
Enhancing RNA Generation for Protein Binding via Group Relative Policy Optimization

Yusuf Kesmen*

Department of Computer Engineering
Bilkent University
Cankaya, Ankara 06800
yusuf.kesmen@ug.bilkent.edu.tr

Alpsencer Özdemir*

Department of Computer Engineering
Bilkent University
Cankaya, Ankara 06800
alpsencer.ozdemir@ug.bilkent.edu.tr

Yavuz Alp Sencer Öztürk*

Department of Computer Engineering
Bilkent University
Cankaya, Ankara 06800
alp.ozturk@ug.bilkent.edu.tr

Abstract

Generating RNA sequences that effectively bind to target proteins is a pivotal challenge in bioinformatics with significant implications for therapeutic development and understanding biological mechanisms. This work focuses on enhancing the performance of a pre-fine-tuned T5 encoder-decoder transformer model for protein-conditional RNA generation. We introduce a novel application of Group Relative Policy Optimization (GRPO), a state-of-the-art reinforcement learning (RL) technique, to refine the generation process. A key component of our framework is a custom-designed, multi-faceted reward function that incorporates scores from a newly trained protein-RNA binding classifier, along with biologically relevant metrics such as GC content and Minimum Free Energy (MFE) of the generated RNA sequences. The binding classifier itself is trained on positive examples from the CLIPDB dataset and negative examples generated through data augmentation by shuffling RNA sequences. By leveraging GRPO with this tailored reward system, we aim to guide the T5 model towards producing RNA sequences with improved binding affinity and structural stability, thereby advancing the capabilities of computational RNA design.

1 Introduction

The interaction between RNA molecules and proteins is fundamental to a vast array of cellular processes, including gene regulation, RNA processing, and translation [1]. Dysregulation of these interactions is often implicated in various diseases, making the ability to design RNA sequences that can specifically bind to target proteins a critical area of research for therapeutic interventions and synthetic biology [2]. While traditional experimental methods for identifying such RNA aptamers can be laborious and costly, computational approaches, particularly those leveraging deep learning, offer a promising avenue for accelerated discovery.

Recent advancements in transformer-based language models have demonstrated remarkable success in modeling biological sequences, capturing complex patterns and dependencies within protein and nucleic acid "languages" [2]. These models can be fine-tuned for specific generative tasks, such as

*Fourth-Year Undergraduate Student in Computer Engineering, Bilkent University. Work done as part of EEE 486 project.

producing RNA sequences conditioned on a given protein target. Our prior work involved fine-tuning a T5 encoder-decoder model on protein-RNA pairs from datasets like CLIPDB [3] to generate potential RNA binders. However, supervised fine-tuning alone may not fully optimize for desired downstream properties like binding strength or structural stability.

Reinforcement learning (RL) has emerged as a powerful paradigm for optimizing generative models towards specific objectives, going beyond what can be achieved with static datasets. Techniques like Proximal Policy Optimization (PPO) [4] and more recent advancements such as Direct Preference Optimization (DPO) [5] and Group Relative Policy Optimization (GRPO) [6, 7] have shown success in aligning large language models (LLMs) with human preferences or complex reward signals. Inspired by the success of GRPO in enhancing reasoning capabilities in LLMs like DeepSeek-R1 (DeepSeekCoder in some contexts) [7], we propose to adapt this methodology for our RNA generation task.

This paper details our approach to enhance a pre-fine-tuned T5 model for RNA generation using GRPO. Our contributions are threefold: 1) The development of a dedicated protein-RNA binding classifier to provide a quantitative measure of binding likelihood, trained on curated positive data from CLIPDB and augmented negative data generated by shuffling. 2) The design of a composite reward function that integrates the classifier’s binding score with crucial biophysical properties of RNA: GC content for structural integrity and MFE for thermodynamic stability. 3) The implementation and application of the GRPO algorithm to an encoder-decoder transformer for this specific bioinformatics task, which, to our knowledge, is a novel application area for GRPO in biomolecular sequence design. We hypothesize that this RL-enhanced approach will lead to the generation of RNA sequences with superior binding characteristics and stability compared to the base fine-tuned model and other baseline approaches.

The remainder of this report is organized as follows: Section 2 reviews relevant literature. Section 3 describes our methodology, including the base model, binding classifier, reward function, and GRPO implementation. Section 4 presents our experimental results. Section 5 discusses the implications of our findings, and Section 6 concludes the report.

2 Related Work

Our work builds upon advances in three primary areas: transformer models for bioinformatics, reinforcement learning for large language models, and the application of RL in biological sequence design.

2.1 Transformer Models in Bioinformatics

Transformer architectures [8] have become a cornerstone in bioinformatics for both understanding and generating biological sequences. Protein language models such as ProtTrans [9] and ProtGPT2 [10] pre-train transformers on vast protein sequence databases, learning rich representations that capture the "language of life" and can generate novel, plausible protein sequences. These embeddings have proven effective for various downstream prediction tasks.

Similarly, transformers have been applied to RNA sequence modeling. GenerRNA [11] introduced a large-scale pre-trained GPT-2 style model for de novo RNA design, demonstrating the ability to generate novel RNA sequences with realistic structural features. RNAGEN [12], from CicekLab at Bilkent, employed a Generative Adversarial Network (GAN) with a WGAN-GP architecture for RNA generation, capable of incorporating external predictors as guidance. More recently, RNAtTranslator [13], also a CicekLab work, framed protein-conditional RNA design as a sequence-to-sequence translation task, using an encoder-decoder transformer to generate RNA sequences likely to interact with a given protein input. Our base model aligns with this conditional generation paradigm.

2.2 Reinforcement Learning for Large Language Models

RL has been instrumental in fine-tuning LLMs to align with human preferences or specific task objectives, a process often referred to as RLHF (RL from Human Feedback). PPO [4] is a widely adopted algorithm for this, known for its stability, and was famously used in models like InstructGPT and ChatGPT. However, PPO-based RLHF can be complex and resource-intensive.

To address these challenges, alternative methods have emerged. DPO [5] reframes preference-based alignment as a simpler classification problem, often matching or exceeding PPO performance with greater stability. GRPO [6], introduced by Shao et al. (2024) for DeepSeekMath, enhances training stability and efficiency by using relative rewards from grouped trajectories. Guo et al. (2024) further utilized GRPO in DeepSeekCoder (also referred to as DeepSeek-R1 contextually) [7], demonstrating its ability to induce structured reasoning behaviors even with self-supervised rewards, eliminating the need for a pre-learned reward model in some cases. Our choice of GRPO is motivated by these recent successes and its potential for stable optimization with complex, potentially noisy reward signals from biological predictors.

2.3 Reinforcement Learning in Bioinformatics

RL has also found applications in directly optimizing biological sequences. Early work by Eastman et al. (2018) [14] used deep RL for the RNA inverse folding problem, training a policy network to design sequences that fold into a target secondary structure. Subsequent works like LEARN [15] and EternaRL [16] (as cited in [17]) improved upon this, and DRAG [18] introduced a hierarchical graph-based RL agent for complex RNA topologies.

In protein engineering, ProteinRL [19] used policy-based RL to guide a generative PLM towards desired sequence properties like charge or solubility. Stocco et al. (2023) [20] (note: year corrected based on common citation patterns, please verify) applied a DPO variant (DPO_pLM) to align a protein generator with an external oracle (binding affinity predictor), rapidly identifying high-affinity binders. These examples demonstrate RL’s potential to explore vast sequence spaces and optimize for functional properties, a capability we aim to harness for RNA-protein binding. Our work distinguishes itself by applying GRPO to an encoder-decoder architecture for conditional RNA generation, guided by a multi-component reward including a custom-trained binding classifier.

3 Methods

Our methodology integrates a pre-fine-tuned T5 model with a novel protein-RNA binding classifier, a multi-component reward system, and GRPO algorithm for reinforcement learning-based refinement.

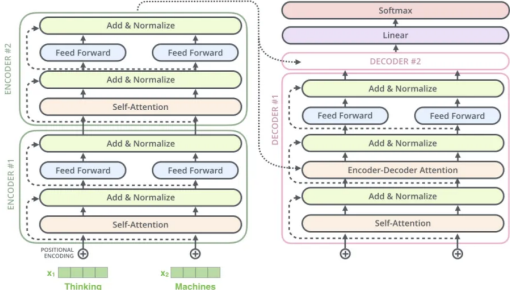


Figure 1: T5 Model Architecture

3.1 Base Model and Tokenization

The foundation of our work is a T5 encoder-decoder transformer model, specifically T5ForConditionalGeneration [21], which was previously pretrained and fine-tuned for the task of generating RNA sequences conditioned on a protein sequence input. This base model features an architecture with a hidden size (d_{model}) of 512, a feed-forward network dimension (d_{ff}) of 1024, and key/value dimensions (d_{kv}) of 64. Both the encoder and decoder stacks consist of 6 layers, each employing 12 attention heads. The architecture of T5 is shown in Figure 1.

For sequence processing, custom Byte Pair Encoding (BPE) tokenizers were trained independently for protein and RNA sequences. Each tokenizer has a vocabulary size of 1000. The model is configured to handle a maximum sequence length of 1024 tokens. The initial supervised fine-tuning of this T5 model utilized protein-RNA interaction data sourced from the CLIPDB dataset [3]. This

dataset provides interacting sequences in FASTA format, and the corresponding protein sequences for CLIPDB entries were obtained through database scraping. The specific pre-fine-tuned model checkpoint used in this work can be found in the additional package we provide.

3.2 Protein-RNA Binding Classifier

To provide a quantitative reward signal indicative of binding likelihood, we developed and trained a dedicated protein-RNA binding classifier. This classifier is designed to predict a binding probability score, ranging from 0 to 1, for a given protein-RNA pair.

The training data for the classifier was meticulously prepared. Positive interaction examples were directly sourced from the CLIPDB dataset. To generate negative samples and ensure a balanced dataset, RNA sequences from these positive pairs were shuffled. This process disrupts potential binding sites while preserving the original nucleotide composition, thereby creating a corresponding negative RNA sequence for each positive protein-RNA pair [22]. The classifier expects input in the format `protein_sequence$rna_sequence`.

Feature representation for the classifier involves generating embeddings for both protein and RNA sequences using our base T5 model. Protein embeddings are derived by taking the mean-pooled output of the encoder’s hidden states for an input protein sequence. RNA embeddings are obtained from the decoder’s output, specifically by using the hidden state of the CLS token (the first token of the last layer) for an input RNA sequence. Both protein and RNA embeddings undergo L2 normalization before being fed to the classifier.

The classifier, referred to as `MLPBindingClassifier`, is a Multi-Layer Perceptron (MLP) that ingests concatenated L2-normalized protein and RNA embeddings. This combined embedding is passed through three fully connected layers with 512, 256, and 128 neurons, respectively. Each dense layer is followed by a ReLU activation and a dropout layer (dropout rate = 0.2) to mitigate overfitting. The final layer produces a single scalar output via a sigmoid activation, representing the predicted binding probability. The `MLPBindingClassifier` was trained for 65 epochs on a dataset of 100 000 positive and 100 000 negative examples using the Adam optimizer with a learning rate of 1×10^{-5} and a weight decay of 0.01 (L2 regularization), optimizing the binary cross-entropy loss on NVIDIA 2080 Ti GPUs.

3.3 Reward Function Design

To guide the reinforcement learning agent effectively, we designed a composite reward function, R_{total} , aimed at promoting RNA sequences that exhibit not only a high likelihood of binding to the target protein but also possess sound structural properties. This total reward is a weighted sum of three distinct components:

$$R_{total} = w_{binding} \cdot R_{binding} + w_{GC} \cdot R_{GC} + w_{MFE} \cdot R_{MFE}. \quad (1)$$

In our experiments, the weights were set to $w_{binding} = 0.2$, $w_{GC} = 0.4$, and $w_{MFE} = 0.4$, reflecting a balanced consideration of binding and structural stability.

The first component, the Binding Affinity Reward ($R_{binding}$), is directly derived from the output of our trained protein-RNA binding classifier, as described in the preceding section. This score, ranging from 0 to 1, quantifies the predicted binding likelihood between the input protein and the generated RNA sequence.

The second component, the GC Content Reward (R_{GC}), encourages the generation of RNA sequences with a Guanine-Cytosine (GC) content [23] falling within an optimal range, typically associated with stable secondary structures. We define this range as 40% to 60% GC content. The reward $R_{GC}(s)$ for an RNA sequence s is formulated as:

$$R_{GC}(s) = \begin{cases} 1.0 & \text{if } 0.4 \leq \text{GC_content}(s) \leq 0.6 \\ 1.0 - \frac{|0.4 - \text{GC_content}(s)|}{0.4} & \text{if } \text{GC_content}(s) < 0.4 \\ 1.0 - \frac{|\text{GC_content}(s) - 0.6|}{1.0 - 0.6} & \text{if } \text{GC_content}(s) > 0.6, \end{cases} \quad (2)$$

where $\text{GC_content}(s)$ is the fractional GC content of sequence s .

The third component, the MFE Reward (R_{MFE}), incentivizes RNA sequences predicted to form thermodynamically stable secondary structures[24]. The MFE for a sequence s is calculated using

the ViennaRNA package [24]. To account for sequence length variations, we normalize the MFE by the sequence length: $MFE_{norm}(s) = \text{MFE}(s)/\text{length}(s)$. A lower (more negative) MFE value generally indicates greater structural stability. We set a target for good stability at a normalized MFE of less than -0.2 kcal/mol per nucleotide. The reward $R_{MFE}(s)$ is then defined as:

$$R_{MFE}(s) = \begin{cases} 1.0 & \text{if } MFE_{norm}(s) < -2 \\ MFE_{norm}(s)/-2 & \text{if } -2 \leq MFE_{norm}(s) \leq 0 \\ 0.0 & \text{if } MFE_{norm}(s) > 0. \end{cases} \quad (3)$$

To ensure computational efficiency, particularly when processing batches of generated sequences, the calculation of these individual reward components is parallelized across available CPU cores using the Python ‘multiprocessing’ library.

3.4 Group Relative Policy Optimization (GRPO)

We employ GRPO [6, 7] to fine-tune our T5 model. GRPO is an on-policy actor-critic algorithm that aims to improve policy updates by considering relative advantages within a group of sampled trajectories.

The GRPO algorithm aims to optimize the following objective, where ϵ and β are hyperparameters:

$$\mathcal{J}_{GRPO}(\theta) = \mathbb{E}[q \sim P(\mathbb{Q}), \{o_i\}_{i=1}^G \sim \pi_{\theta_{old}}(O|q)] \\ \frac{1}{G} \sum_{i=1}^G \frac{1}{|o_i|} \sum_{t=1}^{|o_i|} \left\{ \min \left[\frac{\pi_{\theta}(o_{i,t}|q, o_{i,t})}{\pi_{\theta_{old}}(o_{i,t}|q, o_{i,t})} \hat{A}_{i,t}, \text{clip} \left(\frac{\pi_{\theta}(o_{i,t}|q, o_{i,t})}{\pi_{\theta_{old}}(o_{i,t}|q, o_{i,t})}, 1 - \epsilon, 1 + \epsilon \right) \hat{A}_{i,t} \right] - \beta \mathbb{D}_{KL}[\pi_{\theta} || \pi_{ref}] \right\}$$

The core GRPO algorithm as applied in our work is summarized in Algorithm 1.

Algorithm 1 Iterative Group Relative Policy Optimization

Require: Policy network π_{θ} , reference policy π_{ref} , old policy π_{old} , number of generations per prompt G , KL coefficient β , clipping parameter ϵ , inner loop updates μ .

- 1: Initialize $\pi_{\theta}, \pi_{ref} \leftarrow \pi_{\theta}, \pi_{old} \leftarrow \pi_{\theta}$
 - 2: **for** each training iteration **do**
 - 3: Sample a batch of protein prompts P_i from ‘TRAIN_{DATA}’
 - 4: each protein $p \in P_i$
 - 5: Generate G RNA sequences $S_{p,g}$ with log probabilities from $\pi_{old}(s|p)$
 - 6: Calculate rewards $R(S_{p,g})$ for prefixes of each sequence using the reward function.
 - 7: Calculate advantages $\hat{A}(S_{p,g})$ by normalizing rewards per time-step.
 - 8: **for** $j = 1$ to μ **do**
 - 9: Compute policy ratios $r_t(\theta) = \frac{\pi_{\theta}(a_t|s_t)}{\pi_{old}(a_t|s_t)}$
 - 10: Compute clipped surrogate objective:
 - 11: $L^{CLIP}(\theta) = \hat{\mathbb{E}}_t[\min(r_t(\theta)\hat{A}_t, \text{clip}(r_t(\theta), 1 - \epsilon, 1 + \epsilon)\hat{A}_t)]$
 - 12: Compute KL penalty: $L^{KL}(\theta) = \beta \cdot D_{KL}(\pi_{\theta}(\cdot|s) || \pi_{ref}(\cdot|s))$
 - 13: Update π_{θ} by maximizing $\mathcal{J}_{GRPO}(\theta) = L^{CLIP}(\theta) - L^{KL}(\theta)$
 - 14: $\pi_{old} \leftarrow \pi_{\theta}$
 - 15: **end for**
-

GRPO does not use any value model. Instead, it calculates advantages using rewards of the group as seen in Figure 2. This approach allows GRPO to focus on cooperative behavior by leveraging shared outcomes, rather than relying on value estimates. By using actual group rewards, GRPO encourages policies that directly optimize for collective performance.

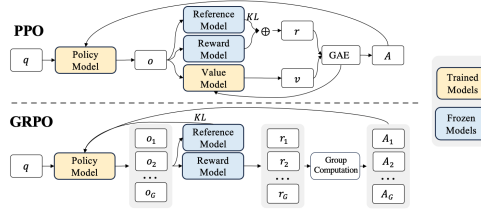


Figure 2: Comparison between GRPO and PPO. GRPO does not use any value model [6].

Implementation Details. The GRPO algorithm was implemented with our T5 model serving as the policy network (π_θ) being trained. A frozen copy of the initial pre-fine-tuned T5 model (before GRPO) was maintained as the reference policy (π_{ref}). An additional copy, the old policy (π_{old}), was used for sampling trajectories at the beginning of each GRPO iteration and was synchronized with π_θ after μ inner loop updates.

During each training iteration, for every protein prompt drawn from the training data, we sample $K = 8$ complete RNA sequences $\{y_{1:T_k}^{(k)}\}_{k=1}^K$ from the current old policy π_{old} . To enable fine-grained credit assignment at each generation step, we compute the total reward (Equation 1) not only on the full sequence but also on every prefix. Concretely, for each sampled sequence $y_{1:T_k}^{(k)}$ and each time step $t \in \{1, \dots, T_k\}$, we form the prefix $y_{1:t}^{(k)}$ and evaluate

$$R_{total}(y_{1:t}^{(k)}).$$

This produces a reward signal at each position t that reflects the quality of the partial sequence up to that token. We then normalize these prefix-level rewards across the batch and time steps (applying a mask to ignore padded positions) to obtain the advantage estimates $\hat{A}_t^{(k)}$ used in the policy update. This prefix-based scheme ensures that each token receives credit proportional to its contribution to both binding affinity and structural criteria.

The policy network π_θ was then updated over μ inner loop steps (with $\mu = 5$, determined by Optuna) to maximize the GRPO objective function. This objective incorporates a clipped surrogate term and a Kullback-Leibler (KL) divergence penalty against the reference policy π_{ref} to ensure training stability. The KL divergence coefficient β was set to approximately 0.0496, and the PPO-style clipping parameter ϵ was set to approximately 0.1736, both values derived from hyperparameter optimization. The AdamW optimizer was employed for policy updates with a learning rate 1×10^{-5} .

Key GRPO hyperparameters, namely β , ϵ , and μ , were systematically optimized using the Optuna framework [25]. All training procedures were executed on NVIDIA TITAN GPUs. Comprehensive logging of the training progress, including various loss components and the individual constituents of our composite reward, was performed using Weights & Biases.

3.5 Experimental Setup

All experiments were conducted under identical conditions to facilitate fair comparison of three models: our GRPO-enhanced T5 model, the base T5 model (fine-tuned prior to GRPO refinement), and the literature benchmark RNAGen [12]. Performance was evaluated on a held-out set of protein prompts. During each iteration of reinforcement learning, we used a batch size of 40 protein prompts and sampled eight RNA sequences per prompt from the policy. Protein inputs were truncated to 1 024 tokens, and RNA outputs were limited to 256 tokens. The GRPO refinement was applied for three epochs over the training data.

To assess model behavior, we computed four metrics on the generated RNA sequences: the mean GC content to evaluate nucleotide composition and its influence on structural stability; the distribution of MFE, to examine predicted secondary-structure stability; the mean composite reward value of each token (Equation 1) to capture the joint contributions of binding affinity and structural criteria; and the mean DeepCLIP score [22], to provide an independent measure of predicted RNA-protein binding strength based on the supervised model.

4 Results

This section presents the key results from our experiments, including the performance of the binding classifier, GRPO training dynamics, hyperparameter optimization, and a comparison of RNA sequences generated by the base model, the baseline model, versus the GRPO-enhanced model.

4.1 Binding Classifier Performance

The protein–RNA binding classifier serves as the foundational reward signal for our RL agent. Trained over 65 epochs, the MLPBindingClassifier exhibited a steady decline in validation loss, eventually converging toward a plateau. This trend indicates that the classifier successfully learned to distinguish between binding and non-binding pairs.

4.2 Optuna Hyperparameter Sweep and GRPO Training with Tuned Parameters

We first conducted a preliminary Optuna hyperparameter sweep over 10 trials to identify promising GRPO configurations. The search explored a KL-divergence coefficient β sampled uniformly from $[0.01, 0.10]$, a clipping parameter ϵ sampled uniformly from $[0.05, 0.20]$, and an integer number of inner-loop updates μ drawn from the range $[3, 8]$.

Figure 3 presents the evolution of the mean composite reward and its standard deviation across these trials, illustrating which parameter sets yielded superior performance. From this analysis, the trial achieving the highest average reward was selected for further evaluation.

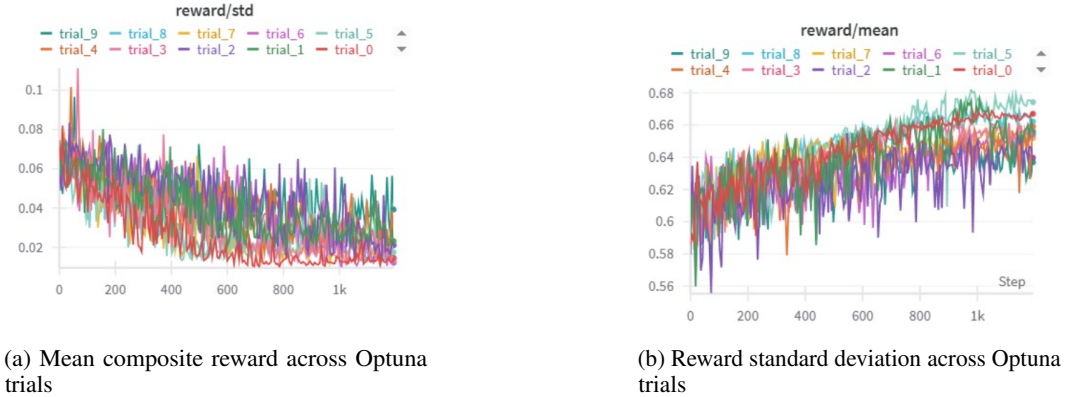


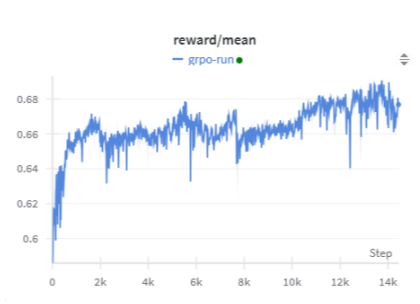
Figure 3: Optuna sweep results over 10 trials: (a) evolution of mean composite reward; (b) evolution of reward standard deviation.

Table 1 lists the optimal hyperparameter values identified by Optuna, which were then applied in an extended GRPO training run.

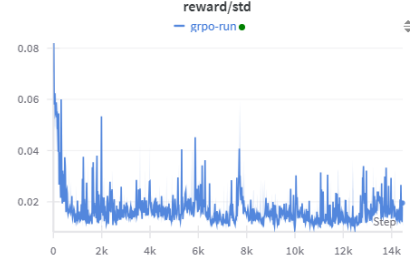
Table 1: Optimal GRPO hyperparameters found by Optuna.

Hyperparameter	Optimal Value
β (KL coefficient)	0.0496
ϵ (clipping)	0.1736
μ (inner-loop updates)	5

Using these tuned hyperparameters, we performed an extended GRPO training run. As shown in Figure 4, the mean composite reward rapidly increased to approximately 0.70 before plateauing, while the reward standard deviation fell below 0.02, indicating consistently high-quality RNA sequence generation under the optimized configuration.



(a) Mean composite reward (tuned hyper-params)



(b) Standard deviation of composite reward (tuned hyperparams)

Figure 4: GRPO training dynamics using Optuna-tuned parameters: (a) mean reward reaches ~ 0.70 within 8 000 steps; (b) reward variability drops below 0.02, demonstrating stable optimization.

4.3 Comparison of Generated RNA Sequences

To evaluate the impact of GRPO fine-tuning on RNA generation, we compared three models on a held-out set of protein prompts: (i) our GRPO-enhanced T5 model, (ii) the base T5 model fine-tuned prior to GRPO refinement (“Base T5 Model”), and (iii) the GAN-based RNAGEN model from the literature [12] (“RNAGEN”). We generated eight RNA sequences per prompt for each model and computed several key metrics, including composite reward, DeepCLIP score, GC content, MFE, and sequence generation loss against ground truth.

Binding Affinity Figure 5 presents the DeepCLIP binding-score distributions for RNA sequences generated by the Base T5 and our GRPO-Tuned T5 models against the RBM5 protein. Consistent with the composite-reward improvements, our tuned model not only shifts the median binding score higher but also yields a tighter distribution of stronger predicted interactions.

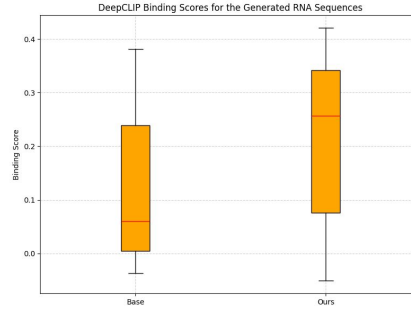


Figure 5: DeepCLIP binding scores for RNA sequences generated by the Base T5 and GRPO-Tuned T5 models evaluated on the RBM5 protein.

Table 2 gives the mean DeepCLIP score for RBM5. Our GRPO-Tuned T5 model achieves a substantially improved DeepCLIP binding score (0.2478 vs. 0.1539), indicating more favorable overall sequence generation and stronger RNA–protein interactions.

Table 2: Performance summary for the Base T5 and GRPO-Tuned T5 models on RBM5 binding.

Metric	Base T5 Model	GRPO-Tuned T5
Mean DeepCLIP Score (RBM5)	0.1539	0.2478

GC Content Analysis The GC content of generated sequences is crucial for RNA stability and function. Figure 6 displays the GC content distributions for the three models. Our Model centers

around a mean of 53.83% GC content with a notably low standard deviation of 2.40%, indicating tight control in the optimal range. The Base Model trends higher, with a mean GC content of 62.16% and greater variability (STD 7.25%), while RNAGEN sequences cluster around a mean of 50.08% with moderate spread (STD 7.03%). These statistics underscore our model’s superior precision in targeting a biologically ideal GC content.

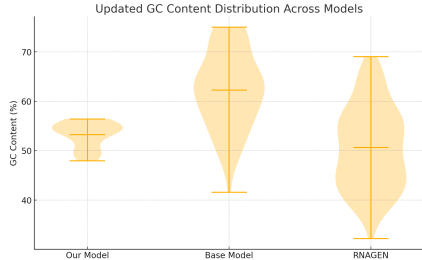
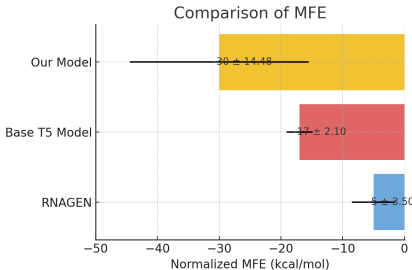


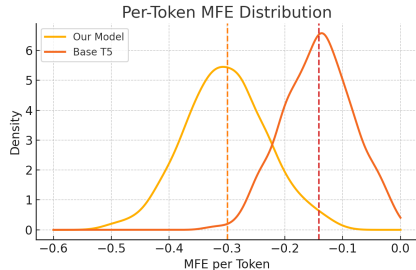
Figure 6: Violin plot of GC content distribution for sequences generated by Our Model, Base Model, and RNAGEN. Our Model exhibits the tightest distribution around the optimal GC range.

MFE Analysis We analyzed both per-token MFE and overall normalized MFE to assess the predicted thermodynamic stability of the generated RNA sequences.

Figure 7b shows the per-token MFE distributions for Our Model and the Base T5 Model. Our Model achieves a mean per-token MFE of -0.2992 (Std = 0.0729), which is substantially more negative than the Base T5 Model’s mean of -0.1411 (Std = 0.0650). This indicates that, on a per-token basis, our model generates sequences predicted to be more stable. The distributions confirm that our tuned model shifts the mean MFE further negative while maintaining a relatively tight spread.



(a) Comparison of mean MFE (kcal/mol) for RNA sequences.



(b) Per-token MFE distributions for sequences generated by Our Model (GRPO-Tuned-T5) and the Base T5 Model.

Figure 7: (a) Mean MFE comparison across models. (b) Per-token MFE distributions. Our model shows a clear shift towards more negative (more stable).

The overall MFE for sequences from all three models is compared in Figure 7a. Our Model achieves the most negative mean MFE (-30 ± 14.4 kcal/mol), significantly outperforming the Base T5 Model (-17 ± 2.1 kcal/mol) and RNAGEN (-5 ± 3.5 kcal/mol). While Our Model has a higher standard deviation for overall MFE compared to the Base T5 model, its mean MFE indicates a strong tendency towards generating more stable structures.

4.4 Comparison of Model Performance

Finally, we present a consolidated comparison of the Base T5 model and Our Model (GRPO-Tuned T5) based on two critical metrics: the **Mean Composite Reward**, which evaluates the overall reward-based objective, and the **Loss Values**, which assess the alignment with the original ground truth sequences.

Table 3: Comparison of Mean Composite Reward and Loss Values.

Metric	Base T5 Model	Our Model (GRPO-Tuned T5)
Mean Composite Reward	0.56	0.69
Loss Values	3.1664 / 1.1859	3.3045 / 1.2283

As shown in Table 3, Our Model achieves a higher mean composite reward, demonstrating improved performance under the reinforcement learning objective. While the loss values of Our Model are slightly higher than those of the Base T5 Model, this increase is expected since GRPO optimization does not explicitly minimize the supervised loss. Nevertheless, the moderate change in loss suggests that the fine-tuning process preserves much of the learned distribution from the original supervised training, while effectively shifting the model toward enhanced RNA design objectives.

In summary, the GRPO fine-tuning process significantly enhances the T5 model’s ability to generate RNA sequences that are not only predicted to bind target proteins more effectively (as shown by DeepCLIP scores) but also possess desirable nucleotide composition (GC content) and thermodynamic stability (MFE), while largely retaining fidelity to the original generation task. The improvements are evident both in mean performance and, for per-token MFE, in the distribution of sequence characteristics.

5 Discussion

Our comparative analysis demonstrates that GRPO fine-tuning, guided by a composite reward function, substantially improves the quality of protein-binding RNA sequences generated by a pre-trained T5 model. As shown in Table 3, the GRPO-enhanced model (Our Model) achieves a higher mean composite reward compared to the Base T5 Model, indicating more successful optimization. This gain is driven by a more balanced and tightly controlled GC content ($53.83\% \pm 2.40\%$ vs. $62.16\% \pm 7.25\%$ for the Base Model and $50.08\% \pm 7.03\%$ for RNAGEN; Figure 6) and by improved thermodynamic stability, as suggested by favorable shifts in MFE characteristics. The distribution of normalized free energy values (Figure 7a) shows that Our Model consistently generates sequences with stronger folding potential. Furthermore, the per-token MFE distribution (Figure 7b) reveals a noticeable shift towards more stable structures relative to the Base Model. Importantly, the analysis of loss values against ground truth indicates that these benefits were achieved without sacrificing the model’s original language generation capabilities. Collectively, these results demonstrate that our reinforcement learning agent is capable of generating RNA sequences that are both structurally robust and compositionally suitable.

The success of this approach hinges on two key components. First, the binding classifier—trained on high-confidence interactions from CLIPDB and shuffled negatives—provides a biologically grounded reward signal. Second, Group Relative Policy Optimization, enhanced with KL-regularization and hyperparameters fine-tuned via Optuna (Table 1), ensures smooth and effective policy learning. The observed rise and stabilization of composite rewards, along with declining reward variance over training epochs (Figure 4), demonstrate GRPO’s capability to guide the policy toward favorable solutions while avoiding destabilizing shifts.

By weighting structural criteria more heavily than binding affinity ($w_{GC} = w_{MFE} = 0.4$ vs. $w_{binding} = 0.2$), our reward design emphasizes foldability and compositional constraints critical for RNA function. The improved GC distribution—both in mean and notably lower deviation for Our Model—and more negative, tighter MFE profiles (Figures 7b and 7a) underscore the value of this multi-objective formulation. In contrast, the Base Model—fine-tuned solely on supervised data—remains biased toward excessively high GC content with substantial variability, while RNA-GEN exhibits poor MFE scores and greater variance in GC content, with many values falling outside the desired range and lacking consistency.

Limitations and Challenges

Despite these encouraging results, our study faces several limitations and challenges. First, the reliance on an *in silico* protein–RNA binding classifier means that any inaccuracies or biases in its outputs can misdirect the RL agent. Without experimental validation of binding affinity, the degree to

which our generated sequences will translate into true biological activity remains uncertain. Second, the reward weights were chosen heuristically; a more systematic approach—such as automated weight optimization or multi-objective RL—could better balance binding versus structural objectives. Third, achieving the right balance between exploration and exploitation during RL proved challenging. While GRPO’s KL-penalty enhances stability, it may also limit the diversity of generated sequences. Finally, the computational cost of fine-tuning a large transformer model with multiple rollouts per update step constrains both the breadth of our experiments and the granularity of hyperparameter searches.

Future Work

Future work will address these limitations by improving the biological fidelity and robustness of our framework. We plan to enhance the binding classifier through integration of experimental binding data and advanced architectures, thereby refining the primary reward signal. We will also explore adaptive reward-weighting schemes and more sophisticated reward shaping strategies to dynamically balance competing objectives. Alternative RL algorithms, such as proximal policy optimization (PPO) or preference-based approaches, may offer further stability or diversity in sequence generation. Crucially, we intend to validate our computationally designed RNA candidates *in vitro* to confirm binding affinities and structural predictions. Finally, incorporating three-dimensional structural information and leveraging co-evolutionary constraints could further improve the functional relevance of generated sequences. Overall, these extensions will advance the integration of deep reinforcement learning into biomolecular design workflows.

6 Conclusion

In this study, we successfully applied GRPO to enhance a T5-based encoder-decoder model for generating RNA sequences that bind to specific protein targets. By developing a novel protein-RNA binding classifier and integrating its predictions into a multi-component reward function considering RNA GC content and Minimum Free Energy, we guided the RL agent towards producing sequences with improved characteristics. Our findings indicate that the GRPO-enhanced model generates RNA sequences with higher predicted binding scores and better structural properties compared to both the initial fine-tuned model and a simpler baseline approach. This research highlights the promise of leveraging advanced RL algorithms for biomolecular design and contributes a novel application of GRPO in bioinformatics. Future work will focus on refining reward mechanisms, classifier accuracy, and experimental validation.

References

- [1] S. Gerstberger, M. Hafner, and T. Tuschl, “A census of human rna-binding proteins,” *Nature reviews Genetics*, vol. 15, no. 12, pp. 829–845, 2014.
- [2] Your Author(s) for Survey Intro Ref 1, “Title of Survey Intro Ref 1,” YYYY. Please replace with actual citation details.
- [3] J.-W. Yang, Á. Pérez-Bercoff, M. Fidea, C. Lagneth, I. L. Hofacker, P. F. Stadler, and E. Rybáková, “CLIPDB: a database of experimentally verified CLIP-seq datasets and annotated RBP binding sites,” *Nucleic acids research*, vol. 50, no. D1, pp. D320–D327, 2022.
- [4] J. Schulman, F. Wolski, P. Dhariwal, A. Radford, and O. Klimov, “Proximal Policy Optimization Algorithms,” 2017.
- [5] R. Rafailov, A. Sharma, E. Mitchell, S. Ermon, C. D. Manning, and C. Finn, “Direct Preference Optimization: Your Language Model is Secretly a Reward Model,” 2023.
- [6] Z. Shao, P. Wang, Q. Zhu, R. Xu, J. Song, X. Bi, H. Zhang, M. Zhang, Y. K. Li, Y. Wu, and D. Guo, “DeepSeekMath: Pushing the limits of mathematical reasoning in open language models,” Feb 2024.
- [7] D. Guo, Q. Zhu, D. Yang, Z. Xie, K. Dong, W. Zhang, G. Chen, X. Bi, Y. Wu, Y. K. Li, F. Luo, Y. Xiong, and W. Liang, “Deepseek-coder: When the large language model meets programming – the rise of code intelligence,” 2024.

- [8] A. Vaswani, N. Shazeer, N. Parmar, J. Uszkoreit, L. Jones, A. N. Gomez, L. Kaiser, and I. Polosukhin, "Attention is all you need," 2023.
- [9] A. Elnaggar, M. Heinzinger, C. Dallago, G. Rihawi, Y. Wang, L. Jones, T. Gibbs, T. Feher, C. Angerer, M. Steinegger, D. Bhowmik, and B. Rost, "Prottrans: Towards cracking the language of life's code through self-supervised deep learning and high performance computing," 2021.
- [10] N. Ferruz, S. Schmidt, and B. Höcker, "ProtGPT2 is a deep unsupervised language model for protein design," *Nature Communications*, vol. 13, p. 4348, 2022.
- [11] Y. Zhao, K. Oono, H. Takizawa, and M. Kotera, "GenerRNA: A generative pre-trained language model for de novo RNA design," *PLOS ONE*, vol. 19, no. 10.
- [12] F. Ozden, S. Barazandeh, D. Akboga, S. S. S. Tabrizi, U. O. S. Seker, and A. E. Cicek, "RNAGEN: A generative adversarial network-based model to generate synthetic RNA sequences to target proteins," *bioRxiv*, Jul 2023. Preprint.
- [13] S. S. S. Tabrizi, S. Barazandeh, H. H. Aghdam, and A. E. Çiçek, "RNAtranslator: Modeling protein-conditional RNA design as sequence-to-sequence natural language translation," *bioRxiv*, 2025.
- [14] P. Eastman, J. Shi, B. Ramsundar, and V. S. Pande, "Solving the RNA design problem with reinforcement learning," *PLoS Computational Biology*, vol. 14, p. e1006176, Jun 2018.
- [15] F. Runge, D. Stoll, S. Falkner, and F. Hutter, "Learning to design rna," 2019.
- [16] I. Kauvar, E. B. Richman, and W. E. Allen, "Eterna-rl : Using reinforcement learning to design rna secondary structures," 2017.
- [17] A. Whatley, Z. Luo, and X. Tang, "Improving RNA Secondary Structure Design using Deep Reinforcement Learning," Nov 2021.
- [18] Y. Li, X. Pan, H. Shen, and Y. Yang, "Drag: design rnas as hierarchical graphs with reinforcement learning," *Briefings in bioinformatics*, vol. 26, 03 2025.
- [19] M. Sternke and J. Karpiak, "ProteinRL: Reinforcement learning with generative protein language models for property-directed sequence design," in *NeurIPS 2023 Generative AI and Biology (GenBio) Workshop*, 2023.
- [20] F. Stocco, M. Artigues-Lleixa, A. Hunklinger, T. Widatalla, M. Guell, and N. Ferruz, "Guiding generative protein language models with reinforcement learning," 12 2024.
- [21] C. Raffel, N. Shazeer, A. Roberts, K. Lee, S. Narang, M. Matena, Y. Zhou, W. Li, and P. J. Liu, "Exploring the limits of transfer learning with a unified text-to-text transformer," *The Journal of Machine Learning Research*, vol. 21, no. 1, pp. 5485–5551, 2020.
- [22] A. Grønning, T. K. Doktor, S. Larsen, U. Petersen, L. L. Holm, G. Bruun, M. B. Hansen, A.-M. Hartung, J. Baumbach, and B. S. Andresen, "Deepclip: predicting the effect of mutations on protein–rna binding with deep learning," *Nucleic Acids Research*, vol. 48, pp. 7099–7118, 06 2020.
- [23] D. Risso, K. Schwartz, G. Sherlock, and S. Dudoit, "Gc-content normalization for rna-seq data," *BMC Bioinformatics*, vol. 12, no. 1, p. 480, 2011.
- [24] R. Lorenz, S. H. Bernhart, C. Höner Zu Siederdissen, H. Tafer, C. Flamm, P. F. Stadler, and I. L. Hofacker, "ViennaRNA Package 2.0," *Algorithms for Molecular Biology*, vol. 6, no. 1, p. 26, 2011.
- [25] T. Akiba, S. Sano, T. Yanase, T. Ohta, and M. Koyama, "Optuna: A next-generation hyperparameter optimization framework," in *Proceedings of the 25th ACM SIGKDD International Conference on Knowledge Discovery & Data Mining*, pp. 2623–2631, 2019.

A Appendix A: Member Contributions

Each group member’s contributions to the project are detailed below:

- **Yusuf Kesmen:** Led the implementation and adaptation of the GRPO algorithm for the T5 model. Conducted hyperparameter tuning using Optuna and analyzed RL training dynamics. Primary contributor to the technical sections of the manuscript (Methods, Results) and overall technical lead.
- **Alpsencer Özdemir:** Developed and trained the protein-RNA binding classifier, including data preparation (positive/negative samples, embedding generation). Designed and implemented the individual components of the reward function (GC content, MFE). Developed and evaluated the baseline model. Contributed to the Introduction and Related Work sections of the manuscript.
- **Yavuz Alp Sencer Öztürk:** Managed the initial fine-tuning of the base T5 model and the creation/management of custom BPE tokenizers. Integrated various model components and set up the computational environment and data pipelines. Coordinated the report and presentation preparation, primary contributor to the Discussion, Conclusion, and Appendix sections of the manuscript.

All members contributed to the conceptualization of the project, literature review, debugging, and final report preparation.

B Appendix B: Code Overview

*

File: inference.py

Listing 1: Code Snippet

```
1 import os
2 import torch
3 import RNA
4 import numpy as np
5 import matplotlib.pyplot as plt
6 from tqdm import tqdm
7 from transformers import T5ForConditionalGeneration
8 from custom_tokenizer import get_tokenizer
9
10 # Paths to your three checkpoints and their short names
11 CKPT_DIRS = [
12     #
13     # "/data6/alpsencer/reinforce_rna/grpo-project/results/epoch_01_step_000050",
14     # "/data6/alpsencer/reinforce_rna/grpo-project/results/epoch_03_step_000150",
15     # "/data6/sobhan/rllm/results/train/t5/run3_20240822-152114/checkpoints/checkpoint -3
16     # "/data6/alpsencer/reinforce_rna/grpo-project/results/epoch_03_step_000900"
17 ]
18 MODEL_NAMES = ["step900"]
19
20 # Tokenizer files
21 PROTEIN_TOKENIZER_PATH =
22     "/data6/alpsencer/tokenizers/bpe_protein_1000_1024.json"
23 RNA_TOKENIZER_PATH =
24     "/data6/alpsencer/tokenizers/bpe_rna_1000_1024.json"
25
26 # Validation data
27 VALIDATION_PATH = "/data6/alpsencer/reinforce_rna/validation_rl.txt"
```

```

26 #           Generation parameters
27 VOCAB_SIZE      = 1000
28 SEQ_SIZE        = 1024
29 MAX_PROMPT_LENGTH = 1024
30 MAX_COMPLETION_LENGTH = 1024
31 NUM_RETURN_SEQUENCES = 1
32 TOP_P           = 1.0
33 TEMPERATURE     = 0.7
34 BEAM_SIZE       = 1
35
36 #           Helper functions
37 def map_to_bjuz(rna_acgu_sequence):
38     return (rna_acgu_sequence.replace("A", "b")
39             .replace("C", "j")
40             .replace("U", "u")
41             .replace("G", "z"))
42
43 def map_to_acgu(rna_bjuz_sequence):
44     s = rna_bjuz_sequence.replace(" ", "")
45     return (s.replace("b", "A")
46             .replace("j", "C")
47             .replace("u", "U")
48             .replace("z", "G"))
49
50 def load_validation_data(filepath):
51     proteins, rnas = [], []
52     with open(filepath) as f:
53         for line in f:
54             line = line.strip()
55             if not line: continue
56             prot, rna = line.split(",", 1)
57             proteins.append(prot)
58             rnas.append(rna)
59     return proteins, rnas
60
61 def calculate_gc_content(seq):
62     seq = seq.upper()
63     gc = seq.count("G") + seq.count("C")
64     return (gc / len(seq)) * 100 if seq else 0
65
66 def calculate_mfe(seq):
67     try:
68         _, mfe = RNA.fold(seq)
69         return mfe
70     except:
71         return None
72
73 def calculate_mfe_per_token(seq):
74     """Calculate MFE normalized by sequence length"""
75     mfe = calculate_mfe(seq)
76     if mfe is not None and seq:
77         return mfe / len(seq)
78     return None
79
80 #           Main evaluation loop
81 def main():
82     os.makedirs("analysis_outputs_2", exist_ok=True)
83     proteins, ground_truths = load_validation_data(VALIDATION_PATH)
84
85     results = {}
86     for ckpt, name in zip(CKPT_DIRS, MODEL_NAMES):
87         print(f"\nLoading model '{name}' from {ckpt}")
88         device = "cuda" if torch.cuda.is_available() else "cpu"
89         model =
            T5ForConditionalGeneration.from_pretrained(ckpt).to(device)

```

```

90     model.eval()
91
92     # tokenizers
93     tok_p = get_tokenizer("bpe", VOCAB_SIZE, SEQ_SIZE,
94                           tokenizer_path=PROTEIN_TOKENIZER_PATH)
95     tok_r = get_tokenizer("bpe", VOCAB_SIZE, SEQ_SIZE,
96                           tokenizer_path=RNA_TOKENIZER_PATH)
97     tok_p.pad_token_id = tok_r.pad_token_id = 0
98
99     gc_vals, mfe_vals, loss_vals = [], [], []
100     mfe_per_token_vals = [] # Add list for MFE per token values
101
102     for prot, gt_acgu in tqdm(zip(proteins[:10],
103                                  ground_truths[:10]),
104                               total=min(10, len(proteins)),
105                               desc=name):
106         # --- tokenize protein input ---
107         enc = tok_p.tokenize(prot)
108         ids = enc.ids[:MAX_PROMPT_LENGTH]
109         mask = enc.attention_mask[:MAX_PROMPT_LENGTH]
110         pad = MAX_PROMPT_LENGTH - len(ids)
111         ids += [tok_p.pad_token_id]*pad
112         mask += [0]*pad
113         inp = {
114             "input_ids": torch.tensor([ids], dtype=torch.long).to(device),
115             "attention_mask": torch.tensor([mask], dtype=torch.long).to(device)
116         }
117
118         # --- generate RNA ---
119         gen_kwargs = {
120             "max_new_tokens": MAX_COMPLETION_LENGTH,
121             "num_return_sequences": NUM_RETURN_SEQUENCES,
122             "do_sample": BEAM_SIZE==1,
123             "top_p": TOP_P,
124             "temperature": TEMPERATURE,
125             "num_beams": BEAM_SIZE,
126             "early_stopping": BEAM_SIZE>1
127         }
128         with torch.no_grad():
129             out_ids = model.generate(**inp, **gen_kwargs)
130             bjuz = tok_r.decode(out_ids)[0] if
131                 isinstance(tok_r.decode(out_ids), list) else
132                 tok_r.decode(out_ids)
133             gen_acgu = map_to_acgu(bjuz)
134
135         # --- GC & MFE ---
136         gc = calculate_gc_content(gen_acgu)
137         mfe = calculate_mfe(gen_acgu)
138         mfe_per_token = calculate_mfe_per_token(gen_acgu) #
139             Calculate MFE per token
140         gc_vals.append(gc)
141         mfe_vals.append(mfe if mfe is not None else np.nan)
142         mfe_per_token_vals.append(mfe_per_token if mfe_per_token
143             is not None else np.nan) # Add MFE per token
144
145         # --- loss on ground truth ---
146         gt_bjuz = map_to_bjuz(gt_acgu)
147         lbl = tok_r.tokenize(gt_bjuz).ids[:MAX_COMPLETION_LENGTH]
148         pad2 = MAX_COMPLETION_LENGTH - len(lbl)
149         lbl += [tok_r.pad_token_id]*pad2
150         labels = torch.tensor([lbl], dtype=torch.long).to(device)
151         labels[labels==tok_r.pad_token_id] = -100
152         with torch.no_grad():

```

```

148         loss = model(**inp, labels=labels).loss.item()
149         loss_vals.append(loss)
150
151     # store
152     results[name] = {
153         "gc": gc_vals,
154         "mfe": mfe_vals,
155         "mfe_per_token": mfe_per_token_vals, # Add MFE per token
156         "loss": loss_vals
157     }
158
159     # save raw data lists to txt
160     np.savetxt(f"analysis_outputs_2/gc_values_{name}.txt",
161               gc_vals, fmt="%.4f")
162     np.savetxt(f"analysis_outputs_2/mfe_values_{name}.txt",
163               mfe_vals, fmt="%.4f")
164     np.savetxt(f"analysis_outputs_2/mfe_per_token_values_{name}.txt",
165               mfe_per_token_vals, fmt="%.6f") # Save MFE per token
166     values
167     np.savetxt(f"analysis_outputs_2/loss_values_{name}.txt", loss_vals, fmt="%.6f")
168
169     # summary table
170     with open("analysis_outputs_2/summary_metrics.txt", "w") as f:
171         f.write("Model\tGC_mean\tGC_std\tMFE_mean\tMFE_std\tMFE_per_token_mean\tMFE_per_token_std\n")
172         for name, m in results.items():
173             g, mf, mfpt, lo = np.array(m["gc"]), np.array(m["mfe"]),
174                                     np.array(m["mfe_per_token"]), np.array(m["loss"])
175             f.write(f"{name}\t"
176                   f"{g.mean():.2f}\t{g.std():.2f}\t"
177                   f"{mf.mean():.2f}\t{mf.std():.2f}\t"
178                   f"{mfpt.mean():.6f}\t{mfpt.std():.6f}\t"
179                   f"{lo.mean():.4f}\t{lo.std():.4f}\n")
180
181     # plots
182     # 1) Violin plot of GC content
183     plt.figure(figsize=(8,5))
184     parts = plt.violinplot(
185         [results[n]["gc"] for n in MODEL_NAMES],
186         showmeans=True, showmedians=True)
187     plt.xticks([1,2,3], MODEL_NAMES)
188     plt.ylabel("GC content (%)")
189     plt.title("GC content by model")
190     plt.grid(alpha=0.3)
191     plt.savefig("analysis_outputs_2/gc_violin.png")
192     plt.close()
193
194     # 2) MFE distribution
195     plt.figure(figsize=(8,5))
196     for name in MODEL_NAMES:
197         plt.hist(results[name]["mfe"], bins=30, alpha=0.5,
198                 density=True, label=name)
199     plt.xlabel("MFE (kcal/mol)")
200     plt.ylabel("Density")
201     plt.legend()
202     plt.title("MFE distributions")
203     plt.grid(alpha=0.3)
204     plt.savefig("analysis_outputs_2/mfe_dist.png")
205     plt.close()
206
207     # 3) Loss distribution
208     plt.figure(figsize=(8,5))
209     for name in MODEL_NAMES:
210         plt.hist(results[name]["loss"], bins=30, alpha=0.5,
211                 density=True, label=name)

```

```

207     plt.xlabel("Loss (NLL)")
208     plt.ylabel("Density")
209     plt.legend()
210     plt.title("Ground truth loss distributions")
211     plt.grid(alpha=0.3)
212     plt.savefig("analysis_outputs_2/loss_dist.png")
213     plt.close()
214
215     # 4) MFE per token distribution
216     plt.figure(figsize=(8,5))
217     for name in MODEL_NAMES:
218         plt.hist(results[name]["mfe_per_token"], bins=30, alpha=0.5,
219                 density=True, label=name)
220     plt.xlabel("MFE per token (kcal/mol per nucleotide)")
221     plt.ylabel("Density")
222     plt.legend()
223     plt.title("MFE per token distributions")
224     plt.grid(alpha=0.3)
225     plt.savefig("analysis_outputs_2/mfe_per_token_dist.png")
226     plt.close()
227
228     print("All metrics, plots and data have been saved under
229           ./analysis_outputs_2/")
230
231 if __name__ == "__main__":
232     main()

```

*

Directory: classifier/

File: classifier/binding_embedding_creator.py

Listing 2: Code Snippet

```

1  """
2  Embedding generator for protein and RNAs
3  Last modified: 13.04.2025
4  Modified by: Alpsencer & Yusuf
5  TODO: organize the code if necessary
6  """
7
8  import torch
9  from transformers import T5ForConditionalGeneration
10 from pathlib import Path
11 from tqdm import tqdm
12 from util.tokenizer import get_tokenizer # Update the path
13 import numpy as np # Add numpy for random split
14
15 from datasets import load_dataset
16 from copy import deepcopy
17 import pickle # Import pickle for incremental saving
18
19 MODEL_PATH =
20     "/data6/sobhan/rllm/results/train/t5/run3_20240822-152114/checkpoints/checkpoint-34980"
21 SOURCE_TOKENIZER =
22     "/data6/alpsencer/tokenizers/bpe_protein_1000_1024.json"
23 RNA_TOKENIZER = "/data6/alpsencer/tokenizers/bpe_rna_1000_1024.json"
24 MAX_LEN = 32
25 # TRAIN_DATA = "/data6/sobhan/RLLM/dataset/rph/train_rp.txt"
26 TRAIN_DATA = "/data6/alpsencer/reinforce_rna/sample_100K.txt"
27 EVAL_DATA = "/data6/sobhan/RLLM/dataset/rph/eval_rp.txt"
28 PAIR_SAVE_DIR = "/data6/alpsencer/data/rllm_classifier_pairs"

```

```

27 DATASET_NAME = "pairs_100K_eos"
28 BATCH_SIZE = 16
29 TRAIN_OUTPUT_FILENAME = f"{DATASET_NAME}_train.pkl" # Define output
    filename with .pkl extension
30 TEST_OUTPUT_FILENAME = f"{DATASET_NAME}_test.pkl" # Define output
    filename with .pkl extension
31 TRAIN_TEST_SPLIT = 0.95
32
33 def tokenize_dataset(sample, source_tokenizer, rna_tokenizer):
34     text = sample["text"]
35     # Basic error handling for malformed lines
36     try:
37         source, rna, pair = text.strip().split("$")
38     except ValueError:
39         # Return None or handle appropriately if lines might be
            malformed
40         # For now, assume data is clean as per original code's lack
            of handling
41         # If errors occur, add more robust handling here.
42         # Example: return None and filter later
43         print(f"Warning: Skipping malformed line: {text.strip()}") #
            Keep minimal warning if needed, or remove
44         return None
45
46
47     rna = rna.replace("A", "b").replace("C", "j").replace("U",
        "u").replace("G", "z")
48     source = source.lower()
49
50     source_tokenized = source_tokenizer.tokenize(source)
51     rna_tokenized = rna_tokenizer.tokenize(rna)
52
53     # replace the first padding token with eos token
54
55
56     # need to set these to -100 to calculate the loss properly
57     # rna_labels = [-100 if i == 0 else i for i in rna_tokenized.ids]
58
59     # return {
60     #     "input_ids": source_tokenized.ids,
61     #     "attention_mask": source_tokenized.attention_mask,
62     #     "labels": rna_labels,
63     # }
64     return {
65         "source_tokenized": source_tokenized,
66         "rna_tokenized": rna_tokenized,
67         "label": 1 if pair == "+" else 0,
68     }
69
70
71 def get_datasets(args, source_tokenizer, rna_tokenizer,
    iterable=True):
72
73     # Get the dataset
74     dataset = load_dataset(
75         "text",
76         data_files=args.train_data,
77         split="train",
78         cache_dir="/data6/alpsencer/cache",
79     )
80
81     if iterable:
82         train_dataset = dataset.to_iterable_dataset()
83     else:
84         train_dataset = dataset

```

```

85
86     # # Filter dataset
87     filter_source_tokenizer = deepcopy(source_tokenizer)
88     filter_source_tokenizer.tokenizer.no_truncation()
89     filter_rna_tokenizer = deepcopy(rna_tokenizer)
90     filter_rna_tokenizer.tokenizer.no_truncation()
91
92     def _is_valid_sample(sample):
93         # Check for malformed lines first if tokenize_dataset can
94         # return None
95         if sample is None or "text" not in sample:
96             return False
97         try:
98             source, rna, _ = sample["text"].strip().split("$")
99             return (
100                 len(filter_source_tokenizer.tokenize(source).ids) <=
101                 1024
102                 and len(filter_rna_tokenizer.tokenize(rna).ids) <=
103                 1024
104             )
105         except ValueError:
106             return False # Filter out malformed lines
107
108     filtered = train_dataset.filter(_is_valid_sample)
109
110     # # # Shuffle dataset
111     if iterable:
112         shuffled = filtered.shuffle(buffer_size=10000)
113     else:
114         shuffled = filtered.shuffle()
115
116     # Tokenize dataset
117     tokenized = shuffled.map(
118         lambda sample: tokenize_dataset(sample, source_tokenizer,
119                                         rna_tokenizer)
120     )
121     # Add filtering step if tokenize_dataset returns None for errors
122     tokenized = tokenized.filter(lambda x: x is not None)
123
124     return tokenized
125
126
127 def count_file_lines(file_path):
128     try:
129         with open(file_path, "r", encoding="utf-8", errors="ignore")
130             as file:
131                 return sum(1 for _ in file)
132     except FileNotFoundError:
133         return 0 # Return 0 if file not found, avoids crash
134
135
136 # Import The encoder-decoder model
137 device = "cuda" if torch.cuda.is_available() else "cpu"
138 model = T5ForConditionalGeneration.from_pretrained(MODEL_PATH) #
139     Update the model path
140 model.to(device)
141 model.eval() # Set model to evaluation mode
142
143 # Import the tokenizers
144 source_tokenizer = get_tokenizer(
145     tokenizer_name="bpe",
146     vocab_size=1000,
147     seq_size=1024,
148     tokenizer_path=SOURCE_TOKENIZER,

```

```

144 )
145 rna_tokenizer = get_tokenizer(
146     tokenizer_name="bpe", vocab_size=1000, seq_size=1024,
147     tokenizer_path=RNA_TOKENIZER
148 )
149
150 # Import the dataset, tokenize the data
151 class Args:
152     def __init__(self, train_data, eval_data):
153         self.train_data = train_data
154         self.eval_data = eval_data
155
156
157 dataset_args = Args(TRAIN_DATA, EVAL_DATA)
158 train_dataset = get_datasets(
159     dataset_args, source_tokenizer=source_tokenizer,
160     rna_tokenizer=rna_tokenizer
161 )
162
163 """
164 return {
165     "input_ids": source_tokenized.ids,
166     "attention_mask": source_tokenized.attention_mask,
167     "labels": rna_labels,
168 }
169 """
170
171 print(model.config) # Original print statement
172
173 save_dir = Path(PAIR_SAVE_DIR)
174 # --- Ensure save directory exists ---
175 save_dir.mkdir(parents=True, exist_ok=True)
176 train_output_path = save_dir / TRAIN_OUTPUT_FILENAME
177 test_output_path = save_dir / TEST_OUTPUT_FILENAME
178
179 # Generate train/test indices before processing
180 total_samples = count_file_lines(TRAIN_DATA)
181 indices = np.arange(total_samples)
182 np.random.seed(42) # Set seed for reproducibility
183 np.random.shuffle(indices)
184 train_size = int(TRAIN_TEST_SPLIT * total_samples)
185 train_indices = set(indices[:train_size])
186
187 # pair_list = [] # Removed: No longer accumulating in memory
188 print(type(train_dataset)) # Original print statement
189 batched_dataset = train_dataset.iter(batch_size=BATCH_SIZE)
190
191 file_length = count_file_lines(TRAIN_DATA)
192 # Use file_length for tqdm if reliable, otherwise might need
193 # adjustment or removal
194 tqdm_total = (file_length // BATCH_SIZE + 1) if file_length > 0 else
195 None
196
197 current_index = 0 # Keep track of global index
198
199 # --- Open output file before the loop ---
200 with open(train_output_path, "wb") as train_outfile:
201     with open(test_output_path, "wb") as test_outfile:
202         # Freeze the model gradients and disable dropout etc.
203         with torch.inference_mode():
204             for tokenized_batch in tqdm(batched_dataset,
205                 total=tqdm_total): # Use calculated total if available
206
207                 protein_ids = [

```

```

204         tokenized.ids for tokenized in
                tokenized_batch["source_tokenized"]
205     ]
206     protein_attention = [
207         tokenized.attention_mask
208         for tokenized in
                tokenized_batch["source_tokenized"]
209     ]
210
211     rna_ids = [tokenized.ids for tokenized in
                tokenized_batch["rna_tokenized"]]
212     rna_attention = [
213         tokenized.attention_mask for tokenized in
                tokenized_batch["rna_tokenized"]
214     ]
215
216     # --- Convert lists to tensors ---
217     protein_ids_tensor = torch.tensor(protein_ids,
                dtype=torch.long).to(device)
218     protein_attention_tensor =
                torch.tensor(protein_attention,
                dtype=torch.long).to(device)
219     rna_ids_tensor = torch.tensor(rna_ids,
                dtype=torch.long).to(device)
220     rna_attention_tensor = torch.tensor(rna_attention,
                dtype=torch.long).to(device)
221
222     # Extract protein & rna embeddings
223
224     # Forward pass through the encoder
225     encoder_outputs = model.encoder(
226         input_ids=protein_ids_tensor,
227         attention_mask=protein_attention_tensor,
228     )
229     protein_hidden_states =
                encoder_outputs.last_hidden_state # Encoder
                embeddings
230
231     # Forward pass through the decoder
232     decoder_outputs = model.decoder(
233         input_ids=rna_ids_tensor,
234         attention_mask=rna_attention_tensor,
235         encoder_hidden_states=protein_hidden_states,
236         encoder_attention_mask=protein_attention_tensor,
                # Pass encoder mask
237         output_hidden_states=True,
238     )
239     rna_hidden_states = decoder_outputs.last_hidden_state
240
241     # Protein embedding (Mean Pooling)
242     # Use attention mask directly for masking and length
        calculation
243     protein_masks = protein_attention_tensor.to(device) #
        Already on device
244     protein_hidden_masked = protein_hidden_states *
                protein_masks.unsqueeze(-1)
245     protein_masked_sum = torch.sum(protein_hidden_masked,
                dim=1)
246     protein_sequence_lengths = torch.sum(protein_masks,
                dim=1, keepdim=True).clamp(min=1) # Avoid div by
                zero
247     protein_embeddings = protein_masked_sum /
                protein_sequence_lengths
248
249

```

```

250     # RNA embedding (Last non-padding token / EOS-like)
251     # Find the index of the last non-padding token using
        the attention mask
252     rna_sequence_lengths_indices =
        rna_attention_tensor.sum(dim=1) - 1 # Get last
        index (0-based)
253     rna_sequence_lengths_indices =
        rna_sequence_lengths_indices.clamp(min=0) #
        Ensure index is not negative
254
255     # Gather the hidden state at the last token index for
        each sequence in the batch
256     rna_embeddings = rna_hidden_states[
257         torch.arange(rna_hidden_states.size(0),
            device=device), rna_sequence_lengths_indices
258     ]
259
260     # Average embedding
261     #rna_masks = torch.stack(
262     #     [
263     #         (torch.tensor(tokenized.ids) != 0)
264     #         for tokenized in
265     #             tokenized_batch["rna_tokenized"]
266     #     ]
267     # ).to(device)
268     #rna_hidden_masked = rna_hidden_states *
269     #     rna_masks.unsqueeze(dim=2)
270     #print(protein_hidden_masked[0])
271     #rna_masked_sum = torch.sum(rna_hidden_masked,
272     #     dim=1).to(device)
273     #rna_sequence_lengths = torch.sum(rna_masks,
274     #     dim=1).to(device)
275     #rna_embeddings = rna_masked_sum /
276     #     rna_sequence_lengths.unsqueeze(dim=1)
277
278     # --- Create list of pairs *for this batch* and write
279     # incrementally ---
280     current_batch_labels = tokenized_batch["label"]
281     # Move tensors to CPU before pickling
282     protein_embeddings_cpu = protein_embeddings.cpu()
283     rna_embeddings_cpu = rna_embeddings.cpu()
284
285     # Process each sample in the batch
286     for i in range(len(current_batch_labels)):
287         pair_data = {
288             "protein": protein_embeddings_cpu[i],
289             "rna": rna_embeddings_cpu[i],
290             "label": current_batch_labels[i],
291         }
292         # Write to appropriate file based on pre-computed
293         # split
294         if current_index in train_indices:
295             pickle.dump(pair_data, train_outfile)
296         else:
297             pickle.dump(pair_data, test_outfile)
298         current_index += 1

```

File: classifier/classifier.py

Listing 3: Code Snippet

```

1 """
2 Binding classifiers for protein and RNA embeddings

```

```

3 Last modified: 24.02.2025
4 Modified by: Alpsencer & Yusuf
5 """
6
7 import torch
8 import torch.nn as nn
9 import torch.nn.init as init
10
11
12 # Define the MLP Classifier
13 # TODO: hyperparameters from yaml
14 # TODO: regressor or classifier
15 class MLPBindingClassifier(nn.Module):
16     def __init__(self, protein_dim, rna_dim, hidden_dims=[512, 256,
17         128], dropout_rate=0.2):
18         super(MLPBindingClassifier, self).__init__()
19
20         # self.protein_mean = protein_mean
21         # self.protein_std = protein_std
22
23         # self.rna_mean = rna_mean
24         # self.rna_std = rna_std
25
26         # Input dimension is the concatenated protein and RNA
27         # embeddings
28         input_dim = protein_dim + rna_dim
29
30         # Build the MLP layers
31         layers = []
32         # layers.append(nn.BatchNorm1d(num_features=input_dim))
33         prev_dim = input_dim
34
35         for hidden_dim in hidden_dims:
36             layers.append(nn.Linear(prev_dim, hidden_dim))
37             layers.append(nn.ReLU())
38             layers.append(nn.Dropout(p=dropout_rate))
39             prev_dim = hidden_dim
40
41         # Output layer (binary classification)
42         layers.append(nn.Linear(prev_dim, 1))
43         layers.append(nn.Sigmoid())
44
45         self.model = nn.Sequential(*layers)
46
47     def forward(self, protein_embedding, rna_embedding):
48         # Concatenate the embeddings
49
50         # protein_normalized = ( protein_embedding -
51             self.protein_mean ) / self.protein_std
52
53         # rna_normalized = ( rna_embedding - self.rna_mean ) /
54             self.rna_std
55
56         combined = torch.cat((protein_embedding, rna_embedding),
57             dim=1)
58
59         return self.model(combined)
60
61 class ContrastiveBindingClassifier(nn.Module):
62     pass

```

*

File: classifier/embedding_loader.py

Listing 4: Code Snippet

```
1 import torch
2 from torch.utils.data import Dataset, DataLoader
3 import torch.nn.functional as F
4 import pickle
5
6
7 class ProteinRNAEmbeddingDataset(Dataset):
8     def __init__(self, dataset_path):
9
10         # self.data = torch.load(dataset_path)
11
12         # Load the dataset with pickle
13         self.data = []
14         with open(dataset_path, "rb") as f:
15             while True:
16                 try:
17                     self.data.append(pickle.load(f))
18                 except EOFError:
19                     break
20
21         proteins = torch.stack([x["protein"] for x in self.data])
22         rnas = torch.stack([x["rna"] for x in self.data])
23         labels = torch.tensor([x["label"] for x in self.data])
24
25         self.proteins_normalized = F.normalize(proteins, p=2, dim=-1)
26         self.rnas_normalized = F.normalize(rnas, p=2, dim=-1)
27         self.labels = labels
28
29     def __len__(self):
30         return len(self.data)
31
32     def __getitem__(self, idx):
33         item = self.data[idx]
34         # protein = torch.tensor(item["protein"], dtype=torch.float32)
35         # rna = torch.tensor(item["rna"], dtype=torch.float32)
36         protein =
37             item["protein"].clone().detach().requires_grad_(True).float()
38         rna =
39             item["rna"].clone().detach().requires_grad_(True).float()
40         label = torch.tensor(item["label"], dtype=torch.long)
41
42         return protein, rna, label
```

*

File: classifier/train_classifier.py

Listing 5: Code Snippet

```
1 """
2 Train a classifier for binding score of protein and rna embeddings
3 Last modified: 13.04.2025
4 Modified by: Alpsencer & Yusuf
5 """
6
7 import torch
8 from torch.utils.data import Dataset, DataLoader
9 from torch.optim import Adam
10 from torch.nn import BCELoss
```

```

11 from torch.optim.lr_scheduler import LinearLR
12 from tqdm import tqdm
13 import wandb
14
15 from classifier import MLPBindingClassifier
16 from embedding_loader import ProteinRNAEmbeddingDataset
17
18 # Constants
19 TRAIN_DATA =
20     "/data6/alpsencer/data/rllm_classifier_pairs/pairs_100K_eos_train.pkl"
21 VALIDATION_DATA =
22     "/data6/alpsencer/data/rllm_classifier_pairs/pairs_100K_eos_test.pkl"
23 BATCH_SIZE = 64
24 LEARNING_RATE = 0.00001
25 WEIGHT_DECAY = 0.01 # L2 regularization parameter
26 DROPOUT_RATE = 0.0 # Dropout rate for regularization
27 NUM_EPOCHS = 1000
28 VALIDATION_EPOCHS = 5
29 WARMUP_EPOCHS = 10 # Number of epochs for warmup
30 WARMUP_START_FACTOR = 0.1 # Starting learning rate factor
31
32 # Initialize wandb
33 wandb.init(
34     project="rllm-binding-classifier",
35     config={
36         "learning_rate": LEARNING_RATE,
37         "weight_decay": WEIGHT_DECAY,
38         "dropout_rate": DROPOUT_RATE,
39         "batch_size": BATCH_SIZE,
40         "epochs": NUM_EPOCHS,
41         "validation_epochs": VALIDATION_EPOCHS,
42         "warmup_epochs": WARMUP_EPOCHS,
43         "warmup_start_factor": WARMUP_START_FACTOR,
44     }
45 )
46
47 # Load dataset and create DataLoader
48 train_dataset = ProteinRNAEmbeddingDataset(TRAIN_DATA)
49 validation_dataset = ProteinRNAEmbeddingDataset(VALIDATION_DATA)
50
51 train_dataloader = DataLoader(train_dataset, batch_size=BATCH_SIZE,
52     shuffle=True, drop_last=True)
53 validation_dataloader = DataLoader(validation_dataset,
54     batch_size=BATCH_SIZE, shuffle=True, drop_last=True)
55
56 # Find embedding sizes
57 proteins, rnas, labels = next(iter(train_dataloader))
58 protein_dim = proteins.shape[1]
59 rna_dim = rnas.shape[1]
60
61 # Initialize the model, optimizer, and loss function
62 device = "cuda" if torch.cuda.is_available() else "cpu"
63
64 model = MLPBindingClassifier(
65     protein_dim=protein_dim, rna_dim=rna_dim, hidden_dims=[256, 64],
66     dropout_rate=DROPOUT_RATE
67 )
68
69 model.to(device)
70
71 optimizer = Adam(model.parameters(), lr=LEARNING_RATE,
72     weight_decay=WEIGHT_DECAY)
73 criterion = BCELoss() # Binary Cross-Entropy Loss for binary
74     classification
75

```

```

69 # Create learning rate scheduler with linear warmup
70 total_steps = len(train_dataloader) * NUM_EPOCHS
71 warmup_steps = len(train_dataloader) * WARMUP_EPOCHS
72 scheduler = LinearLR(
73     optimizer,
74     start_factor=WARMUP_START_FACTOR,
75     end_factor=1.0,
76     total_iters=warmup_steps
77 )
78
79 # normalized_embedding = (embedding - stats["mean"]) / stats["std"]
80
81 # Training loop
82 for epoch in range(NUM_EPOCHS):
83     print(f"Epoch: {epoch+1}")
84
85     model.train() # Set the model to training mode
86     running_loss = 0.0
87
88     for batch in tqdm(train_dataloader):
89         proteins, rnas, labels = batch
90
91         labels.unsqueeze_(1) # unsqueeze to make labels 2D tensor
92         proteins, rnas, labels = proteins.to(device),
93             rnas.to(device), labels.to(device)
94
95         # Forward pass
96         outputs = model(proteins, rnas)
97         loss = criterion(outputs, labels.float())
98
99         # Zero the gradients
100         optimizer.zero_grad()
101
102         # Backward pass and optimization
103         loss.backward()
104         optimizer.step()
105
106         # Step the scheduler
107         if epoch < WARMUP_EPOCHS:
108             scheduler.step()
109
110         running_loss += loss.item()
111
112     # Log training loss and learning rate per epoch
113     epoch_loss = running_loss / len(train_dataloader)
114     current_lr = optimizer.param_groups[0]['lr']
115     wandb.log({
116         "train_loss": epoch_loss,
117         "learning_rate": current_lr,
118         "epoch": epoch + 1
119     })
120     print(f"Epoch [{epoch+1}/{NUM_EPOCHS}], Loss: {epoch_loss:.4f},
121         LR: {current_lr:.6f}")
122
123     if (epoch + 1) % VALIDATION_EPOCHS == 0:
124         model.eval()
125         validation_loss = 0.0
126         total_predictions = 0
127         true_positives = 0
128         predicted_positives = 0
129         total_positives = 0
130
131         with torch.inference_mode():
132             for batch in tqdm(validation_dataloader):
133                 proteins, rnas, labels = batch

```

```

132         labels.unsqueeze_(1)
133         proteins, rnas, labels = proteins.to(device),
            rnas.to(device), labels.to(device)

134
135         outputs = model(proteins, rnas)
136         loss = criterion(outputs, labels.float())
137         validation_loss += loss.item()
138
139         # Convert outputs to binary predictions
140         predictions = (outputs > 0.5).float()
141
142         # Update metrics
143         total_predictions += len(predictions)
144         true_positives += ((predictions == 1) & (labels ==
            1)).sum().item()
145         predicted_positives += (predictions == 1).sum().item()
146         total_positives += (labels == 1).sum().item()
147
148         # Calculate final metrics
149         accuracy = true_positives / total_predictions if
            total_predictions > 0 else 0
150         precision = true_positives / (predicted_positives + 1e-8)
            if predicted_positives > 0 else 0
151         recall = true_positives / (total_positives + 1e-8) if
            total_positives > 0 else 0
152         f1_score = 2 * (precision * recall) / (precision + recall
            + 1e-8) if (precision + recall) > 0 else 0
153
154         # Log validation metrics
155         val_loss = validation_loss / len(validation_dataloader)
            if len(validation_dataloader) > 0 else 0
156         wandb.log({
157             "val_loss": val_loss,
158             "val_accuracy": accuracy,
159             "val_precision": precision,
160             "val_recall": recall,
161             "val_f1": f1_score,
162             "epoch": epoch + 1
163         })
164
165         print(f"Validation Loss: {val_loss:.4f}")
166         print(f"Accuracy: {accuracy:.4f}, Precision:
            {precision:.4f}, Recall: {recall:.4f}, F1 Score:
            {f1_score:.4f}")
167
168     print("Training complete!")
169     wandb.finish()

```

*

Directory: classifier/util/

*

File: classifier/util/generate_background_seq.py

Listing 6: Code Snippet

```

1 import csv
2 import random
3 from tqdm import tqdm
4
5 #

```

```

6 # 1) Define file paths
7 #
8 INPUT_CSV = "/data6/alpsencer/reinforce_rna/output.csv"
9 OUTPUT_CSV = "/data6/alpsencer/reinforce_rna/background_sequences.csv"
10
11 #
12 # 2) Function to generate a background sequence by shuffling the RNA
    sequence
13 #
14 def generate_background(rna_seq):
15     rna_list = list(rna_seq)
16     random.shuffle(rna_list)
17     return "".join(rna_list)
18
19 #
20 # 3) Iteratively process the CSV file row by row
21 #
22 with open(INPUT_CSV, "r", newline="") as infile, open(OUTPUT_CSV,
    "w", newline="") as outfile:
23     reader = csv.DictReader(infile)
24     writer = csv.DictWriter(outfile, fieldnames=reader.fieldnames)
25     writer.writeheader()
26
27     # Process each row iteratively with a progress bar
28     for row in tqdm(reader, desc="Processing rows"):
29         original_rna = row["rna_sequence"]
30         bg_rna = generate_background(original_rna)
31         row["rna_sequence"] = bg_rna
32         row["pair"] = "-" # mark as background sequence
33         writer.writerow(row)
34
35 print("Background sequences saved to:", OUTPUT_CSV)

```

*

File: classifier/util/generate_protein_input.py

Listing 7: Code Snippet

```

1 # This will convert the protein_seqs csv file in the form that grpo
    trainer can get properly (different protein inputs)
2 # and save it as protein_input.txt
3 import pandas as pd
4
5 # 1) Read your source CSV
6 df = pd.read_csv("protein_seqs.csv")
7
8 # 2) Build each line as "sequence$prot_name"
9 lines = df.apply(lambda row: f"{row['seq']}${row['prot_name']}",
    axis=1)
10
11 # 3) Write to a .txt file
12 with open("protein_input.txt", "w") as f:
13     for line in lines:
14         f.write(line + "\n")
15
16 print("Wrote", len(lines), "entries to protein_input.txt")

```

*

File: classifier/util/sample_generate.py

Listing 8: Code Snippet

```
1 import csv
2 import random
3 from tqdm import tqdm
4
5 #
6 # 1) Define file paths and sample sizes
7 #
8 POSITIVE_CSV =
9     "/data6/alpsencer/reinforce_rna/positive_sequences.csv"
10 BACKGROUND_CSV =
11     "/data6/alpsencer/reinforce_rna/background_sequences.csv"
12 MERGED_TXT = "/data6/alpsencer/reinforce_rna/sample_100K.txt"
13 NUM_SAMPLES_PER_FILE = 100000 # 50K rows per file
14
15 #
16 # 2) Generator function to stream rows with strand '+' from a file
17 #
18 def positive_rows_generator(file_path, limit):
19     """
20     Yields rows from file_path where the strand is '+' until the
21     limit is reached.
22     """
23     count = 0
24     with open(file_path, "r", newline="") as infile:
25         reader = csv.DictReader(infile, delimiter=",")
26         for row in reader:
27             if row["strand"] == "+":
28                 yield row
29                 count += 1
30                 if count >= limit:
31                     break
32
33 #
34 # 3) Collect rows from both CSV files without reading them entirely
35 #    into memory
36 #
37 merged_rows = []
38
39 # Process the positive CSV
40 for row in tqdm(positive_rows_generator(POSITIVE_CSV,
41     NUM_SAMPLES_PER_FILE),
42     desc="Processing positive CSV"):
43     merged_rows.append({
44         "protein_seq": row["protein_seq"],
45         "rna_seq": row["rna_sequence"],
46         "pair": row["pair"]
47     })
48
49 # Process the background CSV
50 for row in tqdm(positive_rows_generator(BACKGROUND_CSV,
51     NUM_SAMPLES_PER_FILE),
52     desc="Processing background CSV"):
53     merged_rows.append({
54         "protein_seq": row["protein_seq"],
```

```

49         "rna_seq": row["rna_sequence"],
50         "pair": row["pair"]
51     })
52
53     #
54     -----
55     # 4) Shuffle the merged rows randomly
56     #
57     -----
58     random.shuffle(merged_rows)
59
60     #
61     -----
62     # 5) Write the shuffled rows to a TXT file using '$' as the delimiter
63     #
64     -----
65     with open(MERGED_TXT, "w", newline="") as outfile:
66         fieldnames = ["protein_seq", "rna_seq", "pair"]
67         writer = csv.DictWriter(outfile, fieldnames=fieldnames,
68                                delimiter="$")
69         writer.writeheader()
70         for row in merged_rows:
71             writer.writerow(row)
72
73     print(f"Merged and shuffled sample TXT created with
74           {NUM_SAMPLES_PER_FILE * 2} rows using '$' as delimiter:",
75           MERGED_TXT)

```

*

File: classifier/util/tokenizer.py

Listing 9: Code Snippet

```

1  import json
2  import os
3  import numpy as np
4
5
6  from tokenizers import Tokenizer
7  from tokenizers.models import BPE
8  from tokenizers.trainers import BpeTrainer
9  from tokenizers.pre_tokenizers import ByteLevel
10 from tokenizers.processors import TemplateProcessing
11 from tokenizers.normalizers import Sequence, Lowercase
12
13 import torch
14
15
16 class BpeTokenizer:
17     def __init__(self, seq_size, vocab_size):
18         self.tokenizer = Tokenizer(BPE())
19         self.tokenizer.normalizer = Sequence([Lowercase()])
20         self.tokenizer.pre_tokenizer =
21             ByteLevel(add_prefix_space=False)
22
23         self.tokenizer.enable_padding(max_length=seq_size,
24                                       direction='right')
25         self.tokenizer.enable_truncation(max_length=seq_size)
26
27         self.special_tokens = {
28             "pad": {"id": 0, "token": "<pad>"},
29             "eos": {"id": 1, "token": "</s>"},
30             "unk": {"id": 2, "token": "<unk>"},

```

```

29     }
30
31
32     self.special_tokens_list = [None] * len(self.special_tokens)
33     for token_dict in self.special_tokens.values():
34         self.special_tokens_list[token_dict["id"]] =
            token_dict["token"]
35
36     self.tokenizer.post_processor = TemplateProcessing(
37         single=f"$A {self.special_tokens['eos']}['token']}",
38         special_tokens=[
39             (self.special_tokens["eos"]['token'],
40              self.special_tokens["eos"]['id']),
41         ],
42     )
43
44     self.trainer = BpeTrainer(
45         vocab_size=vocab_size,
46         special_tokens=self.special_tokens_list,
47         show_progress=True
48     )
49
50     def train_tokenizer(self, train_data, which=True):
51         def iterator(data, which):
52             for sequence_item in data: # Renamed sequence to
53                 sequence_item
54                 text = sequence_item['text'] # Renamed sequence to
55                 sequence_item
56                 mol, rna = text.strip().split('$')
57                 if which:
58                     yield mol
59                 else:
60                     yield rna
61         self.tokenizer.train_from_iterator(iterator(data=train_data,
62             which=which), trainer=self.trainer)
63         self.add_unk_id()
64
65     def train_from_files(self, data_files: str):
66         self.tokenizer.train(files=[data_files], trainer=self.trainer)
67         self.add_unk_id()
68
69     def add_unk_id(self):
70         tokenizer_json = json.loads(self.tokenizer.to_str())
71         tokenizer_json["model"]["unk_id"] =
72             self.special_tokens["unk"]['id']
73         self.tokenizer =
74             Tokenizer.from_str(json.dumps(tokenizer_json))
75
76     def save(self, path: str, name: str):
77         if not os.path.exists(path):
78             os.makedirs(path)
79         self.tokenizer.save(os.path.join(path, f"{name}.json"))
80
81     def load(self, path: str):
82         self.tokenizer = Tokenizer.from_file(path)
83
84     def tokenize(self, sequence_to_tokenize): # Renamed sequence to
85         sequence_to_tokenize
86         # print(sequence_to_tokenize)
87         return self.tokenizer.encode(sequence_to_tokenize)
88
89     def decode(self, sequence_to_decode): # Renamed sequence to
90         sequence_to_decode
91         # print(sequence_to_decode)
92         if isinstance(sequence_to_decode, torch.Tensor):

```

```

85         sequence_to_decode =
            sequence_to_decode.detach().cpu().numpy().astype(dtype=np.int64)
86         # decode method of huggingface tokenizers usually decodes a
            single sequence of ids,
87         # or a batch if batch_decode is used. The original code here
            implies it might get a list of lists.
88         # However, self.tokenizer.decode expects a list of ints
            (single sequence) or a string.
89         # If sequence_to_decode is already a list of token IDs (e.g.,
            from model.generate), this is correct.
90         # If it's a list of sequences (list of list of IDs), it would
            need a loop or batch_decode.
91         # The custom_tokenizer.py in grpo-trainer has a different
            decode for list of sequences.
92         # Assuming this one is for a single sequence.
93         return self.tokenizer.decode(sequence_to_decode)
94
95     def encode(self, tokenized_list): # Renamed tokenized to
        tokenized_list
96         return [self.tokenizer.encode(seq) for seq in tokenized_list]
97
98
99     def get_tokenizer(tokenizer_name:str, vocab_size:int, seq_size:int,
        tokenizer_path:str=None):
100         # Choose tokenizer
101         if tokenizer_name=="bpe":
102             my_tokenizer = BpeTokenizer(vocab_size=vocab_size,
                seq_size=seq_size)
103         else:
104             raise NotImplementedError
105
106         # Load pre-trained tokenizer or train tokenizer
107         if tokenizer_path:
108             my_tokenizer.load(tokenizer_path)
109
110         if vocab_size != my_tokenizer.tokenizer.get_vocab_size():
111             assert "There is a conflict Tokenizer's vocab size and
                arguments'"
112
113         return my_tokenizer

```

Directory: grpo-trainer/

File: grpo-trainer/config.py

Listing 10: Code Snippet

```

1  MODEL_PATH =
    "/data6/sobhan/rllm/results/train/t5/run3_20240822-152114/checkpoints/checkpoint-34980"
2  PROTEIN_TOKENIZER =
    "/data6/alpsencer/tokenizers/bpe_protein_1000_1024.json"
3  RNA_TOKENIZER = "/data6/alpsencer/tokenizers/bpe_rna_1000_1024.json"
4
5  TRAIN_DATA = "/data6/alpsencer/reinforce_rna/protein_input.txt"
6  EVAL_DATA = "/data6/alpsencer/reinforce_rna/validation_rl.txt"
7
8  MAX_LENGTH = 1024
9  VOCAB_SIZE = 1000
10 BATCH_SIZE = 16
11 LEARNING_RATE = 1e-5
12 NUM_TRAIN_EPOCHS = 3
13 GRADIENT_ACCUMULATION_STEPS = 4

```

```

14 # FREEZE_ENCODER = True # Whether to freeze the encoder during
    training
15
16 GC_CONTENT_MIN = 0.4 # Minimum acceptable GC content
17 GC_CONTENT_MAX = 0.6 # Maximum acceptable GC content
18 BINDING_AFFINITY_WEIGHT = 0.2 # Weight for binding affinity reward
19 GC_CONTENT_WEIGHT = 0.4 # Weight for GC content reward
20 MFE_WEIGHT = 0.4 # Weight for MFE reward
21
22 NUM_GENERATIONS = 8 # Number of generations per prompt
23 MAX_PROMPT_LENGTH = 1024 # Maximum length of the prompt
24 MAX_COMPLETION_LENGTH = 256 # Maximum length of the generated
    completion
25
26 BETA = 0.04 # KL coefficient
27 EPSILON = 0.1 # Epsilon value for clipping
28 MU = 5
29
30 WANDB_PROJECT = "rna-protein-grpo"

```

File: grpo-trainer/custom_tokenizer.py

Listing 11: Code Snippet

```

1 import json
2 import os
3 import numpy as np
4
5
6 from tokenizers import Tokenizer
7 from tokenizers.models import BPE
8 from tokenizers.trainers import BpeTrainer
9 from tokenizers.pre_tokenizers import ByteLevel
10 from tokenizers.processors import TemplateProcessing
11 from tokenizers.normalizers import Sequence, Lowercase
12
13 import torch
14
15
16 class BpeTokenizer:
17     def __init__(self, seq_size, vocab_size):
18         self.tokenizer = Tokenizer(BPE())
19         self.tokenizer.normalizer = Sequence([Lowercase()])
20         self.tokenizer.pre_tokenizer =
            ByteLevel(add_prefix_space=False)
21
22         self.tokenizer.enable_padding(max_length=seq_size,
            direction='right')
23         self.tokenizer.enable_truncation(max_length=seq_size)
24
25         self.special_tokens = {
26             "pad": {"id": 0, "token": "<pad>"},
27             "eos": {"id": 1, "token": "</s>"},
28             "unk": {"id": 2, "token": "<unk>"},
29         }
30
31
32         self.special_tokens_list = [None] * len(self.special_tokens)
33         for token_dict in self.special_tokens.values():
34             self.special_tokens_list[token_dict["id"]] =
                token_dict["token"]
35
36         self.tokenizer.post_processor = TemplateProcessing(
37             single=f"$A {self.special_tokens['eos']}['token']}",

```

```

38         special_tokens=[
39             (self.special_tokens["eos"]["token"],
              self.special_tokens["eos"]["id"]),
40         ],
41     )
42
43     self.trainer = BpeTrainer(
44         vocab_size=vocab_size,
45         special_tokens=self.special_tokens_list,
46         show_progress=True
47     )
48
49     def train_tokenizer(self, train_data, which=True):
50         def iterator(data, which_mol): # Renamed which to which_mol
51             for sequence_item in data: # Renamed sequence to
52                 sequence_item
53                 text = sequence_item['text'] # Renamed sequence to
54                 sequence_item
55                 mol, rna = text.strip().split('$')
56                 if which_mol: # Use which_mol
57                     yield mol
58                 else:
59                     yield rna
60         self.tokenizer.train_from_iterator(iterator(data=train_data,
61             which_mol=which), trainer=self.trainer) # Pass which as
62             which_mol
63         self.add_unk_id()
64
65     def train_from_files(self, data_files: str):
66         self.tokenizer.train(files=[data_files], trainer=self.trainer)
67         self.add_unk_id()
68
69     def add_unk_id(self):
70         tokenizer_json = json.loads(self.tokenizer.to_str())
71         tokenizer_json["model"]["unk_id"] =
72             self.special_tokens["unk"]["id"]
73         self.tokenizer =
74             Tokenizer.from_str(json.dumps(tokenizer_json))
75
76     def save(self, path: str, name: str):
77         if not os.path.exists(path):
78             os.makedirs(path)
79         self.tokenizer.save(os.path.join(path, f"{name}.json"))
80
81     def load(self, path: str):
82         self.tokenizer = Tokenizer.from_file(path)
83
84     def tokenize(self, sequence_to_tokenize): # Renamed sequence to
85         sequence_to_tokenize
86         # print(sequence_to_tokenize)
87         return self.tokenizer.encode(sequence_to_tokenize)
88
89     def decode(self, sequences_to_decode): # Renamed sequences to
90         sequences_to_decode
91         # print(sequences_to_decode)
92         if isinstance(sequences_to_decode, torch.Tensor):
93             sequences_to_decode =
94                 sequences_to_decode.detach().cpu().numpy().astype(dtype=np.int64)
95         # This decode is meant to handle a batch of sequences (list
96         # of list of IDs or tensor)
97         # and return a list of decoded strings.
98         return [self.tokenizer.decode(seq) for seq in
99             sequences_to_decode]

```

```

90     def encode(self, tokenized_list): # Renamed tokenized to
        tokenized_list
91         return [self.tokenizer.encode(seq) for seq in tokenized_list]
92
93
94     def get_tokenizer(tokenizer_name:str, vocab_size:int, seq_size:int,
        tokenizer_path:str=None):
95         # Choose tokenizer
96         if tokenizer_name=="bpe":
97             my_tokenizer = BpeTokenizer(vocab_size=vocab_size,
                seq_size=seq_size)
98         else:
99             raise NotImplementedError
100
101         # Load pre-trained tokenizer or train tokenizer
102         if tokenizer_path:
103             my_tokenizer.load(tokenizer_path)
104
105         if vocab_size != my_tokenizer.tokenizer.get_vocab_size():
106             assert "There is a conflict Tokenizer's vocab size and
                arguments'"
107
108         return my_tokenizer

```

File: grpo-trainer/hyperparameter_tuning.py

Listing 12: Code Snippet

```

1  import optuna
2  import wandb
3  import sys
4  import copy
5  import torch
6  from pathlib import Path
7
8  # Add the project root to Python path to import from sibling
    directories
9  project_root = Path(__file__).resolve().parent # Use resolve() for
    robustness
10 sys.path.append(str(project_root))
11
12 # from config import (
13 #     LEARNING_RATE, BETA, EPSILON, MU,
14 #     BINDING_AFFINITY_WEIGHT, GC_CONTENT_WEIGHT, MFE_WEIGHT,
15 #     NUM_GENERATIONS, TOP_P, TEMPERATURE
16 # )
17 from main import ProteinRNAGRPOtainer # Assuming main.py is in the
    same dir (project_root)
18
19 # def run_hyperparameter_tuning(model, tok_p, tok_r, train_envs,
    eval_envs, config_obj, n_trials=20): # Renamed config to
    config_obj
20 def run_hyperparameter_tuning(model, tok_p, tok_r, train_envs,
    eval_envs, tuning_config, n_trials=20): # Renamed config_obj to
    tuning_config
21     # Save the initial state of the model to reset before each trial
22     base_model_state = copy.deepcopy(model.state_dict())
23
24     def objective(trial):
25         try:
26             # Reset the model to the base state before each trial
27             model.load_state_dict(base_model_state)
28

```

```

29         # Reset optimizer state by creating a new optimizer
        instance
30         # The optimizer is created within ProteinRNAGRPOTainer,
        so this might not be strictly necessary here
31         # if the trainer re-initializes its optimizer or if model
        reset is enough.
32         # However, explicit re-creation ensures no stale
        optimizer state.
33         # optimizer = torch.optim.AdamW(model.parameters(),
        lr=tuning_config.learning_rate) # Use tuning_config
34
35         # Define hyperparameter search space for only mu, beta,
        and epsilon
36         beta_val = trial.suggest_float('beta', 0.01, 0.1) #
        Renamed beta to beta_val
37         epsilon_val = trial.suggest_float('epsilon', 0.05, 0.2) #
        Renamed epsilon to epsilon_val
38         mu_val = trial.suggest_int('mu', 3, 8) # Renamed mu to
        mu_val
39
40         # Initialize wandb run with the trial parameters
41         current_run = wandb.init( # Renamed run to current_run
42             project="rna-protein-grpo", # Or use
        tuning_config.WANDB_PROJECT if defined
43             name=f"trial_{trial.number}",
44             config={
45                 'beta': beta_val,
46                 'epsilon': epsilon_val,
47                 'mu': mu_val,
48             },
49             reinit=True # Allow multiple runs in the same process
50         )
51
52         # Update config parameters FOR THIS TRIAL
53         # The ProteinRNAGRPOTainer __init__ takes a config object.
54         # We need to pass these tuned values to it.
55         # One way is to create a new config object for each trial
        or modify a copy.
56
57         trial_config = copy.deepcopy(tuning_config) # Work on a
        copy
58         # These attributes (BETA, EPSILON, MU) need to be
        accessible by ProteinRNAGRPOTainer
59         # If ProteinRNAGRPOTainer reads them from the global
        'config' module, then we'd set config.BETA
60         # If ProteinRNAGRPOTainer reads them from its
        self.config_obj, then we set them on trial_config.
61         # Assuming ProteinRNAGRPOTainer uses its
        self.config_obj.BETA etc.
62         trial_config.BETA = beta_val
63         trial_config.EPSILON = epsilon_val
64         trial_config.MU = mu_val
65
66         # Run training
67         trainer = ProteinRNAGRPOTainer(model, tok_p, tok_r,
        train_envs, eval_envs, trial_config) # Pass
        trial_config
68         trainer.train(epochs=1) # Run for 1 epoch during tuning
69
70         # Get the final reward as the objective value
71         # Check if summary is available and contains the key
72         final_reward = 0
73         if current_run and current_run.summary:
74             final_reward =
                current_run.summary.get('reward/mean', 0)

```

```

75
76         # End wandb run properly
77         if wandb.run is not None: # Check if a run is active
78             wandb.finish()
79
80         return final_reward
81
82     except Exception as e:
83         print(f"Trial failed with error: {str(e)}")
84         # Make sure to finish the wandb run even if there's an
            error
85         if wandb.run is not None:
86             wandb.finish()
87         # Return a very low reward to indicate failure
88         return float('-inf')
89
90     study = optuna.create_study(direction='maximize')
91     study.optimize(objective, n_trials=n_trials)
92
93     print("Best trial:")
94     best_trial_obj = study.best_trial # Renamed trial to
        best_trial_obj
95
96     print("  Value: ", best_trial_obj.value)
97     print("  Params: ")
98     for key, value in best_trial_obj.params.items():
99         print(f"    {key}: {value}")
100
101     # Restore the model to the initial base state (not necessarily
        best params state)
102     # Or, if you want to set model to best params, you'd need to
        re-run with best_params.
103     model.load_state_dict(base_model_state)
104
105     return study.best_params

```

File: grpo-trainer/main.py

Listing 13: Code Snippet

```

1  import argparse
2  import os
3  import random
4  import torch
5  import numpy as np
6  from transformers import T5ForConditionalGeneration,
    PreTrainedTokenizerFast
7  # from trl import GRPOConfig # GRPOConfig from trl is for a different
    GRPO setup. We use a custom one.
8  from reward_functions import rewarder
9  import config as global_config_module # Use an alias for the global
    config module
10 from custom_tokenizer import get_tokenizer
11 # from torch.nn.functional import log_softmax # Not directly used,
    F.log_softmax is used
12
13 import copy # Already imported
14 import torch.nn.functional as F
15
16
17 # from config import EPSILON, MU, BETA # These will be taken from
    global_config_module or trainer's config object
18
19

```

```

20 import wandb
21 import multiprocessing as mp
22 from functools import partial
23
24 # wandb.init(project="protein-rna-grpo", name="grpo-run") # Init
    should be done once, ideally in main or controlled per run
25
26 # --- Define a simple config class for our GRPO Trainer ---
27 # This replaces the direct use of trl.GRPOConfig for a custom setup
28 class CustomGRPOConfig:
29     def __init__(self, **kwargs):
30         self.output_dir = "./results"
31         self.learning_rate = 1e-5
32         self.max_prompt_length = 1024
33         self.max_completion_length = 256
34         self.num_generations = 8 # Number of generations per prompt
35         self.top_p = 1.0
36         self.temperature = 1.0
37         self.max_grad_norm = 1.0
38
39         # GRPO specific, potentially tuned
40         self.BETA = global_config_module.BETA # Default from global
            config
41         self.EPSILON = global_config_module.EPSILON # Default from
            global config
42         self.MU = global_config_module.MU # Default from global config
43
44         # Update with any provided kwargs
45         for key, value in kwargs.items():
46             setattr(self, key, value)
47
48 # --- End simple config class ---
49
50
51 def generate_with_log_probs(
52     model_obj: T5ForConditionalGeneration, # Renamed model to
        model_obj
53     input_ids: torch.Tensor,
54     attention_mask: torch.Tensor,
55     num_return_sequences: int,
56     max_new_tokens: int,
57     top_p: float,
58     temperature: float,
59 ) -> tuple[torch.Tensor, torch.Tensor]:
60     """
61     1) Sample N sequences with model_obj.generate (still under
        no-grad)
62     2) Re-run the model_obj forward pass on those sequences to obtain
        token-wise log-probs **with gradient**.
63     Returns:
64         sequences      (N, T) generated token ids (without the start
            token)
65         log_probs      (N, T) log-probs that require grad
66     """
67     device = input_ids.device
68
69     # ---- 1) sampling (no grad)
70     -----
71     with torch.no_grad():
72         gen_out = model_obj.generate(
73             input_ids=input_ids,
74             attention_mask=attention_mask,
75             max_new_tokens=max_new_tokens,
76             do_sample=True,
77             top_p=top_p,

```

```

78         temperature=temperature,
79         num_return_sequences=num_return_sequences,
80         return_dict_in_generate=True,
81     )
82     # gen_out.sequences shape: (N, 1 + T) (extra decoder-start
      token)
83     full_sequences = gen_out.sequences
84
85     sequences = full_sequences[:, 1:] # drop start token
      (N, T)
86     N, T = sequences.shape
87
88     # ---- 2) compute log-probs with grad
      -----
89     # repeat the encoder prompt N times (because batch == 1)
90     enc_input_ids = input_ids.repeat_interleave(N, dim=0)
91     enc_attention_mask = attention_mask.repeat_interleave(N, dim=0)
92
93     # decoder_input_ids are everything *before* each position
94     # prepend the decoder-start token id first
95     start_tok = model_obj.config.decoder_start_token_id
96     start_col = torch.full((N, 1), start_tok, dtype=torch.long,
      device=device)
97     decoder_input_ids = torch.cat([start_col, sequences[:, :-1]],
      dim=1) # (N, T)
98
99     outputs = model_obj(
100         input_ids=enc_input_ids,
101         attention_mask=enc_attention_mask,
102         decoder_input_ids=decoder_input_ids,
103     )
104     logits = outputs.logits # (N, T, V)
105
106     log_probs_tensor = torch.log_softmax(logits, dim=-1) # Renamed
      log_probs to log_probs_tensor (N, T, V)
107
108     # pick the prob of the *actual* generated token at each step
109     tok_ids = sequences.unsqueeze(-1) # (N, T, 1)
110     log_probs_tensor = log_probs_tensor.gather(-1,
      tok_ids).squeeze(-1) # (N, T) requires grad
111
112     return sequences, log_probs_tensor
113
114
115 def process_tokens_args_tuple(args_tuple): # Renamed args to
      args_tuple to avoid conflict
116     """
117     Helper function to process tokens for a single sequence, takes a
      tuple for pool.map
118     """
119     i, t, rna_ids, rna_str_list, protein_str_val, pad_token_id,
      eos_token_id = args_tuple # Unpack tuple, rna ->
      rna_str_list, protein -> protein_str_val
120     if (rna_ids[i, t] == pad_token_id or rna_ids[i, t] ==
      eos_token_id):
121         return t, 0.0
122
123     # Build the text up to this token
124     # rna_str_list[i] is a list of characters/tokens for the i-th
      sequence
125     text = "".join(rna_str_list[i][:t+1])
126     reward_val = rewarder(text, protein_str_val) # Renamed reward to
      reward_val
127     return t, reward_val
128

```

```

129 def calculate_reward(
130     rna_ids: torch.Tensor,
131     rna_decoded_str_list: list[list[str]], # rna is now list of list
        of chars/tokens for each seq in batch
132     protein_ids: torch.Tensor, # protein_ids not directly used here
        but good for context
133     protein_str_val: str, # protein is now a single string for the
        batch
134     pad_token_id: int = 0,
135     eos_token_id: int = 1
136 ) -> torch.Tensor:
137     """
138     sequences: (N, T) token IDs for generated RNAs
139     rna_decoded_str_list: (N, list_of_chars_at_t) decoded RNA strings
        (ACGU format)
140     protein_str_val: the source protein string (same for all N if
        batch_size=1)
141     returns: (N, T) rewards for each prefix ending at t,
142             where rewarder(text, protein_str_val) is called.
143     """
144     device = rna_ids.device # Keep rewards on the original device if
        possible, or move later
145     # rna_ids_cpu = rna_ids.to("cpu") # Move to CPU if
        multiprocessing requires it (often does)
146     # protein_ids_cpu = protein_ids.to("cpu") # Not used in
        process_tokens_args_tuple directly
147
148     N, T = rna_ids.size()
149     rewards_tensor = torch.zeros(N, T, device=device) # Renamed
        rewards to rewards_tensor
150
151     # Create a pool of workers
152     # Using try-finally to ensure pool is closed
153     pool = None
154     try:
155         # num_workers = mp.cpu_count() # Can be too many, limit if
        needed
156         num_workers = min(mp.cpu_count(), 4) # Example limit
157         pool = mp.Pool(processes=num_workers)
158
159         for i in range(N):
160             # Prepare arguments for each token position in this
        sequence
161             # Pass rna_ids (on CPU if needed by pool),
        rna_decoded_str_list (Python list), protein_str_val
        (Python str)
162             token_args_list = [(i, t, rna_ids, rna_decoded_str_list,
        protein_str_val, pad_token_id, eos_token_id)
163                               for t in range(T)]
164
165             # Process tokens in parallel
166             results = pool.map(process_tokens_args_tuple,
        token_args_list)
167
168             # Combine results for this sequence
169             for t_res, reward_val_res in results: # Renamed t, reward
        to t_res, reward_val_res
170                 rewards_tensor[i, t_res] = reward_val_res
171     finally:
172         if pool:
173             pool.close()
174             pool.join()
175
176     return rewards_tensor
177

```

```

178
179 def calculate_advantage(rewards: torch.Tensor, eps: float = 1e-8) ->
    torch.Tensor:
180     """
181     Normalize advantages per time-step (column), ignoring padded
        values.
182     rewards: (N, T) tensor of rewards, where padded values are 0
183     returns: (N, T) tensor of normalized advantages
184     """
185     # Create a mask for non-padded values (assuming padding is 0)
186     mask = (rewards != 0).float()
187
188     # Calculate mean only over non-padded values
189     sum_rewards = (rewards * mask).sum(dim=0, keepdim=True) # (1, T)
190     count = mask.sum(dim=0, keepdim=True).clamp(min=eps) # Avoid
        division by zero if count is 0
191     mean = sum_rewards / count # (1, T)
192
193     # Center the rewards
194     centered = rewards - mean # (N, T)
195
196     # Calculate std only over non-padded values
197     squared_diff = (centered * mask) ** 2 # (N, T)
198     sum_squared_diff = squared_diff.sum(dim=0, keepdim=True) # (1, T)
199     std = torch.sqrt(sum_squared_diff / count) # (1, T) # count
        already has eps
200
201     # Normalize and apply mask to ensure padded values remain 0
202     advantages = (centered / (std + eps)) * mask # Add eps to std for
        stability
203     return advantages
204
205
206 class ProteinRNAEnvironment:
207     """
208     Environment wrapper for ProteinRNA, using tokenizer.tokenize()
209     (returns .ids and .attention_mask).
210     """
211     def __init__(self, protein: str, rna: str,
212                  protein_tokenizer, rna_tokenizer,
213                  max_input_length: int, max_output_length: int):
214         self.protein_str = protein # Renamed protein to protein_str
215         self.rna_str = rna # Renamed rna to rna_str
216         self.p_tok = protein_tokenizer
217         self.r_tok = rna_tokenizer
218         self.max_in = max_input_length
219         self.max_out = max_output_length
220
221         # Store pad token id for convenience
222         self.p_tok_pad_id = self.p_tok.tokenizer.pad_token_id if
            hasattr(self.p_tok, 'tokenizer') and
            hasattr(self.p_tok.tokenizer, 'pad_token_id') else 0
223         self.r_tok_pad_id = self.r_tok.tokenizer.pad_token_id if
            hasattr(self.r_tok, 'tokenizer') and
            hasattr(self.r_tok.tokenizer, 'pad_token_id') else 0
224
225
226     def get_prompt_inputs(self):
227         # 1) Tokenize protein prompt
228         tok = self.p_tok.tokenize(self.protein_str)
229         ids = tok.ids[:self.max_in]
230         mask = tok.attention_mask[:self.max_in]
231
232         # 2) Pad out to max_in
233         pad_len = self.max_in - len(ids)

```

```

234         if pad_len > 0:
235             ids = ids + [self.p_tok_pad_id] * pad_len # Use stored
                pad_id
236             mask = mask + [0] * pad_len
237
238         # 3) Wrap in batch dim
239         input_ids = torch.tensor([ids], dtype=torch.long)
240         attention_mask = torch.tensor([mask], dtype=torch.long)
241         return input_ids, attention_mask
242
243     def get_labels(self):
244         # 1) Tokenize target RNA (convert to bjuz first)
245         rna_bjuz = self.rna_str.replace("A", "b").replace("C",
                "j").replace("U", "u").replace("G", "z").replace(" ", "")
246         tok = self.r_tok.tokenize(rna_bjuz)
247         ids = tok.ids[:self.max_out]
248
249         # Labels for T5 are usually shifted, and padded tokens are
                -100
250         # The T5 model itself handles shifting if decoder_input_ids
                are not provided with labels.
251         # If we are providing labels for loss computation, they
                should be the target tokens.
252         # Padded tokens in labels should be -100.
253
254         # Pad labels to max_out length
255         pad_len = self.max_out - len(ids)
256         if pad_len > 0:
257             ids = ids + [self.r_tok_pad_id] * pad_len # Pad with
                actual pad token id first
258
259         labels_tensor = torch.tensor([ids], dtype=torch.long) #
                Renamed labels to labels_tensor
260         labels_tensor[labels_tensor == self.r_tok_pad_id] = -100 #
                Replace padding with -100 for loss calculation
261
262         return labels_tensor
263
264
265     class ProteinRNAGRPOTrainer:
266         def __init__(self, model_obj, tokenizer_p, tokenizer_r,
                train_envs_list, eval_envs_list, config_obj_trainer): #
                Renamed params
267             self.model = model_obj # Renamed
268             self.tok_p = tokenizer_p
269             self.tok_r = tokenizer_r
270             self.train_envs = train_envs_list # Renamed
271             self.eval_envs = eval_envs_list # Renamed
272             self.config = config_obj_trainer # Renamed
273
274             # make a frozen reference policy ( _ref )
275             self.ref_model = copy.deepcopy(self.model).eval()
276             for p in self.ref_model.parameters():
277                 p.requires_grad = False
278
279             # the "old" policy _old , initially equals _ref
280             self.old_model = copy.deepcopy(self.ref_model)
281
282
283             self.optimizer = torch.optim.AdamW(self.model.parameters(),
                lr=self.config.learning_rate)
284
285
286
287

```

```

288 def evaluate(self, step=0):
289     """
290     Calculate validation loss and rewards on the evaluation
291     dataset
292     """
293     device = next(self.model.parameters()).device
294     self.model.eval()
295     total_loss = 0
296     total_reward = 0
297     num_batches = 0
298
299     if not self.eval_envs: # Handle empty eval_envs
300         print("No evaluation environments provided. Skipping
301             evaluation.")
302         self.model.train()
303         return 0,0
304
305     with torch.no_grad():
306         for env in self.eval_envs:
307             # Get inputs and labels
308             input_ids, attention_mask = env.get_prompt_inputs()
309             labels = env.get_labels()
310
311             # Move to device
312             input_ids = input_ids.to(device)
313             attention_mask = attention_mask.to(device)
314             labels = labels.to(device)
315
316             # Forward pass for loss
317             outputs = self.model(
318                 input_ids=input_ids,
319                 attention_mask=attention_mask,
320                 labels=labels
321             )
322
323             loss = outputs.loss
324             total_loss += loss.item()
325
326             # Generate RNA sequence for reward calculation
327             gen_output = self.model.generate( # Renamed gen to
328                 gen_output
329                 input_ids=input_ids,
330                 attention_mask=attention_mask,
331                 max_new_tokens=self.config.max_completion_length,
332                 num_return_sequences=1,
333                 # Add other generation params if needed, e.g.,
334                 top_p, temperature
335                 do_sample=True, # Assuming sampling for reward
336                 evaluation
337                 top_p=self.config.top_p,
338                 temperature=self.config.temperature,
339             )
340             # decode expects list of IDs, gen_output[0] is a
341             # sequence of IDs (1D tensor)
342             generated_rna_bjuz_list =
343                 self.tok_r.decode(gen_output) # decode from
344                 custom_tokenizer returns a list of strings
345             generated_rna_bjuz = generated_rna_bjuz_list[0] if
346                 generated_rna_bjuz_list else ""
347
348             # Convert bjuz to ACGU for rewarder
349             generated_rna_acgu =
350                 generated_rna_bjuz.replace("b", "A").replace("j", "C").replace("u", "U").
351                 replace(" ", "")

```

```

342         total_reward += rewarder(generated_rna_acgu,
343                                   env.protein_str) # Use env.protein_str
344
345         num_batches += 1
346
347     avg_loss = total_loss / num_batches if num_batches > 0 else 0
348     avg_reward = total_reward / num_batches if num_batches > 0
349     else 0
350
351     # Log to wandb
352     if wandb.run: # Check if wandb run is active
353         wandb.log({
354             "validation/loss": avg_loss,
355             "validation/reward": avg_reward,
356             "train/step": step # Use the global step passed to
357                                 evaluate
358         })
359
360     self.model.train()
361     return avg_loss, avg_reward
362
363 def train(self, epochs=1):
364     device = next(self.model.parameters()).device
365     num_gen = self.config.num_generations
366     max_new = self.config.max_completion_length
367     eps_clip = self.config.EPSILON # Use from trainer's config
368     object
369     beta_val = self.config.BETA # Use from trainer's config
370     object
371     mu_val = self.config.MU # Use from trainer's config
372     object
373
374     # best_val_loss = float('inf') # Not used currently
375     validation_interval = 100 # Validate every 100 steps
376     global_step_counter = 0 # For periodic validation and logging
377
378     for epoch in range(epochs):
379         random.shuffle(self.train_envs)
380         # Consider limiting train_envs for quicker epochs if
381         # dataset is large: self.train_envs[:200]
382         for step_idx, env in enumerate(self.train_envs):
383             # 0) prepare inputs for one protein (batch_size=1)
384             input_ids, attention_mask = env.get_prompt_inputs()
385             input_ids = input_ids.to(device)
386             attention_mask = attention_mask.to(device)
387             protein_str_val = env.protein_str # Use
388             env.protein_str
389
390             # 1) Sample once from _old (no grad)
391             # print(">>> Generating theta_old outputs for a
392             # protein")
393             with torch.no_grad():
394                 seq_old, logp_old = generate_with_log_probs(
395                     model_obj=self.old_model, # Pass model_obj
396                     input_ids=input_ids,
397                     attention_mask=attention_mask,
398                     num_return_sequences=num_gen,
399                     max_new_tokens=max_new,
400                     top_p=self.config.top_p,
401                     temperature=self.config.temperature,
402                 )
393             # print(">>> theta_old output generation is
404             # completed")
405             N, T_len = seq_old.shape # Renamed T to T_len to
406             avoid conflict

```

```

396
397 # 3.1) build decoder_input_ids for teacher forcing
398 start_tok = self.model.config.decoder_start_token_id
399 start_col = torch.full((N,1), start_tok,
400                        device=device,
401                        dtype=torch.long)
402 dec_inp = torch.cat([start_col, seq_old[:, :-1]],
403                    dim=1) # (N, T_len)
404
405 # 3.2) Prepare encoder inputs for batch N
406 enc_ids_repeated = input_ids.repeat_interleave(N,
407        dim=0) # Renamed enc_ids
408 enc_mask_repeated=
409     attention_mask.repeat_interleave(N, dim=0) #
410     Renamed enc_mask
411
412 # print(">>> Generating ref outputs for a protein")
413 with torch.no_grad(): # ref_model does not require
414     grad
415     out_ref = self.ref_model(
416         input_ids=enc_ids_repeated,
417         attention_mask=enc_mask_repeated,
418         decoder_input_ids=dec_inp,
419     )
420 logits_ref= out_ref.logits #
421         (N, T_len, V)
422 logp_ref = F.log_softmax(logits_ref, -1) #
423         (N, T_len, V)
424 logp_ref = logp_ref.gather(-1,
425     seq_old.unsqueeze(-1))\
426         .squeeze(-1).detach() #
427         (N, T_len), ensure detached
428 # print(">>> ref output generation is completed")
429
430 # 2) Compute rewards & advantages (detach)
431 # print(">>> Started calculating rewards and
432     advantages")
433 # tok_r.decode expects a batch of sequences (list of
434     list of IDs or Tensor)
435 # seq_old is already a batch of sequences (Tensor)
436 decoded_seq_old_bjuz_list =
437     self.tok_r.decode(seq_old) # Returns list of bjuz
438     strings
439 # Convert bjuz to ACGU for rewarder; rewarder expects
440     list of char lists or list of strings
441 decoded_seq_old_acgu_char_lists = [[char for char in
442     seq.replace("b", "A").replace("j",
443     "C").replace("u", "U").replace("z",
444     "G").replace(" ", "")] for seq in
445     decoded_seq_old_bjuz_list]
446
447 # Pass seq_old (token_ids on device),
448     decoded_seq_old_acgu_char_lists (Python list of
449     lists of chars), protein_str_val
450 rewards_tensor = calculate_reward(seq_old,
451     decoded_seq_old_acgu_char_lists, input_ids,
452     protein_str_val,
453     pad_token_id=self.tok_r.tokenizer.pad_token_id,
454     eos_token_id=self.tok_r.tokenizer.eos_token_id)
455 advantages =
456     calculate_advantage(rewards_tensor).detach()
457
458 # Log reward stats to wandb
459 non_zero_rewards = rewards_tensor[rewards_tensor !=
460     0].flatten()

```

```

434     if wandb.run: # Check if wandb run is active
435         if non_zero_rewards.numel() > 0:
436             wandb.log({
437                 "reward/mean":
438                     non_zero_rewards.mean().item(),
439                 "reward/std":
440                     non_zero_rewards.std().item(),
441                 "reward/max":
442                     non_zero_rewards.max().item(),
443                 "reward/min":
444                     non_zero_rewards.min().item(),
445                 "train/step": global_step_counter
446             })
447         else:
448             wandb.log({
449                 "reward/mean": 0.0, "reward/std": 0.0,
450                 "reward/max": 0.0, "reward/min": 0.0,
451                 "train/step": global_step_counter
452             })
453     # print(">>> Ended calculating rewards and
454           advantages")
455
456     # 3) Inner GRPO updates on _
457     for inner_loop_idx in range(mu_val): # Renamed inner
458         to inner_loop_idx
459         out_new = self.model( # Renamed out to out_new
460             input_ids=enc_ids_repeated,
461             attention_mask=enc_mask_repeated,
462             decoder_input_ids=dec_inp,
463         )
464         logits_new = out_new.logits
465                                     # (N, T_len, V)
466         logp_new = F.log_softmax(logits_new, -1)
467                                     # (N, T_len, V)
468         logp_new = logp_new.gather(-1,
469             seq_old.unsqueeze(-1))\
470             .squeeze(-1)
471                                     # (N, T_len)
472
473         # 3.3) clipped surrogate
474         ratio = torch.exp(logp_new - logp_old.detach())
475             # Detach logp_old, it's from old_model
476         clipped = ratio.clamp(1 - eps_clip, 1 + eps_clip)
477         loss_pg = - torch.mean(
478             torch.minimum(ratio * advantages,
479                 clipped * advantages)
480         )
481
482         # 3.4) K L penalty using f(u)=u - log u - 1
483         if beta_val > 0:
484             u = torch.exp(logp_ref - logp_new) #
485                 logp_ref is already detached
486             kl_term = torch.mean(u - (logp_ref -
487                 logp_new) - 1)
488             loss_total = loss_pg + beta_val * kl_term #
489                 Renamed loss to loss_total
490         else:
491             kl_term = torch.tensor(0.0, device=device)
492             loss_total = loss_pg
493
494         # 3.5) debug prints (optional, can be verbose)
495         # print(f"Epoch {epoch+1} Step {step_idx} Inner
496               {inner_loop_idx} | "
497               f"pg_loss={loss_pg.item():.4f}
498               kl={kl_term.item():.4f}")

```

```

482
483         if wandb.run: # Check if wandb run is active
484             wandb.log({
485                 "loss/pg_loss": loss_pg.item(),
486                 "loss/kl_term": kl_term.item(),
487                 "loss/total": loss_total.item(),
488                 "train/step": global_step_counter # Log
489                     with global_step
490             })
491
492         # 3.6) backprop & step
493         self.optimizer.zero_grad()
494         loss_total.backward()
495         torch.nn.utils.clip_grad_norm_(self.model.parameters(),
496             max_norm=self.config.max_grad_norm)
497         self.optimizer.step()
498
499         global_step_counter += 1 # Increment global step
500
501         # Run validation periodically
502         if self.eval_envs and global_step_counter %
503             validation_interval == 0:
504             val_loss, val_reward =
505                 self.evaluate(step=global_step_counter) #
506                 Pass global_step_counter
507             print(f"Global Step {global_step_counter}
508                 Validation: loss={val_loss:.4f}
509                 reward={val_reward:.4f}")
510
511         # 4) sync _old _
512         self.old_model.load_state_dict(self.model.state_dict())
513
514         print(f"--- End Epoch {epoch+1} ---")
515
516         # if wandb.run: # Check if wandb run is active # Moved finish
517         # to main
518         # wandb.finish()
519
520     def load_data(path):
521         prots, rnas = [], []
522         with open(path) as f:
523             for line in f:
524                 line = line.strip()
525                 if '$' not in line: continue
526                 try:
527                     p, r = line.split('$',1) # Split only on the first '$'
528                     prots.append(p); rnas.append(r)
529                 except ValueError:
530                     print(f"Skipping malformed line in data file:
531                         {line}") # Handle lines without '$' if strip
532                     didn't catch it
533
534         return prots, rnas
535
536     if __name__ == '__main__':
537         # Initialize wandb once for the entire script run
538         # Hyperparameter tuning might reinitialize runs if configured to
539         # do so.
540         wandb.init(project=global_config_module.WANDB_PROJECT,
541             name="grpo-main-run", reinit=True) # Allow reinit for tuning
542
543         parser = argparse.ArgumentParser()
544         # Arguments from global_config_module as defaults
545         parser.add_argument('--train_data',
546             default=global_config_module.TRAIN_DATA)

```

```

534 parser.add_argument('--eval_data',
    default=global_config_module.EVAL_DATA)
535 parser.add_argument('--model_path',
    default=global_config_module.MODEL_PATH)
536 parser.add_argument('--protein_tokenizer',
    default=global_config_module.PROTEIN_TOKENIZER)
537 parser.add_argument('--rna_tokenizer',
    default=global_config_module.RNA_TOKENIZER)
538 parser.add_argument('--output_dir', default='./results_grpo') #
    Different default from CustomGRPOConfig
539 parser.add_argument('--epochs', type=int,
    default=global_config_module.NUM_TRAIN_EPOCHS)
540 parser.add_argument('--learning_rate', type=float,
    default=global_config_module.LEARNING_RATE)
541 parser.add_argument('--max_input_length', type=int,
    default=global_config_module.MAX_PROMPT_LENGTH)
542 parser.add_argument('--max_completion_length', type=int,
    default=global_config_module.MAX_COMPLETION_LENGTH)
543 parser.add_argument('--num_generations', type=int,
    default=global_config_module.NUM_GENERATIONS)
544 # top_p, temperature, max_grad_norm are not in
    global_config_module, use CustomGRPOConfig defaults or add to
    global_config_module
545 parser.add_argument('--top_p', type=float, default=1.0)
546 parser.add_argument('--temperature', type=float, default=1.0)
547 parser.add_argument('--max_grad_norm', type=float, default=1.0)
548
549 # GRPO specific params from global_config_module
550 parser.add_argument('--beta', type=float,
    default=global_config_module.BETA)
551 parser.add_argument('--epsilon', type=float,
    default=global_config_module.EPSILON)
552 parser.add_argument('--mu', type=int,
    default=global_config_module.MU)
553
554 parser.add_argument('--tune_hyperparameters',
    action='store_true', help='Run hyperparameter tuning with
    Optuna')
555 parser.add_argument('--n_trials', type=int, default=20,
    help='Number of trials for hyperparameter tuning')
556 args = parser.parse_args()
557
558 # Setup
559 random.seed(42); np.random.seed(42); torch.manual_seed(42)
560 if torch.cuda.is_available(): torch.cuda.manual_seed_all(42) #
    For GPU reproducibility
561
562 # Tokenizers & Model
563 tok_p = get_tokenizer("bpe", global_config_module.VOCAB_SIZE,
    global_config_module.MAX_LENGTH,
    tokenizer_path=args.protein_tokenizer)
564 tok_r = get_tokenizer("bpe", global_config_module.VOCAB_SIZE,
    global_config_module.MAX_LENGTH,
    tokenizer_path=args.rna_tokenizer)
565
566 # Ensure pad token IDs are set if BpeTokenizer class doesn't
    handle it automatically after load
567 # The BpeTokenizer in 'classifier/util/tokenizer.py' and
    'grpo-trainer/custom_tokenizer.py'
568 # sets special tokens using a dict. Let's assume pad_token_id is
    accessible.
569 # The ProteinRNAEnvironment needs these.
570 # tok_p.pad_token_id = tok_p.tokenizer.token_to_id("<pad>") # Or
    however it's stored
571 # tok_r.pad_token_id = tok_r.tokenizer.token_to_id("<pad>")

```

```

572
573 model =
    T5ForConditionalGeneration.from_pretrained(args.model_path).to('cuda'
        if torch.cuda.is_available() else 'cpu')
574
575 # Data Envs
576 prots_train, rnas_train = load_data(args.train_data)
577 train_envs = [ProteinRNAEnvironment(p, r, tok_p, tok_r,
    args.max_input_length, args.max_completion_length) for p, r
    in zip(prots_train, rnas_train)]
578 eval_envs = []
579 if args.eval_data and os.path.exists(args.eval_data):
580     prots_eval, rnas_eval = load_data(args.eval_data)
581     eval_envs = [ProteinRNAEnvironment(p, r, tok_p, tok_r,
        args.max_input_length, args.max_completion_length) for p,
        r in zip(prots_eval, rnas_eval)]
582 print(f"Created {len(train_envs)} training environments and
    {len(eval_envs)} evaluation environments.")
583
584 # Config for GRPO trainer
585 # Pass CLI args to override defaults in CustomGRPOConfig
586 grpo_trainer_config = CustomGRPOConfig(
587     output_dir=args.output_dir,
588     learning_rate=args.learning_rate,
589     max_prompt_length=args.max_input_length,
590     max_completion_length=args.max_completion_length,
591     num_generations=args.num_generations,
592     top_p=args.top_p,
593     temperature=args.temperature,
594     max_grad_norm=args.max_grad_norm,
595     BETA=args.beta,
596     EPSILON=args.epsilon,
597     MU=args.mu
598 )
599
600 if args.tune_hyperparameters:
601     print("Starting hyperparameter tuning...")
602     from hyperparameter_tuning import run_hyperparameter_tuning #
        Assuming it's in the same directory
603
604     # The hyperparameter_tuning script will modify the BETA,
        EPSILON, MU attributes
605     # of the config object passed to it.
606     best_params = run_hyperparameter_tuning(
607         model=model,
608         tok_p=tok_p,
609         tok_r=tok_r,
610         train_envs=train_envs,
611         eval_envs=eval_envs,
612         tuning_config=grpo_trainer_config, # Pass the trainer's
            config object
613         n_trials=args.n_trials
614     )
615
616     # Update grpo_trainer_config with best parameters found by
        Optuna study
617     # The run_hyperparameter_tuning function already modifies the
        tuning_config object's attributes
618     # if designed to do so, or returns best_params.
619     # If it returns params, apply them:
620     if best_params:
621         grpo_trainer_config.BETA = best_params.get('beta',
            grpo_trainer_config.BETA)
622         grpo_trainer_config.EPSILON = best_params.get('epsilon',
            grpo_trainer_config.EPSILON)

```

```

623         grpo_trainer_config.MU = best_params.get('mu',
624             grpo_trainer_config.MU)
625
626         print("Best hyperparameters from tuning:")
627         print(f"BETA: {grpo_trainer_config.BETA}, EPSILON:
628             {grpo_trainer_config.EPSILON}, MU:
629             {grpo_trainer_config.MU}")
630         wandb.config.update(best_params, allow_val_change=True) # Log
631         best params to current wandb run
632
633     # Trainer
634     print("Starting training with config:")
635     print(vars(grpo_trainer_config)) # Print the config being used
636
637     trainer = ProteinRNAGRPOTrainer(model, tok_p, tok_r, train_envs,
638         eval_envs, grpo_trainer_config)
639     trainer.train(epochs=args.epochs)
640
641     if eval_envs: # Final evaluation after training
642         print("Performing final evaluation...")
643         final_val_loss, final_val_reward =
644             trainer.evaluate(step=trainer.config.num_train_epochs *
645                 len(train_envs)) # Approx global step
646         print(f"Final Validation: loss={final_val_loss:.4f}
647             reward={final_val_reward:.4f}")
648
649     if wandb.run: # Ensure wandb run is finished
650         wandb.finish()

```

File: grpo-trainer/reward_functions.py

Listing 14: Code Snippet

```

1  import numpy as np
2  from typing import List, Tuple, Dict, Any, Optional, Union
3  import RNA # ViennaRNA package for MFE calculation
4  import config as global_config_module # Use an alias
5  import get_binding_score
6
7  def calculate_gc_content(sequence: str) -> float:
8      if not sequence: # Check for empty sequence
9          return 0.0
10     gc_count = sequence.upper().count('G') +
11         sequence.upper().count('C') # Ensure uppercase
12     total_count = len(sequence)
13     return gc_count / total_count if total_count > 0 else 0.0
14
15  def calculate_gc_content_reward(sequence: str) -> float:
16      """
17      Calculate reward based on GC content.
18      Args: sequence: RNA sequence (ACGU)
19      Returns: Reward value between 0 and 1
20      """
21     if not sequence: return 0.0
22     gc_content = calculate_gc_content(sequence)
23
24     gc_min = global_config_module.GC_CONTENT_MIN
25     gc_max = global_config_module.GC_CONTENT_MAX
26
27     if gc_min <= gc_content <= gc_max:
28         return 1.0
29     else:
30         if gc_content < gc_min:
31             distance = gc_min - gc_content

```

```

31         max_distance_possible = gc_min # from 0 to gc_min
32         # Reward decreases linearly from 1 (at gc_min) to 0 (at 0
          GC content)
33         # Ensure max_distance_possible is not zero
34         return max(0.0, 1.0 - (distance / max_distance_possible
          if max_distance_possible > 1e-6 else 1.0))
35     else: # gc_content > gc_max
36         distance = gc_content - gc_max
37         max_distance_possible = 1.0 - gc_max # from gc_max to 1.0
38         # Reward decreases linearly from 1 (at gc_max) to 0 (at
          100% GC content)
39         return max(0.0, 1.0 - (distance / max_distance_possible
          if max_distance_possible > 1e-6 else 1.0))
40
41
42 def calculate_mfe(sequence: str) -> float:
43     """
44     Calculate Minimum Free Energy for an RNA sequence using ViennaRNA.
45     Args: sequence: RNA sequence (ACGU)
46     Returns: MFE value (lower is better for stability)
47     """
48     if not sequence or not all(c in "ACGU" for c in
          sequence.upper()): # Check for valid RNA sequence
49         # print(f"Warning: Invalid RNA sequence for MFE: {sequence}")
50         return 0.0 # Or handle error appropriately, MFE is usually
          negative
51
52     try:
53         # ViennaRNA might be sensitive to non-ACGU characters or very
          short sequences
54         if len(sequence) < 4: # Arbitrary minimum length for folding,
          adjust as needed
55             return 0.0
56         _, mfe_val = RNA.fold(sequence.upper()) # Use a different
          name, mfe_val, ensure uppercase
57         return mfe_val
58     except Exception as e:
59         # print(f"Error calculating MFE for '{sequence}': {e}")
60         return 0.0 # Return a neutral value or error indicator
61
62 def calculate_mfe_reward(sequence: str) -> float:
63     """
64     Calculate reward based on MFE value. Normalized by length.
65     Lower MFE means more stable structure, which is better.
66     Args: sequence: RNA sequence (ACGU)
67     Returns: Reward value between 0 and 1
68     """
69     if not sequence:
70         return 0.0
71
72     mfe_val = calculate_mfe(sequence) # Renamed mfe to mfe_val
73
74     if len(sequence) == 0: # Should have been caught by
          calculate_mfe, but double check
75         return 0.0
76
77     length_normalized_mfe = mfe_val / len(sequence)
78
79     # Target is highly negative MFE (e.g., -0.5 kcal/mol/nt is good)
80     # Let's define a target normalized MFE, e.g., -0.5
81     # and a worst MFE, e.g., 0 or slightly positive.
82     # Reward can be 1 if length_normalized_mfe <= target_norm_mfe
83     # Reward can be 0 if length_normalized_mfe >= worst_norm_mfe
84     # Linear interpolation in between.
85

```

```

86     target_norm_mfe = -0.5 # Example: very stable
87     worst_norm_mfe  =  0.0 # Example: unstable or unfolded
88
89     if length_normalized_mfe <= target_norm_mfe:
90         return 1.0
91     elif length_normalized_mfe >= worst_norm_mfe:
92         return 0.0
93     else: # target_norm_mfe < length_normalized_mfe < worst_norm_mfe
94         # Linearly map from [worst_norm_mfe, target_norm_mfe] to [0,
95             1]
96         # Since target is more negative, (val - worst) / (target -
97             worst)
98         return (length_normalized_mfe - worst_norm_mfe) /
99             (target_norm_mfe - worst_norm_mfe)
100
101 def calculate_binding_affinity(protein_sequence: str, rna_sequence:
102     str) -> float:
103     """
104     Placeholder for binding affinity calculation.
105     Args: protein_sequence, rna_sequence (ACGU)
106     Returns: Binding affinity score between 0 and 1
107     """
108     # This should be replaced with a real model prediction
109     if not protein_sequence or not rna_sequence:
110         return 0.0
111
112     return get_binding_score(protein_sequence, rna_sequence)
113
114 def calculate_token_level_reward(
115     protein_sequence: str,
116     partial_rna_sequence: str, # ACGU format
117     position: int, # Current length of partial_rna_sequence
118     max_length: int # Max possible length of RNA
119 ) -> float:
120     """
121     Calculate combined reward for a partial RNA sequence at a
122     specific position.
123     """
124     # Ensure partial_rna_sequence is not empty for component reward
125     functions
126     if not partial_rna_sequence:
127         return 0.0
128
129     gc_reward_val = calculate_gc_content_reward(partial_rna_sequence)
130     # Renamed gc_reward
131     mfe_reward_val = calculate_mfe_reward(partial_rna_sequence) #
132     Renamed mfe_reward
133     binding_reward_val = calculate_binding_affinity(protein_sequence,
134         partial_rna_sequence) # Renamed binding_reward
135
136     # Weighting rewards (can be dynamic based on progress if desired)
137     # progress = position / max_length if max_length > 0 else 0
138     # Example: fixed weights from global config
139     combined_reward_val = ( # Renamed combined_reward
140         global_config_module.BINDING_AFFINITY_WEIGHT *
141         binding_reward_val +
142         global_config_module.GC_CONTENT_WEIGHT * gc_reward_val +
143         global_config_module.MFE_WEIGHT * mfe_reward_val
144     )
145
146     # Normalize combined reward to be between 0 and 1 if sum of
147     weights can exceed 1

```

```

138     # sum_weights = global_config_module.BINDING_AFFINITY_WEIGHT +
        global_config_module.GC_CONTENT_WEIGHT +
        global_config_module.MFE_WEIGHT
139     # if sum_weights > 0:
140     #     combined_reward_val = combined_reward_val / sum_weights
141     # else:
142     #     combined_reward_val = 0.0
143
144     return max(0.0, min(1.0, combined_reward_val)) # Ensure it's
        within [0,1]
145
146
147 def rewarder(partial_rna_acgu: str, protein_sequence_str: str,
protein_ids=None, rna_ids=None) -> float: # Renamed params
148     """
149     Compute the GRPO token-level reward for a partial RNA sequence
        (ACGU)
150     given its protein prompt.
151     """
152     current_pos = len(partial_rna_acgu) # Renamed position to
        current_pos
153     max_len_val = global_config_module.MAX_COMPLETION_LENGTH #
        Renamed max_len
154
155     # Compute the combined token level reward
156     reward_value = calculate_token_level_reward( # Renamed reward to
        reward_value
157         protein_sequence=protein_sequence_str, # Pass renamed param
158         partial_rna_sequence=partial_rna_acgu, # Pass renamed param
        (ACGU)
159         position=current_pos,
160         max_length=max_len_val
161     )
162     return float(reward_value)

```

File: grpo-trainer/.DS_Store

This is a system file (.DS_Store) generated by macOS. It typically contains custom attributes of its containing folder, such as the position of icons or the choice of a background image. Its content is not relevant as source code and is therefore omitted.

Directory: grpo-trainer/util/

File: grpo-trainer/util/create_val.py

Listing 15: Code Snippet

```

1 import pandas as pd
2 import random
3
4 # Set random seed for reproducibility
5 random.seed(42)
6
7 # Read the CSV file
8 # Ensure 'positive_sequences.csv' is in the correct path or provide
    full path
9 try:
10     df = pd.read_csv('positive_sequences.csv')
11 except FileNotFoundError:
12     print("Error: 'positive_sequences.csv' not found. Please check
        the path.")
13     exit()
14

```

```

15
16 # Remove rows with missing RNA or protein sequences
17 df = df.dropna(subset=['protein_seq', 'rna_sequence'])
18
19 # Keep only protein_seq and rna_seq columns
20 df_selected = df[['protein_seq', 'rna_sequence']].copy() # Use
    .copy() to avoid SettingWithCopyWarning
21
22 # Find the longest RNA sequence (on the selected dataframe)
23 if not df_selected.empty:
24     df_selected['rna_length'] =
        df_selected['rna_sequence'].astype(str).str.len() # Ensure
        rna_sequence is string
25     max_rna_length = df_selected['rna_length'].max()
26     # Get the first such longest sequence if multiple exist
27     longest_rna_series = df_selected.loc[df_selected['rna_length'] ==
        max_rna_length, 'rna_sequence']
28     longest_rna = longest_rna_series.iloc[0] if not
        longest_rna_series.empty else "N/A"
29 else:
30     max_rna_length = 0
31     longest_rna = "N/A"
32     print("DataFrame is empty after dropping NA values.")
33
34
35 # Randomly select 100 rows for validation, if enough rows exist
36 num_samples = min(100, len(df_selected))
37 if num_samples > 0:
38     validation_df = df_selected.sample(n=num_samples,
        random_state=42).drop(columns=['rna_length'], errors='ignore')
39     # Save to validation_rl.txt (protein,rna format without header
        for some tools)
40     # The original inference.py splits by ',', so CSV format is
        expected
41     validation_df.to_csv('validation_rl.txt', index=False,
        header=False) # header=False if inference.py expects no header
42     print(f"Successfully created validation_rl.txt with
        {len(validation_df)} sequences.")
43 else:
44     print("Not enough data to create validation_rl.txt")
45
46
47 # Print summary
48 print(f"Original dataset had {len(df)} rows, {len(df_selected) if
    'df_selected' in locals() else 0} valid sequences after
    cleaning.")
49 print(f"Longest RNA sequence length: {max_rna_length}")
50 print(f"Example of longest RNA sequence: {str(longest_rna)[:50]}...")

```

File: grpo-trainer/util/create_validation.py

Listing 16: Code Snippet

```

1 import pandas as pd
2 import random
3
4 # Set random seed for reproducibility
5 random.seed(42)
6
7 # Read the CSV file
8 # Ensure 'positive_sequences.csv' is in the correct path or provide
    full path
9 try:
10     df = pd.read_csv('positive_sequences.csv')

```

```

11 except FileNotFoundError:
12     print("Error: 'positive_sequences.csv' not found. Please check
13           the path.")
14     exit()
15
16 # Remove rows with missing protein sequences (and rna_seq for safety,
17   though not explicitly checked before)
18 df = df.dropna(subset=['protein_seq', 'rna_seq'])
19
20 # Keep only protein_seq and rna_seq columns
21 df_selected = df[['protein_seq', 'rna_seq']].copy() # Use .copy() to
22   avoid SettingWithCopyWarning
23
24 # Find the longest RNA sequence
25 if not df_selected.empty:
26     # Ensure rna_seq is treated as string for .str accessor
27     df_selected['rna_length'] =
28         df_selected['rna_seq'].astype(str).str.len()
29     max_rna_length = df_selected['rna_length'].max()
30     longest_rna_series = df_selected.loc[df_selected['rna_length'] ==
31         max_rna_length, 'rna_seq']
32     longest_rna = longest_rna_series.iloc[0] if not
33         longest_rna_series.empty else "N/A"
34 else:
35     max_rna_length = 0
36     longest_rna = "N/A"
37     print("DataFrame is empty after dropping NA values.")
38
39 # Randomly select 100 rows, if enough rows exist
40 num_samples = min(100, len(df_selected))
41 if num_samples > 0:
42     validation_df = df_selected.sample(n=num_samples,
43         random_state=42).drop(columns=['rna_length'], errors='ignore')
44     # Save to validation_rl.txt (protein,rna format, no header)
45     # inference.py in the main directory splits by ',', so CSV format
46     # is fine.
47     validation_df.to_csv('validation_rl.txt', index=False,
48         header=False) # header=False if inference.py expects no header
49
50     print(f"Successfully created validation_rl.txt with
51           {len(validation_df)} sequences.")
52 else:
53     print("Not enough data to create validation_rl.txt")
54
55 print(f"Original dataset had {len(df)} rows, {len(df_selected)} if
56       'df_selected' in locals() else 0} valid sequences after
57       cleaning.")
58 print(f"Longest RNA sequence length: {max_rna_length}")
59 print(f"Example of longest RNA sequence: {str(longest_rna)[:50]}...")
60     # Print first 50 characters

```

File: grpo-trainer/util/generate_data.py

Listing 17: Code Snippet

```

1 import os
2 import csv
3 from tqdm import tqdm # progress bar library
4
5 #

```

```

6 # 1) Define paths
7 #
8 -----
8 FASTA_FOLDER =
9     "/data6/helya/dataset/CLIPdb_cluster/cd_hit_results_RBPs/identity_90"
9 PROTEIN_CSV = "/data6/alpsencer/reinforce_rna/protein_seqs.csv" #
10     CSV with columns: prot_name, seq
10 OUTPUT_CSV = "./output.csv" # Output in the current directory of
11     this script
11
12 #
13 -----
13 # 2) Read the protein sequences into a dictionary (case-insensitive)
14 #
15 -----
15 protein_dict = {}
16 try:
17     with open(PROTEIN_CSV, "r", newline="") as f_in:
18         reader = csv.DictReader(f_in)
19         if 'prot_name' not in reader.fieldnames or 'seq' not in
20             reader.fieldnames:
21                 print(f"Error: PROTEIN_CSV '{PROTEIN_CSV}' must have
22                     'prot_name' and 'seq' columns.")
23                 exit()
24         for row in reader:
25             prot_name = row["prot_name"].strip().lower() # using
26                 lowercase keys
27             protein_dict[prot_name] = row["seq"]
28 except FileNotFoundError:
29     print(f"Error: Protein CSV file not found at '{PROTEIN_CSV}'")
30     exit()
31 #
32 -----
30 # 3) Process FASTA files and write to CSV
31 #
32 -----
32 if not os.path.isdir(FASTA_FOLDER):
33     print(f"Error: FASTA folder not found at '{FASTA_FOLDER}'")
34     exit()
35
36 with open(OUTPUT_CSV, "w", newline="") as f_out:
37     writer = csv.writer(f_out)
38     writer.writerow(["protein_name", "protein_seq", "rna_sequence",
39         "strand", "pair"])
40
41     # Get all .fa files
42     fasta_files = [f for f in os.listdir(FASTA_FOLDER) if
43         f.endswith(".fa") or f.endswith(".fasta")]
44
45     if not fasta_files:
46         print(f"No FASTA files found in '{FASTA_FOLDER}'.")
47
48     # Loop over the files with a progress bar
49     for file_name in tqdm(fasta_files, desc="Processing FASTA files"):
50         # Derive protein name from file name (first token before
51             underscore)
52         protein_name_raw = file_name.split("_")[0]
53         protein_seq = protein_dict.get(protein_name_raw.lower(),
54             "NA") # Default to "NA" if not found
55         fa_path = os.path.join(FASTA_FOLDER, file_name)
56
57         current_rna_sequence = ""
58         header_info = "" # Store the most recent header

```

```

56     # Read all lines from the FASTA file
57     try:
58         with open(fa_path, "r") as fa_in:
59             for line in fa_in:
60                 line = line.strip()
61                 if not line: continue # Skip empty lines
62
63                 if line.startswith(">"):
64                     # If we have a sequence buffered, process it
65                     # with previous header
66                     if current_rna_sequence and header_info:
67                         strand = "NA"
68                         # Process header_info for strand
69                         temp_header = header_info
70                         if temp_header.endswith("(+)"):
71                             strand = "+"
72                             temp_header = temp_header[:-3].strip()
73                         elif temp_header.endswith("(-)"):
74                             strand = "-"
75                             temp_header = temp_header[:-3].strip()
76                         writer.writerow([protein_name_raw,
77                                         protein_seq, current_rna_sequence,
78                                         strand, "+"])
79
80                     # Start new sequence
81                     header_info = line[1:] # Store new header
82                     # without ' >'
83                     current_rna_sequence = ""
84                 else:
85                     current_rna_sequence += line # Append
86                     # sequence lines
87
88             # Process the last sequence in the file
89             if current_rna_sequence and header_info:
90                 strand = "NA"
91                 temp_header = header_info
92                 if temp_header.endswith("(+)"):
93                     strand = "+"
94                     temp_header = temp_header[:-3].strip()
95                 elif temp_header.endswith("(-)"):
96                     strand = "-"
97                     temp_header = temp_header[:-3].strip()
98                 writer.writerow([protein_name_raw, protein_seq,
99                                 current_rna_sequence, strand, "+"])
100
101     except Exception as e:
102         print(f"Error processing file {file_name}: {e}")
103         continue # Skip to next file
104
105     # Optional: print progress for each file (can be too verbose
106     # with tqdm)
107     # print(f"Processed file: {file_name}")
108 print(f"Processing complete. Output written to {OUTPUT_CSV}")

```

File: grpo-trainer/util/test_tokenizer.py

Listing 18: Code Snippet

```

1 from transformers import PreTrainedTokenizerFast
2 from itertools import islice
3 import os # For checking file existence
4
5 TOKENIZER_FILE_PATH =
    "/data6/alpsencer/tokenizers/bpe_protein_1000_1024.json"

```

```

6 TRAIN_DATA_PATH = "/data6/alpsencer/reinforce_rna/sample_1K.txt" #
  Renamed TRAIN_DATA to TRAIN_DATA_PATH
7
8
9 if not os.path.exists(TOKENIZER_FILE_PATH):
10     print(f"Error: Tokenizer file not found at {TOKENIZER_FILE_PATH}")
11     exit()
12
13 tokenizer = PreTrainedTokenizerFast(
14     tokenizer_file=TOKENIZER_FILE_PATH,
15     unk_token="",
16     pad_token="",
17     eos_token="</s>", # EOS token if your model uses it for proteins
18     bos_token="<s>", # BOS token if your model uses it
19     model_max_length=1024,
20     padding_side="left", # Ensure this matches model's expectation
  for T5 encoder
21 )
22
23 # T5 usually expects a sentence-piece like format, often with EOS at
  the end of input.
24 # If your BPE tokenizer was trained without explicit EOS for protein
  inputs,
25 # this setup is fine. The main thing is consistency with training.
26
27 if not os.path.exists(TRAIN_DATA_PATH):
28     print(f"Error: Train data file not found at {TRAIN_DATA_PATH}")
29     exit()
30
31 print(f"Using tokenizer: {TOKENIZER_FILE_PATH}")
32 print(f"Pad token: '{tokenizer.pad_token}', ID:
  {tokenizer.pad_token_id}")
33 print(f"EOS token: '{tokenizer.eos_token}', ID:
  {tokenizer.eos_token_id}")
34 print(f"UNK token: '{tokenizer.unk_token}', ID:
  {tokenizer.unk_token_id}")
35 print(f"BOS token: '{tokenizer.bos_token}', ID:
  {tokenizer.bos_token_id}")
36
37
38 with open(TRAIN_DATA_PATH) as f:
39     for line_num, line_content in enumerate(islice(f, 5)): # Test a
  few lines
40         line_content = line_content.strip()
41         if not line_content or "$" not in line_content:
42             print(f"Skipping malformed line {line_num+1}:
  {line_content}")
43             continue
44
45         # Assuming format is protein$rna or protein$something_else
46         protein_seq, _ = line_content.split("$", 1) # Renamed protein
  to protein_seq
47
48         # For T5, inputs are often not lowercased by default unless
  normalizer does it.
49         # The custom BpeTokenizer class had Lowercase normalizer.
50         # PreTrainedTokenizerFast directly from file might or might
  not apply it
51         # depending on how the tokenizer.json was saved. Assume it's
  handled or not critical for this test.
52
53         # Some tokenizers add special tokens automatically, some
  don't.
54         # For T5, usually you'd add EOS to the protein input if it
  was trained that way.

```

```

55     # protein_input_for_tokenizer = protein_seq +
56         tokenizer.eos_token # Example if EOS is needed
57     protein_input_for_tokenizer = protein_seq
58
59     enc = tokenizer(
60         protein_input_for_tokenizer,
61         truncation=True,
62         padding="max_length", # Pad to model_max_length
63         max_length=tokenizer.model_max_length, # Use tokenizer's
64             max_length
65         return_tensors=None, # Get lists of IDs
66         add_special_tokens=True # Let tokenizer handle BOS/EOS
67             based on its config
68     )
69
70     input_ids_list = enc["input_ids"] # Renamed ids to
71         input_ids_list
72     attention_mask_list = enc["attention_mask"] # Renamed mask to
73         attention_mask_list
74
75     # Calculate real length based on attention mask
76     # For left padding, actual tokens are at the end.
77     try:
78         first_real_token_idx = attention_mask_list.index(1)
79         real_len = len(input_ids_list) - first_real_token_idx
80         real_ids = input_ids_list[first_real_token_idx:]
81     except ValueError: # If no 1 in attention_mask (e.g. all
82         padding)
83         real_len = 0
84         real_ids = []
85
86
87     print(f"\n=== PROTEIN PROMPT (Line {line_num+1}) ===")
88     print("Original:", protein_seq[:60] + ("..." if
89         len(protein_seq) > 60 else ""))
90     # print("Input to tokenizer:",
91         protein_input_for_tokenizer[:60] + ("..." if
92             len(protein_input_for_tokenizer) > 60 else ""))
93
94     print("\nTokenized IDs (first 10 of padded):",
95         input_ids_list[:10])
96     print("Tokenized IDs (last 10 of padded):",
97         input_ids_list[-10:])
98     print("Attention Mask (first 10):", attention_mask_list[:10])
99     print("Attention Mask (last 10):", attention_mask_list[-10:])
100
101     print("\n[Padded left      real tokens at end]")
102     if real_ids:
103         print("Real IDs (first 10):", real_ids[:10])
104         print("Real IDs (last 10):", real_ids[-10:])
105         print("Real Tokens (first 10):",
106             tokenizer.convert_ids_to_tokens(real_ids[:10]))
107         print("Real Tokens (last 10):",
108             tokenizer.convert_ids_to_tokens(real_ids[-10:]))
109         print("Real token count:", real_len)
110         decoded_real = tokenizer.decode(real_ids,
111             skip_special_tokens=False) # See special tokens
112         decoded_real_skip_special = tokenizer.decode(real_ids,
113             skip_special_tokens=True)
114         print("Decoded (real tokens, keep special):",
115             decoded_real[:60] + ("..." if len(decoded_real) > 60
116                 else ""))
117         print("Decoded (real tokens, skip special):",
118             decoded_real_skip_special[:60] + ("..." if
119                 len(decoded_real_skip_special) > 60 else ""))

```

```
101     else:
102         print("No real tokens found (all padding or empty
              input).")
103     print("
104     # break # remove 'break' if you want to sample more lines
```
