

Contents

1	Extended Data	2
2	Evaluation Metrics for Protein and RNA Design	4
2.1	Diversity Score	4
2.2	OpenKnot Score	4
2.3	Eterna Score	5
2.4	Template Modeling Score	5
2.5	Root Mean Square Deviation	6

a.

Iteration

[M] [M] [M] [M] [M] [M]

RDdiffusion

[M] [M] [M] [M] [M] [M]

A G C A C U

Recovered sequence

Ribonanza Net

OpenKnot\ Eterna Score

Secondary Structure

AlphaFold3

Tertiary Structure

b.

OpenKnot Score Distribution

OpenKnot Score

Sequence Length (nt)

20 50 100 200 500

41.0 41.4 37.9 32.3 28.7

c.

Eterna Score Distribution

Eterna Score

Sequence Length (nt)

20 50 100 200 500

81.9 86.8 78.2 64.5 57.5

d.

Diversity Score Distribution

Diversity

Sequence Length (nt)

20 50 100 200 500

0.949 0.904 0.869 0.824 0.753

e.

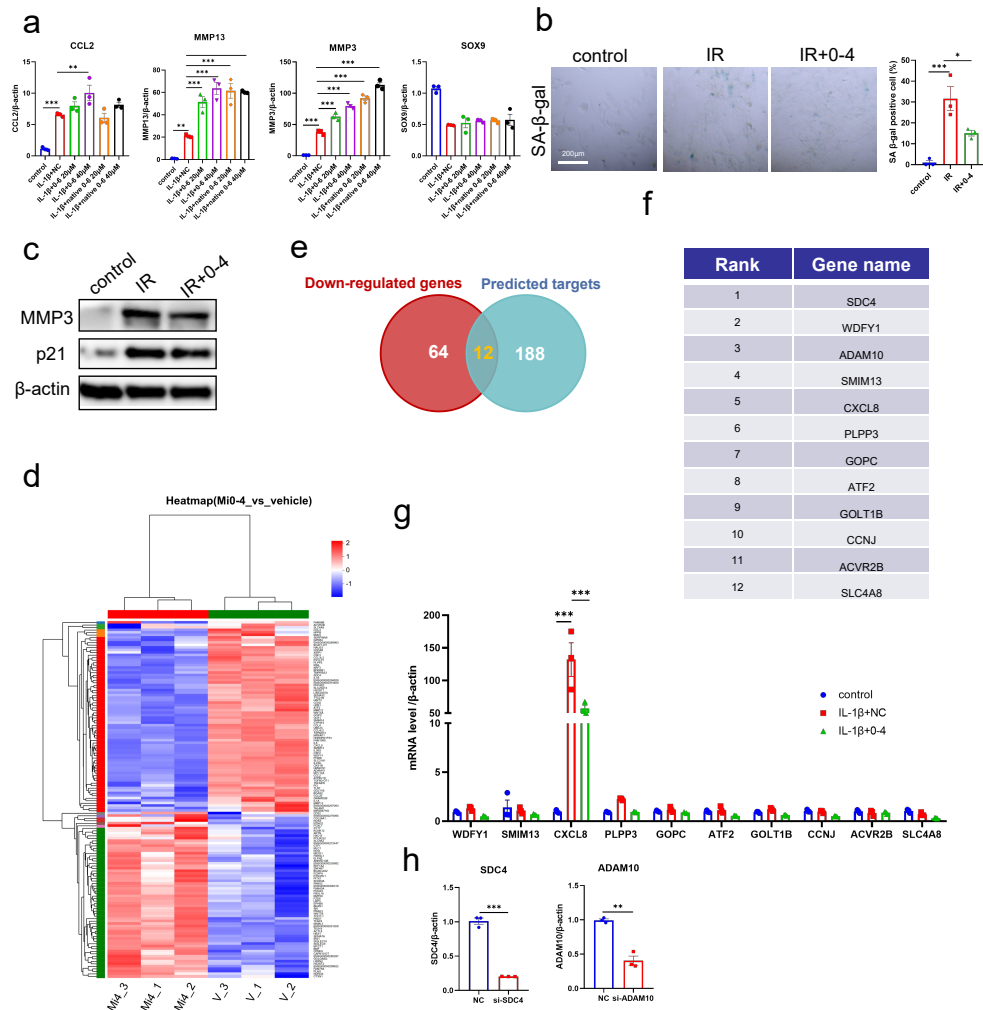
47.9 92.1

53.6 83.4

Length: 50
Avg pLDDT: 74.10

Length: 20
Avg pLDDT: 74.21

Extended Data Fig. 1 Unconditional RNA design of RDiffusion. **a** Workflow of RDiffusion for unconditional RNA design; **b** The OpenKnot score distribution across RDiffusion’s generated RNA sequences in five lengths; **c** The Eterna score distribution across RDiffusion’s generated RNA sequences in five lengths; **d** The diversity distribution across RDiffusion’s generated RNA sequences in five lengths; **e** Two AlphaFold3-predicted tertiary structures for RNA sequences designed by RDiffusion.



Extended Data Fig. 2 De novo RDiffusion-designed miRNA 0-4 attenuates osteoarthritis by suppressing SDC4. **a** RT-qPCR analysis of ECM-related anabolic and catabolic genes and inflammatory cytokines in SW1353 cells transfected with increasing doses of mimic 0-6 or the native miRNA prior to IL-1 β stimulation. **b-c** SW1353 cells were transfected with mimic 0-4 and subsequently subjected to ionizing radiation (IR)-induced senescence. SA- β -gal staining (**b**) and immunoblot analysis of the senescence marker p21 and the cartilage catabolic marker MMP3 (**c**). **d** Hierarchical clustering analysis of differentially expressed genes in the vehicle + IL-1 β and mimic 0-4 + IL-1 β (Mi0-4) groups. **e** Venn diagram showing the overlap between the 76 downregulated differentially expressed genes identified by RNA sequencing and the top 200 predicted target genes from the miRDB database. **f** Identification of 12 candidate target genes based on the overlap analysis shown in (**e**). **g** RT-qPCR validation of the 10 candidate target genes identified in (**e-f**) in SW1353 cells transfected with increasing doses of mimic 0-4 prior to IL-1 β stimulation. **h** RT-qPCR analysis of SDC4 and ADAM10 knockdown efficiency following siRNA transfection. Data are shown as mean $\pm$ SEM. * p < 0.05, ** p < 0.01, *** p < 0.001.

2 Evaluation Metrics for Protein and RNA Design

2.1 Diversity Score

To improve the success rate of protein design, it is important to explore a set of protein sequences rather than placing a bet on a single sequence. In this case, generating diverse sequences is crucial for exploring the reasonable protein sequence space. We define the pairwise diversity as the average Hamming distance between two designed sequences.

Let $r_{i,l}$ indicate the l -th residue (amino acid) of the i -th designed sequence, and let n be the total length of each sequence. The pairwise diversity between sequence i and sequence j is defined as:

$$D_{ij} = \frac{\sum_{l=1}^n \mathbf{1}_{r_{i,l} \neq r_{j,l}}}{n},$$

where $\mathbf{1}_{\text{condition}}$ is the indicator function that equals 1 when the condition is true and 0 otherwise. This measures the proportion of positions where the two sequences differ.

The overall diversity score for a set of m designed sequences is then calculated as the average of all pairwise diversities:

$$\text{Div} = \frac{1}{m^2} \sum_{i=1}^m \sum_{j=1}^m D_{ij},$$

where $i, j \in \{1, 2, 3, \dots, m\}$. By default, we set $m = 10$. A higher Div value indicates a more diverse set of sequences. However, measuring diversity alone without combining it with other metrics may be misleading. For example, a high diversity indicates a low recovery rate, more likely to result in a low structural similarity.

2.2 OpenKnot Score

The OpenKnot score is a metric used to evaluate the complexity and correctness of protein backbone topologies, specifically identifying entangled or knotted structures. Proteins that form knots are rare but functionally significant. The OpenKnot algorithm detects knots by smoothing the protein backbone and analyzing the chain's topological state.

The OpenKnot score typically ranges from 0 to 1, where a score of 0 indicates no knot (an open or trivial knot), and scores approaching 1 indicate complex, deeply embedded knots. The score is calculated based on the minimal number of crossings or the Alexander polynomial coefficients:

$$\text{OpenKnot} = 1 - \frac{\min(\text{Crossings}_{\text{observed}}, \text{Crossings}_{\text{threshold}})}{\text{Crossings}_{\text{threshold}}},$$

or more formally, using the absolute value of the Alexander polynomial $\Delta_K(t)$ evaluated at $t = -1$:

$$\text{OpenKnot} = \frac{|\Delta_K(-1)| - 1}{C_{\text{max}}},$$

where $C_{\max}$ is a normalization constant. A higher OpenKnot score suggests the presence of complex topological features, while a score of zero suggests a knot-free or ideal open-chain conformation. Knots are classified as 3_1 (trefoil), 4_1 (figure-eight), 5_2 , and more complex topologies.

2.3 Eterna Score

The Eterna score is derived from RNA design challenges, evaluating the thermodynamic stability and structural correctness of designed RNA sequences. It integrates multiple measures including ensemble defect, base-pairing probability, and minimum free energy (MFE) distance to the target secondary structure.

For a designed RNA sequence S with target secondary structure T , the Eterna score is defined as:

$$\text{Eterna} = w_1 \cdot (1 - \text{EnsembleDefect}) + w_2 \cdot \text{BasePairPrecision} + w_3 \cdot (1 - \text{MFEDistance}),$$

where the ensemble defect is the fraction of nucleotides that are incorrectly base-paired on average across the Boltzmann ensemble:

$$\text{EnsembleDefect} = \frac{1}{n} \sum_{i=1}^n \left(1 - \sum_{j=1}^n P_{ij} \cdot \mathbf{1}_{(i,j) \in T} \right),$$

with P_{ij} being the base-pairing probability between positions i and j , and n being the sequence length. The base-pair precision measures the fraction of predicted base pairs that match the target. The MFE distance is the normalized number of incorrectly paired nucleotides in the minimum free energy structure. The weights w_1, w_2, w_3 are typically set to $[0.5, 0.3, 0.2]$. A higher Eterna score (ranging from 0 to 1) indicates that the designed sequence is more likely to fold into the desired conformation with minimal off-target structures.

2.4 Template Modeling Score

The TM-score is a measure of structural similarity between two protein structures, typically between a predicted model and a native (ground truth) structure. Unlike RMSD, the TM-score normalizes by the length of the protein and is less sensitive to local deviations, making it more robust for comparing structures of different sizes or with local flexibility.

For two aligned protein structures of length L_{target} (the length of the native protein), the TM-score is defined as:

$$\text{TM-score} = \max \left[\frac{1}{L_{\text{target}}} \sum_{i=1}^{L_{\text{common}}} \frac{1}{1 + \left(\frac{d_i}{d_0(L_{\text{target}})} \right)^2} \right],$$

where:

- L_{common} is the number of aligned residue pairs after optimal superposition,

- d_i is the Euclidean distance between the i -th pair of aligned C_α atoms,
- $d_0(L_{\text{target}})$ is a length-dependent scaling factor defined as:

$$d_0(L_{\text{target}}) = 1.24 \sqrt[3]{L_{\text{target}} - 15} - 1.8.$$

The TM-score ranges from 0 to 1. A TM-score greater than 0.5 generally indicates the same protein fold (significant structural similarity), a score greater than 0.7 indicates very high similarity, and a score greater than 0.9 suggests nearly identical structures. The "max" operator indicates that the score is computed after finding the optimal alignment that maximizes the sum.

2.5 Root Mean Square Deviation

RMSD measures the average distance between the backbone atoms (typically C_α atoms) of two superimposed protein structures. It is the most commonly used metric for quantifying structural similarity, though it is sensitive to local outliers and domain motions.

After performing optimal least-squares superposition of two structures (a predicted model and a native reference), the RMSD is defined as:

$$\text{RMSD} = \sqrt{\frac{1}{N} \sum_{i=1}^N \|\mathbf{x}_i - \mathbf{y}_i\|^2},$$

where:

- N is the number of aligned residue pairs (typically all residues or a selected subset),
- $\mathbf{x}_i = (x_i, y_i, z_i)$ is the coordinate vector of the i -th atom in the predicted model,
- $\mathbf{y}_i = (y_i, y_i, z_i)$ is the coordinate vector of the corresponding atom in the native structure,
- $\|\cdot\|$ denotes the Euclidean distance.

Equivalently, in expanded form:

$$\text{RMSD} = \sqrt{\frac{1}{N} \sum_{i=1}^N [(x_i^{\text{model}} - x_i^{\text{native}})^2 + (y_i^{\text{model}} - y_i^{\text{native}})^2 + (z_i^{\text{model}} - z_i^{\text{native}})^2]}.$$

A lower RMSD indicates higher structural similarity. Typical values: RMSD < 1 Å indicates nearly identical structures, 1 – 2 Å indicates very high similarity, 2 – 5 Å indicates moderate similarity with possible local differences, and > 5 Å indicates significant structural divergence. However, RMSD should be interpreted with caution for proteins of different sizes or with flexible domains.