

MV-CLAM: Multi-View Molecular Interpretation with Cross-Modal Projection via Language Model

Sumin Ha^{1,*}, Jun Hyeon Kim^{2,*}, Yinhua Piao³, Changyun Cho⁴, Sun Kim^{1,3,4,5}

¹Interdisciplinary Program in Artificial Intelligence, Seoul National University

²Bio-MAX/N-Bio, Seoul National University

³Department of Computer Science and Engineering, Seoul National University

⁴AIGENDRUG Co., Ltd.,

⁵Interdisciplinary Program in Bioinformatics, Seoul National University

{suminqw124, tommy0906, 2018-27910, joebrother, sunkim.bioinfo}@snu.ac.kr

Abstract

Deciphering molecular meaning in chemistry and biomedicine depends on context — a capability that large language models (LLMs) can enhance by aligning molecular structures with language. However, existing molecule-text models ignore complementary information in different molecular views and rely on single-view representations, limiting molecule structural understanding. Moreover, naïve multi-view alignment strategies face two challenges: (1) the *aligned spaces differ* across views due to inconsistent molecule-text mappings, and (2) existing loss objectives *fail to preserve complementary information* necessary for fine-grained alignment. To enhance LLM’s ability to understand molecular structure, we propose MV-CLAM, a novel framework that aligns multi-view molecular representations into a unified textual space using a multi-querying transformer (MQ-Former). Our approach ensures cross-view consistency while the proposed token-level contrastive loss preserves diverse molecular features across textual queries. MV-CLAM enhances molecular reasoning, improving retrieval and captioning accuracy. The source code of MV-CLAM is available in <https://github.com/sumin124/mv-clam>.

1 Introduction

A profound contextual understanding of both molecular structures and biomedical text is crucial in chemistry and biomedicine. For large language models to capture these relationships, fine-grained alignment between textual and molecular representations is required to harness their high-context reasoning ability. In vision-language models, researchers have moved beyond coarse image-text matching toward precise region-word alignment, ensuring detailed semantic correspondence between textual descriptions and visual features (Li et al., 2022; Lavoie et al., 2024). Recent studies

*These authors contributed equally to the work

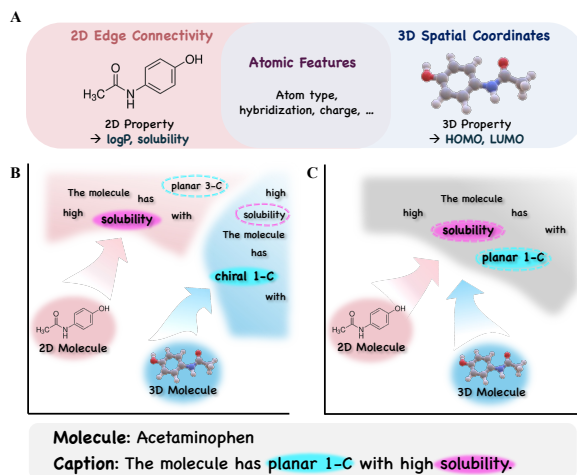


Figure 1: Motivations of MV-CLAM. (A) Complementary molecular information captured by 2D and 3D representations, where 2D graph encodes edge connectivity, and 3D conformers captures spatial coordinate structures. (B) Inconsistent mappings between molecule (2D and 3D) and property tokens (e.g., 2D property token like *solubility* and 3D structural information like *chiral 3-C*) in distinct text spaces. (C) A unified alignment with a Multi-Querying Transformer (MQ-Former) allows all text tokens share a single text space.

have leveraged large language models (LLMs) for molecular understanding by integrating sequential representations (1D SMILES strings) and structural features (2D molecular graphs and 3D conformers) (Edwards et al., 2022; Liu et al., 2023a). This approach mitigates the inherent limitations of LLMs which are primarily trained on textual data, that lacks native reasoning over molecular structures. To enable LLMs to further understand molecule information, Q-former based models (Liu et al., 2023b; Li et al., 2024) align molecular structures into text space (Figure 2B).

Combining multi-view molecular features simultaneously is essential, as their complementary nature provides a more complete understanding of molecular characteristics as illustrated in Figure 1. For example, as shown in Figure 1A, 2D molecular

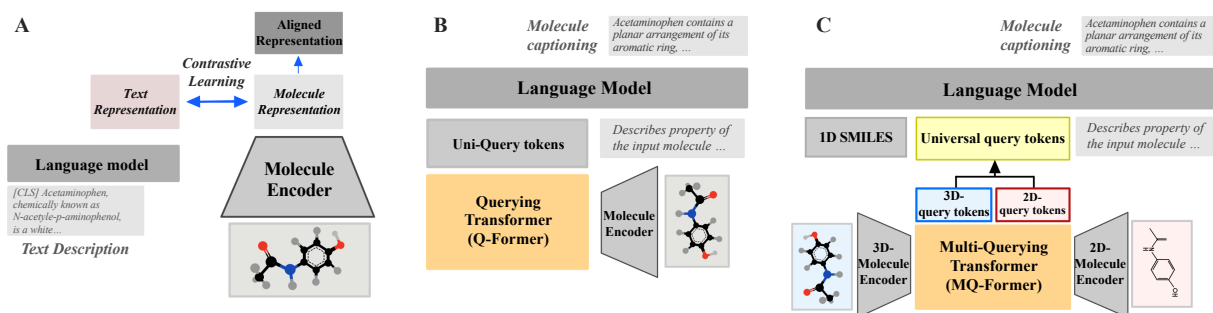


Figure 2: Methods for molecular language modeling. (A) Contrastive learning aligns two modalities via a contrastive objective, excelling in retrieval but lacking generative capabilities. (B) The Q-Former framework uses learnable query tokens for caption generation but is limited to a single molecular representation. (C) MV-CLAM extends this by integrating multiple representations with modality-specific queries, enabling fine-grained knowledge integration.

graphs primarily capture atomic bonding patterns, absent in 3D point clouds. Hence, 2D graphs focus on properties highly affected by atomic bond patterns (e.g., logP, solubility) (Guo et al., 2022) while 3D molecular conformations encode spatial atomic coordinates that influence molecular interactions and quantum properties such as HOMO and LUMO (Kim et al., 2024; Zhou et al., 2023; Du et al., 2023). In the context of molecule understanding, aligning both molecular views into the unified text space of LLMs enables the model to capture all relevant molecular details effectively.

However, existing molecule-text modeling focuses on the alignment of a single molecular view as shown in Figure 2A and 2B (Cao et al., 2023; Li et al., 2024; Liu et al., 2023b,a). Naïve approaches to multi-view alignment might be to independently map each molecular view to text using separate alignment modules. However, this leads to several issues. (1) *Separated aligned spaces*. Aligning 2D and 3D molecular representations separately to text results in distinct aligned spaces for the same molecule. As shown in Figure 1B, “solubility” and “chiral 3-C” correspond to 2D and 3D molecular properties, but each has redundant embeddings in its own space. This inconsistency can prevent the LLM from fully understanding molecular properties, as it lacks a unified representation of 2D and 3D structures. (2) *Insufficient fine-grained molecule-text alignment*. Existing Q-Former-based approaches (Li et al., 2024; Liu et al., 2023b) for aligning molecule queries into a unified text space select the most similar query-to-single token pairs for contrastive learning (Figure 3A). This coarse alignment overlooks structural diversities across molecular views (Appendix Figure 6B), failing to preserve complementary information necessary for

fine-grained alignment and limiting the LLM’s ability to fully understand molecular properties.

To address this, we propose MV-CLAM, a novel framework that aligns multi-view molecule features using a multi-querying transformer, MQ-Former (Figure 2C). Specifically, our approach jointly integrates multi-view molecular representations into a unified textual space, where “solubility” and “chiral 3-C” have unique unified embedding. Such helps generate universal query tokens with more semantic information. Additionally, we propose a multi-token contrastive loss to refine alignment by considering all text tokens within the description, rather than a single CLS token (Figure 3B). Such multi-token contrasting ensures that molecular structures are contextualized with finer, token-level associations, capturing both atomic and functional relevance. MV-CLAM enhances molecular reasoning in LLMs, improving both retrieval and captioning accuracy.

Our main contributions are as follows:

- We propose a novel framework, MV-CLAM, that simultaneously aligns multiple molecular views (1D smiles, 2D graphs, and 3D conformers) to a unified textual space to enhance LLM-based molecular reasoning.
- We introduce a novel contrastive learning loss for molecule–language modeling that achieves fine-grained alignment by comparing all text tokens with enriched molecular query tokens.
- We achieve state-of-the-art performance in molecule-text retrieval and molecule captioning tasks while improving the interpretability of molecular representations.

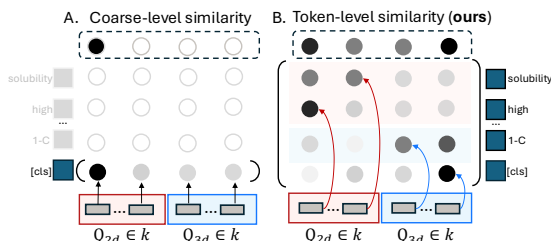


Figure 3: Molecule-text similarity for query-token contrasting. (A) Previous approach compute coarse-level similarity between molecule queries and CLS text token. (B) We propose a new approach to compute token-level similarity between molecule queries and all text tokens, which preserves molecule query diverse information.

2 MV-CLAM

MV-CLAM provides molecule captions given multi-view structural information. 2D and 3D molecular structural information is extracted from specialized encoders and processed through MQ-Former’s cross-attention layers to update learnable query tokens for each dimension. The shared self-attention layer enables information sharing across all modalities. 2D and 3D queries are combined to create a universal query, which is trained with our modified multi-objective loss for fine-grained alignment with textual descriptions. The learned universal query is then passed with the prompt and SMILES strings to the language model for caption generation. The overall framework of MV-CLAM shown in Figure 2C is comprised of three main components: (1) Molecule structural graph encoders for 2D and 3D molecular structures, (2) MQ-Former as a cross-modal projector, and (3) LLaMA2 as the language model.

2.1 Molecular Graph Encoder

To capture structural information from multiple views, we used molecular embeddings from both 3D and 2D structural encoders. For the 3D encoder f_{3d} , we deployed **Uni-Mol** (Zhou et al., 2023), a SE(3)-transformer based model pretrained on 209 million 3D molecular conformations using two tasks: 3D position recovery and masked atom prediction. Input 3D molecule for Uni-Mol is denoted as $m_{3d} = (\mathcal{V}, \mathbf{f}, \mathbf{P})$, where $\mathcal{V}$ and $\mathbf{f}$ each represents atomic nodes and their features, and $\mathbf{P} \in \mathbb{R}^{|\mathcal{V}| \times 3}$ represents 3D coordinates of atoms. Pair representations are initialized by invariant spatial positional encoding from atom coordinates and interact with atom representations. The output atomic representation $H_{3d} \in \mathbb{R}^{|\mathcal{V}| \times d_{3d}}$, where h_i corresponds to the

i -th atom and d_{3d} denotes the hidden dimension size of H_{3d} , updates learnable 3D query tokens through the cross-attention layers in MQ-Former’s 3D molecular transformer block.

$$H_{3d} = [h_1, h_2, \dots, h_{|\mathcal{V}|}] = f_{3d}(m_{3d}) \quad (1)$$

For the 2D molecular encoder f_{2d} , we adopted **Molecule Attention Transformer (MAT)** (Maziarka et al., 2020), pretrained on two million molecule samples from ZINC15 dataset (Irwin et al., 2012). Given 2D molecule $m_{2d} = (\mathcal{V}, \mathbf{f}, \mathbf{A})$ where $\mathbf{A}$ represents edges within the molecule as adjacency matrix, MAT generates atomic representations $H_{2d} \in \mathbb{R}^{|\mathcal{V}| \times d_{2d}}$ using a specialized molecule-specific attention mechanism that considers edges, atomic distances and atomic features. The atomic representations interact with the learnable 2D query tokens via cross-attention layers in 2D molecular transformer block.

$$H_{2d} = [h_1, h_2, \dots, h_{|\mathcal{V}|}] = f_{2d}(m_{2d}) \quad (2)$$

2.2 MQ-Former: Multi-Querying Transformer

Previous studies applying Q-Former to the molecular domain project single-dimensional structural embeddings into the textual space (Li et al., 2024; Zhang et al., 2024). These models consist of a single molecule transformer and a text transformer. However, this approach is inherently limited in preserving molecular information when aligning with text embeddings for two main reasons: (1) separate aligned spaces with inconsistent mappings between molecule and text embeddings, and (2) information loss caused by single-token contrastive learning. MQ-Former addresses this limitation by introducing a novel architecture capable of aligning multiple modalities to a unified aligned space using a refined multi-objective loss for better information preservation (Figure 4).

Our approach combines structural representations of two dimensions, but the architecture can be extended using multiple molecule transformers and a single text transformer. Each molecule transformer, based on the BERT architecture with additional cross-attention layer, processes K learnable query tokens specific to their respective views. Following previous studies (Li et al., 2024; Liu et al., 2023b), we adopt the SciBERT (Beltagy et al., 2019) architecture for the text transformer and initialize all blocks with SciBERT’s pretrained weights. Hence, textual descriptions S of length

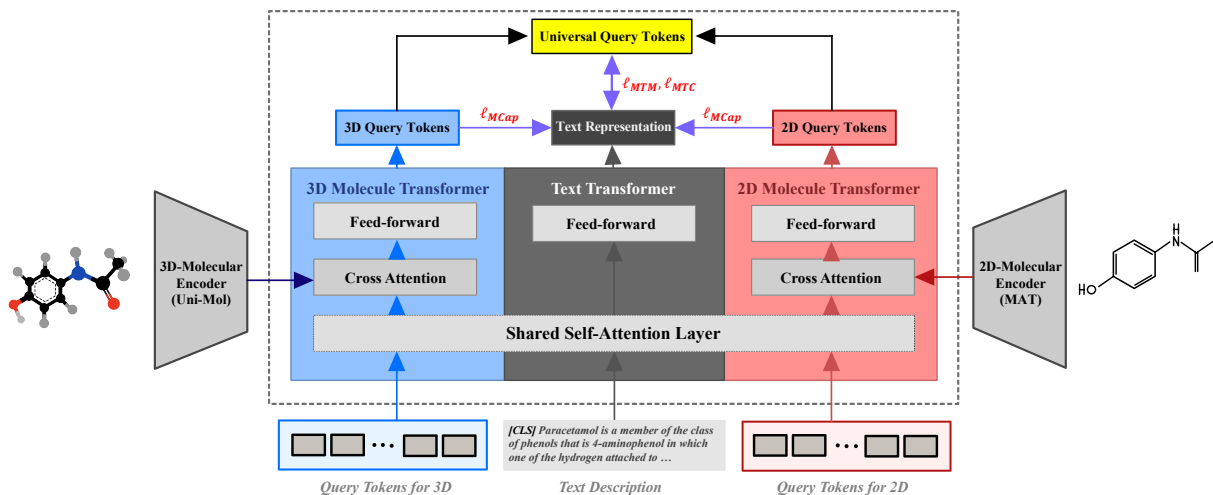


Figure 4: Training scheme of MQ-Former. The proposed MQ-Former enhances molecular language modeling by incorporating multi-token contrasting and amplified molecule captioning losses to the prior multi-objective loss (Li et al., 2023, 2024; Liu et al., 2023b). (1) The novel multi-token contrasting loss ℓ_{MTC} replaces conventional molecule-text contrastive learning, encouraging diverse query-token alignment. (2) The molecule captioning loss ℓ_{MCap} is amplified to improve text generation quality. The molecule-text matching loss ℓ_{MTM} remains unchanged.

L are tokenized with SciBERT’s tokenizer f_{sci} to $X_{\text{text}} = \{x_1, x_2, \dots, x_T\}$ (T denotes the number of tokens in text) before being processed through MQ-Former’s text transformer. The cross-attention mechanism extracts relevant information from embeddings into the query tokens, and shared self-attention layers enable information exchange across all embeddings, overcoming the limitation of separated aligned spaces.

Figure 4 illustrates MQ-Former generating a universal query tokens for a molecule given two different views. Two molecule transformer modules each updates distinct K query tokens $Q_{2d} \in \mathbb{R}^{K \times 768}$ and $Q_{3d} \in \mathbb{R}^{K \times 768}$, which are randomly initialized. The learned query tokens, $\hat{Q}_{2d}$ and $\hat{Q}_{3d}$ of same size, are updated representations of these initial tokens, refined through the alignment of multiple molecule views and textual descriptions $X_{\text{text}} \in \mathbb{R}^{L \times 768}$. Updated query tokens are concatenated to create a single universal query $\hat{Q} \in \mathbb{R}^{2K \times 768}$, containing complementary structural information aligned to textual space. The resulting universal query tokens are then used as inputs for the language model, along with 1D SMILES string and task prompt as depicted in Figure 2C.

$$\begin{aligned} \hat{Q} &= f_{\text{concat}}(\hat{Q}_{2d}, \hat{Q}_{3d}) \\ &= f_{\text{MQformer}}(H_{2d}, H_{3d}, X_{\text{text}}, Q_{2d}, Q_{3d}) \end{aligned} \quad (3)$$

2.3 LLaMA2 & LoRA

The pretraining corpus of LLaMA2 (Touvron et al., 2023) includes a vast amount of biomedical liter-

ature and thereby exerts powerful text generation capability with internal chemistry knowledge. This allows LLaMA2 to effectively interpret 1D molecular sequences and address tasks related to molecular comprehension. Despite its inherent capabilities, the language model necessitates fine-tuning to effectively address the universal queries posed by MQ-Former, particularly due to the modifications in the tokenizer resulting from changes in module processing of textual descriptions. To facilitate efficient fine-tuning, we implemented low-rank adaptation (LoRA, (Hu et al., 2021)).

3 Training MV-CLAM

The training of MV-CLAM consists of two stages. (1) Guiding MQ-Former to align both multi-view molecular representations with a consistent textual space, and (2) refining query tokens to be effectively soft-prompted by LLaMA2. Molecular encoders are frozen during the entire pipeline.

3.1 Stage 1: Training MQ-Former

Two sets of K learnable query tokens are updated by each molecule transformer block in Stage 1. Molecule transformer blocks hold self-attention, cross-attention and feed-forward layers. Specifically, the self attention layers in all blocks of MQ-Former are shared to exchange information between modalities and view. The 2D and 3D query tokens $Q_{2d}(i)$, $Q_{3d}(i)$ for i -th molecule are processed through their respective molecule trans-

formers. Our $2K$ universal query token $\hat{Q}(i)$ is formed by concatenating the learned query sets. The objective is to train MQ-Former to learn a unified latent space for all molecular embeddings and obtain highly informed molecular soft-prompt $\hat{Q}(i)$ without any inconsistencies.

For training, we introduce the following key modifications to the multi-objective loss in previous works inspired by the BLIP-2 framework (Li et al., 2023, 2024), designed to maximize the diversity of queries. In order to preserve complementary chemical aspects embedded in each dimension, we introduce the following key modifications: (1) a novel multi-token contrasting loss ℓ_{MTC} in replacement to single-token (molecule-text) contrasting, and (2) amplification of the molecule captioning loss ℓ_{MCap} . Molecule-text matching is used without further modifications ℓ_{MTM} . This allows our model to capture and preserve both fine-grained atomic interactions and high-level chemical semantics, enhancing interpretability and expressiveness for molecular language modeling. Overall, the total loss for training MQ-Former ℓ_{MQ} in Stage 1 is as follows:

$$\ell_{MQ} = \ell_{MTC} + \ell_{MTM} + \alpha * \ell_{MCap} \quad (4)$$

Multi-Token Contrasting. Unlike the previous approach that retrieved only the maximum similarity between a query token and CLS text token (Figure 3A), we introduce a refined similarity computation where each molecule token is matched against all text tokens, retrieving the maximum similarity for each token against all T text tokens (Figure 3B). The average loss over all $2k$ tokens represents a fine-grained similarity calculation between molecule-text pairs, preventing query collapse, where a single query token with high similarity dominates the training process by aligning only with easily capturable text concepts. By distributing alignment across multiple queries and text tokens, we achieve richer molecule-text representations, improving cross-modal association.

ℓ_{MTC} is measured as the batch mean of the sum of molecule-to-text loss ℓ_{g2t} and text-to-molecule loss ℓ_{t2g} . For each query in the universal query token, we calculate the maximum cosine similarity it has against all text tokens $x(i) \in X_{\text{text}}(i)$ with temperature scaling for precision. The average of the calculated similarity for $2K$ queries represents pairwise similarity in a more precise manner. Similarly, ℓ_{t2g} aligns the text representation with its matching molecular query while contrasting it against

all other queries within the batch. The similarity calculation can be formulated as the following:

$$\begin{aligned} S(i, j) &= \frac{1}{2K} \sum_{2K} \max_t \cos(\hat{Q}_k(i), x_t(j)) \\ S'(i, j) &= \frac{1}{T} \sum_T \max_k \cos(x_t(i), \hat{Q}_k(j)) \end{aligned} \quad (5)$$

Together ℓ_{MTC} form a bidirectional alignment between molecular features and textual descriptions in a detailed token-wise manner. ℓ_{g2t} and ℓ_{t2g} is as written below, where M is the size of the batch and τ is the temperature parameter.

$$\begin{aligned} \ell_{g2t} &= - \sum_{i=1}^M \log \frac{\exp(S(i, i)/\tau)}{\sum_{j=1}^M \exp(S(i, j)/\tau)} \\ \ell_{t2g} &= - \sum_{i=1}^M \log \frac{\exp(S'(i, i)/\tau)}{\sum_{j=1}^M \exp(S'(i, j)/\tau)} \end{aligned} \quad (6)$$

Molecule-text Matching. ℓ_{MTM} is for a binary classification task to predict matching molecule-text pairs. Universal query tokens are obtained then processed through a linear classifier after mean pooling. Let $\rho(\hat{Q}(i), X_{\text{text}}(i))$ denote the predicted probability that universal query $\hat{Q}(i)$ matches its corresponding text description $X_{\text{text}}(i)$. ℓ_{MTM} is calculated as follows:

$$\begin{aligned} \ell_{MTM} &= \frac{1}{M} \sum_{i=1}^M \left(-\log \rho(\hat{Q}(i), X_{\text{text}}(i)) \right. \\ &\quad \left. + \log \rho(\hat{Q}(i), X_{\text{text}}(j)) + \log \rho(\hat{Q}(r), X_{\text{text}}(i)) \right) \end{aligned} \quad (7)$$

where $X_{\text{text}}(j)$, $\hat{Q}(r)$ are randomly selected negative samples from the batch. Overall, ℓ_{MTM} aids MQ-Former to maximize the likelihood of matched pairs and minimize mismatches, enhancing its ability to differentiate between true and false pairs.

Molecule Captioning. ℓ_{MCap} is designed to generate accurate text descriptions based on multi-view query tokens. Text is generated autoregressively, where each token is predicted sequentially based on the corresponding molecular queries. Instead of harnessing universal queries, ℓ_{MCap} sums up separate losses for 2D and 3D query tokens, ensuring that each query token retains its unique dimensional information for high captioning ability. The ℓ_{MCap} is defined as follows:

$$\begin{aligned} \ell_{MCap} &= -\frac{1}{M} \sum_{i=1}^M \log p(X_{\text{text}}(i) | \hat{Q}_{2d}(i)) \\ &\quad -\frac{1}{M} \sum_{i=1}^M \log p(X_{\text{text}}(i) | \hat{Q}_{3d}(i)) \end{aligned} \quad (8)$$

where $p(X_{\text{text}}|\hat{Q}_{2d})$ and $p(X_{\text{text}}|\hat{Q}_{3d})$ represents the probability of generating the text description based independently on 2D or 3D molecular queries, respectively. While the other two losses focus on aligning or matching molecule-text pairs, the ℓ_{MCap} directly impacts the ability to generate new text based on molecular representations, encouraging further diverse feature learning in correspondence to our modified multi-token contrasting loss. Given its critical role, we assigned a greater weight α , guiding MQ-Former to generate quality tokens for text-generation tasks.

3.2 Stage 2: Specializing LLaMA2 for Molecule Captioning

In Stage 2, MQ-Former is further trained alongside LLaMA2 to generate molecular descriptions. The goal is to enhance MQ-Former’s ability to produce universal queries that are not only aligned with the textual space but better interpretable by LLaMA2. In this stage, textual descriptions are tokenized and decoded using LLaMA tokenizer. Universal query tokens, 1D SMILES are given as input with prompt. Autoregressive generation loss of LLaMA2 is used for training the framework with LoRA (Hu et al., 2021). Detailed LoRA settings are in Appendix A.3.

4 Experiments

4.1 Datasets

PubChem324K. For molecule-text alignment and molecule captioning, we collected 324k molecular SMILES-text pairs from PubChem (Kim et al., 2021). 2D graph features were constructed using (Maziarka et al., 2020), and 3D conformers were generated with ETKDG and optimized using the MMFF algorithm in RDKit (Landrum et al., 2013). We follow dataset construction as provided in 3D-MoLM (Li et al., 2024) which also requires 3D molecular conformations. High-quality subset of 15k pairs with text longer than 19 words are sampled for train, valid, test datasets. Shorter pairs are used for pretraining. The statistics for the final PubChem324k dataset used in this study are presented in Appendix Table 4.

4.2 Benchmark models

Baseline models include (1) pretrained language models for science: Sci-BERT (Beltagy et al., 2019), (2) models with molecule-text contrastive learning: KV-PLM (Zeng et al., 2022), MoMu (Su

et al., 2022), MoleculeSTM (Liu et al., 2023a) and (3) models with Q-Former modules: MolCA (Liu et al., 2023b), 3D-MoLM (Li et al., 2024), Uni-MoT (Zhang et al., 2024). For molecule captioning, we also benchmark Llama2-7B and 2D-MoLM, each as a variant of 3D-MoLM using 1D and 2D information along with MolT5 (Edwards et al., 2022) and InstructMol (Cao et al., 2023).

5 Results

5.1 Molecule-Text Retrieval

We evaluate MV-CLAM for molecule-text retrieval on the PubChem324k dataset. We perform two rounds of evaluation on molecule-to-text and text-to-molecule retrieval tasks, using Accuracy and Recall@20 metrics: within batch size of 64 and is across the entire test set. We report baseline performances as written in literature (Li et al., 2024; Zhang et al., 2024).

As shown in Table 1, MV-CLAM outperforms baseline approaches that represent molecules as 1D SMILES strings, 2D graphs, or 3D conformers. We attribute our superior performance to (1) our use of a universal query that aligns both 2D and 3D molecular representations to a consistent text, and (2) a modified multi-objective loss, designed to maximize query diversity and prevent over-reliance on dominant alignment patterns.

5.2 Molecule Captioning

Following previous studies (Li et al., 2024), we use BLEU, ROUGE, METEOR metrics to evaluate captioning results on the PubChem324k dataset. Table 2 shows MV-CLAM consistently outperforms all baselines with notable performance gain from our modified multi-objective loss. We also conducted human evaluation from domain experts (Section A.6), where MV-CLAM was evaluated to generate more accurate captions containing molecular nomenclature, chemical property and clinical usage, compared to its single-view variants. Appendix Table 7 highlights the model’s ability to correctly identify International Union of Pure and Applied Chemistry (IUPAC) nomenclature and generic drug names that differ significantly in language model processing. IUPAC names follow systematic chemical rules, making them complex and highly structured, while generic drug names are more standardized and commonly used in clinical contexts. Despite these differences, MV-CLAM successfully identifies both types of names, show-

Table 1: Molecule-Text retrieval performance in batch and test set for different models. The highest value in each category is indicated in bold, and the second highest value is underlined. For MoleculeSTM* and MolCA*, we report results from UniMoT (Zhang et al., 2024).

Model	Retrieval in batch				Retrieval in test set			
	M2T		T2M		M2T		T2M	
	ACC	R@20	ACC	R@20	ACC	R@20	ACC	R@20
1D SMILES								
Sci-BERT(Beltagy et al., 2019)	85.32	98.74	84.20	98.43	41.67	87.31	40.18	86.77
KV-PLM(Zeng et al., 2022)	86.05	98.63	85.21	98.47	42.80	88.46	41.67	87.80
2D Graph								
MoMu-S(Su et al., 2022)	87.58	99.24	86.44	99.38	47.29	90.77	48.13	89.92
MoMu-K(Su et al., 2022)	88.23	99.41	87.29	99.42	48.47	91.64	49.46	90.73
MoleculeSTM* (Liu et al., 2023a)	90.50	99.60	88.60	99.50	52.70	92.90	53.20	92.50
MolCA* (Liu et al., 2023b)	92.60	99.80	91.30	99.50	67.90	94.40	68.60	93.30
2D Graph + Tokenizer								
UniMoT(Zhang et al., 2024)	93.60	100.0	92.70	99.40	69.50	96.30	69.80	94.40
3D Conformer								
3D-MoLM(Li et al., 2024)	93.50	100.0	92.89	99.59	69.05	95.91	70.13	94.88
2D Graph + 3D Conformer								
MV-CLAM w/ SINGLE-TOKEN CONTRASTING	96.57	99.95	97.03	99.95	76.32	96.57	77.03	96.42
MV-CLAM w/ MULTI-TOKEN CONTRASTING	97.34	<u>99.95</u>	97.19	99.90	78.67	96.98	79.34	96.93

casing its ability to handle a range of linguistic and chemical complexities. Moreover, MV-CLAM demonstrates its capacity to generate literature-matching captions absent in ground truth, as seen in the case of *Rifapentine* (Appendix Table 7), highlighting the ability to produce highly informed outputs.

5.3 Effectiveness of MQ-Former

In this section, we substantiate the effectiveness of incorporating multi-view chemical information within the MQ-Former architecture. We conduct both quantitative and qualitative analysis to compare our superiority to the prior single-view alignment using Q-Former. Molecular encoders are identically set for the ablation studies.

As a quantitative analysis, we conducted ablation by selectively removing each modality from the MV-CLAM framework in different combinations. We also evaluated an alternative setup where multi-view molecular embeddings were pre-combined via concatenation and aligned to text using the naive Q-former module. Overall, Table 3 shows that the simultaneous usage of all 1D, 2D, 3D views lead to a notable synergistic effect, especially under MQ-Former’s architecture and proposed contrastive loss. Although the addition of each view leads to a gradual increase in performance, MV-CLAM’s performance gain does not solely stem from the simple usage of multiple molecule representations. Performance gain was not shown in the pre-combined setup, indicating the value of inter-modality interaction enabled by the shared self-attention layer of MQ-Former, aided with the refined loss.

We further examine the effectiveness of the proposed multi-token contrastive loss in Appendix Table 6, comparing the effect of proposed loss in various modality configurations. While the multi-token loss improves performance in most cases, the gains are relatively modest in ablation settings compared to our full MV-CLAM framework. This suggests that its benefits are most pronounced when combined with multiple modalities via MQ-Former, where the shared self-attention mechanism facilitates richer cross-modal interaction and contextual alignment across 2D, 3D, and text.

We exemplify two case studies to interpret how each transformer module and modality focus on distinct aspects of the molecule and its corresponding text. These qualitative studies provide insight into the alignment process by analyzing how different views contribute to the comprehensive understanding of molecular structures and their textual descriptions.

Case Study 1: Visualizing Attention Maps for 2D and 3D Query Tokens. Embedding grounded on different latent spaces and dimensions differently align molecular information to text. Visualization of the distinct alignment is performed by extracting and comparing the attention maps of the shared self-attention layers when processing 2D and 3D query tokens respectively with text tokens.

With multi-token contrasting loss, each query token attends distinctly to individual tokens in the captioning sentence, exhibiting diverse attention scores (Appendix Figure 6). While query maintaining diversity, 2D query tokens effectively capture 2D-related terms - such as *boiling point* - focusing

Table 2: Molecule captioning performance across models. The highest value in each category is bolded, and the second highest is underlined. Models marked with † were pretrained on larger datasets, as noted in their original papers. Results for InstructMol and MolCA are from UniMoT (Zhang et al., 2024), with MolCA evaluated in two variations using OPT-125M (small) and OPT-1.3B (large) as language models.

	BLEU-2	BLEU-4	ROUGE-1	ROUGE-2	ROUGE-L	METEOR
1D SMILES						
MolT5-Small(Edwards et al., 2022)	22.53	15.23	30.44	13.45	20.30	23.98
MolT5-Base(Edwards et al., 2022)	24.51	16.61	32.19	14.04	21.35	26.10
MolT5-Large(Edwards et al., 2022)	25.87	17.28	34.07	16.42	23.41	28.04
Llama2-7B†(Li et al., 2024)	27.01	20.94	35.76	20.68	28.88	32.11
2D Graph						
MoMu-Small(Su et al., 2022)	22.86	16.01	30.98	13.65	20.75	24.35
MoMu-Base(Su et al., 2022)	24.74	16.77	32.45	14.62	22.09	27.16
MoMu-Large(Su et al., 2022)	26.34	18.01	34.75	16.86	24.76	28.73
2D-MoLM†(Li et al., 2024)	27.15	21.19	36.02	20.76	29.12	32.28
InstructMol*(Cao et al., 2023)	18.90	11.70	27.30	11.80	17.80	21.30
MolCA-Small*(Liu et al., 2023b)	25.90	17.50	34.40	16.60	23.90	28.50
MolCA-Large*(Liu et al., 2023b)	28.60	21.30	36.20	21.40	29.70	32.60
2D Graph + Tokenizer						
UniMoT(Zhang et al., 2024)	31.30	23.80	37.50	23.70	33.60	34.80
3D Conformer						
3D-MoLM(Li et al., 2024)	30.32	22.52	36.84	22.32	31.23	33.06
2D Graph + 3D Conformer						
MV-CLAM w/ SINGLE-TOKEN CONTRASTING	31.75	24.48	40.43	25.72	33.79	36.54
MV-CLAM w/ MULTI-TOKEN CONTRASTING	32.32	25.11	40.87	26.48	34.79	36.87

Table 3: Evaluating contribution of multi-view embeddings to captioning performance. MV-CLAM† denotes our full model using MQ-Former to align multi-view molecular embeddings. Ablation study by removing each molecular modality. Note that the -2D-3D setting refers to the case where only LLaMA2-7B is used with SMILES.

	B-2	B-4	R-1	R-2	R-L	M
Concatenated 3D & 2D	29.80	22.70	39.07	24.92	33.09	35.49
Only 1D (-3D-2D)	27.01	20.94	35.76	20.68	28.88	32.11
Only 3D (-1D-2D)†	27.09	19.43	35.82	20.78	30.17	31.38
Only 2D (-1D-3D)†	29.96	22.37	38.20	23.48	32.36	34.12
1D+2D (-3D)†	30.71	23.32	39.08	24.57	33.01	34.81
1D+3D (-2D)†	29.46	22.19	37.33	23.12	31.53	33.34
2D+3D (-1D)†	31.44	23.85	39.79	25.24	34.06	35.78
MV-CLAM†	32.32	25.11	40.87	26.48	34.79	36.87

on chemical and material properties that may be overlooked in 3D settings. Conversely, 3D query tokens capture 3D-specific structural information, such as *bis (2-dimethylamino)ethyl*, informed by 3D spatial coordinates. In contrast, when MQ-Former is trained with the original contrastive loss, it not only lacks diversity among query tokens but also struggles to properly align with 2D- and 3D-related terms.

Case Study 2: Comparing molecule captions with 2D-Qformer and 3D-Qformer. We illustrate the difference in captioning results between the uni-modal Q-Former ablation models and ours demonstrating the effects of utilizing multi-view molecular understanding in text generation (Appendix Figure 5). The 2D and 3D uni-modal abla-

tions struggle to fully capture complex and large structures like ‘*(R)-3-hydroxytriacontanoyl-CoA*’. The ablation models fail to retain sufficient structural information required to differentiate long carbon chains with their functional groups. However, our model captures not only carboxylic acid but also phosphonate groups, which are often considered bioisosteric replacements for sulfonate acids in medicinal chemistry due to their structural similarity (Macchiarulo and Pellicciari, 2007). In comparison, the ablation models only managed to capture one of these groups, indicating that multi-view approach enables the generation of accurate nomenclature and richer descriptive information.

6 Conclusion

In this paper, we introduce MV-CLAM equipped with MQ-Former, a novel cross-modal projector. The essence of this cross-modal projection lies in aligning both 2D and 3D molecular representation spaces simultaneously to enrich the textual space of language models. Our architecture with its multi-token contrasting loss successfully retains complementary information from multiple dimension into a single universal token easily interpreted by large language models. Extensive experiments demonstrate that MV-CLAM excels in molecule-text retrieval and molecule captioning tasks, with potential for broader applications, being a fine-tuned LLM with expertise in molecular structure.

7 Limitations

For future work, we aim to extend this framework to incorporate additional molecular representations, including other chemical structures, proteomics, and multiomics data. By aligning more views within MV-CLAM’s architecture, we anticipate improved navigation of the drug space and a deeper understanding of molecular interactions across biological contexts. Additionally, curating larger molecule-text datasets is expected to enhance the model’s performance and its ability to generalize to subtle molecular variations.

8 Acknowledgments

This research was supported by the Bio & Medical Technology Development Program of the National Research Foundation (NRF), funded by the Ministry of Science and ICT (MSIT), Republic of Korea (grant number. RS-2022-NR067933), NRF grant funded by the Korea government (MSIT)(RS-2023-00257479), Institute of Information & communications Technology Planning & Evaluation (IITP) grant funded by the Korea government (MSIT) [RS-2021-II211343, Artificial Intelligence Graduate School Program (Seoul National University)]. This study is also funded by AIGENDRUG Co., Ltd. and ICT at Seoul National University provided research facilities.

References

- Iz Beltagy, Kyle Lo, and Arman Cohan. 2019. Scibert: A pretrained language model for scientific text. *arXiv preprint arXiv:1903.10676*.
- He Cao, Zijing Liu, Xingyu Lu, Yuan Yao, and Yu Li. 2023. Instructmol: Multi-modal integration for building a versatile and reliable molecular assistant in drug discovery. *arXiv preprint arXiv:2311.16208*.
- Kirill Degtyarenko, Paula De Matos, Marcus Ennis, Janna Hastings, Martin Zbinden, Alan McNaught, Rafael Alcántara, Michael Darsow, Mickaël Guedj, and Michael Ashburner. 2007. ChEBI: a database and ontology for chemical entities of biological interest. *Nucleic acids research*, 36(suppl_1):D344–D350.
- Wenjie Du, Xiaoting Yang, Di Wu, FenFen Ma, Baicheng Zhang, Chaochao Bao, Yaoyuan Huo, Jun Jiang, Xin Chen, and Yang Wang. 2023. Fusing 2d and 3d molecular graphs as unambiguous molecular descriptors for conformational and chiral stereoisomers. *Briefings in Bioinformatics*, 24(1):bbac560.
- Carl Edwards, Tuan Lai, Kevin Ros, Garrett Honke, Kyunghyun Cho, and Heng Ji. 2022. Translation between molecules and natural language. *arXiv preprint arXiv:2204.11817*.
- Xiaomin Fang, Lihang Liu, Jieqiong Lei, Donglong He, Shanzhuo Zhang, Jingbo Zhou, Fan Wang, Hua Wu, and Haifeng Wang. 2022. Geometry-enhanced molecular representation learning for property prediction. *Nature Machine Intelligence*, 4(2):127–134.
- Zhichun Guo, Kehan Guo, Bozhao Nan, Yijun Tian, Roshni G Iyer, Yihong Ma, Olaf Wiest, Xiangliang Zhang, Wei Wang, Chuxu Zhang, et al. 2022. Graph-based molecular representation learning. *arXiv preprint arXiv:2207.04869*.
- Edward J Hu, Yelong Shen, Phillip Wallis, Zeyuan Allen-Zhu, Yuanzhi Li, Shean Wang, Lu Wang, and Weizhu Chen. 2021. Lora: Low-rank adaptation of large language models. *arXiv preprint arXiv:2106.09685*.
- John J Irwin, Teague Sterling, Michael M Mysinger, Erin S Bolstad, and Ryan G Coleman. 2012. Zinc: a free tool to discover chemistry for biology. *Journal of chemical information and modeling*, 52(7):1757–1768.
- Ross Irwin, Spyridon Dimitriadis, Jiazhen He, and Esben Jannik Bjerrum. 2022. Chemformer: a pre-trained transformer for computational chemistry. *Machine Learning: Science and Technology*, 3(1):015022.
- Seonghwan Kim, Jeheon Woo, and Woo Youn Kim. 2024. Diffusion-based generative ai for exploring transition states from 2d molecular graphs. *Nature Communications*, 15(1):341.
- Sunghwan Kim, Jie Chen, Tiejun Cheng, Asta Gindulyte, Jia He, Siqian He, Qingliang Li, Benjamin A Shoemaker, Paul A Thiessen, Bo Yu, et al. 2021. Pubchem in 2021: new data content and improved web interfaces. *Nucleic acids research*, 49(D1):D1388–D1395.
- Mario Krenn, Florian Häse, AkshatKumar Nigam, Pascal Friederich, and Alan Aspuru-Guzik. 2020. Self-referencing embedded strings (selfies): A 100% robust molecular string representation. *Machine Learning: Science and Technology*, 1(4):045024.
- Greg Landrum et al. 2013. Rdkit: A software suite for cheminformatics, computational chemistry, and predictive modeling. *Greg Landrum*, 8(31.10):5281.
- Samuel Lavoie, Polina Kirichenko, Mark Ibrahim, Mahmoud Assran, Andrew Gordon Wilson, Aaron Courville, and Nicolas Ballas. 2024. Modeling caption diversity in contrastive vision-language pretraining. *arXiv preprint arXiv:2405.00740*.
- Juncai Li and Xiaofei Jiang. 2021. Mol-bert: An effective molecular representation with bert for molecular property prediction. *Wireless Communications and Mobile Computing*, 2021(1):7181815.

- Juncheng Li, Xin He, Longhui Wei, Long Qian, Linchao Zhu, Lingxi Xie, Yueting Zhuang, Qi Tian, and Siliang Tang. 2022. Fine-grained semantically aligned vision-language pre-training. *Advances in neural information processing systems*, 35:7290–7303.
- Junnan Li, Dongxu Li, Silvio Savarese, and Steven Hoi. 2023. Blip-2: Bootstrapping language-image pre-training with frozen image encoders and large language models. In *International conference on machine learning*, pages 19730–19742. PMLR.
- Sihang Li, Zhiyuan Liu, Yanchen Luo, Xiang Wang, Xiangnan He, Kenji Kawaguchi, Tat-Seng Chua, and Qi Tian. 2024. Towards 3d molecule-text interpretation in language models. *arXiv preprint arXiv:2401.13923*.
- Pengfei Liu, Yiming Ren, Jun Tao, and Zhixiang Ren. 2024. Git-mol: A multi-modal large language model for molecular science with graph, image, and text. *Computers in biology and medicine*, 171:108073.
- Shengchao Liu, Weili Nie, Chengpeng Wang, Jiarui Lu, Zhuoran Qiao, Ling Liu, Jian Tang, Chaowei Xiao, and Animashree Anandkumar. 2023a. Multi-modal molecule structure–text model for text-based retrieval and editing. *Nature Machine Intelligence*, 5(12):1447–1457.
- Shengchao Liu, Hanchen Wang, Weiyang Liu, Joan Lasenby, Hongyu Guo, and Jian Tang. 2021. Pre-training molecular graph representation with 3d geometry. *arXiv preprint arXiv:2110.07728*.
- Zhiyuan Liu, Sihang Li, Yanchen Luo, Hao Fei, Yixin Cao, Kenji Kawaguchi, Xiang Wang, and Tat-Seng Chua. 2023b. Molca: Molecular graph-language modeling with cross-modal projector and uni-modal adapter. *arXiv preprint arXiv:2310.12798*.
- Antonio Macchiarulo and Roberto Pellicciari. 2007. Exploring the other side of biologically relevant chemical space: insights into carboxylic, sulfonic and phosphonic acid bioisosteric relationships. *Journal of Molecular Graphics and Modelling*, 26(4):728–739.
- Łukasz Maziarka, Tomasz Danel, Sławomir Mucha, Krzysztof Rataj, Jacek Tabor, and Stanisław Jastrzębski. 2020. Molecule attention transformer. *arXiv preprint arXiv:2002.08264*.
- Colin Raffel, Noam Shazeer, Adam Roberts, Katherine Lee, Sharan Narang, Michael Matena, Yanqi Zhou, Wei Li, and Peter J Liu. 2020. Exploring the limits of transfer learning with a unified text-to-text transformer. *Journal of machine learning research*, 21(140):1–67.
- Bing Su, Dazhao Du, Zhao Yang, Yujie Zhou, Jiangmeng Li, Anyi Rao, Hao Sun, Zhiwu Lu, and Ji-Rong Wen. 2022. A molecular multimodal foundation model associating molecule graphs with natural language. *arXiv preprint arXiv:2209.05481*.
- Jiabin Tang, Yuhao Yang, Wei Wei, Lei Shi, Lixin Su, Suqi Cheng, Dawei Yin, and Chao Huang. 2024. Graphgpt: Graph instruction tuning for large language models. In *Proceedings of the 47th International ACM SIGIR Conference on Research and Development in Information Retrieval*, pages 491–500.
- Hugo Touvron, Louis Martin, Kevin Stone, Peter Albert, Amjad Almahairi, Yasmine Babaei, Nikolay Bashlykov, Soumya Batra, Prajjwal Bhargava, Shruti Bhosale, et al. 2023. Llama 2: Open foundation and fine-tuned chat models. *arXiv preprint arXiv:2307.09288*.
- A Vaswani. 2017. Attention is all you need. *Advances in Neural Information Processing Systems*.
- Sheng Wang, Yuzhi Guo, Yuhong Wang, Hongmao Sun, and Junzhou Huang. 2019. Smiles-bert: large scale unsupervised pre-training for molecular property prediction. In *Proceedings of the 10th ACM international conference on bioinformatics, computational biology and health informatics*, pages 429–436.
- Yuyang Wang, Jianren Wang, Zhonglin Cao, and Amir Barati Farimani. 2022. Molecular contrastive learning of representations via graph neural networks. *Nature Machine Intelligence*, 4(3):279–287.
- Fang Wu, Dragomir Radev, and Stan Z Li. 2023. Molformer: Motif-based transformer on 3d heterogeneous molecular graphs. In *Proceedings of the AAAI Conference on Artificial Intelligence*, volume 37, pages 5312–5320.
- Zheni Zeng, Yuan Yao, Zhiyuan Liu, and Maosong Sun. 2022. A deep-learning system bridging molecule structure and biomedical text with comprehension comparable to human professionals. *Nature communications*, 13(1):862.
- Juzheng Zhang, Yatao Bian, Yongqiang Chen, and Quanming Yao. 2024. Unimot: Unified molecule-text language model with discrete token representation. *arXiv preprint arXiv:2408.00863*.
- Gengmo Zhou, Zhifeng Gao, Qiankun Ding, Hang Zheng, Hongteng Xu, Zhewei Wei, Linfeng Zhang, and Guolin Ke. 2023. Uni-mol: A universal 3d molecular representation learning framework. In *The Eleventh International Conference on Learning Representations*.

A Appendix

A.1 Related Works

Molecule-Text Modeling. Early approaches utilize 1D SMILES molecular sequences to treat molecules as text sequences by adapting Transformer models (Vaswani, 2017) designed for natural language processing (Irwin et al., 2022; Wang et al., 2019). KV-PLM (Zeng et al., 2022) specifically employs a masked language modeling loss to pretrain on biomedical texts with 1D SMILES representation. MolT5 (Edwards et al., 2022) specializes T5 model (Raffel et al., 2020) and tokenizer for SMILES-to-text and text-to-SMILES translations. Further enhancements represent molecules as 2D graphs. In particular, MoMu (Su et al., 2022) and MoleculeSTM (Liu et al., 2023a) leverage cross-modal contrastive learning to align the molecule graph representation to text. GIT-Mol (Liu et al., 2024), which uses molecular images and pretrains on millions of compounds from PubChem. As such, current approaches for leveraging multi-view molecular representations often rely on contrastive learning objectives.

Additionally, aided with the development of vision large language models (VLLMs), molecular large language models with multi-modal learning architectures have been developed. Simple projection layers were used in prior works, InstructMol (Cao et al., 2023) and GraphGPT (Tang et al., 2024), to project molecular graph representations to LLM’s input text token space. Recent works have been concentrated on utilizing Q-Former (Li et al., 2023) suggested in vision domain to bridge the gap between molecule and text modality. MolCA (Liu et al., 2023b) and 3D-MoLM (Li et al., 2024) aligns 2D graph and 3D conformer molecular representations to text in purpose to generate effective soft-prompts for large language models. UniMoT (Zhang et al., 2024) employs a vector quantization-driven tokenizer with a Q-Former. Current methods for utilizing multi-view representations of molecules are limited to contrastive learning or usage of specialized tokenizers, failing to achieve simultaneous alignment across all views and text, thereby neglecting the core principle of cross-modal alignment.

Molecular representation learning. Recent research in representation learning for molecules has seen significant advancements, particularly in leveraging large-scale unlabeled molecular data. SMILES-BERT (Wang et al., 2019), MolBERT (Li

and Jiang, 2021) adapts the BERT architecture on SMILES string for molecular property prediction tasks. To better focus on structural information of molecules, various graph-based representation learning models were presented. MolCLR (Wang et al., 2022) specifically tailored contrastive learning for molecular graphs using data augmentation while MAT (Maziarka et al., 2020) reinterpreted the attention mechanism of transformers to consider distance and edges. More recent works concentrate on employing 3D geometry, mostly to exploit 3D spatial coordinates. GraphMVP (Liu et al., 2021) proposed a contrastive learning framework that bridges 2D topological and 3D geometric views of molecules. GEM (Fang et al., 2022) incorporated 3D geometric information by using bond angles and lengths as additional edge attributes in molecular graphs. Uni-Mol is a SE(3)-transformer based model pretrained via 3D position recovery and masked atom prediction. Additionally, MolFormer (Wu et al., 2023) integrates SMILES, graph, and 3D conformer information in a unified transformer architecture for molecular property prediction. These recent advancements demonstrate a trend towards incorporating more diverse and rich molecular information to improve the quality and applicability of learned representations, validating the approach of our research.

A.2 Datasets Statistics

PubChem. We gathered 324k SMILES-text pairs from PubChem, generating 2D graphs and 3D conformations using existing methods (Maziarka et al., 2020; Landrum et al., 2013). Molecules with valid structures were used, with 15k longer-text pairs for training, and shorter ones for pretraining.

Table 4: PubChem324k dataset statistics

Subset	#Molecule-Text Pairs	#Min Words	#Avg Words
Pretrain	290,507	1	17.84
Train	11,753	20	57.24
Valid	977	20	58.31
Test	1,955	20	55.21

For the molecule captioning task, we chose not to use ChEBI-20 dataset (Degtyarenko et al., 2007) due to two main considerations (Li et al., 2024). First, ChEBI-20 is a curated subset of PubChem, which introduces potential issues of data redundancy and leakage given the overlap between the two datasets. Similarly, PCDes (Zeng et al., 2022) and MoMu (Su et al., 2022) datasets are also sub-

sets of PubChem demonstrating substantial overlap between its test dataset and our pretrain, train dataset. Table 5 shows using these datasets as evaluation benchmarks would not provide a fair assessment of generalization, and may in fact lead to misleading conclusions due to inadvertent data redundancy.

Second, ChEBI-20 replaces molecular names with generic terms like ‘the molecule’, limiting the evaluation of the model’s ability to associate structural features with accurate molecular names. Therefore, we utilized the PubChem dataset, which retains molecular names and offers a broader variety of structures, ensuring a more comprehensive evaluation of our framework in molecule captioning task.

Dataset	Pretrain subset (290,507)	Train subset (11,753)
CheBI-20 Test (3,290)	159 (4.83%)	804 (24.44%)
PCDes Test (2,962)	245 (8.27%)	631 (21.30%)
MoMu dataset (14,991)	6,424 (42.85%)	1,635 (10.91%)

Table 5: Distribution of molecules from external text datasets in the PubChem324k pretrain and train subsets.

A.3 Molecule-Text Retrieval Performance: Recall@10

To provide a more comprehensive assessment of retrieval performance beyond Recall@20, we additionally report results using Recall@10 in Appendix Table 6. The results show a consistent trend across metrics, where models with higher Recall@20 also achieve higher Recall@10. Notably, our framework with the proposed multi-token contrastive loss yields further improvements, with MV-CLAM[†] achieving the best overall performance. These findings highlight the effectiveness of both multi-modal integration and the proposed loss design in enhancing retrieval accuracy.

A.4 Experimental Settings

Stage 1 Molecule-Text Retrieval Pretraining.

Stage 1 serves to effectively transform molecular representations into query tokens interpretable in textual space. Using the PubChem324k pretraining subset with shorter textual descriptions, that is less informative but easier to align, MQ-former is trained for 35 epochs. A total of 301,658 molecules generated valid 2D graphs and 3D conformers, and thereby was used for pretraining. The goal of this stage was to optimize MQ-Former’s universal query generation by multi-objective training (molecule-text contrasting, molecule-text contrast-

ing, and molecule captioning). Pretraining was conducted for 35 epochs using 3 NVIDIA A6000 GPUs with a batch size of 99. Learnable query tokens of each view were set to 12 tokens and were randomly initialized. Both the Uni-Mol and MAT graph encoders were frozen throughout the pipeline to prevent the model from focusing too much on modifying the graph encoders, ensuring the training prioritized aligning representations with the textual space. To put emphasis on the decoding ability given the molecule tokens, we assigned a weight of 2 to the captioning loss. Maximum text length was configured to 256. We used an optimizer with a warmup step of 200 and a learning rate scheduler with a decay rate of 0.9. Gradient accumulation was set to 1 batch per step.

Stage 1 Molecule-Text Retrieval Finetuning.

After 35 epochs of pretraining, we loaded the checkpoint and fine-tuned MQ-Former for an additional 10 epochs on PubChem’s train, validation and test datasets, consisting of 12,000, 1,000, and 2,000 molecules respectively. Training is conducted using our modified multi-token contrastive loss. This serves to raise alignment capability given longer and more complex textual descriptions. The optimizer, learning rate scheduler, batch size and text length settings are identical to the previous phase.

Stage 2 Molecule Captioning Pretraining.

Stage 2 serves to further refine the universal tokens in a manner suited to a specific language model, LLaMA2 (Touvron et al., 2023) available at <https://huggingface.co/baffo32/decapoda-research-llama-7B-hf>. Using the trained model checkpoint from Stage 1 training stage, we conducted 4 epochs of pretraining on the PubChem dataset. The universal query generated by MQ-Former, along with the 1D SMILES string and an instruction prompt were given as input to the language model to generate textual descriptions for the molecules.

To fine-tune LLaMA2 efficiently, we employed LoRA (Hu et al., 2021) with a configuration of $r=8$, $\alpha=32$, and a 0.1 dropout rate. These settings were applied to the $[k_{proj}, v_{proj}, q_{proj}, o_{proj}, gate_{proj}, up_{proj}, down_{proj}]$ modules, adding 19 million trainable parameters, which constituted 0.29% of the total parameters in the LLaMA2-7B model. Unlike Stage 1, we used batch size of 30 with a maximum text length of 320 considering the prompt size. Token length for generation was set to range between 128 and 320. Gradient accumulation was set

Table 6: *Recall@10* and *Recall@20* for variant of MV-CLAM in molecule-text retrieval task. [‡] denotes the use of multi-token contrastive loss.

	Retrieval in batch				Retrieval in test set			
	Molecule-to-Text		Text-to-Molecule		Molecule-to-Text		Text-to-Molecule	
	REC@10	REC@20	REC@10	REC@20	REC@10	REC@20	REC@10	REC@20
1D+2D (-3D)	99.69	99.90	99.80	99.90	95.55	96.78	95.09	96.52
1D+2D (-3D) [‡]	99.64	99.80	99.64	99.80	99.75	97.03	95.19	96.52
1D+3D (-2D)	99.54	99.90	99.64	99.90	94.02	96.42	94.32	96.27
1D+3D (-2D) [‡]	99.80	99.85	99.64	99.74	94.17	96.42	94.32	95.70
MV-CLAM	99.80	99.95	99.69	99.95	95.19	96.62	95.19	96.37
MV-CLAM [‡]	99.85	99.95	99.80	99.90	96.06	96.98	95.70	96.93

to 2. The training was carried out using 3 NVIDIA A6000 GPUs.

Stage 2 Molecule Captioning Fine-tuning.

Stage 2 pre-training checkpoint was further fine-tuned on the train dataset for additional 10 epochs. Experimental settings are same as stage 2 pre-training phase, and validated using valid, test datasets.

A.5 Effectiveness of MQ-Former

In this section, we provide the detailed explanations and figures of Section 5.3. We illustrate the underlying mechanism for MQ-Former, which aligns two representations by providing (1) generated captions with ground truth, (2) caption comparison with Q-former based single-view alignment, and (3) attention map visualization.

A.5.1 Comparison of MV-CLAM Captions with Ground Truth

Appendix Table 7 provides caption examples within the test dataset as specified in Section 5.2. MV-CLAM not only correctly generates IUPAC and generic names but also additional information unavailable in ground truth labels.

A.5.2 Single-View Alignment Captions

Appendix Figure 5 highlights the differences in captioning results between the uni-modal Q-Former ablation models and ours. This demonstrates that the multi-view approach generates richer and more precise molecular descriptions.

A.5.3 Attention Map Visualization

We provide images of the attention maps explained in Section 5.3 (Appendix Figure 6). The attention maps of the shared self-attention layers are visualized to compare the processing of 2D and 3D query tokens with and without the multi-token contrasting loss. With the proposed loss, query tokens exhibit

diverse attention scores for each word in the captioning sentences while effectively distinguishing 2D- and 3D-related terms. Specifically, 2D query tokens focus on chemical and material properties (e.g., *boiling point*, *toxic*, *eye contact*), while 3D query tokens capture structural information (e.g., *bis(2-(dimethylamino)ethyl)*). In contrast, the original contrastive loss reduces query token diversity and weakens MQ-Former’s ability to align with 2D- and 3D-specific terms. This demonstrates that MQ-Former with the revised contrastive loss not only effectively preserves modality-specific information from 2D and 3D while aligning seamlessly with textual semantics but also guarantees query token diversity.

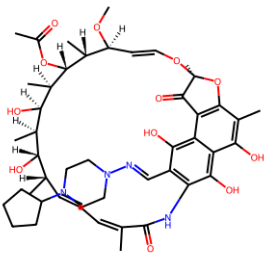
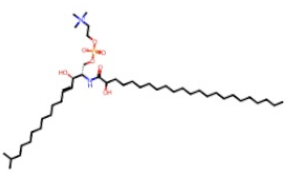
A.6 Human Evaluation by Domain Experts

We conducted a human evaluation by domain experts to complement the limitations of automated metrics and validate the chemical correctness and informativeness of generated captions. This evaluation aimed to assess whether our proposed MV-CLAM model does indeed benefit from its multi-view approach to generate more accurate and informative molecular descriptions than its 2D- and 3D-only ablated variants using the Q-former module.

A.6.1 Molecule Selection for Human Evaluation

To ensure representative and chemically diverse assessment, we selected 15 molecules spanning a range of compound types including peptides, hormones, drugs, aromatic compounds, coenzyme intermediates, vitamins, and environmental chemicals. Selection was based on structural diversity and clarity of model outputs - for fair comparison, we only selected cases where all three models gave reasonably comprehensive explanation about the molecule. Table 8 summarizes the selected

Table 7: Comparison of ground truth and MV-CLAM descriptions. Matching keywords are highlighted in bold, while additional details provided by MV-CLAM are marked in red.

Molecule	Ground Truth	MV-CLAM
	Rifapentine is a rifamycin antibiotic that is similar in structure and activity to rifampin and rifabutin and that is used in combination with other agents as therapy of tuberculosis , particularly in once or twice weekly regimens. Rifapentine is associated with transient and asymptomatic elevations in serum aminotransferase and is a likely cause of clinically apparent acute liver injury .	Rifapentine is a rifamycin antibiotic that is similar in structure and activity to rifampin and rifabutin and that is used in combination with other agents as therapy of tuberculosis , particularly in once or twice weekly regimens. Rifapentine is associated with transient and asymptomatic elevations in serum aminotransferase and is a likely cause of clinically apparent acute liver injury . Rifapentine is a long-acting, cyclopentyl-substituted derivative of rifamycin .
	N-(2-hydroxytricosanoyl)-15-methylhexadecaspHING-4-enine-1-phosphocholine is an N-acyl-15-methylhexadecaspHING-4-enine-1-phosphocholine in which the acyl group has 23 carbons and 0 double bonds and is 2-hydroxylated . It is functionally related to a 15-methylhexadecaspHING-4-enine.	N-(2-hydroxytricosanoyl)-15-methylhexadecaspHING-4-enine-1-phosphocholine is an N-acyl-15-methylhexadecaspHING-4-enine-1-phosphocholine in which the acyl group has 23 carbons and 0 double bonds and is 2-hydroxylated . It is functionally related to a 15-methylhexadecaspHING-4-enine.

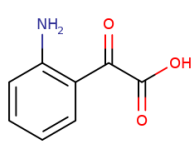
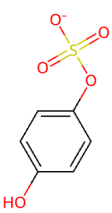
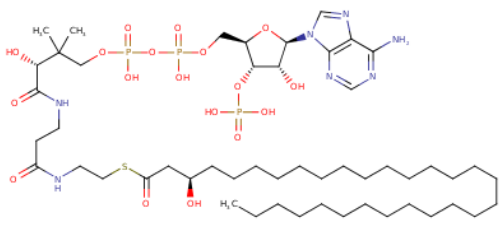
2D Only	3D Only	Original	Ground Truth
isatinic acid is a member of the class of 4-aminobenzoic acids that is anthranilic acid substituted by a hydroxy group at C-5. It has a role as a bacterial metabolite. It is a monohydroxybenzoic acid and a member of 4-aminobenzoic acids. It is functionally related to an anthranilic acid. It is a conjugate acid of an isatinate. Anthranilic acid is a metabolite found in or produced by Escherichia coli (stra	4-hydroxyphenyl sulfate(1-) is a phenyl sulfate oxoanion that is the conjugate base of 4-hydroxyphenyl hydrogen sulfate, obtained by deprotonation of the sulfate group; major species at pH 7.3. It has a role as a human metabolite. It is a conjugate base of a 4-hydroxyphenyl hydrogen sulfate. Phenyl hydrogen sulfate is a metabolite found in or produced by Escherichia coli	(R)-3-hydroxytriacontanoyl-CoA is a 3-hydroxy fatty acyl-CoA that results from the formal condensation of the thiol group of coenzyme A with the carboxy group of (R)-3-hydroxytriacontanoic acid. It is a (R)-3-hydroxyacyl-CoA, a 3-hydroxy fatty acyl-CoA and an ultra-long-chain fatty acyl-CoA. It is a conjugate acid	(R)-3-hydroxytriacontanoyl-CoA is a 3-hydroxy fatty acyl-CoA that results from the formal condensation of the thiol group of coenzyme A with the carboxy group of (R)-3-hydroxytriacontanoic acid [(R)-3-hydroxymelissic acid]. It is a (R)-3-hydroxyacyl-CoA, a 3-hydroxy fatty acyl-CoA and an ultra-long-chain fatty acyl-CoA. It is functionally related to a triacontanoic acid. It is a conjugate acid of a (R)-3-hydroxytriacontanoyl-CoA(4-)
			
2D only	3D only	Original	Ground Truth

Figure 5: Comparison of Uni-modal Q-Former Ablation and Ours

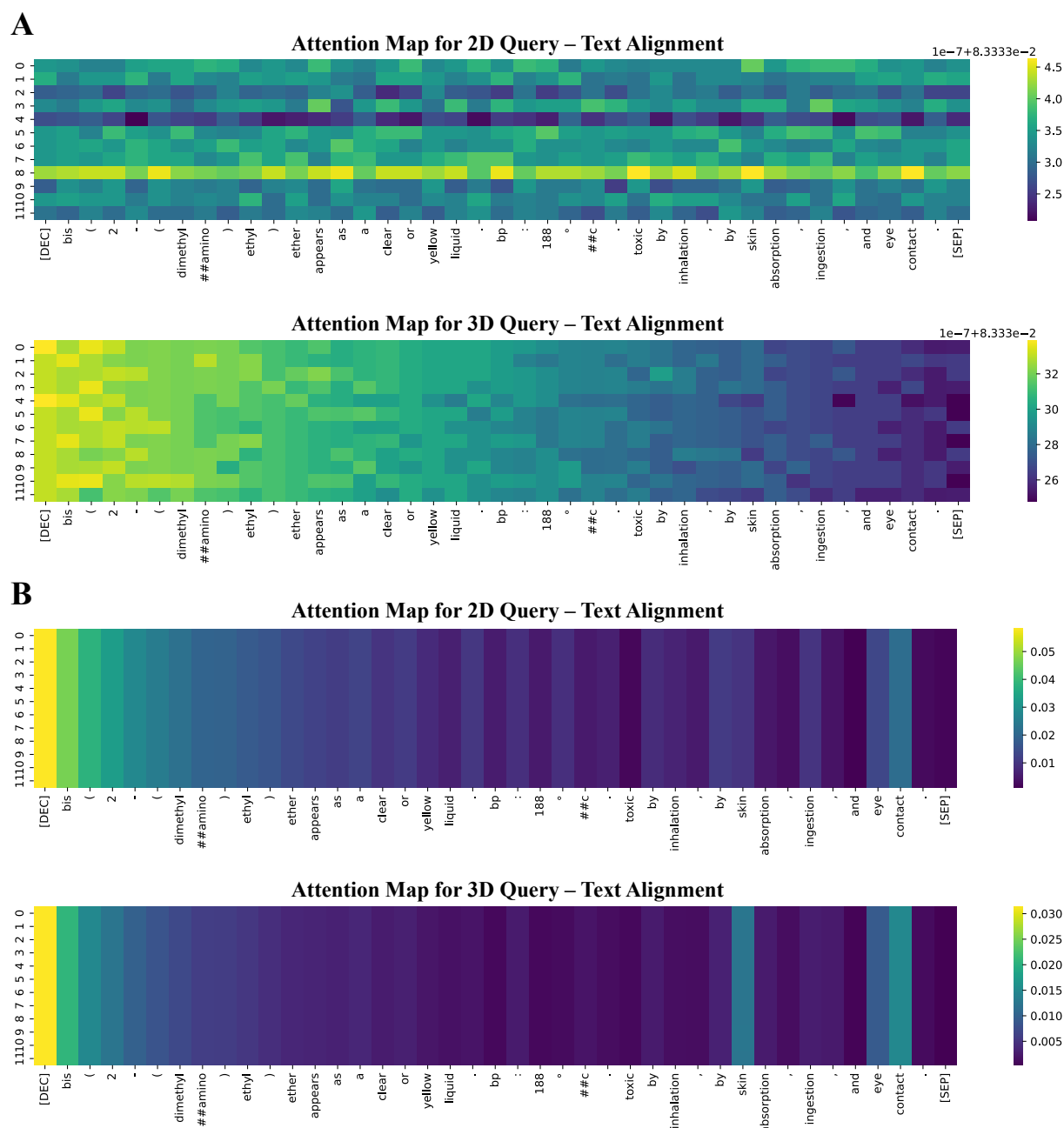


Figure 6: Comparison of attention map visualizations using different contrasting losses. The x-axis represents the word tokens in the sentence: *[DEC] bis (2 - (dimethylamino) ethyl) ether appears as a clear or yellow liquid . bp : 188 ° c . toxic by inhalation , by skin absorption , ingestion , and eye contact . [SEP]*., while the y-axis corresponds to the query tokens representing the molecule. (A) shows each query exhibiting different attention weights across the textual descriptions. Additionally, 2D query tokens focus on chemical and material properties (e.g., boiling point, toxic, eye contact), while 3D query tokens capture structural information (e.g., bis(2-(dimethylamino)ethyl)). Comparatively in (B), all query tokens have consistent attention distributions for all text tokens and lack word specificity for each dimension.

molecules and their associated chemical class.

Molecule	Chemical Class
Lanreotide	Peptide / Drug
5alpha-dihydrotestosterone	Steroid / Hormone
N(4)-acetylsulfathiazole	Sulfonamide Antibiotic
2,3',4,6-Tetrachlorobiphenyl	Aromatic / Pollutant
l-serine zwitterion	Amino Acid
DY-701(1-)	Fluorescent Dye
4-fluorobenzoic acid	Aromatic Acid
(S)-nicotine	Alkaloid / Stimulant
3-oxoadipyl-CoA	Coenzyme Intermediate
Ipratropium Bromide	Anticholinergic Drug
Riboflavin sodium	Vitamin B Derivative
(S)-verapamil hydrochloride	Calcium Channel Blocker
Soyasapogenol B 3-O-beta-glucuronide	Saponin / Glycoside
Betamethasone	Corticosteroid

Table 8: Selected molecules for human evaluation and their chemical categories.

A.6.2 Selection of Experts and Rater Instructions

We recruited a total of 11 domain experts to participate in our human evaluation. The expert pool consisted of licensed pharmacists, as well as MS and PhD graduates or candidates in organic chemistry, pharmaceutical sciences and chemical and biomolecular engineering. All participants possessed foundational knowledge in molecular structure and nomenclature, ensuring an informed assessment of caption quality.

To thank them for their time, each participant was provided with a small token of appreciation in the form of a Starbucks gift card after completing the evaluation. Participation was voluntary, and compensation was not tied to the correctness of their responses. This token was deemed appropriate given the minimal time commitment and the demographic of the participant pool (graduate-level or professional experts).

The evaluation was conducted using an anonymized Google Form online. For each molecule, participants were presented with:

- A high-resolution 2D structure image of the molecule and
- Three captions (corresponding to predictions from 2D-only, 3D-only, and MV-CLAM models), labeled as options A, B, and C in randomized order.

Raters were instructed to select the caption that best accurately and informatively described the molecule - considering structural fidelity, naming accuracy and chemical detail (Figure 7). They were

also encouraged to optionally explain their rationale or note any issues with the captions. At the end of the survey, experts were asked to rate the informativeness of the captions within a range of 1-5 and report recurring model weaknesses. A disclaimer clarified that all captions were AI-generated and may contain factual or chemical errors to encourage critical evaluation.

A.6.3 Ethics Review Board Approval

We did not seek IRB approval, as:

The evaluation posed minimal risk;

- Participants were informed, anonymous, and voluntarily participated.
- The only collected data were expert assessments of AI-generated content.
- No personally identifiable information was collected.

However, we acknowledge the inclusion of a small thank-you gift and were careful to avoid any form of coercion or undue influence.

A.6.4 Evaluation of MV-CLAM Captions

Across 15 molecule examples and 11 expert raters (165 total judgments), captions generated by MV-CLAM were preferred in 64% of cases. When aggregating preferences per example, MV-CLAM received the majority vote in 11 out of 15 molecules. The remaining cases were split between the 2D-only and 3D-only baselines.

This consistent preference indicates that integrating both 2D and 3D views contributes meaningfully to the chemical correctness and informativeness of captions. MV-CLAM was frequently favored for its ability to:

- Correctly identify core functional backbones and moieties (e.g., peptide bonds, phytosphingosine, soyasapogenol structure),
- Provide more specific structural terminology (e.g., aware of stereochemistry),
- Capture molecule class or role with chemically accurate language.

A.6.5 Limitations of LLM-generated Molecule Captions

Despite MV-CLAM’s improvements, experts noted several recurring limitations across all models. We provide the categories, the number of raters raising the issue and several comment examples:

Human Evaluation of Molecule Captioning Models

This survey is part of an academic research project evaluating the quality of AI-generated molecule descriptions. Your responses will help validate model performance beyond automated metrics.

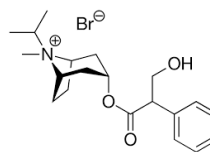
You will be shown 15 molecular structures, each accompanied by three AI-generated captions describing the molecule. Your task is to:

1. Carefully observe the molecular structure image.
2. Read all three caption options (labeled A, B, and C).
3. Choose the caption that best describes the molecule in terms of (1) structural correctness (atoms, bonds, stereochemistry), (2) chemical naming accuracy, (3) informativeness and clarity
4. Optionally, explain (1) why you chose your preferred caption and (2) any issues you noticed with the captions.

You do **not** need to identify which caption came from which model. The caption order is randomized to prevent bias.

▲ **Note:** The captions shown are automatically generated by AI models. Some may contain errors, be vague, or not fully reflect the chemical structure. You are encouraged to use your expert judgment to identify such issues. Also, there may be cases where **none of the captions are fully correct**. Please still select the best one and feel free to comment on limitations you notice.

Thank you for your time and participation.



Hyoscyamine methobromide is a natural product found in *Brugmansia arborea* and *Brugmansia suaveolens* with data available. Hyoscyamine Bromide is the bromide salt form of hyoscyamine, a belladonna alkaloid derivative and anticholinergic drug. Hyoscyamine bromide blocks the muscarinic cholinergic receptors in the smooth muscles of the gastrointestinal, urinary, and reproductive tracts as well as vascular smooth muscle.

[[[(4S)-3-exo-sec-butyl-8-isopropyl-8-methyl-8-azoniabicyclo[3.2.1]octan-3-yl]oxy]-3-hydroxy-2-phenylpropanoate is an alkylbenzene. Ipratropium is a muscarinic antagonist used in the treatment of chronic obstructive pulmonary disease (COPD). Ipratropium has not been associated with serum enzyme elevations.

Ipratropium bromide is the anhydrous form of the bromide salt of ipratropium. An anticholinergic drug, ipratropium bromide blocks the muscarinic cholinergic receptors in the smooth muscles of the bronchi in the lungs. This opens the bronchi, so providing relief in chronic obstructive pulmonary disease and acute asthma. It has a role as a bronchodilator agent, a muscarinic antagonist and an antispasmodic drug.

A

B

C

- ☐ A
☐ B
☐ C

Reasons for your selection or additional comments.

Figure 7: Evaluation interface presented to human raters. Experts were asked to assess the accuracy and informativeness of three AI-generated captions describing the shown molecule.

- Chemical names were incorrect or inconsistent - 11: *Wrong IUPAC name, Mis-identified to benzofurans, Wrong number of chlorides in chemical name*
- Functional groups were missing or hallucinated - 9: *Functional group is not glucoside, Does not capture acetaminophen structure*
- Descriptions were too generic - 5: *Do not have useful information about the chemical*
- Stereochemistry was ignored - 1: *Wrong stereochemistry for (L)-serine*
- Hallucinated properties/nomenclature - 1: *Androsphinganine does not exist, Molecule is not trichlorinated*

These observations highlight the need for more precise structure-to-text alignment mechanisms in molecular captioning systems, particularly to capture small yet chemically significant functional groups. To ensure practical deployment of large language models in molecular domains, further refinement informed by expert evaluation is essential and remains an important direction for future work.

A.7 Downstream Task 1. Question Answering

A.7.1 Dataset: 3D-MoIT

A total of 18439K molecule-instruction text pairs are employed using the dataset split as given in the original paper (Li et al., 2024). The dataset consists of two types of molecular property prediction tasks: (1) Computed property prediction including 3D-dependent properties (e.g. HOMO) and (2) descriptive property prediction.

Table 9: Statistics of the PubChemQC and PubChem datasets across different subsets.

Subset	PubChemQC			PubChem		
	#Mol	#Comp.	QA	#Mol	#Comp.	QA
Pretrain	3,119,717	12,478,868		301,658	1,199,066	1,508,290
Train	623,944	2,495,776		12,000	46,680	60,000
Valid	77,993	311,972		1,000	3,898	5,000
Test	77,993	311,972		2,000	7,785	10,000

A.7.2 Experimental Settings

For the molecular question-answering task, we utilized the 3D-MoIT (Li et al., 2024) dataset, which includes question-prompt and text-answer pairs derived from the same PubChem data we used in prior. Dataset statistics are in Appendix Table 9. The dataset consists of three distinct subsets: (1) Question-answering about non-3D properties, (2) Question-answering about 3D properties, and (3) Descriptive molecular properties.

For robust guidance into instruction tuning, the

three sub-datasets of 3D-MoIT (Li et al., 2024) were used in combination for training a single epoch. To ensure a fair comparison with single-view methods, we initialized the instruction-tuning process using the pretrained MV-CLAM checkpoints from the molecule captioning stage, employing the original loss function rather than the multi-token contrasting loss. Given the dataset size, the model was further fine-tuned for 5 epochs on non-3D, descriptive property tasks and 1 epoch on 3D property tasks. For computed property prediction, we evaluated performance using mean absolute error (MAE). For descriptive property prediction, we measured BLEU, ROUGE, and METEOR scores.

A.7.3 Results

For baselines, we reproduced results for 3D-MoLM and 2D-MoLM (with MAT (Maziarka et al., 2020) graph encoder). These baselines represent single-modal alignment using Q-Former, and provides a fair point of comparison to demonstrate the efficacy of our multi-view cross-modal alignment. Appendix Tables 10, 11 and 12 show that MV-CLAM consistently outperformed the single-modal models.

Table 10: Comparison of Descriptive Property Generation Performance

Model	B-2	B-4	R-1	R-2	R-L	M
2D-MoLM	31.24	25.13	39.30	25.16	34.11	49.88
3D-MoLM	29.22	22.82	37.38	22.54	31.47	27.29
MV-CLAM [‡]	31.70	25.60	39.61	25.46	34.51	50.61

Table 11: Q&A performance on 3D properties

Model	HOMO	LUMO	HOMO-LUMO	SCF Energy
2D-MoLM	0.78 (0.99)	0.47 (0.99)	0.39 (0.90)	0.98 (1.00)
3D-MoLM	0.42 (0.99)	0.44 (0.98)	1.26 (0.99)	1.22 (0.98)
MV-CLAM [‡]	0.35 (0.98)	0.42 (0.93)	0.35 (0.99)	0.32 (0.99)

Table 12: Q&A performance on non-3D properties. MW, TPSA denotes molecular weight and topological surface area.

Model	MW	LogP	Complexity	TPSA
2D-MoLM	47.51 (0.98)	0.89 (0.99)	110.78 (0.99)	16.65 (0.99)
3D-MoLM	42.76 (0.96)	1.25 (0.96)	105.03 (0.96)	20.97 (0.92)
MV-CLAM [‡]	21.35 (0.92)	0.69 (0.94)	55.14 (0.91)	9.65 (0.91)

A.8 Downstream Task 2: Zero-shot Molecule Editing

Unlike conventional natural languages, SMILES encode molecular topology and properties demanding a specialized understanding of its notation system. Thereby, previous efforts in text-based de-novo molecule generation with large language

models typically involves training or developing tokenizers that account for the unique grammar of SMILES (Edwards et al., 2022). By fine-tuning MV-CLAM, we enabled the model to output SMILES strings without additional tokenizer training.

A.8.1 Dataset: ZINC20

Following the experiment settings of (Liu et al., 2023a), 200 molecules randomly selected from the ZINC20 dataset are given 6 single-objective molecule editing instructions. The 200 molecules follow the property distribution of the entire dataset, and do not overlap with the PubChem324k training dataset in previous stages. The six instructions are the following. (1) The molecule is soluble in water. (2) The molecule is insoluble in water. (3) The molecule has high permeability. (4) The molecule has low permeability. (5) The molecule is like a drug. (6) The molecule is not like a drug. (7) The molecule has more hydrogen bond donors. (8) The molecule has more hydrogen bond acceptors.

A.8.2 Experimental Settings

Zero-shot molecule editing was conducted on the curated dataset presented in (Liu et al., 2023a) which consists of 200 randomly sampled molecules from the ZINC dataset. Each molecule was paired with molecule editing prompts (chemical instructions such as "*The molecule is more soluble in water*") and their corresponding SMILES. The dataset included molecular structures that were unseen during training. Starting with the original SMILES, the universal molecular token generated by the trained MQ-Former, and the editing prompt, we generated SMILES of the edited molecule. Using the pretrained MV-CLAM checkpoints from the molecule captioning stage, we conducted zero-shot molecule editing, utilizing the model’s pre-existing multi-view molecular understanding from prior stages. The model was further fine-tuned for 4 epochs on the PubChem 324k pretraining and training datasets. This fine-tuning enabled MV-CLAM to directly generate SMILES from molecular universal tokens and was crucial to produce valid SMILES, considering the nature of LLaMA’s general-purpose tokenizer which was not explicitly trained for SMILES generation. We evaluate the edited results by computing desired chemical properties using RDKit (Landrum et al., 2013), and classify whether the modification was valid shot.

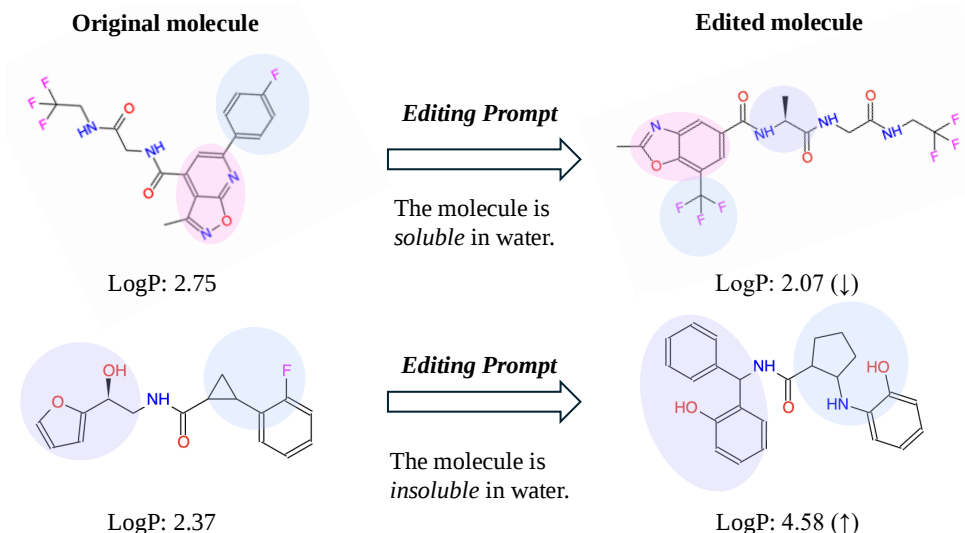


Figure 8: Zero-shot editing with chemical properties

A.8.3 Results

In this section we show successful case studies of the language model generating valid SMILES strings with adequate property modifications. Compared to previous works which mostly generate mere modifications of a single functional group, MV-CLAM generates diversified chemical structure modifications that may not be immediately obvious. This ability to generate more complex modifications is particularly advantageous for domain experts, as simple functional group changes are typically easy to perform manually. We attribute this diversity to the model’s robust understanding of molecules within the textual space. The alignment between molecules and text is achieved by focusing on distinct substructures and molecular properties through the multi-view approach.

The values presented indicate the predicted LogP (octanol-water partition coefficient), topological surface area (TPSA), quantitative estimate of drug-likeness (QED) and number of hydrogen bond and acceptors (Appendix Figure 8, 9,10,11,12). Each figure showcases original molecules alongside their modified counterparts with numerical indicators representing the chemical properties before and after the zero-shot editing. LogP values reflect solubility in water, while topological surface area relates to molecular permeability. QED reflects drug likeliness. The modifications are aligned with targeted property-based editing prompt, demonstrating the flexibility and chemical expertise of MV-CLAM.

Table 13: Comparison of cases with and without multi-token contrastive loss in both the full MV-CLAM framework and its ablation settings. ‡ denotes the use of multi-token contrastive loss.

Model	BLEU-2	BLEU-4	ROUGE-1	ROUGE-2	ROUGE-L	METEOR
1D+3D (-2D)	29.45	22.03	37.86	23.11	31.83	33.79
1D+3D (-2D) [‡]	29.46	22.19	37.33	23.12	31.53	33.34
1D+2D (-3D)	29.72	22.26	38.22	23.45	31.61	34.22
1D+2D (-3D) [‡]	30.71	23.32	39.08	24.57	33.01	34.81
MV-CLAM	31.75	24.48	40.43	25.72	33.79	36.54
MV-CLAM [‡]	32.32	25.11	40.87	26.48	34.79	36.87

A.9 Ablation Studies for Stage 2. Specializing LLaMA2 for Molecule Captioning

Effect of Multi-token contrastive loss We conducted an ablation study to examine the effect of multi-token loss not only in our full MV-CLAM framework, but also in ablation settings that use the pre-existing Q-Former with either the 2D or 3D modality removed (Appendix Table 13). While the multi-token contrastive loss improves performance in most cases, the gains are comparatively marginal in ablation settings. These results suggest that the benefits of multi-token loss are most pronounced when combined with MQ-Former and multiple modalities, where the shared self-attention layer enables effective cross-modal interaction and contextual alignment across 2D, 3D, and text. This interpretation is further supported by the attention map (Appendix Figure 6), which shows that multi-token contrastive loss encourages different query tokens to enhance fine-grained semantic alignment, with 2D- and 3D-based query tokens focusing on modality-specific segments of the caption—interacting through the shared self-attention layer within MQ-Former.

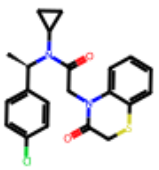
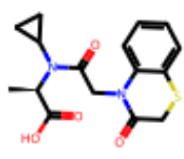
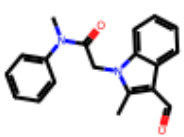
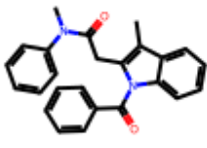
The molecule is soluble in water.		The molecule is insoluble in water. (LogP)	
	Original (4.53)		Modified (1.59)
	Original (3.43)		Modified (4.84)

Figure 9: Editing Solubility (LogP Adjustments): Smaller LogP indicates higher solubility in water. Molecules were successfully modified given the prompt "The molecule is soluble/insoluble in water".

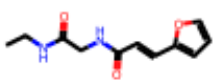
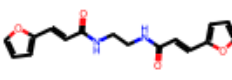
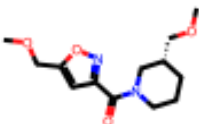
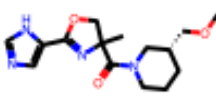
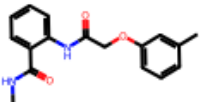
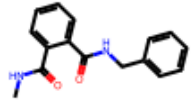
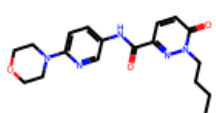
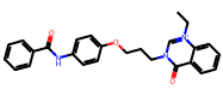
The molecule has high permeability.		The molecule has low permeability. (Topological Surface Area)	
	Original (71.34)		Modified (84.48)
	Original (64.80)		Modified (79.81)
	Original (67.43)		Modified (58.20)
	Original (89.35)		Modified (64.21)

Figure 10: Editing Permeability (Topological Surface Area, TPSA Adjustments): A higher TPSA implies lower permeability, while a lower TPSA suggests higher permeability. Molecules were successfully modified given the prompt "The molecule has high/low permeability".

1D Molecular Representations We conducted an ablation study to compare the use of SELFIES (Krenn et al., 2020) with SMILES as input representations (Appendix Table 14). Using the pretrained Stage 2 checkpoint, the model was further trained for captioning under identical settings. After 10 stages of training with SELFIES, SMILES consistently demonstrated superior performance across metrics such as BLEU, METEOR, and ROUGE, validating the effectiveness of our selection.

Table 14: Captioning performance comparison for 1D molecular representations

Model	B-2	B-4	R-1	R-2	R-L	M
SELFIES	28.39	20.89	33.25	37.58	22.49	31.37
SMILES	31.75	24.48	40.43	25.72	33.79	36.54

A.10 Failure Case Study

Appendix Table 15 showcases two instances where MV-CLAM fails to differentiate structurally similar molecules. First, the model misclassifies lactoyl-CoA as oleoyl-CoA despite the key difference being the length of the carbon chain. This indicates a limitation in the model’s capacity to capture subtle variations in carbon chain lengths. Second, the model misidentifies Ajugaciliatin B as subtypes E and C, demonstrating that while it successfully recognizes the molecule’s primary backbone, it struggles to distinguish the small functional groups that define each subtype. This suggests that the model is not sufficiently sensitive to minor structural modifications. Both errors appear to stem from the model’s difficulty in perceiving refined differences in chemical properties and spatial structure between the ground truth and its predictions. This underscores a broader challenge in molecular

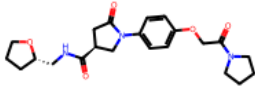
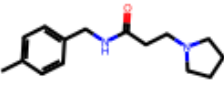
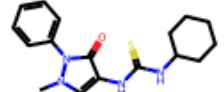
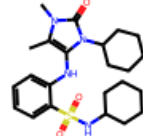
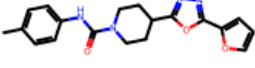
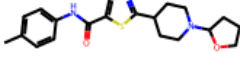
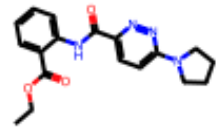
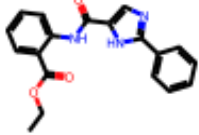
The molecule is like a drug.		The molecule is not like a drug.	
		(QED)	
			
Original (0.73)	Modified (0.86)	Original (0.84)	Modified (0.69)
			
Original (0.77)	Modified (0.88)	Original (0.84)	Modified (0.70)

Figure 11: Editing Drug Likelihood (Quantitative Estimate of Drug-likeness, QED): A higher QED suggests a compound is more likely to possess favorable pharmacokinetic and ADMET (absorption, distribution, metabolism, excretion, and toxicity) properties, being more drug-like. Molecules were successfully modified given the prompt "The molecule is/is not like a drug".

captioning: capturing subtle yet critical molecular features that may not greatly impact the primary structure but are crucial contributors for property.

To overcome these limitations, we propose several future studies. First, expanding our MQFormer to align additional views or modalities, along with finer-grained molecular or related biological embeddings, could offer complementary insights to enhance the model's ability to differentiate between similar molecules. This multi-view alignment could offer a more holistic understanding of the molecule's structure and properties. In addition, curating larger molecule datasets would enhance the model's capacity to generalize, ensuring it has sufficient exposure to a wide range of molecular variations during training. These developments will address the current shortcomings and pave the way for more accurate molecular identification in future iterations of the model.

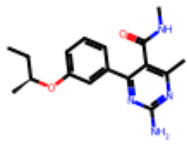
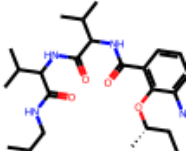
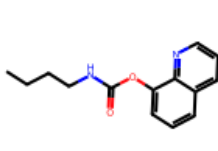
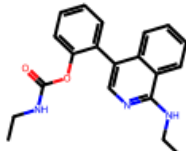
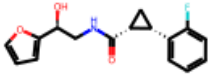
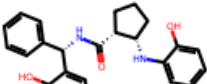
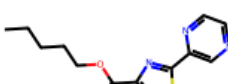
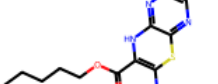
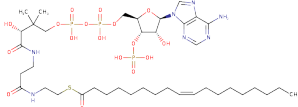
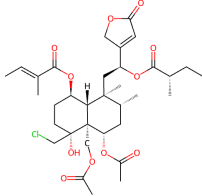
The molecule has more hydrogen bond donors.		The molecule has more hydrogen bond acceptors. (number of donors/acceptor)	
			
Original (2)	Modified (4)	Original (3)	Modified (4)
			
Original (2)	Modified (4)	Original (6)	Modified (7)

Figure 12: Editing Hydrogen Bond Acceptor/Donors: The number of hydrogen bond acceptors and donors in the molecule were given for evaluation. Molecules were successfully modified given the prompt "The molecule has more hydrogen bond donors/acceptors".

Table 15: Failure Case Study: Comparison of Ground Truth and MV-CLAM Descriptions

Molecule	Ground Truth	MV-CLAM
	Lactoyl-CoA(4-) is an acyl-CoA(4-) that is the tetraanion of lactoyl-CoA arising from deprotonation of phosphate and diphosphate functions. It is a conjugate base of a lactoyl-CoA.	oleoyl-CoA(4-) is an acyl-CoA(4-) that is the tetraanion of oleoyl-CoA, arising from deprotonation of phosphate and diphosphate functions. It has a role as a human metabolite and a <i>Saccharomyces cerevisiae</i> metabolite. It is a saturated fatty acyl-CoA(4-) and a 3-substituted propionyl-CoA(4-). It is a conjugate base of an oleoyl-CoA
	Ajugaciliatin B is a diterpene lactone isolated from the whole plants of <i>Ajuga ciliata</i> and has been shown to exhibit neuroprotective activity. It has a role as a plant metabolite and a neuroprotective agent. It is a diterpene lactone, an acetate ester, a butenolide, a carbobicyclic compound, an organochlorine compound and a tertiary alcohol. Ajugaciliatin B is a natural product found in <i>Ajuga ciliata</i> with data available.	ajugaciliatin C is a natural product found in <i>Ajuga ciliata</i> with data available. Ajugaciliatin E is a diterpene lactone isolated from the whole plants of <i>Ajuga ciliata</i> . It has a role as a plant metabolite. It is a butenolide, an acetate ester, a diterpene lactone and a organochlorine compound. It is functionally related to a tiglic acid. Ajugaciliatin E is a natural product found in <i>Ajuga ciliata</i>