the most typical unsupervised AAM tool, which has a significant accuracy advantage for the criteria of unbalanced reactions. From the perspective of graphs, Nikitin *et al.* [63] considered the complete reaction as an unconnected graph and proposed a disconnected graph neural network (DRACON), which can annotate AAM for reaction centers but fails to solve the AAM for the complete molecule. Considering that manual labeling of AAM is intensely time-consuming, Chen *et al.* [77] recently proposed an AAM model based on active learning named LocalMapper. LocalMapper achieves state-of-the-art AAM prediction accuracy by only labeling 2% of the reactions, and outputs entirely correct AAM for more than 97% reactions in the USPTO-50K dataset.

Compared with methods based on data-driven models, methods based on structural changes are more intuitive, interpretable and easier to check, but are also more challenging because such a subgraph isomorphism problem has been known to be an NP-hard problem [78,79].

AAM helps computers to understand and learn reactions. Many researchers [149,77] have found by comparison that AAM can significantly affect the feasibility and efficiency of the output

pathway of retrosynthesis planning, which is one of the reasons why USPTO-50K with good AAM is widely used in this field. Fortunately, existing open-source tools such as Indigo [80] and RXNMapper [71] can provide relatively accurate results for simple, balanced, and reaction-center-unique reactions, so that the retrosynthesis planning method will not be completely ineffective in datasets missing AAM.

4. Pathway Generation and Evaluation

Pathway generation, through the employment of pre-processed reaction data, trains ML models to realize the output of synthetic pathways, and can be divided into two parts: single-step pathway extension and multi-step pathway planning. Pathway evaluation, on the other hand, guides the decisions during pathway generation based on the ranks of candidate synthetic pathways or reactions. Currently, research on single-step, multi-step, and pathway evaluation is relatively independent. This neglect of coordination leads to the fact that current ML-assisted retrosynthesis planning still facing challenges in presenting novel and practically advantageous synthetic pathways.

4.1. Single-step pathway extension

Single-step pathway extension refers to the prediction of reactants (or precursors) that could be transformed into the target compound through chemical reactions, thus extending the synthetic pathway. With the development of ML, researchers have introduced a series of data-driven single-step pathway extension methods, which can be divided into three categories: templated-based methods, semi-template-based methods and template-free methods. Fig. 5 shows the workflows of three types of single-step pathway extension methods and their respective representative models, with colors indicating the reported Top-1 accuracy of each method on the USPTO-50K dataset (for an introduction to Top-1 or Top-k accuracy, see Section 4.1.4).

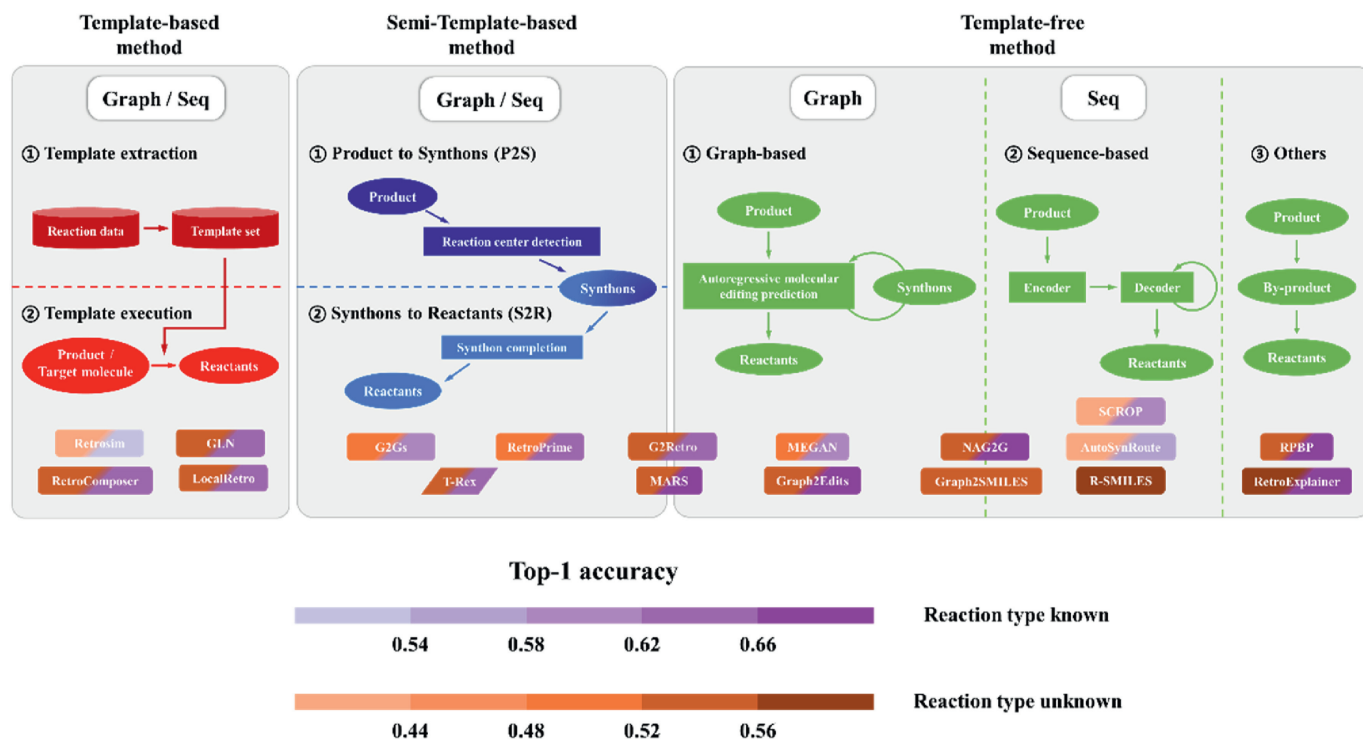


Fig. 5. Workflows of single-step pathway extension methods.

4.1.1. Template-based method

Template-based single-step pathway extension methods encompass two steps: (1) the extraction of a series of reaction templates from a collection of real chemical reactions; (2) the prediction of reactants through the execution of the extracted reaction templates. The former relies on the AAM and the identification of reactive centers, while the latter depends upon substructure matching. In particular, a reaction template refers to the group transformation occurring at the reaction center during a chemical reaction [28,81,82]. Given the diversity of reaction templates, and considering that the product molecules involved in retrosynthesis planning often contain multiple reactive functional groups, a target molecule can typically match multiple reaction templates, resulting in various potential reactants. Fig. 6 presents an example wherein various functional groups of the target product may match with distinct templates, consequently generating an array of diverse precursor candidates.

Template-based methods are intuitive, interpretable, and can provide reaction conditions. The extraction of reaction templates is the cornerstone of template-based methods. The accuracy and comprehensiveness of the reaction templates directly determine the upper limit of the quality of candidate pathways generated by template-based approaches. Template-based methods leverage ML to automatically and efficiently extract reaction templates, addressing the issues of low efficiency, error-proneness, and non-uniform formatting associated with manual extraction [17,66,74,81–85]. There are mainly two approaches to extracting reaction templates based on ML models, each relying on different representations of reactions. The first method uses SMILES or SMARTS to represent reactions, employing sequence-based methods to identify reaction centers and extract templates. The second approach represents reactions as graphs, utilizing GNNs to classify atoms within molecules, thereby achieving both the identification of reaction centers and reaction template acquisition simultaneously. Upon extracting reaction templates, template-based methods require only the substructure matching of the products to these templates and the

execution of the group transformations indicated by the template to achieve the prediction of reactants and the expansion of the pathway.

The open-source tool RDChiral introduced by Coley *et al.* [8] can achieve rapid reaction template extraction on reactions that have well-annotated AAMs. Additionally, it possesses the capability to process stereochemical information, which is vital for accurately reaction prediction and designing synthetic pathways. RDChiral enables researchers to process large-scale reaction data to extract a comprehensive set of reaction templates.

Concerning the poor generalization of template-based methods, Levin *et al.* [17] combined reaction templates from enzymatic reactions with those from synthetic reactions and found that the resulting pathways had advantages in terms of pathway length, selectivity, and efficiency. This demonstrates the impact of the scale and diversity of reaction templates on the outcomes of reaction pathways, which also indicates that using only the USPTO dataset is insufficient to encompass the practical demands such as synthesizing natural products.

Focusing on the efficiency of computationally expensive substructure matching [86], scientists have attempted to convert the simple enumeration into a strategy for template selection. Segler *et al.* [28] pioneered the use of ML models for the selection of reaction templates with Neursym, introduced in 2017, which employs the highway network to predict the most probable template based on the molecular fingerprint (ECFP4) of the product. It is noteworthy that Neursym currently stands as the most frequently utilized single-step model within the research of multi-step pathway planning. Coley *et al.* [87] proposed Retrosim, which frames the selection of reaction templates as a similarity problem, positing that structurally similar molecules can be synthesized using similar reactions. Based on the GNN, Dai *et al.* [88] have also proposed the conditional Graph Logic Network (GLN) to assess the compatibility of compounds with reaction templates, quickly providing template recommendations while ensuring the feasibility of the reactions.

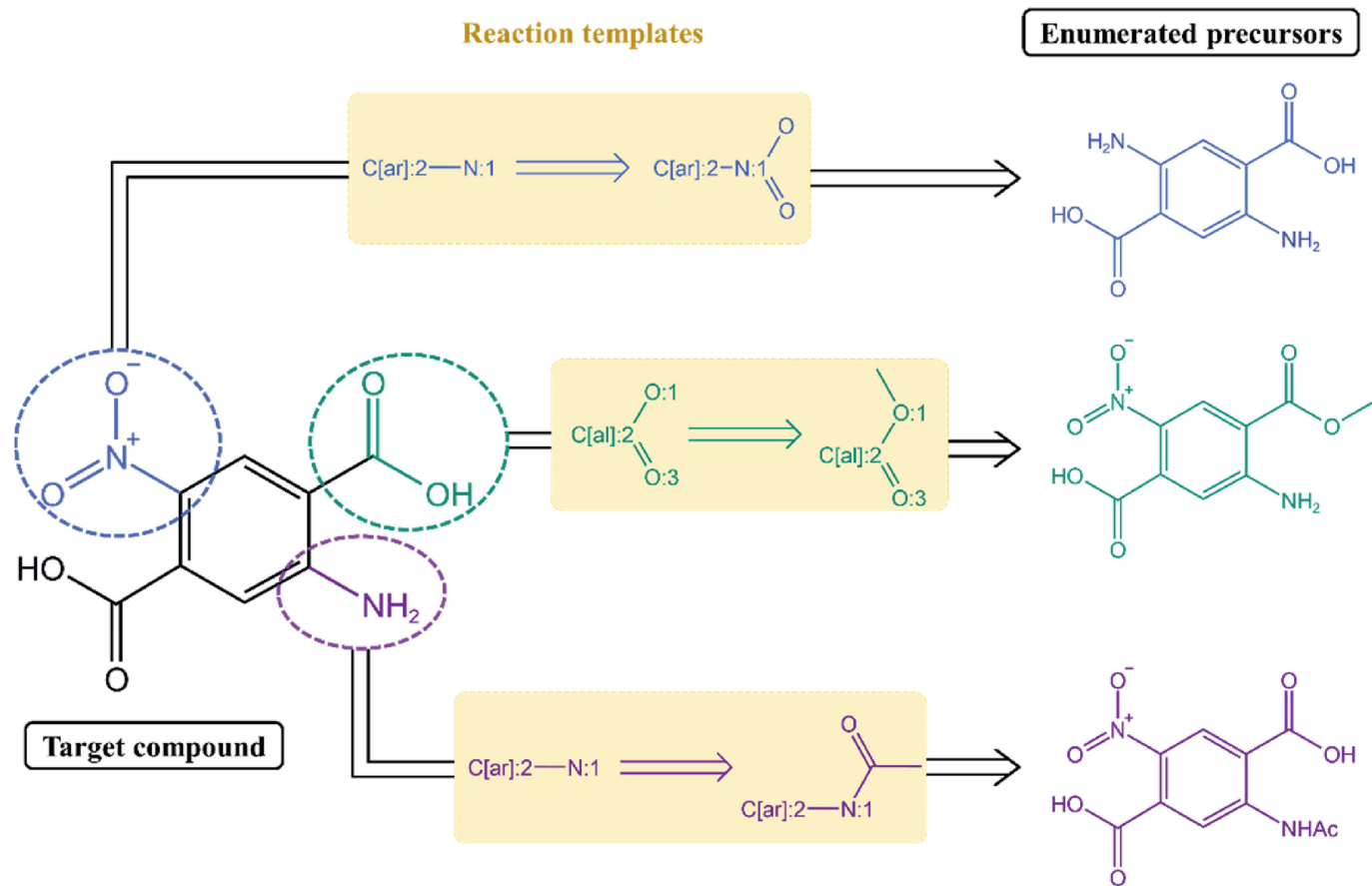


Fig. 6. An illustration of substructure matching between a target product and reaction templates. C[ar] signifies an aromatic carbon with undefined bonds, while C[al] represents an aliphatic carbon with undefined bonds.

Considering that reactions are essentially local structural changes in molecules, Chen *et al.* [89] developed LocalRetro, which outputs locally applicable reaction templates for all potential reaction centers, based only on local features, thereby enabling reaction prediction. RetroComposer, introduced by Yan *et al.* [90], transforms the template selection stereotype into template composition. It employs ML models and a subgraph dictionary to compose several candidate templates that are most suitable for the current product. RetroComposer is the best-performing template-based model on the USPTO-50K dataset, with the Top-1 accuracy of 54.5% and 65.9% on without and with reaction types, respectively.

It is noteworthy that template-based methods typically presume that the molecular structure outside the reaction center remains chemically inert and irrelevant to the reaction. However, under certain circumstances, the stereochemistry beyond the reaction center might influence the binding between the catalyst and the reaction site. Presently, research on this disadvantage remains scant, and we believe that a more robust experimental verification and counterexample data are imperative to leverage the complete structure of the product for a more precise prediction of the template-product compatibility.

4.1.2. Semi-template-based method

Considering the difficulties faced by template-based methods in computational efficiency and generalization, scientists have begun to ponder the use of ML to identify the transformation of atoms and chemical bonds during chemical reactions. Under this line of thought, the semi-template-based method divides the

extension of a single-step pathway into two stages: (1) identifying the reaction center and breaking the bonds at the reaction center to obtain synthons (from product to synthons, P2S), and (2) synthon completion (from synthons to reactants, S2R) [91,92]. Specifically, P2S is similar to the extraction of reaction templates in template-based methods, which is also the reason why this type of two-step approach is named as semi-template-based. The reason intermediate molecules are called synthons is that the synthons are generally incomplete local molecular structures, which may have chemical bonds with missing atoms and are not stable, viable independent chemical structures. Both Seq2Seq models (especially Transformer) and GNNs can both be used to construct P2S and S2R models, depending on the type of reaction representation. Theoretically, models for P2S and S2R can be combined in any manner, although the advanced semi-template-based approaches typically address both tasks simultaneously.

Following the naming convention of Seq2Seq, Shi *et al.* [93] proposed the first semi-template-based model based on graph convolutional neural networks in 2020, which they named Graph to Graphs (G2Gs). G2Gs first use a GNN to identify the most probable reaction centers in the product, breaking the chemical bonds between the atom pairs to generate synthons. Then, each synthon is translated into reactants using a Markov decision process, thus completing the prediction of reactants.

RetroPrime proposed by Wang *et al.* [91] utilizes two Transformers for P2S and S2R, respectively. They used a “mix and match” strategy to improve the diversity of the model, and a “label and align” strategy to reduce the proportion of chemically implausible

prediction results. It's worth mentioning that the introduction of this alignment strategy predates Root-aligned SMILES [49].

Considering that chemical reactions are essentially local group transfers within molecules, Liu *et al.* [94] proposed MARS which viewed S2R as the connection of attachment atoms with motifs. They first obtained a motif vocabulary of size 210 from the USPTO-50K dataset through molecular fragmentation, and then trained a model to predict the attachment atoms for the current synthon as well as the motifs added, autoregressively achieved synthon completion. MARS is the state-of-the-art semi-templated based model with the best prediction performance on the USPTO-50K dataset.

Semi-template-based methods combine the advantages of both the aforementioned template-based and the subsequent template-free methods. On one hand, semi-template-based methods describe a single-step pathway into two steps, better complying with how chemical reactions are understood, and offer a certain level of interpretability for the predicted reaction centers and derived reactants. On the other hand, semi-template-based methods can leverage powerful graph representation learning paradigms and attention mechanisms to better capture molecular structural features, thereby ensuring the accuracy of the model and the ability to predict new reactions.

4.1.3. Template-free method

Template-based methods are unable to handle novel or unseen reactions outside of the dataset, while semi-template-based methods train two independent models to complete the reaction prediction, ignoring the connection between P2S and S2R. Considering the aforementioned limitations and the development of ML, scientists have begun to use end-to-end data-driven models to directly predict reactions, without the need to individually extract and store reaction templates or reaction centers. Template-free methods include two approaches: sequence-based methods and graph-based methods [63,95,96].

• Sequence-based methods

The concept of sequence-based methods is quite intuitive, as the process of predicting the SMILES of possible reactants from the product's SMILES can be naturally likened to a translation task between two strings, which allows for the transfer of a series of Seq2Seq models [91,97,98]. However, due to the fact that a single product may correspond to multiple reactants, this one-to-many mapping relationship makes retrosynthesis planning based on the translation very challenging.

Liu *et al.* [9] proposed the first LSTM-based template-free model in 2017, which has comparable predictive capabilities for simple reactions that do not involve ring formation to the template-based models available at the time. With the subsequent introduction of the Transformer and attention mechanisms [99,100], sequence-based methods have almost entirely adopted this new structure, which can better extract global and local features of the input sequence.

Although sequence-based methods are conceptually intuitive, they still possess pronounced drawbacks. (1) Sequence-based models treat SMILES as sentences, yet their training often lacks prior knowledge of SMILES syntax, frequently resulting in the output of invalid SMILES. (2) Sequence-based models focus on only one-dimensional character and fail to utilize the rich structural information contained in the two-dimensional structure of molecules.

To address these disadvantages, Zheng *et al.* [101] designed a self-correcting module to correct grammatical errors in the SMILES output by Transformer framework, reducing the probability of the model outputting unfeasible SMILES. Zhong *et al.* [49] proposed a

reaction representation method called Root-aligned SMILES (R-SMILES), which aligned the same atoms in reactants and products based on AAM, and demonstrated through comparison that reducing the edit distance between the SMILES of reactants and products can effectively improve the reaction prediction performance of Transformer-based models. R-SMILES stands as one of the most advanced template-free models when reaction type is unknown, achieving a Top-1 accuracy of 56.3%.

Recently, Yao *et al.* [102] proposed a new template-free model called NAG2G that employs the Transformer encoder-decoder framework and utilizes molecular 1D, 2D, and 3D structural information. NAG2G achieves alignment between product and reactant SMILES through molecular graphs and AAM, requiring the model to prioritize outputting atoms that are present in both reactants and products in sequence, and finally predicting atoms that are only in the reactants. This strategy effectively avoids the output of unfeasible SMILES. NAG2G achieved the best performance on the USPTO-50K dataset when given reaction type, reaching a Top-1 accuracy of 67.2%.

• Graph-based methods

Graph-based methods employ a series of GNNs for single-step pathway extension and take two-dimensional molecular graphs as input. Any subgraph within a molecular graph carries structural significance, representing the molecular local environment, an aspect that linear symbol sequences, such as SMILES, fail to capture. Specifically, graph-based methods view chemical reactions as combinations of a series of predefined graph edits (including all possible bond and atom changes), implementing reactant prediction with an autoregressive strategy [96]. Semi-template-based models like G²Retro [92] in 2023 and MARS [94] in 2022 also use an autoregressive strategy during the S2R process.

MEGAN [96] is the first graph-based template-free model proposed in 2021 and the pioneer and baseline of the autoregressive molecule edit framework. MEGAN utilized Graph Attention Networks [103] (GATs) to extract features from molecular graphs, defining five types of graph actions at the atomic level (EditAtom, EditBond, AddAtom, AddBenzene and Stop). These actions allow MEGAN to modify molecular structures step-by-step in a controlled manner. Based on this baseline, Zhong *et al.* [95] proposed an end-to-end graph generative architecture termed Graph2Edits. Graph2Edits employed directed message passing neural network (D-MPNN) [104] to encode both local and global features of molecules, and modified the bond generation action to the addition of leaving groups with the consideration of stereochemistry. These designs can significantly reduce the length of graph edits and enhance predictive performance.

Graph-based methods can predict edits of arbitrary length in multi-step generation, which enable them to more effectively explore the potential space of reasonable reactions and enhance the diversity of the predictions [105]. However, current methods of this type heavily rely on AAM. The absence or errors in AAM can lead to misleading edit sequences, thereby affecting the accuracy of pathway extension.

• Other methods

Aside from the aforementioned two typical approaches, scientists have also conducted some other interesting template-free ideas.

On the basis of sequence-based methods, Tu and Coley [106] introduced a novel Graph2SMILES model, which changes the model's input to a molecular graph. The permutation invariance of D-MPNN and the graph-aware positional embeddings in

Graph2SMILES eliminates the need for any SMILES augmentation, while achieving 8% improvements in Top-1 accuracy over the Transformer baseline [107].

Yan *et al.* [108] innovatively shifted the prediction target from reactants to byproducts, realizing the prediction of reactants based on both the byproducts and target product with Transformer. Their RPBP first constructed a byproduct library containing 217 types of byproducts by comparing the atomic differences between reactants and products in the real reaction set based on R-SMILES [49], and then used a GNN to build a classification model for predicting byproducts.

Template-free methods are efficient, demonstrate strong learning capabilities with reaction data, and possess a degree of generalization. However, they still face limitation with poor interpretability and the inability to provide reaction conditions and catalyst information. The improvement the reliability and feasibility of the output pathways of template-free methods remains a challenge.

4.1.4. Evaluation for single-step pathway extension

Pathway evaluation is used to quantitatively evaluates a single-step model's performance and facilitate rigorous comparisons between models with several appropriate evaluation metrics. A deep understanding of these evaluation metrics helps to more accurately comprehend model's performance.

- Top-*k* accuracy

Top-*k* accuracy refers to the percentage of test target products for which the ground-truth precursor is included among the top *k* recommendations given by a single-step pathway extension model. Common values for *k* include 1, 3, 5, 10, 50, *etc.* In order to compare with previous work, Top-*k* accuracy is currently the most commonly used and sometimes the only evaluation metric used in various single-step models.

Intuitively, scientists believe that a higher Top-*k* accuracy indicates a better single-step model. However, researchers have gradually become aware of the limitations of Top-*k* accuracy [109,110]. Top-*k* accuracy pursues a match between the predicted results and the actual reactions recorded in the dataset, neglecting the value of diversity in model outputs. Since current reaction data rarely contain infeasible counterexamples, it is difficult to accurately assess the feasibility of new, unseen reactions. This leads scientists to arbitrarily hope that a model capable of predicting reactions within a dataset is also capable of predicting feasible reactions outside the dataset. However, it is challenging to determine whether an increase in Top-*k* accuracy is due to the model's improved learning of reaction behavior or simply due to overfitting. It is undoubtedly valuable to provide multiple options for each step in finding the globally optimal solution for multi-step pathway planning. Therefore, in addition to focusing on the Top-1 accuracy, we should also pay attention to Top-10 and Top-50 accuracy to obtain a more comprehensive evaluation.

- Maximal fragment accuracy

Inspired by classical retrosynthesis planning, the maximal fragment (MaxFrag) accuracy [31] is an extension of Top-*k* accuracy proposed by Tetko. MaxFrag degrades the prediction target from all reactants to the mean (largest) reactant, aligning more closely with the way chemists manually conduct retrosynthesis planning. With the largest reactant determined, there may be different ways to achieve the transformation, and the reactions within the dataset may not necessarily represent the optimal among these candidate options.

The MaxFrag accuracy is higher compared to traditional Top-*k* accuracy. Since its proposal, some models have started to consider using MaxFrag as an evaluation metric [27,49]. However, due to the fact that most baseline models still only report Top-*k* accuracy, MaxFrag has not received as much attention.

- Round-trip accuracy

Round-trip accuracy [111] quantifies the effective percentage of model suggestions by forward reaction predictions for candidate reactants. A higher round-trip accuracy implies that the model can offer a greater number of valid suggestions. In other words, the reactions predicted by the model are more likely to be feasible in practice [112,113]. It's worth mentioning that round-trip accuracy is highly related to the size of the beam in each single-step pathway extension. A larger beam means recommending more candidate reactions during each pathway extension, which may lead to a decrease in quality.

- Diversity

Diversity measures the number of reaction categories covered by the single-step predictions. Generally, a prediction of reactants with higher diversity is considered capable of covering a larger search space, which is more advantageous in finding the optimal synthetic pathway in further multi-step pathway planning [111].

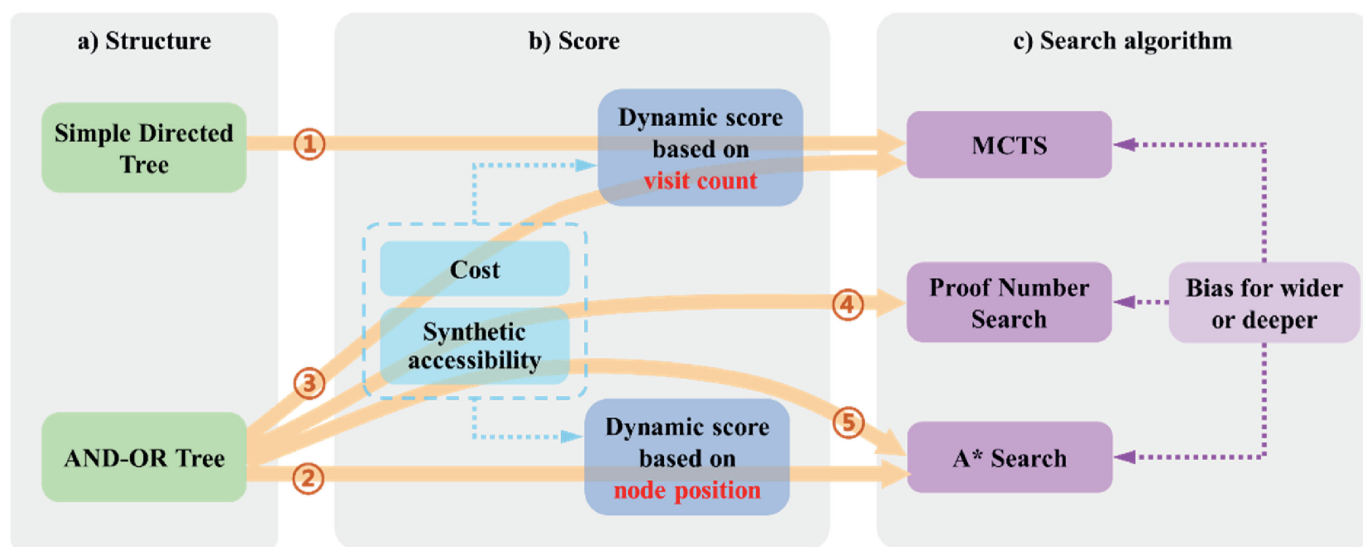
- Inference efficiency

Inference efficiency generally refers to the time it takes for the model to output recommended reactions. Since the single-step pathway extension model will be continuously invoked during subsequent pathway planning, its efficiency is crucial for finding a sufficiently good synthetic pathway within an acceptable time-frame. Unfortunately, most current single-step pathway extension models do not report their efficiency, possibly because model efficiency is highly related to the hardware performance of the computer. Zhong *et al.* [25] discovered through experiments conducted under identical hardware conditions that NeuralSym is the single-step model with the highest inference efficiency. This is attributed to its exceptionally simple network architecture, which may be the reason NeuralSym is favored as the preferred single-step model in multi-step pathway planning.

In addition to the aforementioned metrics, single-step evaluation also utilizes metrics such as coverage [111], Jensen-Shannon divergence [111], *etc.* Currently, these metrics are generally used independently to evaluate a model, which may be vulnerable to adversarial strategies. How to reasonably integrate these metrics is an issue for pathway evaluation that needs consideration. Methods such as the analytic hierarchy process(AHP) [114] might be very useful in this context.

4.2. Multi-step pathway planning

Chemical synthesis is a complex process involving multiple steps of consecutive reactions. The generation of multi-step synthetic pathways with practical significance is not a simple iteration of single-step pathway extensions, which requires the design of specialized planning methods with the guide of evaluation. On one hand, any unfeasible reaction within a multi-step pathway directly compromises the feasibility of the entire pathway. On the other hand, regardless of the single-step pathway extension method used, each extension typically yields multiple possible reactant outcomes. If one chooses to extend the pathway for all reactants at



Typical: ① 3N-MCTS, ② Retro*, ③/④ DFPN-E, ⑤ Chematica.

Fig. 7. The framework of multi-step pathway planning.

each iteration, the pathway size would explode rapidly due to the vast chemical search space.

Multi-step pathway planning aims to ensure both the efficiency of the output synthetic pathway and the feasibility of the pathway. Under this idea, multi-step pathway planning utilizes various search algorithms to start from the target product and gradually extend the pathway to more molecules until obtain until reasonable synthetic pathways are obtained, employing the single-step pathway extension methods previously introduced. Fig. 7 shows the three parts of multi-step pathway planning we divided: pathway structure, evaluation score calculation, and search algorithm.

4.2.1. Pathway structure

Pathway structure refers to the type of structure used to represent reaction pathways during the planning and expansion process, which includes two main categories: simple directed tree and AND-OR tree, as shown in Fig. 8. The former contains only molecule nodes, while the latter includes both reaction nodes and molecule nodes.

In simple directed trees, directed edges between nodes are used to guide the reaction trees (generally point from reactants to product). A simple directed tree is natural and intuitive, allowing for easy evaluation of each molecule node. However, this type of pathway structure may lead to misunderstandings of the reaction process. For example, in the simple reaction process shown in Fig. 8, it might be misunderstood that C can be obtained from either A or B, rather than requiring both A and B simultaneously. This

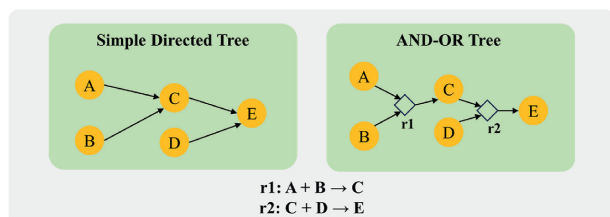


Fig. 8. Two different tree structures used in multi-step pathway planning.

ambiguity means that unless all reactions are “one-to-one”, additional information is needed to accurately reconstruct the reactions and correctly evaluate the pathways [115].

The AND-OR tree can effectively avoid this ambiguity by introduce a new type of nodes in the graph, which is more suitable for expressing pathways. The AND-OR tree is a type of bipartite graph that includes reaction nodes and molecule nodes, where nodes of the same type do not directly connect to each other in the graph. In an AND-OR tree, molecule nodes (OR nodes) and reaction nodes (AND nodes) alternate, with product nodes and reactant nodes connected to their corresponding reaction nodes, respectively. The AND-OR tree preserves all the causal relationships between reactants and products, making it easy to reconstruct the entire reaction process.

For both pathway structures mentioned above, there is a contradiction in how to handle the repeated appearance of molecule nodes in the graph. If such repetition is allowed, it may reduce the efficiency of the search. Conversely, if all identical nodes in the graph are merged directly, directed cycles may appear in the graph, which means the synthetic distance between two compound nodes may not be unique. This is also not conducive to the search for the optimal pathway. To address the issue of repeated molecule nodes and cycles in pathway graphs, scientists have experimented with using directed acyclic graphs (DAGs) and hypergraphs to represent reaction pathways [84,115]. A hypergraph is an extension of a simple directed tree that allows each molecule node to contain multiple compounds. Each expansion takes into account all substrates and other reactants required in this pathway and includes them in the current leaf node. However, because these structures require more careful handling during expansion and scoring, they have not yet been widely adopted.

4.2.2. Evaluation score calculation

The calculation of evaluation scores is part of pathway evaluation, and is central to pathway planning, guiding search algorithms to find search direction. Regardless of the pathway structure used, assessing which molecule node in the current pathway graph is more likely to lead to a viable synthetic pathway is crucial for both search efficiency and the feasibility of the final pathway obtained.

Based on whether the scores are dependent on the current state of the graph, the evaluation scores for molecule nodes can be categorized into two types: (1) predefined scores and (2) dynamic scores.

• Predefined scores

Predefined scores refer to the scores of molecule nodes calculated directly based on predefined formulas or pre-trained models. Common predefined scores include two categories: (1) quantified molecular properties [116] and (2) the synthetic accessibility (SA) of molecules [29].

Predefined scores based on molecular properties are calculated directly from existing molecular property data (such as compound cost, toxicity, etc.). Considering that the cost of a molecule is the easiest to obtain, it is frequently employed directly to score nodes, subsequently informing the prioritization of the nodes.

SA quantifies the difficulty of synthesizing a particular molecule. Ertl *et al.* [62] first proposed a representative SAScore ranged from 1 to 10, by analyzing the frequency of molecular structure fragments in the PubChem database. SAScore considers that the difficulty of synthesizing a molecule is determined by the groups it contains, and it statistically analyzes the frequency of molecular substructures. The higher the SAScore, the more difficult the molecule is to synthesize. Based on the similar idea, Voršilák *et al.* [117] proposed the synthetic Bayesian accessibility (SYBA) of molecules by performing Bayesian analysis on the frequency of ECFP8 fragments [60] of molecules. The SYBA score of a molecule is the sum of the contributions of all its fragments, and the SYBA score of each fragment represents the odds of its posterior probabilities of belonging to the easy-to-synthesize and hard-to-synthesize categories. Compared with SAScore, SYBA used Nonpher [118] to generate some molecules that are currently impossible to synthesize during training, which improved the generalization of the model.

Coley *et al.* [29] believed that molecules requiring more steps to reach are more difficult to synthesize, and they trained a neural network to calculate the synthetic complexity score (SCScore), based on real reactant-product pairs in Reaxys. They believe that the synthetic pathway should start from the molecule with the lowest synthetic complexity and gradually transform into the target product with the highest synthetic complexity.

In order to consider the global environment of the molecule and the interaction between different functional groups, Yu *et al.* [119] proposed a molecular SA evaluation model, GASA, based on graph attention mechanism. Utilizing GNNs, GASA automatically captures important structural features related to molecular SA. Compared with SAScore and SCScore, GASA can better distinguish structurally similar molecules.

The aforementioned models can quickly provide the SA of molecules, but since they are all trained on datasets of limited size, they may encounter difficulties when facing new molecules. Based on SMILES, SMALLER [12,115] is a simple and universal method for calculating molecular synthetic accessibility. It directly uses the length of the molecular SMILES to calculate the score. $SMALLER = \sum_{\text{substrings}} \text{SMILES_LEN}_i^n$. n is a hyperparameter that represents the preference for the size of the molecule. The preference of central disconnections and the preference of multicomponent reactions increases with increasing exponent n . SMALLER has the advantages of simplicity and excellent generalization.

Sometimes, scientists have preferences for the substances in a pathway, such as a tendency to avoid toxic compounds. By penalizing toxic compounds [120] with predefined scores (for example, a large negative value), we can effectively bias the pathway toward our preferences (*i.e.*, avoiding toxic intermediates) in subsequent

search processes [116]. This flexibility allows to freely incorporate knowledge into the ideal choice of pathway search through the design of scores.

• Dynamic scores

Dynamic scores are combined with the search process, and are based on both the evaluation of the molecule node (such as predefined scores) and the current state of the graph. Based on the source of dynamism involved in the score, dynamic scores can also be divided into two categories: (1) depend on the visit counts [84,116,121] and (2) depend on the “distance” of a node from the starting node (target product) [115,122].

The general calculation formulas for the 2 categories of dynamic scores on a molecule node m_i are shown in Eqs. (1) and (2), respectively.

$$\text{Score}(m_i) = \frac{f(m_i)}{g[n(m_i)]} + C(m_i) \times \frac{\sqrt{\ln(N)}}{G[n(m_i)]} \quad (1)$$

$$\text{Score}(m_i) = f(m_i) + h(m_i) \quad (2)$$

Dynamic scores that depend on visit counts are the sum of two components. $f(m_i)$ represents the static estimation score of the node, such as the predicted cost or SA. $n(m_i)$ is the visit count for node m_i , sometimes it is used as the denominator of the first term, that is, $g[n(m_i)] = n(m_i)$ or $g[n(m_i)] = 1$. The second component is dynamic, where N is the total visit count for the current state of the graph. Notably, the coefficient $C(m_i)$ could be a constant or a value related to the current node, such as the likelihood of the associated reaction [84]. Since the initial visit count for each node is 0, directly using it in the scoring could result in all unvisited leaf nodes being indistinguishable. Therefore, a function G is sometimes used to ensure the denominator of the dynamic term is not zero, for example, $G[n(m_i)] = n(m_i) + 1$. This type of dynamic score is highly flexible, allowing for the embedded functions or hyperparameters to be changed at any time to accommodate various retrosynthesis planning tasks.

Dynamic scores that depend on “distance” are also the sum of two terms, but their form is simpler. The first term $f(m_i)$ represents the estimated cost of the node, typically calculated using a pre-trained model. The second term $h(m_i)$ is the “distance” from the current node to the root node, which, for retrosynthesis planning tasks, is the sum of the costs of all reactions from the target product to the current node [122]. The cost of a reaction can be calculated based on the scores of its reactants and product. Since the scores involves calculations for both nodes and reactions, this type of dynamic score is generally applied only to AND-OR trees.

Dynamic scores, by considering the state of the current pathway graph, allow for a more reasonable evaluation of nodes. Specifically, in addition to utilizing a heuristic function unrelated to the search process as the static term, one may also refine a value network which is more apt to delineate the search space corresponding to the current target product by introducing a reinforcement learning strategy [121,123]. Theoretically, by iteratively alternating between planning and updating, the value network can be incrementally optimized, thereby steering the algorithm towards a more favorable trajectory.

4.2.3. Search algorithms and representative methods

During the process of pathway extension, due to the “combinatorial explosion” problem, the time complexity of expanding all nodes is unacceptable [115]. On one hand, scientists hope to prioritize the expansion and exploration of the currently most promising nodes (search direction), and on the other hand, they

also wish to consider to some currently suboptimal but unexplored nodes, because these seemingly “disadvantaged” nodes could potentially lead to globally optimal synthetic pathways.

To obtain synthetic pathways with foresight, a series of advanced heuristic search algorithms have been applied to multi-step pathway planning [122,124]. The search algorithms encompass two key processes: (1) determining the nodes in the current graph that will be expanded, and (2) updating the scores of nodes in the graph after each expansion. In fact, search algorithms act like a framework, while the aforementioned pathway structures and calculation methods of the score are flexible components within this framework. The same search algorithm can be paired with different pathway structures and scoring calculation methods. In this section, we will introduce three typical search algorithms, and present representative multi-step pathway planning models.

- Monte Carlo Tree Search

Monte Carlo Tree Search (MCTS) [125] is a search algorithm that iteratively learns the value of child nodes to find the optimal decision and guide the extension towards a specific search direction. MCTS is well-suited for tasks like synthetic pathway planning, which have a huge search space, as it can find the currently relatively optimal single-step pathway extension choices with less computational cost. Each iteration of MCTS includes four steps: selection, expansion, rollout (fast reasoning), and backpropagation. In multi-step pathway planning, the most important is the rollout step, which quickly obtains a synthetic pathway under the current leaf node through a certain sampling method. This rollout strategy can effectively improve the computational efficiency of MCTS iterations.

MCTS can be applied to both simple directed tree or AND-OR tree, but generally uses dynamic scores that rely on visit counts, thereby balancing exploration and exploitation. Segler *et al.* [84] combined three neural networks with MCTS to form 3N-MCTS. Based on the original MCTS, the three neural networks are used for template-based single-step pathway extension, the score calculation and selection of the nodes, and sampling in the rollout, respectively. The introduction of neural networks improves the speed and quality of the model.

- Proof-Number Search

Proof-Number Search (PNS) [126] is a search algorithm that can only be applied to AND-OR trees. During the search process, it maintains the proof number and disproof number for each node in the tree. The proof number and disproof number of a new leaf node is both set to 1, or estimated using a heuristic function. For retrosynthesis planning, the proof number represents the number of expansion steps required to prove that a node can be synthesized, while the disproof number represents the number of expansion steps indicating that a node cannot be synthesized. The goal of PNS is to determine whether a node can be proven (whether it can be synthesized) with the fewest number of expansions, that is, how to quickly ascertain the true state of a node.

Although PNS can be directly applied to retrosynthesis planning, there is a significant imbalance in the search space when using PNS and an AND-OR tree. A molecule may be synthesized through multiple candidate reactions, but most reactions only involve one or two reactants. This means that the branching factor (*i.e.*, the number of children each node can produce) of OR nodes is usually much greater than that of AND nodes [25,127]. To address this issue, Kishimoto *et al.* [127] proposed a depth-first proof-number search algorithm (DFPN-E), which uses a heuristic method to initialize the estimation of the cost of edges from OR nodes to AND nodes.

Compared with MCTS and traditional PNS, DFPN-E can provide almost the same number of pathway solutions while significantly reducing the number of node expansions, thereby increasing the efficiency of the search.

- A* Search

A* search [128] is a search algorithm that is commonly paired with AND-OR trees, using a heuristic dynamic score based on “distance” to measure the synthetic cost of each node. Similar to PNS, A* search also includes three steps: selection, expansion, and backup. After expansion, the score of each node between the new leaf node and the root node is updated bottom-up. The score of each AND node is calculated as the sum of the scores of all its child OR nodes. Conversely, the score of each OR node is determined by taking the maximum score of its child AND nodes. Since the heuristic function estimates the expected cost of nodes after expansion, the node score estimation in A* search is much faster than the rollout step in MCTS. Considering that the score includes the informed cost (edge cost function) related to the root node, A* search is expected to efficiently find better synthetic pathways.

Retro* proposed by Chen *et al.* [122] combines A* search with AND-OR trees, using a heuristic pre-trained model to calculate the expected cost of nodes, guiding the search towards more likely directions. To further improve the efficiency of node score calculation, Chematica [115,129] uses two predefined formulas, chemicals’ scoring function (CSF) and reaction scoring function (RSF), to directly calculate the expected cost and informed cost of nodes [130], respectively, based on their sum (CSF + RSF) to perform node selection.

The three search algorithms mentioned are essentially best-first searches, which may tend toward depth-first or breadth-first, so scientists sometimes artificially balance the two. One typical method is the beam search strategy, which expands the best *k* nodes at the same time and updates the scores of nodes in the graph based on the results of the expansion [115]. This strategy brings a forced breadth balance to the search. Another method is to alternate the use of score functions with different tendencies. For example, Chematica [115] uses scoring functions with *n* values of 1.5 and 3 in *SMALLER*. During expansion, it alternates between two queues to select nodes for expansion, ensuring a balance between the breadth and depth of the search.

Although existing multi-step pathway planning algorithms can provide some meaningful pathway results, they often fail in complex scenarios and tend to overlook reagents and chemical reaction conditions during the search process [122]. The ultimate goal of multi-step pathway planning is to hope that the model can function like human scientists by: (1) having the ability to plan with foresight (discarding the node that currently seems the best), exploring the most promising options in depth while also broadly exploring many different synthetic options; (2) exchanging and learning from the experience of different search strategies [115].

4.2.4. Evaluation for multi-step pathway planning

For multi-step pathway planning, scientists are concerned with their efficiency of generating multi-step synthetic pathways, as well as the quality of these pathways. Therefore, pathway evaluation utilizes a series of model-centric and pathway-centric metrics to evaluate multi-step pathway planning models and multi-step pathways, respectively. Given the unclear definition of a ‘good synthetic pathway’ and the limited research focus on multi-step pathway planning, there is currently no widely accepted metric for evaluating multi-step pathways. In this section, we will introduce several metrics that can be used to evaluate intuitively. The first three metrics are oriented towards evaluating a multi-step

model, while the last is designed to evaluate a multi-step synthetic pathway. The efficacy of multi-step pathway planning approaches depends not only on search algorithms and scoring functions, but also on single-step methods employed.

- Search efficiency

Due to the vast chemical search space, the efficiency of multi-step pathway planning models is crucial. The search efficiency can be measured in two ways [127,131]: (1) the number of times a single-step pathway extension model is called within the same time limit, i.e., the number of extensions in the search; (2) the total search time when calling the single-step pathway extension model the same number of times. As we previously emphasized, multi-step models and single-step models can be freely combined, meaning that we should choose the same single-step model (e.g., NeuralSym [28]) for different multi-step models to accurately compare their search efficiency [25].

- Success rate

Success rate [122,124] refers to the percentage of target molecules for which a synthetic pathway starting from building block molecules can be found by the multi-step pathway planning model within a limited time or number of expansions. Success rate provides an intuitive comparison of the ability of multi-step models to discover pathways. Without considering the quality and diversity, scientists hope that multi-step models should at least provide one pathway for each target molecule.

- Top- n accuracy

Similar to Top- k accuracy for single-step models, Top- n accuracy [131] is a metric to show how well a search model is in recovering the reference pathways. Compared to Top- k , multi-step models generally take a larger value for n in Top- n . This arises from the existing evaluation of multi-step pathways being somewhat indiscriminate, meaning that the top few pathways provided may not prove to be of significant utility. For the top n optimal pathways by the multi-step model, a tree edit distance (TED) [132] method or deep learning models [133] can be used to quickly determine if they contain the reference pathway. In order to better evaluate different multi-step models, one should select the same set of target compounds and corresponding reference pathways for each target compound. Genheden *et al.* [131] extracted two sets of 10000 pathways with maximum diversity (maximum inter-pathway distance) through sampling from reaction subnetworks in the USPTO patents and calculated the Top- n accuracy of different multi-step models on these two datasets.

- Pathway diversity

Pathway diversity [131] refers to the variety of different synthetic pathways identified by a multi-step model for synthesizing a target molecule. It is worth mentioning that because different search algorithms may find varying numbers and types of pathways for the same target, pathway diversity typically needs to be used in conjunction with other metrics to evaluate multi-step pathway planning models. Genheden *et al.* [131] proposed the PaRoutes framework, which utilizes the optimal number of clusters of pathways with the Silhouette method [134] to calculate pathway diversity. The diversity of the pathways found by a multi-step model indeed suggests a greater number of choices and opportunities for optimization, which can be particularly useful when facing practical constraints such as reagent availability, cost, or the need for specific reaction conditions.

In addition to the aforementioned metrics, scientists also use some other metrics (pathway length [122,124], average chemical complexity [107,135,136]) to evaluate synthetic pathways. Although the double-blind validation by Segler *et al.* [84] and the experimental validation by Klucznik *et al.* [129], which introduce human judgment, are very intuitive and practical, they are too time-consuming and costly to be extended to large-scale data.

In fact, to date, there is still no consensus on how to combine the metrics and what a good multi-step pathway planning model should look like. Some researchers believe that search speed is not as important as quality. Some researchers may believe that the algorithm should not find as many as possible pathways, as many pathways may not be reliable. As a pioneer, PaRoutes provides researchers with an open and fair framework for evaluating multi-step pathway planning models, including a set of reference pathways and a trained single-step pathway extension model.

5. Tools and Platforms

Compared to manual operations, computers possess powerful computational capabilities and can work continuously. Utilizing computers to automate the prediction of synthetic pathways can greatly facilitate the work of chemists. Although most data-driven models established by researchers using ML are targeted at specific problems, there are already some mature, readily useable retrosynthesis planning platforms that can address both single-step pathway extension and multi-step pathway planning, thereby achieving the generation of complete pathways. Table 2 organizes some tools/platforms related to retrosynthesis planning. These tools/platforms come in both open-source and commercial varieties, with commercial platforms generally able to offer better pathway recommendations with higher feasibility. Two major factors contribute to this difference. Firstly, open-source platforms primarily rely on reaction data from patents (such as USPTO datasets), whereas commercial platforms utilize reaction data from both literature and patents. Furthermore, these reaction data used by commercial platforms are typically subjected to meticulous manual refinement, resulting in better quantity and quality. Secondly, reaction data utilized by commercial platforms can contain additional information such as reaction conditions. This implies that commercial platforms are better equipped to evaluate the theoretical feasibility of pathways. Inevitably, although current platforms can perform the fundamental function of providing pathways, neither open-source nor commercial tools provides sufficiently precise and interpretable pathway evaluation and ranking. Consequently, they have difficulty in ensuring the practical feasibility of the outputted pathways.

6. Large Language Model in Retrosynthesis Planning

With the advancement of AI and the emergence of large models, ML models have gained the ability to process information in more complex modal. Large language models (LLMs), as the most typical and widely used category of large models, are capable of processing and analyzing vast amounts of data and have the ability to directly handle tasks for which they have not been explicitly trained [143]. LLMs can integrate knowledge from different domains and provide explanations in natural language, which, once encoded, can be applied to more specific domains, such as retrosynthesis planning. Although current LLMs are more general-purpose, and are less focused on performance in specific scientific domains, scientists have also begun to use some LLMs in retrosynthesis planning for tasks such as data preparation and preprocessing, as well as single-step pathway extension.

Table 2

Typical tools/platforms related to retrosynthesis planning.

Name	Developer	Data sources	Approaches	Availability	Characteristics
Pathway planning tools/platforms					
Chematica/SYNTHIA™ [24,115]	Merck	UnKnown	Template-based + A*	Commercial	Provide synthetic pathway planning and experimental validation for natural products.
ASKCOS [110]	MIT	USPTO, Pistachio	Template-based + MCTS	Free with registration	Provide various services related to synthetic pathway.
RXN/RoboRXN for chemistry [111]	IBM	USPTO, Pistachio	Template-free + tree search	Free with registration	Integrate ML, process automation, and cloud computing.
Reaxys Predictive Retrosynthesis [84]	Elsevier	Reaxys, ZINC	Template-based + MCTS	Commercial	Directly utilize the continuously updated, high-quality experimental reaction data in Reaxys.
SciFinder ⁿ (Chemplanner)	CAS	Scifinder	Template-based + MCTS	Commercial	Utilize the CAS dataset for synthetic pathway planning.
Retro*	Le Song	USPTO, eMolecules	Template-based + A*	Free	Propose a novel learning-based retrosynthesis planning algorithm to learn from previous planning experience and construct synthesis pathways.
RetroSynX [137]	Lei Zhang	Available literature in March's advanced organic chemistry	Template-based + breadth-first search	Free	Establish group contribution-based thermodynamic models to ensure the thermodynamic feasibility of proposed pathways.
RetroPath [116,138]	Jean-Loup Faulon	Rhea, MetaNetX	Template-based + MCTS	Free	Provide a full-service workflow for biosynthetic pathway planning.
XTMS [120]	Jean-Loup Faulon	MetaCyc, KEGG	Template-free	Free	Focus on bio-retrosynthesis planning, with the consideration pathway theoretical yield, node toxicity, and enzyme efficiency.
ATLAS [139]	EPFL	KEGG	Template-based	Free with registration	Focus on bio-retrosynthesis planning, with over 130000 virtual reactions which were generated using reaction templates.
Transform-Miner [140]	EMBL-EBI	KEGG	Template-based	Free for academics	Only use enzyme-catalyzed reactions to compose the synthetic pathway.
AiZynthFinder [141,142]	AstraZeneca	USPTO	Template-based + template-free + MCTS	Free	Integrate various single-step and multi-step planning models for retrosynthesis planning, while supporting extension.
Visualization tools					
Networkx	Networkx group	/	/	Open-source	Implement the drawing of various types of network diagrams.
Code bases and toolkits					
RDKit	Greg Landrum et al.	/	/	Open-source	A widely used open-source toolkit for chemoinformatics with functions include molecular descriptor generation and molecular operations. It is continuously updated and maintained by the community.
Indigo	EPAM	/	Unique algorithms developed by EPAM + standard algorithms	Open-source	A universal molecular toolkit [80] that can be used for molecular tasks including molecular fingerprinting, substructure search, AAM, and molecular visualization.

Before large models like ChatGPT received adequate attention, scientists had already started to use deep learning models to quickly, automatically, and accurately extract relevant information about chemical reactions from the literature. They utilized data-driven models to identify details provided in reaction information, such as reactants, solvents, reaction types, catalysts, and reaction temperatures [144]. However, these models faced challenges due to the unstructured format and diverse writing styles of chemical literature, as well as the scarcity of annotated reactions. With the introduction of LLMs, scientists have decomposed the task of reaction extraction and identification into multiple sub-tasks, with each task being handled by a different LLM, and allowing for information exchange between these models. Under this framework, AI can build a comprehensive, standardized format chemical reaction database through automated data collection, organization, and annotation, facilitating data preparation and preprocessing for retrosynthesis planning. Of course, prompt engineering can also be introduced to enhance the performance of LLMs within the framework.

LLMs primarily focus on the textual aspect of reaction information processing and extraction, considering less the chemical behavior corresponding to the reactions. For single-step pathway

extension in retrosynthesis planning, Liu *et al.* [145] recently proposed T-Rex, a text-assisted approach that incorporates chemical information provided by a LLM (ChatGPT) as part of the model input to assist in the prediction of reactants. T-Rex is semi-template-based, with the LLM being used in both P2S and S2R processes. In P2S, T-Rex integrates textual information from the LLM with graphical information from the molecular graph as input to a GNN, jointly identifying potential reaction centers. In S2R, T-Rex uses replies from the LLM to rank potential reaction centers and achieves synthon completion based on GNNs, thus reaches reaction prediction. Although T-Rex has not achieved Top-k accuracy close to state-of-the-art single-step models, this method of integrating textual embeddings from LLMs into the input of GNNs provides insights into the deep integration of large models with retrosynthesis planning.

The integration of retrosynthesis planning and large models is in its nascent stage. Considering the impressive capabilities and the challenge of relatively low operational efficiency associated with large models, we believe that in the future, large models could act as coordinators. They would be capable of grasping the overall extension strategy of the synthetic pathway, and orchestrating the models responsible for each module, thereby

achieving superior performance in synthesis of complex molecules.

7. Challenges and Opportunities

Retrosynthesis planning has become one of the most important tasks for the design and synthesis of functional molecules. With the rapid development of AI technology in recent years, the cross-integration of ML techniques and big data algorithms has brought new changes and opportunities to retrosynthesis planning. ML-assisted retrosynthesis planning dedicates to using computer learning of chemical knowledge to automate and efficiently assist chemists in designing more rational and effective synthetic pathways. Despite the significant progress made in the scale of reaction data that can be processed, the efficiency, quantity, and quality of pathways generated by ML-assisted retrosynthesis planning, there are still many challenges to be addressed.

7.1. The scale and quality of reaction data

Reaction data in data preparation module, as the fundamental part of the entire framework, is the basis of retrosynthesis planning and directly determines the quality of the candidate pathways output by retrosynthesis planning models [109].

Undoubtedly, the public datasets currently used in retrosynthesis planning research are far from sufficient for planning synthetic pathways for compounds across the entire chemical space. The most widely used USPTO series datasets contain only millions of reactions, while molecular databases contain over a billion annotated molecules. Although commercial datasets include more scale of carefully annotated reaction data, their requirement for subscription or purchase limits their widespread use in scientific research. In addition to insufficient quantity, the reaction datasets represented by USPTO also suffer from some quality issues, including insufficient diversity, lack of negative reaction data, unrecorded reaction conditions, and risks of data leakage. Ucak *et al.* [146] analyzed the atomic environments (AEs) covered by different datasets and found that the USPTO datasets encompass less than 3% of AE types found in the chemical space. This suggests that relying only on the USPTO datasets could hinder the training of a general retrosynthesis planning model, as many reaction space transformations are not represented in these datasets. Additionally, since failed reactions are almost never reported in literature or patents, the datasets rarely record reactions that could serve as negatives, which is not conducive to accurate assessment of reaction feasibility [119,147]. Finally, an interesting contradiction is that having AAM is generally considered an advantage of USPTO-50K, yet Wang *et al.* [148] found that USPTO-50K has risks of data leakage because the reaction centers tend to have a lower AAM number, leading to potential “cheating” by single-step pathway extension models, especially those that are semi-template-based. These deficiencies imply that future retrosynthesis planning needs to collect more reaction data with better diversity and ensure data availability, which requires a concerted effort by researchers in the community.

7.2. The development of evaluation system

Compared to the progress made by ML-assisted retrosynthesis planning in replicating references within datasets, there is still less attention paid to the evaluation system for retrosynthesis planning. In our view, a comprehensive, automated evaluation system is equally important for retrosynthesis planning, and there is currently a lack of a reasonable, universal, and practical application-oriented evaluation system. It should be emphasized

that an automated evaluation system can quickly provide us with rankings of pathways and reactions, but the final confirmation of feasibility still depends on experiments in actual laboratories.

A well-developed evaluation system includes two parts: (1) the selection of reasonable and comprehensive evaluation metrics, and (2) the sensible integration of these evaluation metrics. Those currently widely-used model-centric metrics are not sufficient to fully assess the feasibility of the reactions and pathways output by models. Pathway length, reaction cost [149,150] (not just the simple sum of reactant costs, but also considering the influence of various factors such as reaction kinetics and conditions), compatibility of reactions within the pathway, yield, product distribution, novelty and other metrics [149,151] based on the attributes of the pathway itself should all be considered in pathway evaluation. After determining metrics for evaluation, it is necessary to integrate these evaluation metrics sensibly, including the identification and handling of correlations between different metrics (some metrics may be highly correlated), quantitative processing (converting qualitative metrics into quantitative values), and normalization (fixing all indicators within the same range), as well as the calculation of weights and weighted summation for different metrics. Hyperparameters in some evaluation metrics [115,151], as well as the calculation of weights in the integration of metrics, bring flexibility to the evaluation system, allowing us to manually introduce different preferences for reactions or pathways.

Another interesting discussion is about the fairness of the evaluation for single-step models. In fact, although all works have chosen the same data set division ratio (*i.e.*, training/validation/test set as 80%/10%/10%), the division for most models is random. This means that comparisons between different models might not be simply achieved by the reported Top-*k* accuracy, and using a consistent experimental environment to evaluate all models is more repeatable and unbiased [25]. At the same time, due to the potential scaffold evaluation bias [152], we should avoid having molecules that are highly similar in both the test set and the training set. Some researchers have already noticed challenges brought by dataset division and data leakage. Zhong *et al.* [25] evaluated some models under the same dataset division and hardware environment, fairly comparing the performance of these typical models. To avoid data leakage, Wang *et al.* [148] used the Tanimoto similarity splitting method and shuffled the mapping numbers.

Of course, a more practical idea is to hope that ML-assisted retrosynthesis planning tools can think and judge like humans, but a fixed, well-developed evaluation system might weaken the foresight planning ability that retrosynthesis planning models possess. Future work should, on one hand, involve the collective efforts of all researchers to establish and maintain a well-developed evaluation system, with the PaRoutes framework being a pioneer in this aspect. On the other hand, there should be attempts to endow machines with a certain level of “intelligence”, perhaps the introduction of large models would be very intriguing.

7.3. The interpretability of models

Interpretability is a key dimension in evaluating ML models, explaining how models make reaction predictions and pathway planning, and is the basis for chemists to understand the synthetic pathways output by the model, helping to inspire chemists. Currently, almost all ML-assisted retrosynthesis planning research aims to replicate real synthetic pathways and lacks focus on model interpretability. Because the templates used are extracted from real reactions, template-based single-step models are generally considered to have a certain degree of interpretability. In addition to the interpretability problems posed by the complexity of the

model itself, the lack of interpretability in retrosynthesis planning models mainly stems from three aspects: (1) the inability to predict reaction conditions (catalysts, solvents, reaction temperature, etc.) while providing reaction predictions [1,84,148]; (2) the provision of unbalanced, incorrect, or incomplete reactions due to unbalanced reaction training data; (3) the lack of a reasonable and effective mature evaluation system [1,119,153]. The first two aspects are interrelated, as different reaction corrections or completions generally determine different reaction conditions. For complete reactions, scientists have already made some attempts to predict reaction conditions [154]. Regarding the third aspect, beyond SA of molecules, scientists are actively exploring the use of quantum chemistry and transition state theory, augmented by ML models, to predict the activation energy of reactions [65,155,156]. Decision-making is informed by reaction kinetics models, which not only enhance the kinetic performance of the pathways but also provide interpretability [157]. Recently, Wang *et al.* [148] proposed a single-step pathway planning model for molecular assembly guided by energy scores. The presence of energy scores provides interpretability to the single-step pathway planning decisions.

Improving the interpretability of retrosynthesis planning models remains a pressing issue. Employing clearer molecular representations, more transparent model architectures, and more comprehensive pathway evaluations are potential viable directions.

7.4. The computational complexity of models

Computational complexity, which includes both time complexity and space complexity, is generally related to the design and architecture of ML models and is an unavoidable issue for ML models beyond performance. Currently, ML-assisted retrosynthesis planning may encounter computational difficulties in certain situations, such as traversing data in large search spaces, running templates based on subgraph matching in single-step pathway extension, and iteration in multi-step pathway planning. It is worth mentioning that compared to the performance of the model, the speed (time complexity) advantage of the model may not be as important, because even if the model takes some time to output results, its efficiency is higher than manual pathway design. This may also be the reason why most models do not focus on and report computational efficiency. At the same time, although deep learning algorithms can better learn reaction behavior than simple ML algorithms, they have already encountered space complexity issues (insufficient memory) on large-scale reaction datasets [25]. The immense scale of the chemical search space brings great potential for development to retrosynthesis planning, but it also means that careful study of data-driven models is still needed to avoid excessive computational complexity.

7.5. Practical verification of pathway feasibility

Undoubtedly, with the development of ML, ML-assisted retrosynthesis planning has made significant progress in the past few years. However, current research focuses on only the first three modules of the framework, pursuing how to use data-driven models to efficiently output intuitively better synthetic pathways, while neglecting the experimental verification of these pathways. The number of novel pathways proposed entirely by ML and successfully synthesized in the laboratory remains limited, failing to meet the demands of chemists. Although the experimental verification of pathways is a costly and time-consuming process, it is indispensable for retrosynthesis planning. In the future, it is necessary to establish high-throughput automatic evaluation devices to verify the feasibility of pathways. The results of the experimental verification of pathways will also be used for model

correction, thereby achieving continuous iteration and optimization of the entire ML-assisted retrosynthesis planning framework. Ultimately, ML-assisted retrosynthesis planning tools aim to provide truly experimentally feasible new synthetic pathways for target compounds, with advantages in cost, yield, and other aspects, thereby becoming useful assistants for chemists. Of course, there is still a long way to go to achieve this goal. It requires the close cooperation of chemists and computer scientists. The emergence of automated experimental robots makes this goal no longer completely out of reach [158–160].

8. Conclusions

ML has brought rapid development and numerous opportunities to the field of retrosynthesis planning. In this review, we have organized the framework from the perspective of the ML-assisted retrosynthesis planning task, reviewed the current status of each module in the framework, and analyzed the pros and cons of different methods within each module. Furthermore, we address current challenges in ML-assisted retrosynthesis planning and propose potential research directions to enhance future efforts in this field.

Currently, research in the field of ML-assisted retrosynthesis planning is in an early stage of rapid development. We believe that in the future, with the efforts, communication, and cooperation of chemists, computer scientists, and all researchers, ML-assisted retrosynthesis planning will achieve the goals of being community-based, automated, and generalized. The framework of retrosynthesis planning is a very flexible and modular structure, the methods and ML models responsible for different modules can be freely combined. We look forward to advancing the discovery of superior pathways by focusing on improved coordination between modules, thereby fulfilling the original goal envisioned by chemists for retrosynthesis planning. Considering the tremendous potential of large models, large models will integrate more closely with retrosynthesis planning, just like current ML models, to promote the intelligent evolution of retrosynthesis planning.

CRedit Authorship Contribution Statements

Yixin Wei: Writing – review & editing, Writing – original draft, Project administration, Methodology. Leyu Shan: Writing – review & editing, Methodology, Investigation. Tong Qiu: Supervision, Funding acquisition. Diannan Lu: Supervision. Zheng Liu: Supervision, Conceptualization.

Declaration of Competing Interest

The authors declare that they have no known competing financial interests or personal relationships that could have appeared to influence the work reported in this paper.

Acknowledgements

This work is supported by the National Key Research and Development Program of China (2022ZD0117501).

References

- [1] J.X. Dong, M.Y. Zhao, Y.S. Liu, Y.S. Su, X.X. Zeng, Deep learning in retrosynthesis planning: datasets, models and tools, *Brief. Bioinform.* 23 (1) (2022) bbab391.
- [2] J.A. DiMasi, H.G. Grabowski, R.W. Hansen, Innovation in the pharmaceutical industry: new estimates of R&D costs, *J. Health Econ.* 47 (2016) 20–33.
- [3] T.J. Struble, J.C. Alvarez, S.P. Brown, M. Chytil, J. Cisar, R.L. Desjarlais, O. Engkvist, S.A. Frank, D.R. Greve, D.J. Griffin, X.J. Hou, J.W. Johannes, C.

- Kreatsoulas, B. Lahue, M. Mathea, G. Mogk, C.A. Nicolaou, A.D. Palmer, D.J. Price, R.I. Robinson, S. Salentin, L. Xing, T. Jaakkola, W.H. Green, R. Barzilay, C. W. Coley, K.F. Jensen, Current and future roles of artificial intelligence in medicinal chemistry synthesis, *J. Med. Chem.* 63 (16) (2020) 8667–8682.
- [4] A. Thakkar, S. Johansson, K. Jorner, D. Buttar, J.L. Reymond, O. Engkvist, Artificial intelligence and automation in computer aided synthesis planning, *React. Chem. Eng.* 6 (1) (2021) 27–51.
- [5] H.M. Chen, O. Engkvist, Y.H. Wang, M. Olivecrona, T. Blaschke, The rise of deep learning in drug discovery, *Drug Discov. Today* 23 (6) (2018) 1241–1250.
- [6] A. Kadurin, S. Nikolenko, K. Khrabrov, A. Aliper, A. Zhavoronkov, druGAN: an advanced generative adversarial autoencoder model for *de novo* generation of new molecules with desired molecular properties in silico, *Mol. Pharm.* 14 (9) (2017) 3098–3104.
- [7] T. Blaschke, M. Olivecrona, O. Engkvist, J. Bajorath, H.M. Chen, Application of generative autoencoder in *de novo* molecular design, *Mol. Inform.* 37 (1–2) (2018) 1700123.
- [8] C.W. Coley, W.H. Green, K.F. Jensen, RDChiral: an RDKit wrapper for handling stereochemistry in retrosynthetic template extraction and application, *J. Chem. Inf. Model.* 59 (6) (2019) 2529–2537.
- [9] B. Liu, B. Ramsundar, P. Kawthekar, J. Shi, J. Gomes, Q. Luu Nguyen, S. Ho, J. Sloane, P. Wender, V. Pande, Retrosynthetic reaction prediction using neural sequence-to-sequence models, *ACS Cent. Sci.* 3 (10) (2017) 1103–1113.
- [10] M. Ragoza, J. Hochuli, E. Idrobo, J. Sunseri, D.R. Koes, Protein-ligand scoring with convolutional neural networks, *J. Chem. Inf. Model.* 57 (4) (2017) 942–957.
- [11] A.A. Lee, Q.Y. Yang, V. Sresht, P. Bolgar, X.J. Hou, J.L. Klug-McLeod, C.R. Butler, Molecular Transformer unifies reaction prediction and retrosynthesis across pharma chemical space, *Chem. Commun.* 55 (81) (2019) 12152–12155.
- [12] E.J. Corey, The logic of chemical synthesis: multistep synthesis of complex carbogenic molecules (Nobel lecture), *Angew. Chem. Int. Ed.* 30 (5) (1991) 455–465.
- [13] E.J. Corey, W.T. Wipke, Computer-assisted design of complex organic syntheses, *Science* 166 (3902) (1969) 178–192.
- [14] S.Z. Ding, X.Q. Jiang, C. Meng, L.X. Sun, Z.Q. Wang, H.B. Yang, G.W. Shen, N. Xia, Application of artificial intelligence and big data technology in synthesis planning, *Sci. Sin.-Chim.* 53 (1) (2023) 66–78.
- [15] Y.J. Jiang, Y.M. Yu, M. Kong, Y. Mei, L.T. Yuan, Z.X. Huang, K. Kuang, Z.H. Wang, H.X. Yao, J. Zou, C.W. Coley, Y. Wei, Artificial intelligence for retrosynthesis prediction, *Engineering* 25 (2023) 32–50.
- [16] J. Goodman, Computer software review: reaxys, *J. Chem. Inf. Model.* 49 (12) (2009) 2897–2898.
- [17] I. Levin, M.J. Liu, C.A. Voigt, C.W. Coley, Merging enzymatic and synthetic chemistry with computational synthesis planning, *Nat. Commun.* 13 (2022) 7747.
- [18] X. Zhang, E. King-Smith, L.B. Dong, L.C. Yang, J.D. Rudolf, B. Shen, H. Renata, Divergent synthesis of complex diterpenes through a hybrid oxidative approach, *Science* 369 (6505) (2020) 799–806.
- [19] J. Li, A. Amatuni, H. Renata, Recent advances in the chemoenzymatic synthesis of bioactive natural products, *Curr. Opin. Chem. Biol.* 55 (2020) 111–118.
- [20] N.R. Patel, C.C. Nawrat, M. McLaughlin, Y.J. Xu, M.A. Huffman, H. Yang, H.M. Li, A.M. Whittaker, T. Andreani, F. Lévesque, A. Fryszkowska, A. Brunskill, D. M. Tschäen, K.M. Maloney, Synthesis of islatravir enabled by a catalytic, enantioselective alkynylation of a ketone, *Org. Lett.* 22 (12) (2020) 4659–4664.
- [21] Y. Li, W. Liang, L. Peng, D. Zhang, C. Yang, K.C. Li, Predicting drug-target interactions via dual-stream graph neural network, *IEEE ACM Trans. Comput. Biol. Bioinf.* 21 (4) (2024) 948–958.
- [22] A. Mayr, G. Klambauer, T. Unterthiner, S. Hochreiter, DeepTox: toxicity prediction using deep learning, *Front. Environ. Sci.* 3 (2016) 80.
- [23] S. Szymkuć, E.P. Gajewska, T. Klucznik, K. Molga, P. Dittwald, M. Startek, M. Bajczyk, B.A. Grzybowski, Computer-assisted synthetic planning: the end of the beginning, *Angew. Chem. Int. Ed.* 55 (20) (2016) 5904–5937.
- [24] B. Mikulak-Klucznik, P. Gołębiewska, A.A. Bayly, O. Popik, T. Klucznik, S. Szymkuć, E.P. Gajewska, P. Dittwald, O. Staszewska-Krajewska, W. Beker, T. Badowski, K.A. Scheidt, K. Molga, J. Mlynarski, M. Mrksich, B.A. Grzybowski, Computational planning of the synthesis of complex natural products, *Nature* 588 (7836) (2020) 83–88.
- [25] Z.P. Zhong, J. Song, Z.L. Feng, T.T. Liu, L.X. Jia, S.L. Yao, T.J. Hou, M.L. Song, Recent advances in deep learning for retrosynthesis, *Wires Comput. Mol. Sci.* 14 (1) (2024) e1694.
- [26] Q. Zhang, J. Liu, W. Zhang, F. Yang, Z.H. Yang, X.L. Zhang, A multi-stream network for retrosynthesis prediction, *Front. Comput. Sci.* 18 (2) (2023) 182906.
- [27] K. Zhang, V. Mann, V. Venkatasubramanian, G-MATT: single-step retrosynthesis prediction using molecular grammar tree transformer, *AIChE J.* 70 (1) (2024) e18244.
- [28] M.H.S. Segler, M.P. Waller, Neural-symbolic machine learning for retrosynthesis and reaction prediction, *Chemistry* 23 (25) (2017) 5966–5971.
- [29] C.W. Coley, L. Rogers, W.H. Green, K.F. Jensen, SCScore: synthetic complexity learned from a reaction corpus, *J. Chem. Inf. Model.* 58 (2) (2018) 252–261.
- [30] L. Fang, J.R. Li, M. Zhao, L. Tan, J.G. Lou, Single-step retrosynthesis prediction by leveraging commonly preserved substructures, *Nat. Commun.* 14 (1) (2023) 2446.
- [31] I.V. Tetko, P. Karpov, R. Van Deursen, G. Godin, State-of-the-art augmented NLP transformer models for direct and single-step retrosynthesis, *Nat. Commun.* 11 (1) (2020) 5575.
- [32] P. Reiser, M. Neubert, A. Eberhard, L. Torresi, C. Zhou, C. Shao, H. Metni, C. van Hoesel, H. Schopmans, T. Sommer, P. Friederich, Graph neural networks for materials science and chemistry, *Commun. Mater.* 3 (1) (2022) 93.
- [33] A. Toniato, P. Schwaller, A. Cardinale, J. Geluykens, T. Laino, Unassisted noise reduction of chemical reaction datasets, *Nat. Mach. Intell.* 3 (2021) 485–494.
- [34] L. Daniel, Chemical reactions from US patents (1976–Sep2016), figshare (2017), <https://doi.org/10.6084/m9.figshare.5104873.v1>. Accessed Nov 2024.
- [35] W. Jin, C.W. Coley, R. Barzilay, T. Jaakkola, Predicting organic reaction outcomes with Weisfeiler-Lehman network, in: Proceedings of the 31st International Conference on Neural Information Processing Systems, US, Long Beach, CA, 2017.
- [36] J.Y. He, D.Q. Nguyen, S. Akhondi, C. Druckenbrodt, C. Thorne, R. Hoessel, Z. Afzal, Z.N. Zhai, B.Y. Fang, H. Yoshikawa, A. Albahem, L. Cavedon, T. Cohn, T. Baldwin, K.M. Verspoor, Overview of ChEMU 2020: named entity recognition and event extraction of chemical reactions from patents, *Conf. Lab. Evaluat. Forum, Thessaloniki, Greece* 12260 (2020) 237–254.
- [37] A. Morgat, T. Lombardot, K.B. Axelsen, L. Aimo, A. Niknejad, N. Hyka-Nouspikel, E. Coudert, M. Pozzato, M. Pagni, S. Moretti, S. Rosanoff, J. Onwubiko, L. Bougueleret, I. Xenarios, N. Redaschi, A. Bridge, Updates in Rhea – an expert curated resource of biochemical reactions, *Nucleic Acids Res.* 45 (D1) (2017) D415–D418.
- [38] R. Mercado, S.M. Kearnes, C.W. Coley, Data sharing in chemistry: lessons learned and a case for mandating structured reaction data, *J. Chem. Inf. Model.* 63 (14) (2023) 4253–4265.
- [39] S.M. Kearnes, M.R. Maser, M. Wlekliński, A. Kast, A.G. Doyle, S.D. Dreher, J.M. Hawkins, K.F. Jensen, C.W. Coley, The open reaction database, *J. Am. Chem. Soc.* 143 (45) (2021) 18820–18826.
- [40] S. Kim, J. Chen, T.J. Cheng, A. Gindulyte, J. He, S.Q. He, Q.L. Li, B.A. Shoemaker, P.A. Thiessen, B. Yu, L. Zaslavsky, J. Zhang, E.E. Bolton, PubChem 2019 update: improved access to chemical data, *Nucleic Acids Res.* 47 (D1) (2019) D1102–D1109.
- [41] M. Kanehisa, S. Goto, KEGG Kyoto encyclopedia of genes and genomes, *Nucleic Acids Res.* 28 (1) (2000) 27–30.
- [42] J. Hastings, G. Owen, A. Dekker, M. Ennis, N. Kale, V. Muthukrishnan, S. Turner, N. Swainston, P. Mendes, C. Steinbeck, ChEBI in 2016: improved services and an expanding collection of metabolites, *Nucleic Acids Res.* 44 (D1) (2016) D1214–D1219.
- [43] L. Jeske, S. Placzek, I. Schomburg, A. Chang, D. Schomburg, BRENDA in 2019: a European ELIXIR core data resource, *Nucleic Acids Res.* 47 (D1) (2019) D542–D549.
- [44] T.U. Consortium, UniProt: the universal protein knowledgebase in 2023, *Nucleic Acids Res.* 51 (D1) (2022) D523–D531.
- [45] A. Jain, S.P. Ong, G. Hautier, W. Chen, W.D. Richards, S. Dacek, S. Cholia, D. Gunter, D. Skinner, G. Ceder, K.A. Persson, Commentary: the Materials Project: a materials genome approach to accelerating materials innovation, *APL Mater.* 1 (1) (2013) 011002.
- [46] K.T. Winther, M.J. Hoffmann, J.R. Boes, O. Mamun, M. Bajdich, T. Bligaard, Catalysis-Hub.org, an open electronic structure database for surface reactions, *Sci. Data* 6 (2019) 75.
- [47] N. Schneider, N. Stiefl, G.A. Landrum, What's what: the (nearly) definitive guide to reaction role assignment, *J. Chem. Inf. Model.* 56 (12) (2016) 2336–2346.
- [48] Y.H. Ding, B. Qiang, Q.X. Chen, Y.Q. Liu, L.R. Zhang, Z.M. Liu, Exploring chemical reaction space with machine learning models: representation and feature perspective, *J. Chem. Inf. Model.* 64 (8) (2024) 2955–2970.
- [49] Z.P. Zhong, J. Song, Z.L. Feng, T.T. Liu, L.X. Jia, S.L. Yao, M. Wu, T.J. Hou, M.L. Song, Root-aligned SMILES: a tight representation for chemical reaction prediction, *Chem. Sci.* 13 (31) (2022) 9023–9034.
- [50] Y.X. Wei, T. Qiu, MDs-NP: a property prediction model construction procedure for naphtha based on molecular dynamics simulation, *J. Phys. Condens. Matter* 36 (31) (2024) 315402.
- [51] RDKit: open-source cheminformatics, <https://www.rdkit.org/>. Accessed Nov 2024.
- [52] A.A. Toropov, A.P. Toropova, QSPR/QSAR: state-of-art, weirdness, the future, *Molecules* 25 (6) (2020) 1292.
- [53] D.S. Wigh, J.M. Goodman, A.A. Lapkin, A review of molecular representation in the age of machine learning, *Wires Comput. Mol. Sci.* 12 (5) (2022) e1603.
- [54] A.M. Moore, A line-formula chemical notation, *J. Am. Chem. Soc.* 77 (7) (1955) 2032.
- [55] D. Weininger, SMILES, a chemical language and information system. 1. Introduction to methodology and encoding rules, *J. Chem. Inf. Comput. Sci.* 28 (1) (1988) 31–36.
- [56] I. Daylight Chemical Information Systems, Daylight theory manual. <https://www.daylight.com/dayhtml/doc/theory/>, 2008. Accessed Nov 2024.
- [57] N.M. O'Boyle, A. Dalke, DeepSMILES: an adaptation of SMILES for use in machine-learning of chemical structures, *ChemRxiv* (2018). <https://chemrxiv.org/engage/chemrxiv/article-details/60c73ed6567dfe7e5fec388d>.
- [58] A. Lo, R. Pollice, A. Nigam, A.D. White, M. Krenn, A. Aspuru-Guzik, Recent advances in the self-referencing embedded strings (SELFIES) library, *Dig. Dis.* 2 (4) (2023) 897–908.
- [59] V.D. Hähnke, E.E. Bolton, S.H. Bryant, PubChem atom environments, *J. Cheminf.* 7 (1) (2015) 41.

- [60] D. Rogers, M. Hahn, Extended-connectivity fingerprints, *J. Chem. Inf. Model.* 50 (5) (2010) 742–754.
- [61] J.Y. Deng, Z.B. Yang, H.H. Wang, I. Ojima, D. Samaras, F.S. Wang, A systematic study of key elements underlying molecular property prediction, *Nat. Commun.* 14 (1) (2023) 6395.
- [62] P. Ertl, A. Schuffenhauer, Estimation of synthetic accessibility score of drug-like molecules based on molecular complexity and fragment contributions, *J. Cheminf.* 1 (1) (2009) 8.
- [63] F. Nikitin, O. Isayev, V. Strijov, DRACON: disconnected graph neural network for atom mapping in chemical reactions, *Phys. Chem. Chem. Phys.* 22 (45) (2020) 26478–26486.
- [64] A.I. Lin, T.I. Madzhidov, O. Klimchuk, R.I. Nugmanov, I.S. Antipin, A. Varnek, Automatized assessment of protective group reactivity: a step toward big reaction data analysis, *J. Chem. Inf. Model.* 56 (11) (2016) 2140–2148.
- [65] E. Heid, W.H. Green, Machine learning of reaction properties via learned representations of the condensed graph of reaction, *J. Chem. Inf. Model.* 62 (9) (2022) 2101–2110.
- [66] C.W. Coley, R. Barzilay, T.S. Jaakkola, W.H. Green, K.F. Jensen, Prediction of organic reaction outcomes using machine learning, *ACS Cent. Sci.* 3 (5) (2017) 434–443.
- [67] A. Varnek, D. Fourches, F. Hoonakker, V.P. Solov'ev, Substructural fragments: an universal language to encode reactions, molecular and supramolecular structures, *J. Comput. Aided Mol. Des.* 19 (9–10) (2005) 693–703.
- [68] W. Jaworski, S. Szymkuć, B. Mikulak-Klucznik, K. Piecuch, T. Klucznik, M. Kaźmierowski, J. Rydzewski, A. Gambin, B.A. Grzybowski, Automatic mapping of atoms across both simple and complex chemical reactions, *Nat. Commun.* 10 (1) (2019) 1434.
- [69] M. Latendresse, J.P. Malerich, M. Travers, P.D. Karp, Accurate atom-mapping computation for biochemical reactions, *J. Chem. Inf. Model.* 52 (11) (2012) 2970–2982.
- [70] R. Körner, J. Apostolakis, Automatic determination of reaction mappings and reaction center information. 1. The imaginary transition state energy approach, *J. Chem. Inf. Model.* 48 (6) (2008) 1181–1189.
- [71] P. Schwaller, B. Hoover, J.L. Reymond, H. Strobelt, T. Laino, Extraction of organic chemistry grammar from unsupervised learning of chemical reactions, *Sci. Adv.* 7 (15) (2021) eabe4166.
- [72] R. Nugmanov, N. Dyubankova, A. Gedich, J.K. Wegner, Bidirectional graphormer for reactivity understanding: neural network trained to reaction atom-to-atom mapping task, *J. Chem. Inf. Model.* 62 (14) (2022) 3307–3315.
- [73] S.A. Rahman, G. Torrance, L. Baldacci, S. Martínez Cuesta, F. Fenninger, N. Gopal, S. Choudhary, J.W. May, G.L. Holliday, C. Steinbeck, J.M. Thornton, Reaction decoder tool (RDT): extracting features from chemical reactions, *Bioinformatics* 32 (13) (2016) 2065–2066.
- [74] H. Kraut, J. Eiblmaier, G. Grethe, P. Löw, H. Matuszczyk, H. Saller, Algorithm for reaction classification, *J. Chem. Inf. Model.* 53 (11) (2013) 2884–2895.
- [75] E.L. First, C.E. Gounaris, C.A. Floudas, Stereochemically consistent reaction mapping and identification of multiple reaction mechanisms through integer linear optimization, *J. Chem. Inf. Model.* 52 (1) (2012) 84–92.
- [76] D. Fooshee, A. Andronico, P. Baldi, ReactionMap: an efficient atom-mapping algorithm for chemical reactions, *J. Chem. Inf. Model.* 53 (11) (2013) 2812–2819.
- [77] S. Chen, S. An, R. Babazade, Y. Jung, Precise atom-to-atom mapping for organic reactions via human-in-the-loop machine learning, *Nat. Commun.* 15 (1) (2024) 2250.
- [78] W.L. Chen, D.Z. Chen, K.T. Taylor, Automatic reaction mapping and reaction center detection, *Wires Comput. Mol. Sci.* 3 (6) (2013) 560–593.
- [79] G.A. Preciat Gonzalez, L.R.P. El Assal, A. Noronha, I. Thiele, H.S. Haraldsdóttir, R.M.T. Fleming, Comparative evaluation of atom mapping algorithms for balanced metabolic reactions: application to recon 3D, *J. Cheminf.* 9 (1) (2017) 39.
- [80] Indigo Toolkit, <https://lifescience.opensource.epam.com/indigo/>. Accessed Nov 2024.
- [81] J. Law, Z. Zsoldos, A. Simon, D. Reid, Y. Liu, S.Y. Khew, A.P. Johnson, S. Major, R.A. Wade, H.Y. Ando, Route Designer: a retrosynthetic analysis tool utilizing automated retrosynthetic rule generation, *J. Chem. Inf. Model.* 49 (3) (2009) 593–602.
- [82] J.L. Baylon, N.A. Cilfone, J.R. Gulcher, T.W. Chittenden, Enhancing retrosynthetic reaction prediction with deep learning using multiscale reaction classification, *J. Chem. Inf. Model.* 59 (2) (2019) 673–688.
- [83] C.D. Christ, M. Zentgraf, J.M. Kriegl, Mining electronic laboratory notebooks: analysis, retrosynthesis, and reaction based enumeration, *J. Chem. Inf. Model.* 52 (7) (2012) 1745–1756.
- [84] M.H.S. Segler, M. Preuss, M.P. Waller, Planning chemical syntheses with deep neural networks and symbolic AI, *Nature* 555 (7698) (2018) 604–610.
- [85] M.H.S. Segler, M.P. Waller, Modelling chemical reasoning to predict and invent reactions, *Chemistry* 23 (25) (2017) 6118–6128.
- [86] C.W. Coley, W.H. Green, K.F. Jensen, Machine learning in computer-aided synthesis planning, *Acc. Chem. Res.* 51 (5) (2018) 1281–1289.
- [87] C.W. Coley, L. Rogers, W.H. Green, K.F. Jensen, Computer-assisted retrosynthesis based on molecular similarity, *ACS Cent. Sci.* 3 (12) (2017) 1237–1245.
- [88] H. Dai, C. Li, C.W. Coley, B. Dai, L. Song, Retrosynthesis prediction with conditional graph logic network, in: Proceedings of the 33rd International Conference on Neural Information Processing Systems, Vancouver, Canada, 2019.
- [89] S. Chen, Y. Jung, Deep retrosynthetic reaction prediction using local reactivity and global attention, *JACS Au.* 1 (10) (2021) 1612–1620.
- [90] C.C. Yan, P.L. Zhao, C. Lu, Y. Yu, J.Z. Huang, RetroComposer: composing templates for template-based retrosynthesis prediction, *Biomolecules* 12 (9) (2022) 1325.
- [91] X.R. Wang, Y.Q. Li, J.Z. Qiu, G.Y. Chen, H.X. Liu, B.B. Liao, C.Y. Hsieh, X.J. Yao, RetroPrime: a Diverse, plausible and Transformer-based method for Single-Step retrosynthesis predictions, *Chem. Eng. J.* 420 (2021) 129845.
- [92] Z.Q. Chen, O.R. Ayinde, J.R. Fuchs, H. Sun, X. Ning, G²Retro as a two-step graph generative models for retrosynthesis prediction, *Commun. Chem.* 6 (1) (2023) 102.
- [93] C. Shi, M. Xu, H. Guo, M. Zhang, J. Tang, A graph to graphs framework for retrosynthesis prediction, in: Proceedings of the 37th International Conference on Machine Learning, Vancouver, Canada, 2020.
- [94] J.H. Liu, C.C. Yan, Y. Yu, C. Lu, J.Z. Huang, L. Ou-Yang, P.L. Zhao, MARS: a motif-based autoregressive model for retrosynthesis prediction, *Bioinformatics* 40 (3) (2024) btac115.
- [95] W. Zhong, Z. Yang, C.Y. Chen, Retrosynthesis prediction using an end-to-end graph generative architecture for molecular graph editing, *Nat. Commun.* 14 (1) (2023) 3009.
- [96] M. Sacha, M. Biał, P. Byrski, P. Dąbrowski-Tumański, M. Chromiński, R. Loska, P. Włodarczyk-Pruszyński, S. Jastrzębski, Molecule edit graph attention network: modeling chemical reactions as sequences of graph edits, *J. Chem. Inf. Model.* 61 (7) (2021) 3273–3284.
- [97] R.X. Sun, H. Dai, L. Li, S. Kearnes, B. Dai, Towards understanding retrosynthesis by energy-based models, Proceedings of the 35th International Conference on Neural Information Processing Systems, San Diego, CA, USA, 2021.
- [98] Y. Wan, C.-Y. Hsieh, B. Liao, S. Zhang, Retroformer: pushing the limits of end-to-end retrosynthesis transformer, in: Proceedings of the 39th International Conference on Machine Learning, California, USA, 2022.
- [99] P. Karpov, G. Godin, I.V. Tetko, A transformer model for retrosynthesis, artificial neural networks and machine learning – ICANN 2019: workshop and special sessions, in: Proceedings of 28th International Conference on Artificial Neural Networks, Munich, Germany, 2019.
- [100] A. Vaswani, N. Shazeer, N. Parmar, J. Uszkoreit, L. Jones, A.N. Gomez, L. Kaiser, I. Polosukhin, Attention is all you need, in: 31st Annual Conference on Neural Information Processing Systems, NIPS, Long Beach, CA, USA, 2017.
- [101] S.J. Zheng, J.H. Rao, Z.Y. Zhang, J. Xu, Y.D. Yang, Predicting retrosynthetic reactions using self-corrected transformer neural networks, *J. Chem. Inf. Model.* 60 (1) (2020) 47–55.
- [102] L. Yao, W.T. Guo, Z. Wang, S. Xiang, W.T. Liu, G.L. Ke, Node-aligned graph-to-graph: elevating template-free deep learning approaches in single-step retrosynthesis, *JACS Au* 4 (3) (2024) 992–1003.
- [103] P. Veličković, G. Cucurull, A. Casanova, A. Romero, P. Liò, Y. Bengio, Graph attention networks, arXiv:1710.10903 (2017), <https://doi.org/10.48550/arXiv.1710.10903>.
- [104] K. Yang, K. Swanson, W.G. Jin, C. Coley, P. Eiden, H. Gao, A. Guzman-Perez, T. Hopper, B. Kelley, M. Mathea, A. Palmer, V. Settels, T. Jaakkola, K. Jensen, R. Barzilay, Analyzing learned molecular representations for property prediction, *J. Chem. Inf. Model.* 59 (8) (2019) 3370–3388.
- [105] B. Chen, T.X. Shen, T.S. Jaakkola, R. Barzilay, Learning to make generalizable and diverse predictions for retrosynthesis, arXiv:1910.09688 (2019), <https://doi.org/10.48550/arXiv.1910.09688>.
- [106] Z.K. Tu, C.W. Coley, Permutation invariant graph-to-sequence model for template-free retrosynthesis and reaction prediction, *J. Chem. Inf. Model.* 62 (15) (2022) 3503–3513.
- [107] K.J. Lin, Y.J. Xu, J.F. Pei, L.H. Lai, Automatic retrosynthetic route planning using template-free models, *Chem. Sci.* 11 (12) (2020) 3355–3364.
- [108] Y.C. Yan, Y. Zhao, H.F. Yao, J. Feng, L. Liang, W.J. Han, X.H. Xu, C.T. Pu, C.D. Zang, L.F. Chen, Y.Y. Li, H.C. Liu, T. Lu, Y.D. Chen, Y.M. Zhang, RBPB: deep retrosynthesis reaction prediction based on byproducts, *J. Chem. Inf. Model.* 63 (19) (2023) 5956–5970.
- [109] A. Thakkar, T. Kogej, J.L. Reymond, O. Engkvist, E.J. Bjerrum, Datasets and their influence on the development of computer assisted synthesis planning tools in the pharmaceutical domain, *Chem. Sci.* 11 (1) (2020) 154–168.
- [110] M.E. Fortunato, C.W. Coley, B.C. Barnes, K.F. Jensen, Data augmentation and pretraining for template-based retrosynthetic prediction in computer-aided synthesis planning, *J. Chem. Inf. Model.* 60 (7) (2020) 3398–3407.
- [111] P. Schwaller, R. Petraglia, V. Zullo, V.H. Nair, R.A. Haeuselmann, R. Pisoni, C. Bekas, A. Iuliano, T. Laino, Predicting retrosynthetic pathways using transformer-based models and a hyper-graph exploration strategy, *Chem. Sci.* 11 (12) (2020) 3316–3325.
- [112] H. Lee, S. Ahn, S. Seo, Y. Song, E. Yang, S.J. Hwang, J. Shin, RetCL: a selection-based approach for retrosynthesis via contrastive learning, in: Proceedings of the Thirtieth International Joint Conference on Artificial Intelligence, Montreal, Canada, 2021.
- [113] Y. Xia, T. He, X. Tan, F. Tian, D. He, T. Qin, Tied transformers: neural machine translation with shared encoder and decoder, in: Proceedings of the Thirty-Third AAAI Conference on Artificial Intelligence and Thirty-First Innovative Applications of Artificial Intelligence Conference and Ninth AAAI Symposium on Educational Advances in Artificial Intelligence, Honolulu, Hawaii, USA, 2019.
- [114] O.S. Vaidya, S. Kumar, Analytic hierarchy process: an overview of applications, *Eur. J. Oper. Res.* 169 (1) (2006) 1–29.

- [115] B.A. Grzybowski, T. Badowski, K. Molga, S. Szymkuć, Network search algorithms and scoring functions for advanced-level computerized synthesis planning, *Wiley Interdiscip. Rev. Comput. Mol. Sci.* 13 (1) (2023) e1630.
- [116] M. Koch, T. Duigou, J.L. Faulon, Reinforcement learning for bioretrosynthesis, *ACS Synth. Biol.* 9 (1) (2020) 157–168.
- [117] M. Voršilák, M. Kolár, I. Čmelo, D. Svozil, SYBA: Bayesian estimation of synthetic accessibility of organic compounds, *J. Cheminf.* 12 (1) (2020) 35.
- [118] M. Voršilák, D. Svozil, Nonpher: computational method for design of hard-to-synthesize structures, *J. Cheminf.* 9 (1) (2017) 20.
- [119] J.H. Yu, J.K. Wang, H. Zhao, J.B. Gao, Y. Kang, D.S. Cao, Z. Wang, T.J. Hou, Organic compound synthetic accessibility prediction based on the graph attention mechanism, *J. Chem. Inf. Model.* 62 (12) (2022) 2973–2986.
- [120] P. Carbonell, P. Parutto, J. Herisson, S.B. Pandit, J.L. Faulon, XTMS: pathway design in an eXTended metabolic space, *Nucleic Acids Res.* 42 (Web Server issue) (2014) W389–W394.
- [121] X.X. Wang, Y.J. Qian, H.Y. Gao, C. Coley, Y.M. Mo, R. Barzilay, K.F. Jensen, Towards efficient discovery of green synthetic pathways with Monte Carlo tree search and reinforcement learning, *Chem. Sci.* 11 (40) (2020) 10959–10972.
- [122] B. Chen, C. Li, H. Dai, L. Song, Retro*: learning retrosynthetic planning with neural guided a* search, in: Proceedings of the 37th International Conference on Machine Learning, Long Beach, California, USA, 2020.
- [123] G. Liu, D. Xue, S. Xie, Y. Xia, A. Tripp, K. Maziarz, M. Segler, T. Qin, Z. Zhang, T.-Y. Liu, Retrosynthetic planning with dual value networks, in: Proceedings of the 40th International Conference on Machine Learning, San Francisco, USA, 2023.
- [124] J. Kim, S. Ahn, H. Lee, J. Shin, Self-improved retrosynthetic planning, in: Proceedings of the 38th International Conference on Machine Learning, Los Angeles, California, USA, 2021.
- [125] T.X. Ou, Y.N. Lu, X.S. Wu, J.W. Cao, Monte Carlo tree search: a survey of theories and applications, in: 2022 3rd International Conference on Big Data, Artificial Intelligence and Internet of Things Engineering (ICBAIE), IEEE, Xi'an, China, 2022.
- [126] L.V. Allis, M. van der Meulen, H.J. van den Herik, Proof-number search, *Artif. Intell.* 66 (1) (1994) 91–124.
- [127] A. Kishimoto, B. Buesser, B. Chen, A. Botea, Depth-first proof-number search with heuristic edge cost and application to chemical synthesis planning, in: Proceedings of the 33rd International Conference on Neural Information Processing Systems, ACM, Vancouver, 2019.
- [128] P.E. Hart, N.J. Nilsson, B. Raphael, A formal basis for the heuristic determination of minimum cost paths, *IEEE Trans. Syst. Sci. Cybern.* 4 (2) (1968) 100–107.
- [129] T. Klucznik, B. Mikulak-Klucznik, M.P. McCormack, H. Lima, S. Szymkuć, M. Bhowmick, K. Molga, Y.B. Zhou, L. Rickershauser, E.P. Gajewska, A. Touchkine, P. Dittwald, M.P. Startek, G.J. Kirkovits, R. Roszak, A. Adamski, B. Sieredzińska, M. Mrksich, S.L.J. Trice, B.A. Grzybowski, Efficient syntheses of diverse, medically relevant targets planned by computer and executed in the laboratory, *Chem* 4 (3) (2018) 522–532.
- [130] E.W. Dijkstra, A note on two problems in connexion with graphs, *Numer. Math.* 1 (1) (1959) 269–271.
- [131] S. Genheden, E. Bjerrum, PaRoutes: towards a framework for benchmarking retrosynthesis route predictions, *Dig. Dis.* 1 (4) (2022) 527–539.
- [132] S. Genheden, O. Engkvist, E. Bjerrum, Clustering of synthetic routes using tree edit distance, *J. Chem. Inf. Model.* 61 (8) (2021) 3899–3907.
- [133] S. Genheden, O. Engkvist, E. Bjerrum, Fast prediction of distances between synthetic routes with deep learning, *Mach. Learn.: Sci. Technol.* 3 (1) (2022) 015018.
- [134] P.J. Rousseeuw, Silhouettes: a graphical aid to the interpretation and validation of cluster analysis, *J. Comput. Appl. Math.* 20 (1987) 53–65.
- [135] J.S. Schreck, C.W. Coley, K.J.M. Bishop, Learning retrosynthetic planning through simulated experience, *ACS Cent. Sci.* 5 (6) (2019) 970–981.
- [136] R. Shibukawa, S. Ishida, K. Yoshizoe, K. Wasa, K. Takasu, Y. Okuno, K. Terayama, K. Tsuda, CompRet: a comprehensive recommendation framework for chemical synthesis planning with algorithmic enumeration, *J. Cheminf.* 12 (1) (2020) 52.
- [137] W.L. Wang, Q.L. Liu, L. Zhang, Y.C. Dong, J. Du, RetroSynX: a retrosynthetic analysis framework using hybrid reaction templates and group contribution-based thermodynamic models, *Chem. Eng. Sci.* 248 (2022) 117208.
- [138] B. Delépine, T. Duigou, P. Carbonell, J.L. Faulon, RetroPath2.0: a retrosynthesis workflow for metabolic engineers, *Metab. Eng.* 45 (2018) 158–170.
- [139] N. Hadadi, J. Hafner, A. Shajkofci, A. Zisaki, V. Hatzimanikatis, ATLAS of biochemistry: a repository of all possible biochemical reactions for synthetic biology and metabolic engineering studies, *ACS Synth. Biol.* 5 (10) (2016) 1155–1166.
- [140] J.D. Tyzack, A.J.M. Ribeiro, N. Borkakoti, J.M. Thornton, Transform-MinER: transforming molecules in enzyme reactions, *Bioinformatics* 34 (20) (2018) 3597–3599.
- [141] S. Genheden, A. Thakkar, V. Chadimová, J.L. Reymond, O. Engkvist, E. Bjerrum, AiZynthFinder: a fast, robust and flexible open-source software for retrosynthetic planning, *J. Cheminf.* 12 (1) (2020) 70.
- [142] P. Torren-Peraire, A.K. Hassen, S. Genheden, J. Verhoeven, D.A. Clevert, M. Preuss, I.V. Tetko, Models Matter: the impact of single-step retrosynthesis on synthesis planning, *Dig. Dis.* 3 (3) (2024) 558–572.
- [143] H.X. Liu, H.Y. Yin, Z.Y. Luo, X.N. Wang, Integrating chemistry knowledge in large language models via prompt engineering, *Synth. Syst. Biotechnol.* 10 (2024) 23–38.
- [144] J. Guo, A.S. Ibanez-Lopez, H. Gao, V. Quach, C.W. Coley, K.F. Jensen, R. Barzilay, Automated chemical reaction extraction from scientific literature, *J. Chem. Inf. Model.* 62 (9) (2022) 2035–2045.
- [145] Y. Liu, H. Xu, T. Fang, H. Xi, Z. Liu, S. Zhang, H. Poon, S. Wang, T-rex: text-assisted retrosynthesis prediction, arXiv:2401.14637 (2024), <https://doi.org/10.48550/arXiv.2401.14637>.
- [146] U.V. Ucak, I. Ashyrmamatov, J. Ko, J.Y. Lee, Retrosynthetic reaction pathway prediction through neural machine translation of atomic environments, *Nat. Commun.* 13 (1) (2022) 1186.
- [147] R. Kurczab, S. Smusz, A.J. Bojarski, The influence of negative training set size on machine learning-based virtual screening, *J. Cheminf.* 6 (2014) 32.
- [148] Y. Wang, C. Pang, Y.Z. Wang, J.R. Jin, J.J. Zhang, X.X. Zeng, R. Su, Q. Zou, L.Y. Wei, Retrosynthesis prediction with an interpretable deep-learning framework based on molecular assembly tasks, *Nat. Commun.* 14 (1) (2023) 6155.
- [149] J.R. Li, L. Fang, J.G. Lou, Retro-BLEU: quantifying chemical plausibility of retrosynthesis routes through reaction template sequence analysis, *Dig. Dis.* 3 (3) (2024) 482–490.
- [150] T. Badowski, K. Molga, B.A. Grzybowski, Selection of cost-effective yet chemically diverse pathways from the networks of computer-generated retrosynthetic plans, *Chem. Sci.* 10 (17) (2019) 4640–4651.
- [151] K. Papineni, S. Roukos, T. Ward, W.-J. Zhu, BLEU: a method for automatic evaluation of machine translation, in: Proceedings of the 40th Annual Meeting on Association for Computational Linguistics, Stroudsburg, PA, USA, 2002.
- [152] A. Mayr, G. Klambauer, T. Unterthiner, M. Steijaert, J.K. Wegner, H. Ceulemans, D.A. Clevert, S. Hochreiter, Large-scale comparison of machine learning methods for drug target prediction on ChEMBL, *Chem. Sci.* 9 (24) (2018) 5441–5451.
- [153] G. Bhisetti, C. Fang, Artificial intelligence-enabled de novo design of novel compounds that are synthesizable, *Method. Mol. Biol.* 2390 (2022) 409–419.
- [154] H.Y. Gao, T.J. Struble, C.W. Coley, Y.R. Wang, W.H. Green, K.F. Jensen, Using machine learning to predict suitable conditions for organic reactions, *ACS Cent. Sci.* 4 (11) (2018) 1465–1476.
- [155] T. Stuyver, C.W. Coley, Quantum chemistry-augmented neural networks for reactivity prediction: performance, generalizability, and explainability, *J. Chem. Phys.* 156 (8) (2022) 084104.
- [156] I. Ismail, C. Robertson, S. Habershon, Successes and challenges in using machine-learned activation energies in kinetic simulations, *J. Chem. Phys.* 157 (1) (2022) 014109.
- [157] Q.L. Liu, K. Tang, L. Zhang, J. Du, Q.W. Meng, Computer-assisted synthetic planning considering reaction kinetics based on transition state automated generation method, *AIChE J.* 69 (7) (2023) e18092.
- [158] C.W. Coley, D.A. Thomas 3rd, J.A.M. Lummiss, J.N. Jaworski, C.P. Breen, V. Schultz, T. Hart, J.S. Fishman, L. Rogers, H.Y. Gao, R.W. Hicklin, P.P. Plehiers, J. Byington, J.S. Piotti, W.H. Green, A.J. Hart, T.F. Jamison, K.F. Jensen, A robotic platform for flow synthesis of organic compounds informed by AI planning, *Science* 365 (6453) (2019) eaax1566.
- [159] A.M. Bran, S. Cox, O. Schilter, C. Baldassari, A.D. White, P. Schwaller, Augmenting large language models with chemistry tools, *Nat. Mach. Intell.* 6 (2024) 525–535.
- [160] D.A. Boiko, R. MacKnight, B. Kline, G. Gomes, Autonomous chemical research with large language models, *Nature* 624 (2023) 570–578.