

Cite this: *Chem. Sci.*, 2026, 17, 1666

All publication charges for this article have been paid for by the Royal Society of Chemistry

Grammar-driven SMILES standardization with *TokenSMILES*

Luis Armando Gonzalez-Ortiz,^a Lisset Noriega,^a Filiberto Ortiz-Chi,^b Gabriela Vidales-Ayala,^a Emmanuel Soberanis-Cáceres,^a Amilcar Meneses-Viveros,^b Alan Aspuru-Guzik^c and Gabriel Merino^a

The redundancy of *SMILES* notation, where multiple strings can describe the same molecule, remains a challenge in computational chemistry and cheminformatics. To mitigate this issue, we introduce *TokenSMILES*, a grammatical framework that standardizes *SMILES* into structured sentences composed of context-free words. By applying five syntactic constraints (including branch limitations, balanced parentheses, and aromaticity exclusion), *TokenSMILES* minimizes redundant *SMILES* enumerations for alkanes while maintaining valence and octet compliance through semantic parsing rules. *TokenSMILES* does not replace *SMILES* but rather formalizes its syntax into a standardized, machine-interpretable form. This grammatical structure enables controlled generation and manipulation of valid *SMILES* strings, ensuring syntactic and semantic consistency while substantially reducing redundancy. Implemented into *SmilX*, an open-source tool, *TokenSMILES* generates valid *SMILES* with accuracy comparable to existing computational implementations for molecules with low hydrogen deficiency ($\text{HDI} \leq 4$). Its applicability extends beyond alkanes through stoichiometric modifications such as bond insertion, cyclization, and heteroatom substitution. Nevertheless, challenges remain for highly unsaturated systems, where canonicalization artifacts highlight the need for dynamic feasibility checks. By integrating linguistic principles with cheminformatics, *TokenSMILES* establishes a scalable framework for systematic chemical space exploration, supporting applications in drug discovery, materials design, and machine learning-driven molecular innovation.

Received 7th July 2025
Accepted 12th November 2025

DOI: 10.1039/d5sc05004a
rsc.li/chemical-science

Introduction

The analysis of chemical space using artificial intelligence relies on comprehensive and standardized molecular representations. The Simplified Molecular Input Line Entry System (*SMILES*), introduced by Weininger and co-workers in the 1980s, encodes molecular structures *via* a context-sensitive grammar using ASCII characters.^{1–3} Initially, *SMILES* was designed to represent atoms, bonds, branching, cycles, aromaticity, charges, and hydrogen counts. Over time, its syntax was expanded to include stereochemistry, isotopic labeling, and hybridization. Despite its versatility, *SMILES* is not unique:

several distinct strings can describe the same molecule (Fig. 1). This redundancy arises from permissible syntactic variations within the language. Furthermore, grammatical extensions have been proposed to include more complex systems as polymers and crystal structures.^{4,5}

Table 1 shows syntactic variations,⁶ including Kekulé, aromatic, branching, and ring number/dot bond syntax, using 2-(aminomethyl)benzoic acid ($\text{C}_8\text{H}_9\text{NO}_2$) as an example.

Recent advances have sought to overcome these limitations by introducing alternative representations with improved

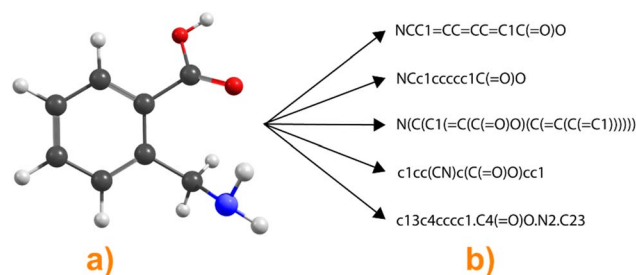


Fig. 1 2-Aminomethylbenzoic acid: (a) molecular model and (b) selected *SMILES* strings representing the same molecule, illustrating the standardization problem addressed by *TokenSMILES*.

^aDepartamento de Física Aplicada, Centro de Investigación y de Estudios Avanzados, Unidad Mérida, km 6 Antigua Carretera a Progreso, Apdo. Postal 73, Cordemex 97310, Mérida, Yucatán, Mexico

^bSecihti-Departamento de Física Aplicada, Cinvestav-IPN, Antigua Carretera a Progreso km 6, Mérida, Yucatán, 97310, Mexico

^cDepartamento de Computación, Centro de Investigación y de Estudios Avanzados, Unidad Zacatenco, Av. IPN No. 2508, Apdo. Postal 07000, Col. San Pedro Zacatenco, CDMX, Mexico

^dDepartment of Chemistry, University of Toronto, DB 421D, Lash Miller Chemical Laboratories, 80 St. George Street, Toronto, ON, M5S 3H6, Canada. E-mail: gmerino@cinvestav.mx; luis.gonzalez@cinvestav.mx; alan@aspuru.com; amilcar.meneses@cinvestav.mx

Table 1 Syntactic variations in SMILES notation

Syntax style	Representative SMILES
Kekulé syntax	<chem>NCC1=CC=CC=C1C(=O)O</chem>
Aromatic syntax	<chem>NCc1ccccc1C(=O)O</chem>
Branching syntax	<chem>N(C(C1(=C(C(C(=O)O)(C(=C(C(=C1)))))))</chem>
Ring numbers and dot bond syntax	<chem>c13c4cccc1.C4(=O)O.N2.C23</chem>

syntactic control. For instance, *DeepSMILES* simplifies parenthesis handling by adopting postfix ring-numbering rules,⁷ while *t-SMILES* encodes functional groups explicitly, avoiding parentheses and ring indices.⁸ *BigSMILES* extends the notation to polymers through the use of braces to denote stochastic binding patterns,^{8,9} whereas *CurlySMILES* embeds annotations within braces { } to describe noncovalent or coordinated structures, while preserving the core SMILES grammar.¹⁰

Beyond these structural adaptations, new languages such as *SELFIES*,¹¹ *Group SELFIES*,¹² and *JAM*^{13,14} have emerged. *SELFIES* guarantees that every token sequence corresponds to a chemically valid structures, thereby minimizing parsing and valency errors. *Group SELFIES* refines this idea by representing rings or

functional groups as single tokens, simplifying substructure encoding. In contrast, *JAM* adapts SMILES-like syntax to describe stacking sequences in crystalline or layered materials, combining chemical and geometric information. Table 2 summarizes these languages using 2-(aminomethyl)benzoic acid as an example, highlighting improvements in grammatical clarity and robustness.

In this work, we introduce *TokenSMILES*, a grammatical and graph-theoretical framework that formalizes the SMILES language, together with *SmilX*, its open-source implementation. *SmilX* applies the *TokenSMILES* grammar to generate and validate molecular structures based on explicitly defined syntactic

Table 2 Grammatical representations of 2-(aminomethyl)benzoic acid in different notations

Notation	Representation
SMILES	<chem>NCc1ccccc1C(=O)O</chem>
DeepSMILES	<chem>NCcccccc6C=O)O</chem>
t-SMILES	<chem>c1([1*])c([2*])cccc1^[1*]C(=O)O^[2*]CN</chem>
SELFIES	<chem>[N][C][C][=C][C][=C][C][=C][Ring1][=Branch1][C][=Branch1][C][=O][O]</chem>
Group SELFIES	<chem>[;benzene][Branch][;CH2NH2][pop][Branch][;COOH][pop]</chem>

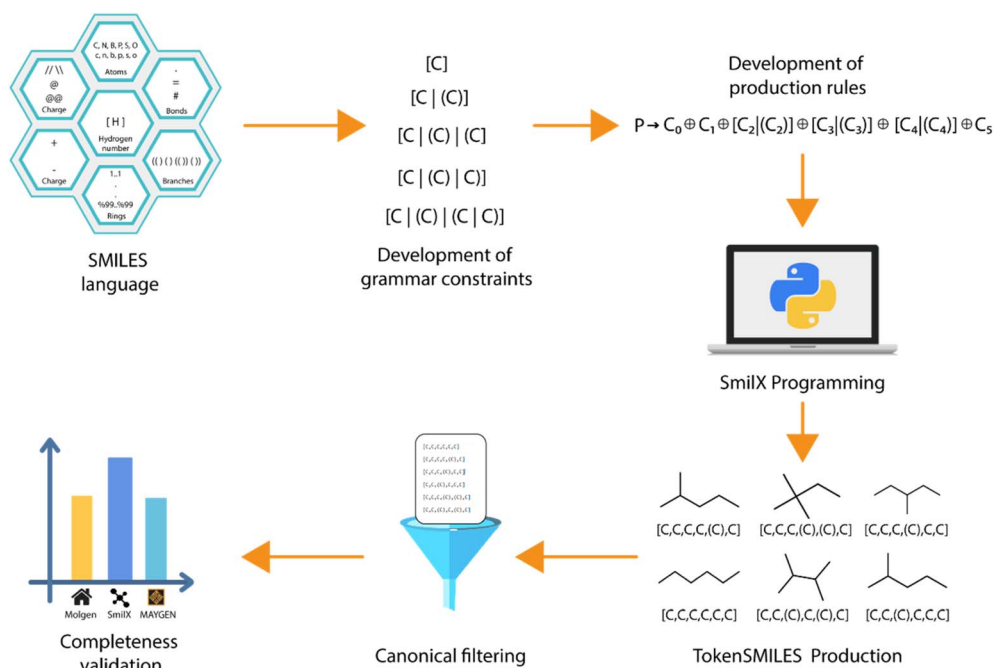


Fig. 2 General workflow diagram for the development of the *TokenSMILES* grammatical framework and the validation of the number of structures.



and semantic rules. Both the conceptual framework and its software implementation are presented here for the first time.

The proposed grammar provides a standardized representation for organic molecules that retains the descriptive capacity of *SMILES* while minimizing ambiguity. By treating molecular strings as complete sentences governed by grammatical rules rather than collections of discrete symbols, *TokenSMILES* enables systematic computational analysis and exhaustive isomer enumeration (Fig. 2). Implemented in *SmilX*, this framework departs from matrix-based approaches (e.g., *MOLGEN*,¹⁵ *MAYGEN*¹⁶) and block-based methods (e.g., *SMILIB*^{17,18}), offering a more structured and linguistically consistent paradigm for molecular representation.

Theoretical development

Kekulé syntax for *SMILES* generation

Let us start by constructing *SMILES* strings for saturated hydrocarbons. The procedure defines transversal paths across all atoms and bonds in a molecule, forming the basis for grammatical constraints in the Kekulé *SMILES* syntax. Using 2,3-dimethylbutane as example (Fig. 3a), the process follows these steps:

- (1) Hydrogen removal. Generate a hydrogen-free molecular graph (Fig. 3b).
- (2) Atom labeling. Assign unique numerical labels to non-hydrogen atoms (Fig. 3b).
- (3) Path definition. Define an ordered set W representing the transversal path. In Fig. 3c, the path $P = \{(C_3, C_2), (C_2, C_4), (C_4, C_6)\}$ corresponds to $W = \{C_3, C_2, C_4, C_6\}$.

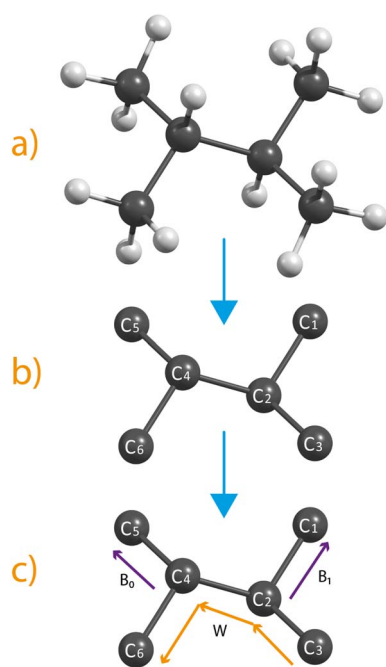


Fig. 3 (a) Molecular model for 2,3-dimethylbutane; (b) hydrogen-free molecular graph with labeled carbon atoms; (c) traversal pathway (orange and purple arrows) covering all atoms.

(4) Branch identification. Identify branches B_m , where each branch contains the maximum possible number of connected atoms. The ordered set of branches is $B = \{B_0, B_1, \dots, B_m\}$. For the system in Fig. 3c, $B_0 = \{C_4, C_5\}$ and $B_1 = \{C_2, C_1\}$.

(5) Branch insertion. Insert the branches $B = \{B_0, B_1\}$ into W , omitting the first atom of each branch since it already appears in W . Parentheses mark the branch boundaries:

$$\{C_3, C_2, (, C_1,), C_4, (, C_5,), C_6\} \quad (1)$$

(6) Symbol replacement and concatenation. Replace the atomic labels in (1) with the corresponding atomic symbols and concatenate them to obtain the string CC(C)C(C)C.

The resulting *SMILES* string, CC(C)C(C)C, represents 2,3-dimethylbutane with implicit hydrogen atoms. Since *SMILES* grammar is non-commutative, the steps must be performed in the specified order.

Tokenization of *SMILES* into *TokenSMILES*

Our method transforms the Kekulé syntax into a standardized form that equalizes string lengths and isolates chemical information by assigning individual tokens to each atom and symbol. For example, the *SMILES* representation of 2,3-dimethylbutane, CC(C)C(C)C, can be rewritten as [C, C, (, C, (, C,)], where “(” and “)” denote the beginning and end of branches, respectively. So, in this tokenized form, branches occur at positions 1 and 3 (zero-indexed). The resulting sequence of tokens constitutes a standardized-length representation referred to as *TokenSMILES*.

This tokenization follows two sequential rules: First, the original string is parsed into individual characters, each enclosed in square brackets. For CC(C)C(C)C, this yields [C, C, (, C,), C, (, C,), C], maintaining the exact order of symbols in the original notation. In this representation, all symbols are enclosed within square brackets to form an ordered sequence. Although conventional set notation implies unique elements, here repeated tokens are intentionally preserved to retain positional information.

Second, the tokens are categorized according to their syntactic context. Left-context symbols, [, (, =, and #, are placed immediately before atomic symbols, while right-context symbols, @,), %,], and a digit n , are placed immediately after them. Applying these rules to the example yields the standardized *TokenSMILES* form [C, C, (, C, (, C,)].

Grammar constraints for *TokenSMILES*

The production rules used to generate *SMILES* strings from molecular paths allow multiple valid representations for a single structure, making exhaustive *SMILES* enumeration a non-deterministic polynomial-time hard (NP-hard) problem.^{19,20} To mitigate this redundancy, and following previous work on grammatical constraints in formal languages,²¹ we introduce five grammar rules to reduce the number of equivalent strings representing organic isomers.

Constraint 1. For a molecular model G without ring number and dot bonds, *TokenSMILES* strings are constructed from



Chemical context definition

Semantic parsing^{22,23} was applied to determine the connectivity encoded in each *TokenSMILES* string. Following the conventions of Weininger *et al.*, each atom (α) is represented by its atomic symbol, and the symbol is linked to a set of chemical properties collectively referred to as the context of the atom, $\text{Context}(\alpha)$. For our analysis, the context is defined as:

$$\text{Context}(\alpha) = (p_0, p_1, p_2)$$

where p_0 is the atomic symbol, p_1 corresponds to the valence, and p_2 is the atomic connectivity.

As an example, the *TokenSMILES* [C, C, (C), C, (C), C] can be indexed as [C₀, C₁, (C₂), C₃, (C₄), C₅]. The chemical context for each substring is summarized in Table 3, where the first element identifies the atomic symbol ("C"), the second indicates the valence (4), and the third lists the corresponding connectivity set.

Connectivity extraction from *TokenSMILES*

In *SMILES* notation, connectivity is often implicit and must be derived through semantic parsing to reconstruct chemical relationships. Fig. 5 shows the procedure used to extract connectivity from the *TokenSMILES* sequence [C₀, C₁, (C₂), C₃, (C₄), C₅]:

(a) "C₀" is not concatenated with any other token, so its connectivity set remains empty.

(b) "C₁" is concatenated with "C₀", creating a new bond represented by the tuple (0, 1).

(c) The same procedure is applied to the remaining tokens. Fig. 5c shows how (C₂) is concatenated, using parentheses to indicate the start and end of a branch emerging from "C₁", adding the tuple (1, 2).

(d) "C₃" is then concatenated, forming the tuple (1, 3) rather than (2, 3), since "(C₂)" is a closed branch.

(e) "(C₄)" is concatenated next, adding the tuple (3, 4).

(f) Finally, "C₅" is concatenated, producing the tuple (3, 5).

After processing all tokens, the resulting bond set is {(0, 1), (1, 2), (1, 3), (3, 4), (3, 5)}, corresponding to [C, C, (C), C, (C), C]. Each tuple represents the connectivity in the *TokenSMILES* notation.

As shown in Fig. 5a, the bond contexts initially consist of empty sets. As the strings are concatenated according to the *SMILES* grammar and defined constraints, their contexts expands as new bonds are introduced. From this stage onward, the strings listed in Table 3 are treated as context-free strings or

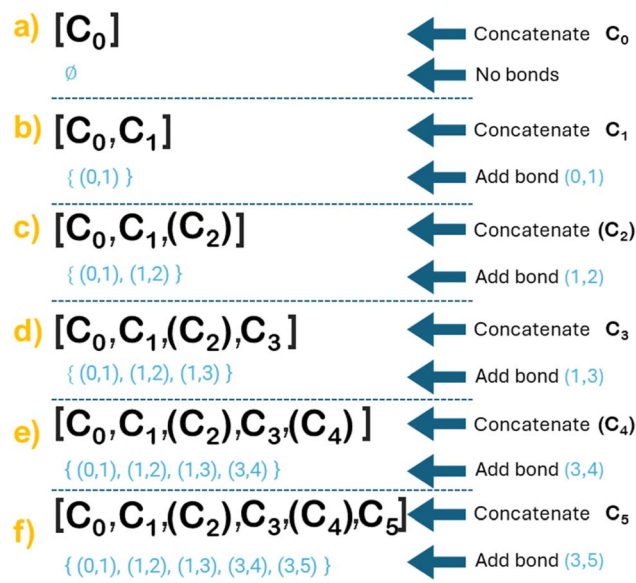


Fig. 5 Procedure for extracting bonds from the *TokenSMILES* of 2,3-dimethylbutane. (a)–(f) Depict the sequential steps of the method.

simply words.²² These words lack explicit connectivity to other atoms, excluding hydrogen. Consequently, a *TokenSMILES* can be interpreted as a sequence of such words forming a chemically meaningful sentence.

To show this process, we return to the example of the C₆H₁₄ isomers. Using production rule (7), connectivity was assigned to each *TokenSMILES* string, and those that violated the octet rule were removed, yielding seven valid *SMILES* candidates (Fig. 6). To eliminate duplicates, each *TokenSMILES* was converted to a canonical *SMILES* form using the canonicalization algorithm implemented in the *RDKit* module.²⁴ Strings sharing the same canonical form were identified, and only unique *SMILES* were retained. After filtering, five distinct *SMILES* remained, corresponding precisely to the five constitutional isomers of C₆H₁₄.

Integrating grammatical constraints into the Kekulé syntax transforms variable-length isomer strings into standardized, fixed-length representations, reducing redundancy and ensuring syntactic coherence. For example, classical *SMILES* enumeration of C₆H₁₄ yields 125 valid strings, whereas *TokenSMILES* generates only seven normalized candidates: [C, C, C, C, C, C], [C, C, C, C, (C), C], [C, C, C, (C), C, C], [C, C, C, (C), (C), C], [C, C, (C), C, C, C], [C, C, (C), C, (C), C], and [C, C, (C), (C), C, C], each conforming to a consistent six-word structure (Table 4).

Modifying *TokenSMILES* stoichiometry

This section defines the rules for modifying *TokenSMILES* syntax to change stoichiometry and to generate *SMILES* for systems beyond C_nH_{2n+2}. Three operations are introduced: two insertion operations, which add new symbols into a word, and one replacement operation, which alters atomic symbols. Each operation produces a copy of the entire sentence while recording the applied modifications. This procedure enables the systematic generation of isomers with stoichiometries different from those of alkanes.

Table 3 Chemical context of [C₀, C₁, (C₂), C₃, (C₄), C₅]

String	Chemical context
C ₀	(C, 4, {(0,1)})
C ₁	(C, 4, {(0,1), (1,2), (1,3)})
(C ₂)	(C, 4, {(1,2)})
C ₃	(C, 4, {(1,3), (3,4), (3,5)})
(C ₄)	(C, 4, {(3,4)})
C ₅	(C, 4, {(3,5)})



insertion of rings or double bonds, from the *TokenSMILES* representation without relying on IUPAC nomenclature or pre-defined templates (see SI for details).

To illustrate the process, we begin with 2,3-dimethylbutane (C_6H_{14}) and modify it to generate 1-cyclopropylideneethanol (C_5H_8O , Fig. 7). First, the connectivity set is extracted as $\{(0, 1), (1, 2), (1, 3), (3, 4), (3, 5)\}$. Using these bonds, the algorithm identifies adjacent words representing bonds in *TokenSMILES*. Words at positions 1 and 3 (Fig. 7a) are selected for double-bond insertion. To evaluate feasibility, the valence excess (Δ) is computed for each atom, defined as the difference between its valence and its degree (the number of incident bonds). When $\Delta \geq 1$, a “=” symbol is inserted into the word at the higher position in the sentences (Fig. 7b).

To prevent excessive “=” insertions, the hydrogen deficiency index (HDI) of the initial system (C_6H_{14}) is used as a control variable. Each additional double bond increases the HDI by one, ensuring that the connectivity modifications remain chemically valid and stoichiometrically consistent.

Next, to insert a cycle, two non-adjacent words are selected based on the same connectivity $\{(0, 1), (1, 2), (1, 3), (3, 4), (3, 5)\}$. If no bond exists between the chosen positions, the algorithm proceeds with cycle insertion. For the example in Fig. 7b, words at positions 0 and 2 are chosen, and the absence of bond (0, 2) is confirmed. The valence excess for both atoms is then evaluated; if $\Delta > 0$ for each, cycle symbols are inserted (Fig. 7c). A random ring number between 1 and 9 is assigned and placed to the right of the atomic symbol (e.g., “C1”). For numbers 10–99, the cycle symbol is prefixed with “%” (e.g., “C%10”). Each cycle insertion increases the HDI by one (Fig. 7c).

Finally, to substitute a carbon atom with oxygen, the condition $\text{degree}(\text{C}) \leq \text{valence}(\text{O})$ must be satisfied before replacement. The operation reduces the carbon and hydrogen count by one and two, respectively (Fig. 7d). For heteroatoms, the allowed are B, Br, Cl, F, I, N, O, P, and S. This approach generates all structural isomers of C_5H_8O (Table S11). Once all transformations are complete, equivalent *SMILES* are filtered using the canonicalization algorithm in *RDKit*.

SmilX software

To automate the generation of *SMILES* under grammatical constraints, the *SmilX* program was developed as an open-

source Python tool, available at <https://github.com/LuisOrz/SmilX.git>. *SmilX* constructs *SMILES* representations of isomers while maintaining compliance with the specified stoichiometry. A user-oriented web interface was also implemented using *Streamlit*, which provides both interactive functionality and server infrastructure. The interface is included with the package and accessible at <https://smilx-isogenerator.streamlit.app/>.

The workflow begins with a molecular formula provided as a string, such as C_nH_{2n+2} . Based on this input, *SmilX* generates all possible *TokenSMILES* corresponding to the defined stoichiometry and adjusts their syntax accordingly. The resulting words in each *TokenSMILES* are concatenated to form complete *SMILES* strings, which are then processed using a canonicalization algorithm to eliminate duplicates and retain unique structures. Finally, the software uses the *RDKit* module to generate stereoisomeric variants, returning a curated list of *SMILES* strings that satisfy the input molecular formula.

Results and discussion

Two experiments were performed to evaluate the isomer-generation capabilities of *SmilX* and to validate the *TokenSMILES* framework. The first reproduced the data reported by Elyashberg *et al.*²⁵ for C–H systems (Fig. 7), which serve as reference structures for more complex compositions. The resulting isomer counts from *SmilX* were compared with those obtained using *MOLGEN* (a matrix-based generator) and *MAYGEN* (an open-source, *Java*-based tool). The second experiment assessed *SmilX*'s performance in systems containing heavier elements (O, N, Cl) across a range of HDI and atom counts.

First experiment

Isomer enumeration followed the molecular formulas reported by Elyashberg *et al.*²⁵ to allow direct comparison between *SmilX* and *MAYGEN*. Systems containing at least two hydrogen atoms were prioritized, while hydrogen-free cases (C_6 or C_{10}) were excluded by omitting the “C = n” notation.

SmilX reproduced Elyashberg's isomer counts for most systems (Fig. 8). Minor deviations occurred when the HDI approached the total heavy-atom count: *SmilX* occasionally yielded one missing structure (false negative, e.g., C_8H_4) or 1–7

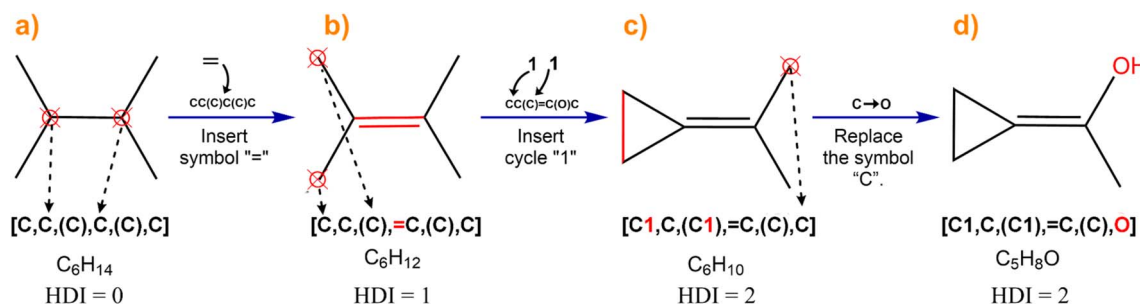


Fig. 7 Syntax transformation from 2,3-dimethylbutane *TokenSMILES* to 1-cyclopropylideneethanol *TokenSMILES*. (a)–(d) Show the sequential operations: double-bond insertion, cycle formation, and heteroatom substitution. Changes are highlighted in red, and dotted arrows indicate the modified word.



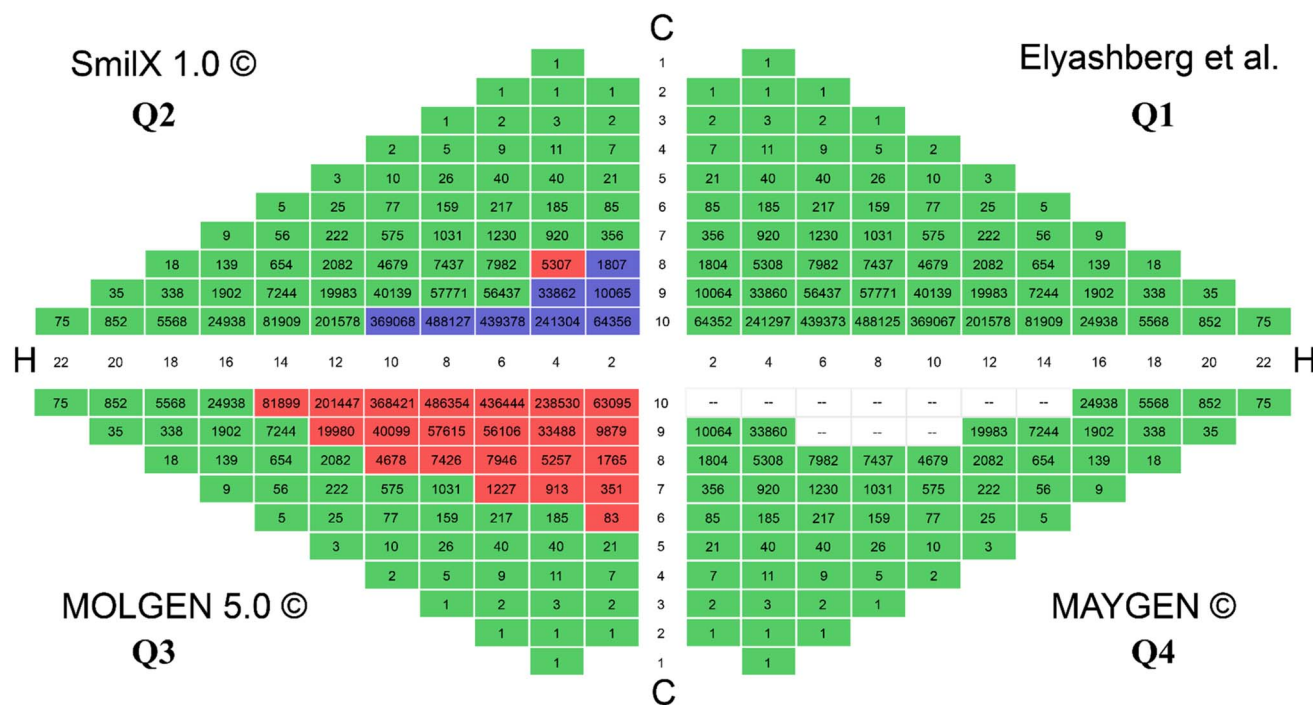


Fig. 8 Isomers composed of carbon and hydrogen as reported by Elyashberg *et al.* (Quadrant 1) compared with those obtained using *SmilX* (Q2), *MOLGEN* (Q3), and *MAYGEN* (Q4). Green cells indicate matching isomers; blue cells represent false positives; red cells correspond to false negatives; and white cells denote missing data due to resource limitations. All results are shown relative to the reference data from Elyashberg *et al.*

additional isomers (false positives) in systems as C_8H_2 , C_9H_2 , $C_{10}H_2$, C_9H_4 , $C_{10}H_4$, $C_{10}H_6$, C_9H_8 , and $C_{10}H_{10}$ (Fig. 8).

MOLGEN accurately enumerated isomers at low HDI values but showed decreased accuracy as HDI increased (Fig. 8). In contrast, *MAYGEN*'s online version encountered memory saturation in larger systems (C_9H_6 , C_9H_8 , C_9H_{10} , $C_{10}H_2$, $C_{10}H_4$, $C_{10}H_6$, $C_{10}H_8$, $C_{10}H_{10}$, $C_{10}H_{12}$, $C_{10}H_{14}$, and $C_{10}H_{16}$) but reproduced Elyashberg's data for smaller, computationally feasible cases.

MAYGEN employs a canonicalization procedure distinct from *RDKit*'s implementation of the Weininger algorithm.¹⁶ Consequently, canonical *SMILES* can struggle to distinguish equivalent atoms in molecules containing multiple nested rings, particularly when HDI is equal to or greater than half the number of heavy atoms. *SmilX*'s reliance on *RDKit* likely accounts for the few observed false positives and negatives.

Second experiment

Building on these results, the second experiment examined *SmilX*'s robustness in systems of increasing compositional complexity: (C, H, O), (C, H, O, N), and (C, H, O, N, Cl). These compositions typically produce more isomers than pure C–H systems. For each composition, six molecular formulas were generated for HDI values ranging from 0 to 5, and the number of heavy atoms was limited to 3–10 to avoid cases where $HDI \geq \frac{1}{2}$ the heavy-atom count, previously associated with enumeration errors. Results are summarized in Tables SI1–SI3.

SmilX showed efficient performance, aided by disk-caching optimization. All three tools produced identical isomer counts

for $HDI \leq 4$, with small discrepancies emerging at higher HDI values. These deviations further support the hypothesis that *RDKit*'s canonicalization algorithm faces difficulties in handling highly nested ring topologies.

The *TokenSMILES* framework effectively reduced both string redundancy and computational overhead through grammatical constraints and caching. While *SmilX* correctly enumerated the majority of organic systems, boundary cases where HDI approached the heavy-atom count remained problematic. These discrepancies are consistent with the theoretical limitations of canonicalization algorithms in systems containing high symmetry or multiple fused rings, rather than with specific software errors. Despite these edge-case issues, *SmilX* maintained low execution times, and the reuse of cached structures enabled scalable exploration of extensive chemical spaces.

Classical *SMILES* representations provide a foundation for cheminformatics but suffer from significant redundancy, as shown in Table 4. To overcome this limitation, the present work redefines *SMILES* not merely as atomic sequences but as grammatically structured sentences, a conceptual framework rooted in formal language theory. The *TokenSMILES* approach implements this through hierarchical syntax (word- and sentence-level organization), enforced grammatical constraints, and standardized string lengths. This reformulation reduces redundancy, ensures systematic chemical-space coverage, and facilitates new computational applications.

Unlike conventional structure generators such as *MOLGEN* or *MAYGEN*, *TokenSMILES* emphasizes formal representation rather than speed or memory efficiency, which justifies the



- 17 A. Schüller, G. Schneider and E. Byvatov, *QSAR Comb. Sci.*, 2003, **22**, 719–721, DOI: [10.1002/qsar.200310008](https://doi.org/10.1002/qsar.200310008).
- 18 A. Schüller, V. Hähnke and G. Schneider, *QSAR Comb. Sci.*, 2007, **26**, 407–410, DOI: [10.1002/qsar.200630101](https://doi.org/10.1002/qsar.200630101).
- 19 D. S. Hochbaum, *Approximation Algorithms for NP-hard Problems*, ACM Sigact News, 1997, vol. 28, pp. 447–476, ISBN 978-0534949681.
- 20 R. M. Wharton, *Inf. Contr.*, 1977, **33**, 253–272, DOI: [10.1016/S0019-9958\(77\)80005-3](https://doi.org/10.1016/S0019-9958(77)80005-3).
- 21 S. Kadioglu and M. Sellmann, *Constraints*, 2010, **15**, 117–144, DOI: [10.1007/s10601-009-9073-4](https://doi.org/10.1007/s10601-009-9073-4).
- 22 S. L. Graham and M. A. Harrison, in *Advances in Computers*, ed. M. Rubinoff and M. C. Yovits, Elsevier, 1976, vol. 14, pp. 77–185, DOI: [10.1016/S0065-2458\(08\)60451-9](https://doi.org/10.1016/S0065-2458(08)60451-9).
- 23 D.-Q. Zhang, K. Zhang and J. Cao, *Comput. J.*, 2001, **44**, 186–200, DOI: [10.1093/comjnl/44.3.186](https://doi.org/10.1093/comjnl/44.3.186).
- 24 G. Landrum, P. Tosco, B. Kelley, R. Rodriguez, D. Cosgrove, R. Vianello, Sriniker, P. Gedeck, G. Jones, N. Schneider, E. Kawashima, D. Nealschneider, A. Dalke, M. Swain, B. Cole, S. Turk, A. Savelev, T. Hurst, A. Vaucher, M. Wójcikowski, I. Take, V. F. Scalfani, R. Walker, K. Ujihara, D. Probst, J. Lehtivarjo, H. Faara, G. Godin, A. Pahl and J. Monat, *rdkit/rdkit: 2024_09_5 (Q3 2024) Release*, Zenodo, 2025, DOI: [10.5281/zenodo.28640](https://doi.org/10.5281/zenodo.28640).
- 25 M. Elyashberg, A. Williams and K. Blinov, *Contemporary Computer-assisted Approaches to Molecular Structure Elucidation*, Royal Society of Chemistry, Cambridge, U.K, 2012, vol. 1, pp. 28–30, ISBN 978-1-84973-432-5.

