

Sequence analysis

NEFFy: a versatile tool for computing the number of effective sequences

Maryam Haghani¹ , Debswapna Bhattacharya^{1,*} , T.M. Murali^{1,*} 

¹Department of Computer Science, Virginia Tech, Blacksburg, VA 24061, United States

*Corresponding author. Debswapna Bhattacharya, Department of Computer Science, Virginia Tech, Blacksburg, VA 24061, United States. E-mail: dbhattacharya@vt.edu; T.M. Murali, Department of Computer Science, Virginia Tech, Blacksburg, VA 24061, United States. E-mail: murali@cs.vt.edu.

Associate Editor: Jianlin Cheng

Abstract

Motivation: A Multiple Sequence Alignment (MSA) contains fundamental evolutionary information that is useful in the prediction of structure and function of proteins and nucleic acids. The “Number of Effective Sequences” (NEFF) quantifies the diversity of sequences of an MSA. While several tools embed NEFF calculation with various options, none are standalone tools for this purpose, and they do not offer all the available options.

Results: We developed NEFFy, the first software package to integrate all these options and calculate NEFF across diverse MSA formats for proteins, RNAs, and DNAs. It surpasses existing tools in functionality without compromising computational efficiency and scalability. NEFFy also offers per-residue NEFF calculation and supports NEFF computation for MSAs of multimeric proteins, with the capability to be extended to DNAs and RNAs.

Availability and implementation: NEFFy is released as open-source software under the GNU Public License v3.0. The source code in C++ and a Python wrapper are available at <https://github.com/Maryam-Haghani/NEFFy>. To ensure users can fully leverage these capabilities, comprehensive documentation and examples are provided at <https://Maryam-Haghani.github.io/NEFFy>.

1 Introduction

A Multiple Sequence Alignment (MSA) organizes a set of similar sequences by introducing gaps to ensure that all sequences are of the same length, l . The MSA contains each sequence in a row with l columns or positions, where each residue or gap occupies a distinct position. Computing an MSA involves maximizing the similarity in each position across the rows while minimizing the number of gaps. By uncovering evolutionarily conserved sequence patterns, regions of similarity that cannot be identified from a single sequence alone, MSAs are used in applications including contact map prediction (He *et al.* 2017, Wang *et al.* 2017, Adhikari *et al.* 2018, Liu *et al.* 2021), RNA and protein structure prediction (Baek *et al.* 2021, Jumper *et al.* 2021, Wang *et al.* 2023, Baek *et al.* 2024, Zheng *et al.* 2024), and protein function annotations (Zhang *et al.* 2017, Hu *et al.* 2022, Shao *et al.* 2025).

Classical MSA generation involves aligning a set of user-provided homologous sequences (Notredame *et al.* 2000, Chenna *et al.* 2003, Edgar 2004). Recent applications construct an MSA starting from a single query sequence by searching large databases with iterative methods such as PSI-BLAST (Altschul *et al.* 1997) or HHblits (Remmert *et al.* 2012) to retrieve similar sequences and align them with the query. Recent advancements in DNA/RNA sequencing technology have expanded public databases, enabling the generation of MSAs with high sequence diversity (Wilke *et al.* 2016,

Zhang *et al.* 2020). Such MSAs are generally believed to provide richer evolutionary and coevolutionary insights, and they can thereby enhance the effectiveness of models utilizing them for downstream tasks (Zheng *et al.* 2024). However, since MSAs can contain redundant sequences, the number of sequences by itself may not be an accurate reflection of their diversity. The concept of “Number of Effective Sequences,” NEFF, addresses this redundancy and assesses the quality of an MSA. Higher NEFF values often indicate a more diverse and informative MSA, leading to enhanced accuracy in predicting contact maps and the tertiary structures of proteins or RNA molecules (Wu *et al.* 2020, Pearce and Zhang 2021). For example, the accuracy of AlphaFold declines substantially when the NEFF value is below approximately 30 (Jumper *et al.* 2021). Additionally, for RNA structure prediction models such as trRosettaRNA, which utilize MSAs of RNAs as their input, prediction accuracy is correlated with NEFF (Wang *et al.* 2023), and for high-quality MSAs, these models can outperform other methods (Tarafder *et al.* 2024).

We introduce NEFFy, a fast and dedicated standalone tool for NEFF calculation. NEFFy is uniquely equipped to parse MSAs and calculate NEFF across a wide range of MSA formats for protein and nucleic acid sequences. It integrates all the features from earlier NEFF tools (see Table 1) and offers a set of new functionalities. NEFFy is developed in C++ for optimal performance and is also provided as a Python library that wraps the C++ executables. This approach enables

Table 1. An overview of NEFFy and other tools that incorporate NEFF calculation, highlighting their respective features (refer to [Supplementary Section S2.1](#) for detailed information on each feature).^a

Tool	Language	MSA format						Similarity		Query gaps	Gappy position	Multi alphabet	Non-standard residues		Validation
		a2m	a3m	sto	aln	fasta	clustal	pfam	sym.	asym.					
RaptorX (Källberg et al. 2012)	Python	✓	✓	×	×	×	×	×	×	✓	×	×	×	×	×
Conkit (Simkovic et al. 2017)	Cython	✓	✓	✓	×	✓	✓	×	×	×	×	×	×	×	×
DeepMSA (Zhang et al. 2020)	C++	×	×	×	✓	×	×	×	✓	✓	×	×	×	×	×
Gremlin (Kamisetty et al. 2013)	C++	×	×	×	✓	✓	×	×	✓	×	✓	✓	×	×	×
rMSA (Zhang et al. 2023)	C++	✓	×	×	✓	✓	×	×	✓	×	×	×	×	×	×
NEFFy	C++ & Python	✓	✓	✓	✓	✓	✓	✓	✓	✓	✓	✓	✓	✓	✓

^a **Language:** The programming language used to implement each tool. **MSA Format:** Formats used to represent aligned sequences in an MSA. **Similarity:** Methods for determining sequence similarity between pairs of sequences within the MSA, with two options: symmetric (sym.) and asymmetric (asym.). **Query Gaps:** Handling query gaps—defined as “gaps aligned to insertions”—can be customized during NEFF calculation based on user preference: either by removing them along with the corresponding positions in the aligned sequences or by keeping them intact. **Gappy position:** Inspired by Gremlin, this option addresses positions with a gap frequency exceeding the desired gap threshold. **Alphabet:** Alphabet used to represent a biological sequence. Of the tools mentioned, RaptorX and Conkit do not explicitly define an alphabet, while DeepMSA and rMSA use the protein alphabet. In contrast, Gremlin supports both protein and RNA sequences (implicitly including DNA). NEFFy accommodates proteins, RNAs, and DNAs. **Non-standard Residues:** Residues outside the standard set of biological sequence residues are handled differently across tools. RaptorX and Conkit do not support these residues, while rMSA and Gremlin treat them as gaps. DeepMSA considers them as standard symbols in its symmetric version but treats them as gaps when calculating similarity cutoff in its asymmetric version. NEFFy offers users the flexibility to customize the handling of these residues according to their preference. **Validation:** Indicates if validation is provided for the MSA file before the NEFF calculation.

* Conkit can manage query gaps for the a3m format by offering two distinct options: a3m-inserts and a3m.

seamless integration into Python-based workflows, simplifying use for a wider audience while preserving efficiency.

2 Materials and methods

The NEFF of an MSA M can be formulated as

$$\text{NEFF}(M) = \left(\frac{1}{\sqrt{L}} \right) \frac{1}{1 + \sum_{m=1, m \neq n}^N I[S_{m,n} \geq \theta]},$$

where L represents the length of the query sequence, N denotes the number of sequences in M , $S_{m,n}$ indicates the amount of sequence similarity between m th and n th sequences, θ is the similarity threshold (typically set to 0.8), and $I[\cdot]$ is the Iverson bracket, meaning that $I[S_{m,n} \geq \theta]$ equals 1 if $S_{m,n} \geq \theta$, and 0 otherwise.

This formula defines NEFF as the normalized sum of the weights of the sequences in the MSA. Specifically, each sequence i is assigned a weight of $\frac{1}{n_i}$, where n_i is the number of sequences (including itself) similar to it. This approach for calculating NEFF is widely used in many contact prediction and MSA generation methods ([Wu et al. 2020](#), [Zhang et al. 2020](#), [Li et al. 2021](#), [Liu et al. 2021](#)). Note that $\frac{1}{\sqrt{L}}$ is a normalization factor, which is a commonly used approach ([Zhang et al. 2020](#), [Li et al. 2021](#), [Liu et al. 2021](#)), although NEFFy offers other normalization options. We discuss alternative formulations of NEFF in [Supplementary Section S5](#). In this work, we adopt the specified NEFF formula because of its broad acceptance in the structure prediction community and its proven ability to capture the evolutionary diversity present in MSAs.

Aligning with the given formulation, several tools incorporate NEFF calculation as part of their functionality ([Källberg et al. 2012](#), [Kamisetty et al. 2013](#), [Simkovic et al. 2017](#), [Zhang et al. 2020, 2023](#)). Additional details about each are provided in [Supplementary Section S3](#) of the Supplementary File. However, none of these tools are exclusively designed for NEFF calculation and they typically offer NEFF as

required for their primary tasks. Consequently, they do not offer the full range of NEFF-related capabilities. [Table 1](#) compares these tools with NEFFy in terms of the features they provide. It highlights that NEFFy is the only solution compatible with a wide range of MSA formats and supports all the features provided by the other tools. Moreover, NEFFy introduces several new features:

- 1) *NEFF calculation for multiple MSAs:* Accepts multiple MSA input files, combines their sequences in the specified order, removes duplicates, and calculates the NEFF for the resulting merged MSA.
- 2) *Per-residue (column-wise) NEFF calculation:* Calculates NEFF for each position in the MSA by summing the weights of sequences containing a residue (i.e. non-gap character) at that position.
- 3) *NEFF for multimeric MSAs:* Computes NEFF for multimeric MSAs, which comprising multiple chains. NEFFy can process both homomers and heteromers. The user can specify the multimer format to enable NEFFy to automatically identify the appropriate blocks within the MSA and calculate NEFF values for each.
- 4) *MSA format conversion:* Converts MSA formats while preserving sequence integrity and annotations, with no need for user intervention.
- 5) *MSA validation:* Ensures that the input follows the specified format and contains only residues from the permitted character set for the given alphabet.

[Supplementary Section S2](#) details each feature.

3 Results

We conducted an experiment using the CASP15 dataset, which consists of 93 targets, to compare the reliability and efficiency of NEFFy with other tools. To generate MSA files for these targets, we ran AlphaFold 2.3 locally, utilizing its default MSA generation pipeline. This process resulted in three MSA files per target, sourced from the Uniref90, Mgnify, and BFD datasets. In total, we generated 279 MSA files in STO

and A3M formats. Given that certain tools do not support NEFF calculation in these formats, we used NEFFy's built-in converter to convert the files into formats compatible with each tool. We then calculated the NEFF value for each of these MSAs. To do this, we ran NEFFy using the options specified by each tool and also used each tool directly to calculate NEFF (Supplementary Section S4.1). The results showed that NEFFy consistently produced NEFF values identical to those of DeepMSA and highly similar to those from other tools (Supplementary Figs S2–S6).

To assess the computational efficiency of NEFFy relative to other tools, we used the MSA files from the previous analysis and recorded the execution time for each tool (Supplementary Section S4.2). NEFFy achieves the same efficiency as DeepMSA and Gremlin and significantly outperforms other tools. This makes it particularly ideal for processing deep MSAs of long query sequences (Supplementary Fig. S7).

To evaluate the scalability of NEFFy with respect to MSA depth, we conducted an analysis by progressively increasing the MSA depth and measuring the corresponding execution times (Supplementary Section S4.3). The results show that NEFFy, along with DeepMSA and Gremlin, exhibits relatively constant execution times across all depths, demonstrating superior scalability and minimal sensitivity to increasing MSA depth. This indicates these methods are well-suited for efficiently handling deeper MSAs. In contrast, RaptorX and Conkit show a substantial increase in execution time as the depth increases, indicating that their performance is more sensitive to larger MSA depths and may struggle with higher computational loads at higher depths (Supplementary Fig. S8). These findings provide valuable insights for selecting appropriate tools for MSA-based workflows depending on the computational resources available and the required scalability.

As a case study on multi-domain proteins, we explored the relationship between the NEFF values of individual domains and those of entire protein chains. Using 19 multi-domain proteins from the CASP15 dataset, we calculated NEFF values for the MSA of each domain and the entire target, using two separate sets of MSAs—one generated by AlphaFold (Jumper *et al.* 2021) and the other by RoseTTAFold (Baek *et al.* 2021) (Supplementary Section S4.4). To compare the relative NEFF values, we generated Grishin plots (Kinch *et al.* 2011), which display the correlation between the weighted sum of NEFF values for individual domains (y -axis) and the NEFF values for the full chain (x -axis). Our results showed that individual domains tend to have higher NEFF values than the complete protein chains in both the AlphaFold and RoseTTAFold-generated MSAs (Supplementary Figs S9 and S10).

In a recent study, Moussad *et al.* (2023) examined the prediction accuracies of individual domains and their overall packing in tertiary structures generated by AlphaFold and RoseTTAFold models for these 19 multi-domain targets using the Global Distance Test (GDT-TS) metric designed to compare the similarity between predicted and reference structures (Zemla 2003). They found that individual domains were predicted with high accuracy, while the overall packing of these domains was less accurate. Our NEFF analysis for the MSAs used as input for these multi-domain targets reflects similar patterns, suggesting that the higher NEFF

values observed for MSAs of individual domains may contribute to the improved accuracy of their predicted structures.

4 Conclusion

We present NEFFy, a fast and comprehensive tool for calculating the “Number of Effective Sequences” (NEFF) for a Multiple Sequence Alignment (MSA). By merging established features with innovative functionalities, NEFFy delivers an efficient and versatile solution for NEFF calculation, enhanced by its robust capabilities to convert and validate a wide range of MSA formats. Its implementation as a Python library further increases its accessibility and ease of use. Our experiments demonstrated that NEFFy is highly consistent with existing tools and is efficient, making it a valuable addition to the bioinformatics toolkit for processing MSAs of proteins and nucleic acids. As foundational models continue to evolve (Sumi *et al.* 2024) and generate synthetic biological sequences, NEFFy can also be used to evaluate the effectiveness of evolutionary sequence search and alignment algorithms for proteins and nucleic acids (DNA and RNA) in comparison to synthetic sequences generated by these models. For future work, we will integrate pairwise NEFF calculations into NEFFy for heteromeric MSAs, enabling nuanced analysis of cases where some monomer pairs are accurately predicted while others are not.

Author contributions

Maryam Haghani (Conceptualization [equal], Methodology [equal], Software [equal], Writing—original draft [equal], Writing—review & editing [equal]), Debswapna Bhattacharya (Conceptualization [equal], Methodology [equal], Supervision [equal], Writing—review & editing [equal]), and T. M. Murali (Funding acquisition [equal], Supervision [equal], Writing—review & editing [equal])

Supplementary data

Supplementary data are available at *Bioinformatics* online.

Conflict of interest: None declared.

Funding

This work was partially supported by awards from the National Institute of General Medical Sciences [R35GM138146 to D.B.]; the National Science Foundation [DBI 2208679 to D.B., CCF 2200045 to T.M.M.].

Data availability

The MSA files used for NEFF analysis of reliability and scalability are hosted on Zenodo at <https://doi.org/10.5281/zenodo.14210949>, while those for multi-domain analysis can be found at <https://doi.org/10.5281/zenodo.7682977>. The source code is archived at Zenodo: <https://doi.org/10.5281/zenodo.14908219>.

References

- Adhikari B, Hou J, Cheng J. DNCON2: improved protein contact prediction using two-level deep convolutional neural networks. *Bioinformatics* 2018;34:1466–72.

- Altschul S, Madden T, Schäffer A *et al.* Gapped BLAST and PSI-BLAST: a new generation of protein database search programs. *Nucleic Acids Res* 1997;25:3389–402.
- Baek M, DiMaio F, Anishchenko I *et al.* Accurate prediction of protein structures and interactions using a three-track neural network. *Science* 2021;373:871–6.
- Baek M, McHugh R, Anishchenko I *et al.* Accurate prediction of protein–nucleic acid complexes using RoseTTAFoldNA. *Nat Methods* 2024;21:117–21.
- Chenna R, Sugawara H, Koike T *et al.* Multiple sequence alignment with the Clustal series of programs. *Nucleic Acids Res* 2003;31:3497–500.
- Edgar R. MUSCLE: multiple sequence alignment with high accuracy and high throughput. *Nucleic Acids Res* 2004;32:1792–7.
- He B, Mortuza S, Wang Y *et al.* NeBcon: protein contact map prediction using neural network training coupled with naïve Bayes classifiers. *Bioinformatics* 2017;33:2296–306.
- Hu M, Yuan F, Yang K *et al.* Exploring evolution-aware &-free protein language models as protein function predictors. *Adv Neural Inf Process Syst* 2022;35:38873–84.
- Jumper J, Evans R, Pritzel A *et al.* Highly accurate protein structure prediction with AlphaFold. *Nature* 2021;596:583–9.
- Källberg M, Wang H, Wang S *et al.* Template-based protein structure modeling using the RaptorX web server. *Nat Protoc* 2012;7:1511–22.
- Kamisetty H, Ovchinnikov S, Baker D. Assessing the utility of coevolution-based residue–residue contact predictions in a sequence-and structure-rich era. *Proc Natl Acad Sci USA* 2013;110:15674–9.
- Kinch L, Shi S, Cheng H *et al.* CASP9 target classification. *Proteins* 2011;79:21–36.
- Li Y, Zhang C, Bell E *et al.* Deducing high-accuracy protein contact-maps from a triplet of coevolutionary matrices through deep residual convolutional networks. *PLoS Comput Biol* 2021;17:e1008865.
- Liu X, Jin L, Gao S *et al.* Protein contact map prediction using multiple sequence alignment dropout and consistency learning for sequences with less homologs. bioRxiv, 2021, preprint: not peer reviewed.
- Moussad B, Roche R, Bhattacharya D. The transformative power of transformers in protein structure prediction. *Proc Natl Acad Sci USA* 2023;120:e2303499120.
- Notredame C, Higgins D, Heringa J. T-Coffee: a novel method for fast and accurate multiple sequence alignment. *J Mol Biol* 2000;302:205–17.
- Pearce R, Zhang Y. Toward the solution of the protein structure prediction problem. *J Biol Chem* 2021;297:100870.
- Remmert M, Biegert A, Hauser A *et al.* HHblits: lightning-fast iterative protein sequence searching by HMM-HMM alignment. *Nat Methods* 2012;9:173–5.
- Shao J, Chen J, Liu B. ProFun-SOM: protein function prediction for specific ontology based on multiple sequence alignment reconstruction. *IEEE Trans Neural Netw Learn Syst* 2025;36:8060–71.
- Simkovic F, Thomas J, Rigden D. ConKit: a python interface to contact predictions. *Bioinformatics* 2017;33:2209–11.
- Sumi S, Hamada M, Saito H. Deep generative design of RNA family sequences. *Nat Methods* 2024;21:435–43.
- Tarafder S, Roche R, Bhattacharya D. The landscape of RNA 3D structure modeling with transformer networks. *Biol Methods Protoc* 2024;9:bpae047.
- Wang S, Sun S, Li Z *et al.* Accurate de novo prediction of protein contact map by ultra-deep learning model. *PLoS Comput Biol* 2017;13:e1005324.
- Wang W, Feng C, Han R *et al.* trRosettaRNA: automated prediction of RNA 3D structure with transformer network. *Nat Commun* 2023;14:7266.
- Wilke A, Bischof J, Gerlach W *et al.* The MG-RAST metagenomics database and portal in 2015. *Nucleic Acids Res* 2016;44:D590–4.
- Wu T, Hou J, Adhikari B *et al.* Analysis of several key factors influencing deep learning-based inter-residue contact prediction. *Bioinformatics* 2020;36:1091–8.
- Zemla A. LGA: a method for finding 3D similarities in protein structures. *Nucleic Acids Res* 2003;31:3370–4.
- Zhang C, Freddolino P, Zhang Y. COFACTOR: improved protein function prediction by combining structure, sequence and protein–protein interaction information. *Nucleic Acids Res* 2017;45:W291–9.
- Zhang C, Zhang Y, Pyle A. rMSA: a sequence search and alignment algorithm to improve RNA structure modeling. *J Mol Biol* 2023;435:167904.
- Zhang C, Zheng W, Mortuza S *et al.* DeepMSA: constructing deep multiple sequence alignment to improve contact prediction and fold-recognition for distant-homology proteins. *Bioinformatics* 2020;36:2105–12.
- Zheng W, Wuyun Q, Li Y *et al.* Improving deep learning protein monomer and complex structure prediction using DeepMSA2 with huge metagenomics data. *Nat Methods* 2024;21:279–89.