

Evaluating deep learning based structure prediction methods on antibody–antigen complexes

Samuel Fromm, Marko Ludaic, Arne Elofsson*, 

Science for Life Laboratory and Department of Biochemistry and Biophysics, Stockholm University, Solna 171 21, Sweden

*Corresponding author. Science for Life Laboratory and Department of Biochemistry and Biophysics, Stockholm University, Box 1031, Solna 171 21, Sweden. E-mail: arne@bioinfo.se.

Associate Editor: Pier Luigi Martelli

Abstract

Motivation: AlphaFold2 significantly improved the prediction of protein complex structures. However, its accuracy is lower for interactions without coevolutionary signals, such as host–pathogen and antibody–antigen interactions. Two strategies have been developed to address this limitation: massive sampling and replacing the evoformer with the pairformer, which does not rely on coevolution, as introduced in AlphaFold3, thereby enabling more structural reasoning by the network.

Results: In this study, we benchmark structure prediction methods on unseen antibody–antigen complexes. We found that increased sampling improves the chances of generating a correct protein model, roughly in a log-linear manner. However, the internal quality estimates by AlphaFold often cannot identify the best predicted structures for each target, resulting in a significant loss of performance for the top-ranked protein model compared with the best model. For all methods, a significant challenge remains the identification of the best model. We also show that AlphaFold3 outperforms AlphaFold2, Boltz-1, and Chai-1. Furthermore, AlphaFold3 performance declines significantly for complexes that lack structural similarity to the training set, indicating that it has to some extent learned to detect remote structural similarities.

Availability and implementation: All code is available from github.com/samuelfromm/abag-benchmark-set/ and all data from DOI: 10.5281/zenodo.17978681. The latter repository also contains the code.

1 Introduction

Antibody–antigen interactions play a crucial role in immunology and various biotechnological applications. Furthermore, the use of therapeutic antibodies has increased over the past few decades (Lu *et al.* 2020). Therefore, accurate prediction of these interactions is essential. It can accelerate discovery by identifying potential immunotherapy candidates and enhancing our ability to design targeted treatments for various diseases, including cancers and autoimmune disorders.

The prediction of protein complexes has advanced considerably with the advent of AlphaFold2 (Jumper *et al.* 2021), originally designed to predict the structures of individual proteins (Bryant *et al.* 2022). However, AlphaFold2 accuracy decreases for interactions lacking coevolutionary signals, particularly in complex biological systems such as host–pathogen (Saluri *et al.* 2025), and for antibody–antigens (Yin and Pierce 2024). Recognizing this limitation, two primary strategies have been established to improve prediction accuracy.

The first strategy involves extensive sampling methods demonstrated during the CASP15 competition (Lu *et al.* 2020,

Wallner 2023), where multiple structures were generated, increasing the likelihood of identifying accurate conformations. Several alternative methods have been proposed to increase sampling efficiency, including using dropout (Lu *et al.* 2020, Wallner 2023, Raouraoua *et al.* 2024), masking columns in the MSA (Kalakoti and Wallner 2025), and subsampling the MSA (Del Alamo *et al.* 2022).

Several scoring functions have been developed to estimate the accuracy of predicted structures, starting with the introduction of pDockQ, used to predict the DockQ from AlphaFold2.0 (Bryant *et al.* 2022). With the introduction of AlphaFold-multimer (Evans *et al.* 2022), an internal score (ipTM) can be used to estimate the quality of the predicted interaction. Internally, AlphaFold2 and the other methods use the ranking confidence, a linear combination of 80% ipTM and 20% pTM. Several attempts have been made to improve the ranking. For instance, pDockQ2 (Zhu *et al.* 2023) can be used to predict the quality of interactions of each chain, actipTM (Varga *et al.* 2024) is tuned to rank interactions involving flexible regions, and ipSAE (Dunbrack 2025) is developed to better separate true from false interactions than ipTM. All of these methods might provide

an improvement over AlphaFold2 (i) pTM scores. However, it has not been clearly demonstrated that they significantly improve the per-target ranking of antibody–antigen complexes.

In this study, our objective was to evaluate new structure prediction methods for predicting antibody–antigen complexes using a new benchmark. We assess whether larger sampling, selection, and confidence metrics can improve the identification of accurately predicted structures.

2 Material and methods

2.1 Dataset

The dataset consists of 110 antibody–antigen structures with <80% sequence identity to any protein in the training set. The dataset was created using a slightly modified version of the Antibody Antigen Dataset Maker (AADaM) (McCoy *et al.* 2024), available from <https://github.com/samuelfromm/AADaM-fork>. We initially used all entries deposited in SAbDab (Schneider *et al.* 2022) as of 23 October 2024. These structures are then divided into two subsets: a “before” dataset (30 September 2021) and an “after” (benchmark) dataset (30 September 2021). The chosen date, 30 September 2021, serves as the cut-off for the training data used by all benchmarked methods. The “before” dataset was filtered to retain only structures in which at least one antigen chain is a peptide or protein.

Antibody–peptide structures were excluded from the benchmark dataset, which was further filtered to include only structures determined by X-ray diffraction or electron microscopy, with a resolution of 3.5 Å or better. Structures with non-standard residues were excluded unless at the sequence ends, in which case they were trimmed. Structures missing >10% of residues were also discarded. Structures are discarded from the benchmark dataset if the sequence identity between the benchmark and “before” dataset exceeded 80%. Sequence identity between the two datasets was calculated individually for the heavy (H), light (L), and antigen chains using a global alignment. The benchmark set itself is further clustered by sequence identity using the same procedure and threshold. During clustering, the structure with the fewest missing residues in the H, L, and antigen chains is retained. If two structures have the same number of missing residues, the one with the shorter antigen sequence is selected. One of the now 111 complexes failed due to an HHblits-related issue and was therefore excluded, leaving the final dataset with 110 proteins. In the final set, no two antibodies are identical, while one antigen is present twice, bound to two different antibodies.

We also quantified how the benchmark size would change if redundancy was assessed using concatenated complementarity-determining regions (CDRs) sequences. Applying the same 80% identity threshold retained 416 complexes, reflecting the much higher diversity of CDR regions compared to full antibody sequences. Furthermore, we examined CDR redundancy in 110 benchmark complexes relative to the training set and found that although some loops are identical to those in the training set, no single antibody has all its loops identical to any antibody in the training set. The maximum sequence identity for the concatenated loops is 72%. For antigens, 15 sequences were

found in the training set, while the other 110 are unique to the benchmark (note that some complexes have more than one antigen chain). In Fig. S1, available as supplementary data at *Bioinformatics* online, it can be seen that most antibodies share a maximum identity to a member of the training set of ~50% for the concatenated loops, although individual loops have higher identities.

2.2 Structure generation

2.2.1 AlphaFold2.3 methods to enhance sampling

To increase the sampling diversity, we examined previously described methods (Del Alamo *et al.* 2022, Wallner 2023, Kalakoti and Wallner 2025) by modifying the AlphaFold2.3.2 code. All methods were run without relaxation. The maximum template date is set to 30 September 2021. We used the latest version of the weights (v3) and for each of the five available network models, 40 predictions were generated, resulting in a total of 200 models per method.

In total, five different AlphaFold2-based methods were evaluated:

- AlphaFold2.3: AF2.3 with default parameters.
- AF2.3-SubSampl-32-64: AF2.3 with MSA subsampling using num_msa=32 and num_extra_msa=64 compared to the default values of 508 and 2048, respectively.
- AF2.3-SubSampl-16-32: AF2.3 with MSA subsampling using num_msa=16 and num_extra_msa=32.
- AF2.3-Dropout: using the ideas in AFsample (Wallner 2023), i.e. we turn on dropout during interference.
- AF2.3-ColMask: using the ideas in AFsample2 (Kalakoti and Wallner 2025), i.e. we randomly mask 15% of the columns in the MSA.

The code for our implementation of the above methods is available at <https://github.com/samuelfromm/alphafold-clone>. It should be noted that we used the default parameters for all methods. Other hyperparameters, such as dropout rates, might yield better performance. In CASP, up to five models were allowed, and it is often good to submit a diverse set of models (Elofsson 2025), so if such a strategy were applied here, the difference to the best possible model would decrease. However, the best approach for selecting five models has not been well studied; therefore, we have focused only on selecting the top-ranked model.

2.2.2 AlphaFold 3, Chai-1, and Boltz-1 models

For each complex in the dataset, 40 random seeds (1–40) were used to generate five models each, resulting in a total of 200 models per complex.

2.2.3 MSA generation

Multiple sequence alignments (MSAs) were generated using the AlphaFold 2.3.2 pipeline (specifically commit f251de6). For running AF3 and Boltz-1, we specifically used the `bfd_uniref_hits.a3m` file produced by this pipeline, while for Chai-1, MSAs were generated using the `—use-msa-server` option via the ColabFold (Mirdita *et al.* 2022) server.

2.3 Evaluation

2.3.1 Scoring

To calculate scores, we created a Snakemake workflow. For each prediction, the corresponding predicted model (query) and the PDB ground truth structure of the corresponding PDBID (reference) were used as inputs.

AADaM (McCoy *et al.* 2024), the pipeline used to generate our dataset, outputs a FASTA sequence for each antibody–antigen structure. However, these structures do not always align fully with their corresponding PDB entries. Since AADaM focuses on the antibody–antigen interface, some chains present in the original PDB file may be excluded from the interface structures used in our dataset. Additionally, the sequence in the reference structure from the PDB is often shorter than the SEQRES sequence used to generate the protein structures. To address both issues, we first used MMalign (Mukherjee and Zhang 2009) to calculate the alignment between the query and reference structures. We then take each pair of aligned chains and calculate a sequence alignment between them. Based on this sequence alignment, we select the residues present in both the query and the reference structure.

Using the cropped query and reference structures as input, we calculate several metrics:

- 1) We used MMalign to calculate the TM-score (Zhang and Sjolnick 2004).
- 2) We calculate the DockQ score (Basu and Wallner 2016, Mirabello and Wallner 2024) for the antibody–antigen interface (“merging” multi-chain antigens into a “single chain”). This results in a single DockQ score for each complex.
- 3) Similarly, we calculate the ipSAE score (Dunbrack 2025) with a PAE and distance cut-off of 10 Å for the antibody–antigen interface, resulting in a single score for each complex.
- 4) We calculate pDockQ version 2 (Zhu *et al.* 2023).
- 5) We extract the confidence scores from the predictions as output by the respective method: `ranking_confidence` for AlphaFold-based methods, `confidence_score` for Boltz-1, and `aggregate_score` for Chai-1.
- 6) We calculate the aligned error between the cropped query and reference structure as well as the aligned error ranking confidence (see Section 2.7).

The workflow is the same for all samples and methods, except for the last step: the aligned error analysis was performed only for a selection of methods. The complete workflow is included in the provided code.

In some cases, MMalign cannot properly align the chains, indicating that the query and reference structures are dissimilar and that the predicted structure is of very low quality. We still run the rest of the workflow even when the number of aligned residues is extremely low. When the number of aligned residues, normalised by the query length, is below 0.9, the resulting DockQ scores typically fall between 0 and 0.05. We consider these low scores to accurately reflect the poor quality of the predicted interfaces in such cases.

2.4 Sampling

For each method, we generated 200 models for each complex and treated them as independent predicted structures, although

in AF3’s inference, each random seed generates five different models. These models use the same information as the pair-former, but differ in the built-in randomness of the diffusion process for training. To assess the impact of sampling, we randomly select k models with replacement, identify the best (highest DockQ) and top-ranked (highest confidence) models, and record their DockQ scores. For each sampling size k , we average the recorded DockQ score over 50 iterations. We then average the resulting score over all IDs and sample sizes.

2.5 Structural comparison

We used both global and interface similarity to investigate whether structural similarity could be detected between the proteins in the test and training sets. First, we used FoldSeek-Multimer (van Kempen *et al.* 2024, Kim *et al.* 2025) to identify global similarity between a complex in the test set and any protein present in PDB prior to the split date (30 September 2021). For each complex, the maximum TM-score (Zhang and Sjolnick 2004) was normalized by the length of the query protein. Next, we created a smaller dataset to examine the interface similarity using USalign (Zhang *et al.* 2022). This dataset included all 1998 PDB entries with a TM-score > 0.5 to any complex in the test set. We first extracted all interface residues from the proteins in the test set to evaluate interface similarity. Interface residues were defined as any antigen residue containing an atom within 12 Å of the antibody or *vice versa*. Next, the interface was aligned to the 1998 PDB entries mentioned above using USalign with the `-mm 1` flag, and the maximum TM-score was recorded.

2.6 Model diversity, PconsDock

To estimate the diversity of models for a given target, we used PconsDock (Pozzati *et al.* 2022), i.e. a pairwise comparison of all the models for a target using DockQ. Due to computational constraints, we did not pairwise compare all 200 models. Instead, for each model, we randomly selected 20 other models for comparison. We then report the average DockQ per target or per model.

2.7 Aligned error

AlphaFold is trained to estimate two types of errors: the pLDDT, which estimates the error for each residue, and the PAE, which estimates the *Aligned Error* for each pair of residues. The predicted aligned error (PAE) is then used to calculate pTM and ipTM, and hence the confidence score. We want to examine if it is the quality of these estimates or the functional form of ipTM/pTM itself that hampers discrimination among model predictions. To test this, we evaluate discrimination performance with oracle-based AEs (aeTMs) for a given prediction, described below.

We follow the notation in Jumper *et al.* (2021) [§1.9.7]. Assume that we are given two structures—a reference and query structure—for the same sequence. Let $Y = \{\vec{y}_j\}$ and $X = \{\vec{x}_j\}$ denote the C^α -positions of the reference and query structure, respectively, let $T_{Y,i}$ and $T_{X,i}$ be the corresponding backbone frames, and let $e_{ij}(X, Y) = \|T_{X,i}^{-1} \circ \vec{x}_j - T_{Y,i}^{-1} \circ \vec{y}_j\|$ be the error in the C^α -position of residue j when the query and reference

structures are aligned using the backbone frame of residue i . The non-symmetric matrix $(e_{ij})_{i,j=1,\dots,N_{\text{res}}}$, where N_{res} is the number of residues, is called the aligned error (AE) between the two structures. Let

$$f(d, n) = \frac{1}{1 + \left(\frac{d}{d_0(n)}\right)^2}, d_0(n) = 1.24 \sqrt[3]{\max(n, 19) - 15} - 1.8.$$

We define the aligned error TM-score (aeTM) and aligned error interface TM-score (aeiTM) as

$$\text{aeTM}(X, Y) := \max_{i \in [1, \dots, N_{\text{res}}]} \frac{1}{N_{\text{res}}} \sum_{j=1}^{N_{\text{res}}} f(e_{ij}(X, Y), N_{\text{res}}) \quad (1)$$

and

$$\text{aeiTM}(X, Y) := \max_{i \in [1, \dots, N_{\text{res}}]} \frac{1}{|\mathcal{D}_i|} \sum_{j=1}^{N_{\text{res}}} f(e_{ij}(X, Y), |\mathcal{D}_i|), \quad (2)$$

respectively. Here, $\mathcal{D}_i$ denotes the set of all residues except residue i .

Note that the aeTM score and the aeiTM scores can be viewed as idealized versions of the predicted TM (pTM) score and interface predicted TM (ipTM) score, respectively. Indeed, let us replace X with a predicted structure and assume that the ground truth structure X^{true} is unknown. Assuming that we have a random variable $Y(\zeta)$, sampling structures from a distribution of probable ground truth structures. Then the pTM and ipTM scores are given by

$$\text{pTM}(X) = \max_{i \in [1, \dots, N_{\text{res}}]} \frac{1}{N_{\text{res}}} \sum_{j=1}^{N_{\text{res}}} \mathbb{E}[f(e_{ij}(X, Y(\zeta)), N_{\text{res}})] \quad (3)$$

and

$$\text{ipTM}(X) = \max_{i \in [1, \dots, N_{\text{res}}]} \frac{1}{|\mathcal{D}_i|} \sum_{j=1}^{N_{\text{res}}} \mathbb{E}[f(e_{ij}(X, Y(\zeta)), |\mathcal{D}_i|)], \quad (4)$$

respectively (see AF2, §1.9.7 and AF2-multimer, §7.9). Here, the expectation is taken over the probability distribution of $e_{ij}(X, Y(\zeta))$. Now, in an ideal situation, $\mathbb{P}(Y(\zeta) = X^{\text{true}}) = 1$, i.e. the random variable Y takes on the values X^{true} with probability 1, in which case $\mathbb{E}[f(e_{ij}(X, Y(\zeta)))] = f(e_{ij}(X, X^{\text{true}}))$ so that $\text{pTM}(X) = \text{aeTM}(X, X^{\text{true}})$ and similarly $\text{ipTM}(X) = \text{aeiTM}(X, X^{\text{true}})$.

In practice, AlphaFold learns to predict the probability distribution of the elements of the aligned error $e_{ij}(X, Y(\zeta))$. To this end, the distribution of e_{ij} is discretized into 64 bins, ranging from 0 to 31.5 Å with a bin width of 0.5 Å, where the final bin captures any errors larger than 31.5 Å. When calculating the aligned error between two structures X and Y , we do, however, allow unlimited error values.

While the pTM score is a predicted metric, the aeTM takes as input any two structures of the same sequence, similarly to the TM-score (Zhang and Sjolnick 2004). Note also that the aeTM score is bounded from above by the TM-score (see AF2, §1.9.7).

Finally, recall that AlphaFold's ranking confidence score is defined as $0.2 \cdot \text{pTM} + 0.8 \cdot \text{ipTM}$. In a similar fashion, we define the aligned error ae ranking confidence (aeRankConf) as $0.2 \cdot \text{aeTM} + 0.8 \cdot \text{aeiTM}$.

3 Results and discussion

Here, we evaluated the performance of AlphaFold2.3 (Evans *et al.* 2022), AlphaFold3 (Abramson *et al.* 2024), Boltz-1 (Wohlwend *et al.* 2025), and Chai-1 (Discovery *et al.* 2024), using a dataset of 110 unseen antibody-antigen structures with limited identity to the training set, for varying numbers of generated protein models. For AlphaFold2.3, we also examined other strategies reported to increase sampling efficiency.

Several studies have reported that generating more samples improves the quality of AlphaFold2 predictions (Wallner 2023, Raouraoua *et al.* 2024, 2026). Further, in the AlphaFold3 (Abramson *et al.* 2024) paper, it is claimed that generating thousands of samples improved the predictions of antibody-antigen complexes. We generated 200 models for each method and examined their performance, measured by the DockQ score (Basu and Wallner 2016, Mirabello and Wallner 2024) between the antigen and the antibody. The performance was measured for each sampling size by using the mean DockQ of the best model for each target and the mean DockQ of the top-ranked models, see Section 2 for details. A summary of the results is presented in Fig. 1.

3.1 AlphaFold3 makes better predictions than the other methods

First, we will discuss the ability to generate a good model, assuming an oracle exists to identify the best model among all models generated by a method. Figure 1A shows that AlphaFold3 outperforms all other methods by some margin, while AlphaFold2, Chai-1, and Boltz-1 show similar performance when a single model is generated. One example where AF3 performed considerably better than other methods is for the target 8TQ7, see Fig. 5. Here, the AF3 prediction is spot-on, while all other methods predicted binding to different regions of the antigen.

Secondly, it can be observed that all methods produce better protein models as sampling increases. A roughly linear increase is observed with the logarithm of the sample size. This means that the average DockQ for AF3 of the best model increases from below 0.3 to above 0.5 when 200 samples are generated. A similar increase is observed for AlphaFold2 and Chai-1 (Discovery *et al.* 2024), while a smaller improvement is observed for Boltz-1.

In AlphaFold3, additional models can be generated by rerunning the entire pipeline with a different random seed or by letting the diffusion model produce multiple models. Figure 1C shows that using different seeds yields greater structural diversity than generating more samples from the diffusion step alone.

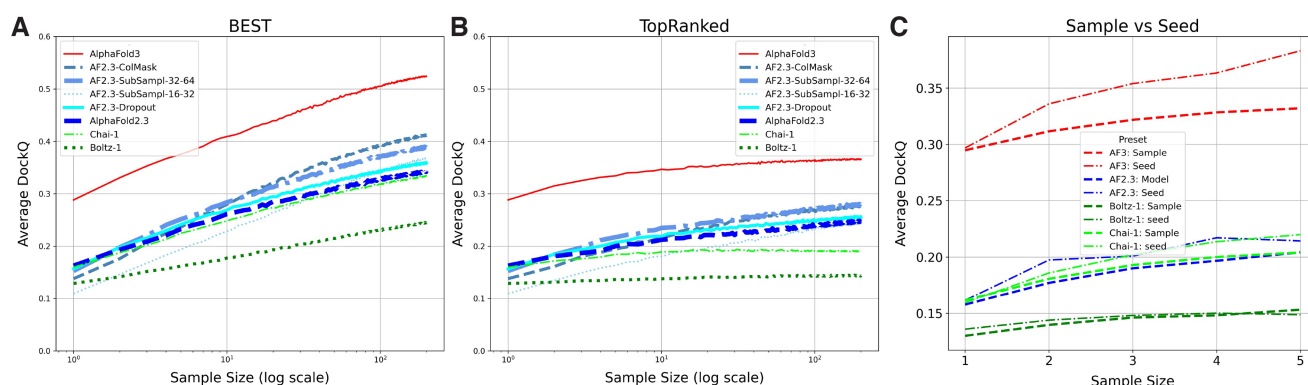


Figure 1 For each method and sample size n we randomly select n proteins out of the 200 structures we created (with replacement). Out of these n selected models, we select the model with the highest DockQ score (A) or highest-ranking confidence (B), respectively, and record its DockQ score on the y-axis. (C) Identification of the best, i.e. highest DockQ score, model when changing either the seed or only the sample for up to five models.

3.2 AF2.3-ColMask (AF-Sample2) is the most efficient strategy to increase sampling

As shown before, it is possible to increase the diversity of the generated proteins using several different strategies. We examined these strategies by modifying AlphaFold2.3, as the AlphaFold3 licence does not allow such modifications and because this study was initiated before the release of AlphaFold3. In CASP15, it was shown that increased sampling via dropouts improved prediction accuracy, particularly for antigen-antibody pairs (Wallner 2023) using AFsample. We refer to it here as AF2.3-Dropout, as the implementation is not identical. More recently, the Wallner group developed AFsample2 (hereafter referred to as AF2.3-ColMask), an even more aggressive method to add noise by masking entire columns in the MSA (Kalakoti and Wallner 2025). It is also possible to generate alternative conformations by subsampling the MSA (Del Alamo *et al.* 2022).

Figure 1A shows that, on average, all methods generate slightly worse models than the default AlphaFold2.3 setting when only one model is generated, see *SampleSize* = 1 (10^0). The difference is negligible for all methods except for AF2.3-ColMask (which is very similar to AFsample2), which produces significantly worse average models, DockQ of 0.138 versus 0.163. However, as more models are generated, all these methods produce better models than AlphaFold2.3 does natively. Overall, AF2.3-ColMask generated the best protein models, likely due to its more aggressive noise addition.

3.3 A large gap between the best and top-ranked models

It is necessary to identify good models as well as generate them. Figure 1B shows the average DockQ score for the top-ranked model against sample size. For all methods, the increase in average DockQ with sample size is much smaller than what is observed for the best model (Fig. 1A). This means that the internal ranking of the models cannot identify the best generated model. For AlphaFold3, it means the average DockQ increases only from 0.29 to 0.37, rather than 0.52. For Boltz-1 and Chai-1, there is hardly any improvement with larger sample sizes (from 0.12 to

0.14). For the AlphaFold2.3-based models, a similar trend to that observed for AlphaFold3 is evident. Here, AF2.3-SubSample performs as well as AF2.3-ColMask when 30 or more models are generated. It is clear that better detection of top-ranked models would significantly improve the prediction of antibody-antigen structures.

In Fig. 2, we take a more detailed look at the distribution of the scores for each target. It can be seen that for all methods, the targets can, in general, be divided into three types. To the left in the plots are the easy-to-predict targets, where most models are correct, though a few are incorrect. To the right in the plots are challenging targets, where most models are wrong, but a few good models exist for some targets. In the middle are models with many good and bad models. Further, the main difference between the four methods is the number of targets within each category, i.e. AlphaFold3 has many more easy-to-predict targets.

In Fig. 3, it can be seen that for both top-ranked and best models, some targets where AlphaFold2.3 outperforms AlphaFold3 exist. This indicates that there is no unique set of easier models; instead, there is a larger set of models where AlphaFold3 produces significantly better structures than the other methods.

3.4 Has AlphaFold3 learned to predict structure beyond its training set?

Next, we investigated how AlphaFold3 can predict many targets almost perfectly, while other methods failed for the same targets. There should be no coevolutionary signal between the antibody and the antigen; therefore, AlphaFold3 may have learned something about the physics governing the interactions. Alternatively, it may have identified patterns in the database that other methods did not. To answer this, we compared each target in our test set with the training data, i.e. all of PDB until September 2021. The search was made either by examining just the interface (Fig. 4A) or the entire complex (Fig. 4B). For each complex in our test set, we noted the maximum TM similarity to any complex in the training set for either the entire complex or just the interface region.

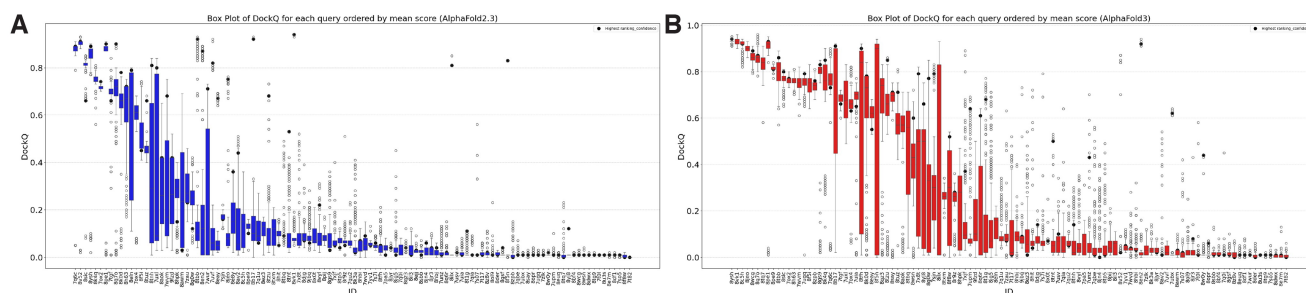


Figure 2 Each figure shows a box plot for the DockQ score for each target ID, sorted by mean DockQ score for that target ID and that method. The highest-ranked model is highlighted with a thick dot. (A) AlphaFold2.3, (B) AlphaFold3, for Boltz-1 and Chai-1 see [Fig. S2, available as supplementary data at Bioinformatics online](#). It can be seen that the variance of DockQ and ranking confidence is highest for targets with intermediate difficulty (see [Figs S3 and S4](#), and that these two variations correlate for AlphaFold3, [Figs S5 and S6, available as supplementary data at Bioinformatics online](#)).

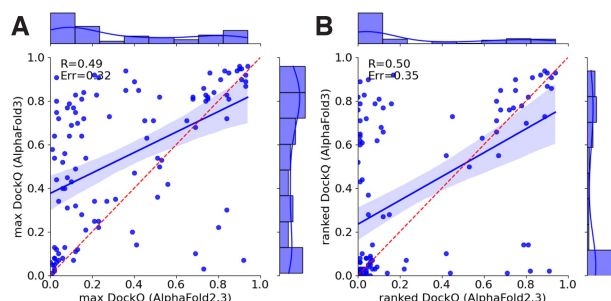


Figure 3 DockQ scores for AlphaFold3 versus AlphaFold2.3 for best models (A) and top-ranked models (B).

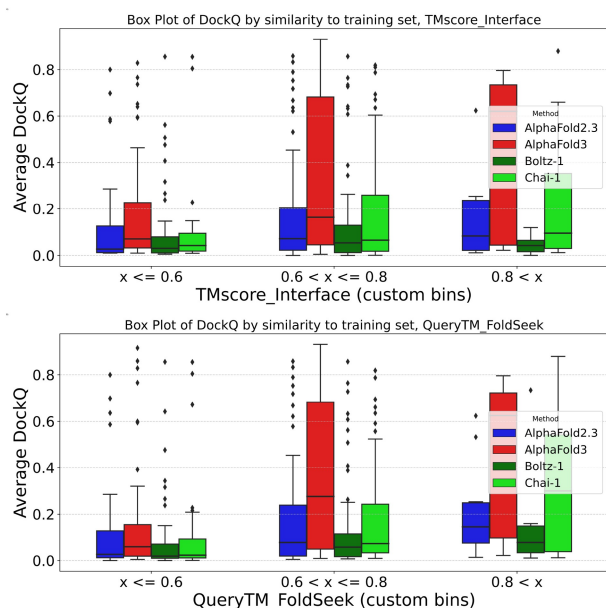


Figure 4 Prediction performance versus similarity to the training data for AlphaFold2.3 and AlphaFold2 using interface TM-score similarity (A) or query TM-score similarity (B) versus the mean AbAG-score for each target. For Max or first ranked scores, see [Fig. S7, available as supplementary data at Bioinformatics online](#).

As shown in [Fig. 4](#), the mean DockQ score significantly increases when a structurally similar antibody–antigen complex exists in the training set of AlphaFold3, and to a lesser extent

for Chai-1. When the best model for each target is used, a similar trend can be seen for all methods, [Fig. S7, available as supplementary data at Bioinformatics online](#). However, it should be noted that well-predicted models exist in which no structural similarity can be detected, and that the prediction is often dire, even if a structurally similar complex is present in the training set, as seen in [Fig. S8, available as supplementary data at Bioinformatics online](#). In general, the correlation between the AbAG-DockQ score and the similarity to the training set is low ($Cc=0.32$ for complex search and 0.17 for interface) ([Fig. S9, available as supplementary data at Bioinformatics online](#)), further supporting the idea that similarity with the training set is neither a requirement nor a guarantee for a correct prediction. We could not detect any trends for better performance for targets with high sequence identity in either the CDR regions or for the antigens, [Figs S10 and S11, available as supplementary data at Bioinformatics online](#).

We also noted that it is not trivial to understand what AlphaFold3 might have used from the training set to make excellent predictions. [Figure 5B](#) shows an example of the most similar hit in the training set to a model that AlphaFold3 (but not other methods) predicted well (8TQ7). It can be seen (in magenta) that the best hit (6DDM) has a similar beta-sheet domain that binds to the antibody. However, it is also clear that the interface is not identical [the interface TM-score ([Zhang and Sjolnick 2004](#)) is 0.63], indicating that, in this case, AlphaFold3 may have learned the orientation of the domains rather than the exact packing of the interface.

3.5 Ranking of AlphaFold3 models

[Figure 1](#) clearly shows that an improved ranking of the models for each target would significantly improve antibody–antigen predictions. The average DockQ score would increase from 0.37 to 0.52 . While the correlation between ranking confidence and DockQ across all targets is reasonably good ($Cc=0.80$ for ranking confidence), it is quite poor when considered on a per-target basis (mean $CC=0.28$), as can be seen in [Fig. 6](#). Although for some targets ranking confidence and DockQ correlate, for many they do not. Some targets have almost identical ranking confidences, but the DockQ varies widely, while others, including 7WVM, have virtually identical (bad) DockQ scores, while the ranking confidence differs widely. Alternative measures, such as

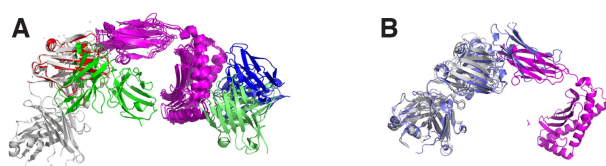


Figure 5 (A) Example (8TQ7) of a prediction where AlphaFold3 is superior to all other methods. The antigen in the center is shown in magenta, grey (to the left)—PDB structure, red (to the upper left) AlphaFold3, green (in the lower center) Boltz-1, lime (to the bottom right) Chai-1, blue (to the right) AlphaFold2.3. (B) Superposition of 8TQ7, in magenta and grey, with the most similar protein complex in the training set (6DDM) in blue.

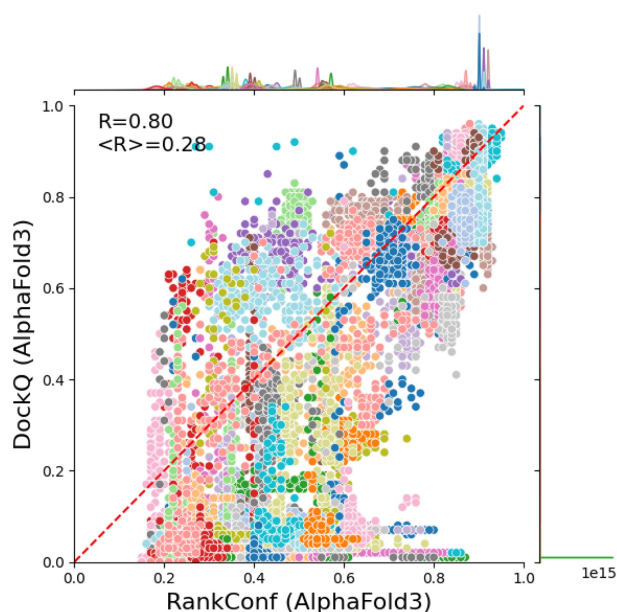


Figure 6 DockQ scores versus ranking confidence for all AlphaFold3 models. All models are shown in a single colour for each target. The overall correlation coefficient is 0.8, while the mean correlation coefficient per target ($\langle R \rangle$) is only 0.28. A plot with a maximum of 10 targets per plot is shown in Fig. S15, available as supplementary data at Bioinformatics online.

pDockQ, which are also based on predicted qualities (pLDDT and PAEs), show very similar behaviour (data not shown).

An alternative method to estimate the quality of a model is to use what we earlier referred to as PconsDock (Pozzati *et al.* 2022), i.e. to pairwise compare all the models for a target using DockQ. In Figs S12 and S13, available as supplementary data at Bioinformatics online, we show that the PconsDock score correlates as well as ipTM with the real DockQ scores, but is no better at ranking the targets. The reason is, obviously, because for a hard target, the models do not agree on how to locate the antibody on the antigen. In contrast, for an easy target, all antibodies are placed at the same location, as shown in Fig. S14, available as supplementary data at Bioinformatics online. Interestingly, there is one target, 8U3S, where both ipTM and PconsDock scores are high, although the DockQ score is low; see Fig. S14, available as supplementary data at Bioinformatics online.

In this case, SabDab is likely wrong, as the asymmetric unit does not correspond to the biological assembly according to PDB.

3.6 Ranking analysis

One of AlphaFold2's most innovative features was its ability to accurately estimate the quality of the generated protein structures. The estimates stem from two predicted features: the prediction of the accuracy of an individual residue (pLDDT) and the PAE, which estimates the error in the distance between pairs of residues (Jumper *et al.* 2021). Since the introduction of AlphaFold-multimer, these estimated distances have been used to calculate ipTM, specifically estimating the accuracy of interactions between chains.

Attempts to improve these scores have been proposed since the introduction of pDockQ (Bryant *et al.* 2022), which predicts the DockQ score for protein models generated by AlphaFold2.0. pDockQ bases its predictions on the number of interacting residues and their pLDDT values. Several other methods based on similar ideas have been introduced later. Recently, attempts have been made to improve AlphaFold's ipTM and pTM scores (and thus the ranking confidence) by recalculating them from PAE values (Varga *et al.* 2024, Dunbrack 2025). These methods modify functions (3) and (4) used to calculate the pTM and ipTM scores, respectively, while still using the same features used to calculate the ipTM and pTM scores.

To demonstrate the importance of accurate PAE predictions, we have also included these measures calculated from perfect PAE values, aeTM and aeITM. One example of how this affects predictions is shown in Fig. 7A–D. The two models are of 8HIT, one is almost perfect (DockQ: 0.690), while the other is rather poor (DockQ: 0.040). However, the PAE matrices for the two cases are almost identical.

To better understand the impact of feature quality versus function accuracy, we investigate how well the ideal aligned error scores, aeITM and aeTM (see Equations (1) and (2)), perform. We compared these measures on their ability to identify top-ranked model, as well as their correlation with DockQ, both overall (R) and per target ($\langle R \rangle$), see Fig. 7E and Table 1. For measures based on predicted PAE values—regardless of the specific choice of measure—(i) the average DockQ score of the top-ranked model is significantly lower than the best possible (DockQ=0.35, versus 0.54), (ii) the overall correlation with DockQ is acceptable ($R > 0.7$), while (iii) the per-target correlation is much worse ($\langle R \rangle < 0.30$). At the same time, using any of the measures with idealized PAE predictions solves these problems. Figure 7E shows the average DockQ score for each ranking method and sample size. Selecting by aeITM yields results much closer to the best possible outcome compared to selection by ipTM, indicating that the limiting factor when it comes to the ipTM score, at least for antibody–antigen complexes, is the accuracy of the PAE (and not necessarily a problem with the way the ipTM score is calculated itself). In Fig. 7F, it can be seen that the average per target correlation is almost as high as the overall correlation ($R=0.70$ versus 0.89).

3.7 Analysis of the PAE matrices

As can be seen in Fig. 7 the PAE matrices often have some sort of patterns. To examine whether this could have a rational

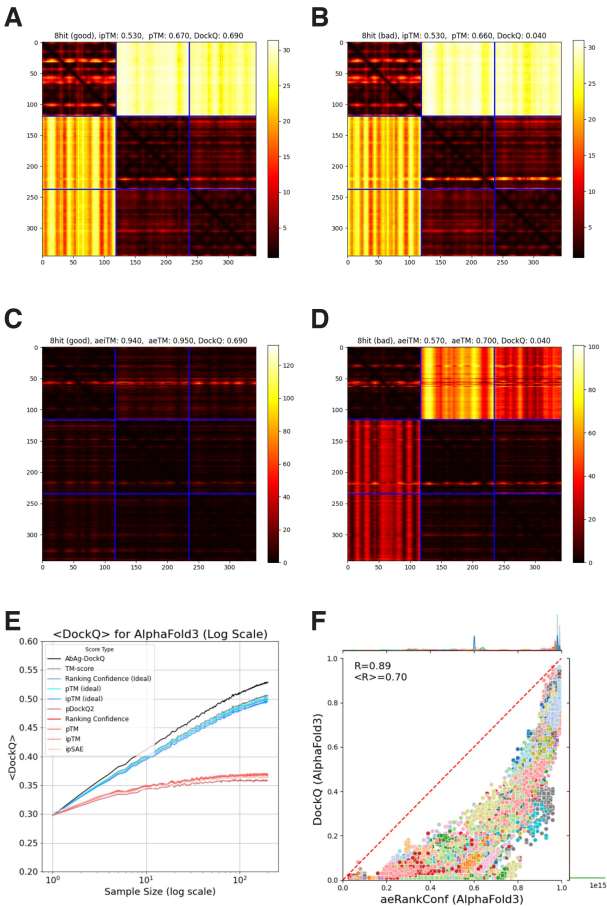


Figure 7 (A and B) PAE plots for two models of 8HIT. The chain breakers are marked in red (darker). The mean absolute error (MAE) between the two PAE matrices is 1.03 across the whole matrix and 1.37 across the inter-chain region. (C and D) AE plots for the same two models 8HIT. The mean average error between the two AE matrices is 13.24 across the whole matrix and 19.53 across the inter-chain region. (E) Ranking using different ways to use the PAEs to calculate the scores, in red methods using predicted PAEs (lower curves) and in blue (top curves) using AE (i.e. idealized PAE) scores. The primary issue is that the ground-truth aligned error values are not sufficiently accurate, resulting in a significant performance gap. (F) the aligned error ranking confidence values (compare with Fig. 6).

Table 1 Summary for different methods to rank AlphaFold3 models, average DockQ (<DockQ>) and Spearman correlation for all targets (R) or average per target <R>.

Method	<DockQ>	R	<R>
DockQ	0.544	1.000	1.000
aeTM	0.516	0.814	0.665
aeiTM	0.508	0.873	0.562
aeRankConf	0.511	0.879	0.595
pTM	0.369	0.670	0.200
ipTM	0.370	0.707	0.220
RankConf	0.375	0.799	0.214
pDockQ2	0.352	0.603	0.184
ipSAE	0.353	0.750	0.247

explanation, we compared all PAE matrices across models. First, we defined two regions of the PAE matrix. Region1 is the region to the top right related to the antibody–antigen contacts, while Region2 is the corresponding area in the bottom left part of the matrix. Region1 is related to the PAEs when the models are aligned on the antigen, while Region2 is related to the PAEs when the alignment occurs on the antibody.

First, it can be noted that the scores are almost always lower (better PAEs) in Region2 than in Region1, this is likely just an effect of size. However, this results in the fact that in almost all cases pTM and ipTM are calculated from one line in this region. Secondly, the correlation between different targets within the same region is much higher than that between Region1 and Region2 (Fig. S16, available as supplementary data at Bioinformatics online), indicating that it might be of important to develop a score that takes both regions into account.

Next, we analysed the correlation between rows and columns in the two regions, see Fig. S17, available as supplementary data at Bioinformatics online, this refers to how consistent the PAEs are when aligning on a specific residue in the other chain. It can be seen that the strongest correlation is for rows in Region2. This means that the strongest pattern occurs when the antigen is aligned with antibody residues, i.e. it is not particular residues (loops) in the antibody that dominate the patterns. We observed that the lowest PAEs were often found on surface exposed residues (data not shown).

4 Conclusion

In this study, we benchmarked different sampling and ranking strategies against a selected ensemble of unseen antibody–antigen complexes.

Our findings reveal that (i) AlphaFold3, but not Chai-1 and Boltz-1, consistently offers improved predictions over AlphaFold2, (ii) larger sampling sizes correlate with an increased probability of generating at least one good model, (iii) a crucial bottleneck persists in accurately identifying the best model among the generated models, and (iv) AlphaFold3’s performance drops considerably for complexes without prior structural similarities to the training dataset, underscoring the model’s dependence on previously encountered data for accurate predictions, the same is not seen for analysing sequence similarity to the training set. However, even if we see an increased success rate for targets with greater structural similarity to the training set, it is clear that this is neither a sufficient nor a necessary condition, as there are successful predictions with low similarity and unsuccessful predictions with high similarity.

Furthermore, to analyse the confidence metric used by AlphaFold (ipTM score), we introduced an idealized version of the ipTM score, showing that if the underlying estimates of accuracy (the PAE) are sufficiently accurate, the default ipTM formula is sufficient to achieve close to ideal per-target ranking and thereby significantly increase the DockQ score of the top-ranked model. This indicates that the limitation lies not in the ipTM formula itself but in AlphaFold’s ability to accurately predict aligned errors.

Author contributions

Samuel Fromm (Formal analysis [equal], Investigation [equal], Resources [equal], Software [equal], Writing—original draft [equal], Writing—review & editing [equal]), Marko Ludaic (Data curation [supporting], Formal analysis [supporting], Investigation [equal], Methodology [supporting], Resources [equal], Writing—review & editing [supporting]), and Arne Elofsson (Formal analysis [equal], Funding acquisition [lead], Investigation [supporting], Supervision [lead], Writing—review & editing [equal])

Supplementary material

[Supplementary material](#) is available at *Bioinformatics* online.

Conflicts of interest

None declared.

Funding

A.E. was funded by Vetenskapsrådet, Grant no. 2021–03979, and the Knut and Alice Wallenberg Foundation, Grant no. 2022.0032. The computations and data handling were enabled by the supercomputing resource Berzelius, provided by the National Supercomputer Centre at Linköping University, the Knut and Alice Wallenberg Foundation, and SNIC, grant numbers SNIC 2021/5–297 and Berzelius-2021–29.

References

- Abramson J, Adler J, Dunger J *et al.* Addendum: accurate structure prediction of biomolecular interactions with AlphaFold 3. *Nature* 2024;**636**:E4.
- Basu S, Wallner B. DockQ: a quality measure for protein–protein docking models. *PLoS One* 2016;**11**:e0161879.
- Bryant P, Pozzati G, Elofsson A. Improved prediction of protein–protein interactions using AlphaFold2. *Nat Commun* 2022;**13**:1265.
- Del Alamo D, Sala D, Mchaourab HS *et al.* Sampling alternative conformational states of transporters and receptors with AlphaFold2. *Elife* 2022;**11**.
- Discovery C, Boitreaud J, Dent J *et al.* Chai-1: decoding the molecular interactions of life. *bioRxiv*, 2024.10.10.615955, 2024, preprint: not peer reviewed.
- Dunbrack RL. 2025. Rēs ipSAE loquunt: what’s wrong with AlphaFold’s ipTM score and how to fix it. *bioRxiv*, 2025.02.10.637595, 2025, preprint: not peer reviewed.
- Elofsson A. 2025. AlphaFold3 at CASP16. *bioRxiv*, 2025.04.10.648174, 2025, preprint: not peer reviewed.
- Evans R, O’Neill M, Pritzel A *et al.* 2022. Protein complex prediction with AlphaFold-multimer. *bioRxiv*, 2021.10.04.463034, 2022, preprint: not peer reviewed.
- Jumper J, Evans R, Pritzel A *et al.* Highly accurate protein structure prediction with AlphaFold. *Nature* 2021;**596**:583–9.
- Kalakoti Y, Wallner B. AFsample2 predicts multiple conformations and ensembles with AlphaFold2. *Commun Biol* 2025;**8**:373.
- Kim W, Mirdita M, Levy Karin E *et al.* Rapid and sensitive protein complex alignment with foldseek-multimer. *Nat Methods* 2025;**22**:469–72.
- Lu R-M, Hwang Y-C, Liu I-J *et al.* Development of therapeutic antibodies for the treatment of diseases. *J Biomed Sci* 2020;**27**:1–30.
- McCoy KM, Ackerman ME, Grigoryan G. A comparison of antibody–antigen complex sequence-to-structure prediction methods and their systematic biases. *Protein Sci* 2024;**33**:e5127.
- Mirabello C, Wallner B. DockQ v2: improved automatic quality measure for protein multimers, nucleic acids, and small molecules. *Bioinformatics* 2024;**40**:btac586.
- Mirdita M, Schütze K, Moriaki Y *et al.* Colabfold: making protein folding accessible to all. *Nat Methods* 2022;**19**:679–82.
- Mukherjee S, Zhang Y. MM-align: a quick algorithm for aligning multiple-chain protein complex structures using iterative dynamic programming. *Nucleic Acids Res* 2009;**37**:e83.
- Pozzati G, Zhu W, Bassot C *et al.* Limits and potential of combined folding and docking. *Bioinformatics* 2022;**38**:954–61.
- Raouraoua N, Lensink MF, Brysbaert G. Massive sampling strategy for antibody–antigen targets in CAPRI round 55 with MassiveFold. *Proteins* 2026 (in press).
- Raouraoua N, Mirabello C, Véry T *et al.* MassiveFold: unveiling AlphaFold’s hidden potential with optimized and parallelized massive sampling. *Nat Comput Sci* 2024;**4**:824–8.
- Saluri M, Landreh M, Bryant P. AI-first structural identification of pathogenic protein target interfaces. *PLoS Comput Biol* 2025;**21**:e1013168.
- Schneider C, Raybould MIJ, Deane CM. SABDab in the age of biotherapeutics: updates including SABDab-nano, the nanobody structure tracker. *Nucleic Acids Res* 2022;**50**:D1368–72.
- van Kempen M, Kim SS, Tumescheit C *et al.* Fast and accurate protein structure search with foldseek. *Nat Biotechnol* 2024;**42**:243–6.
- Varga JK, Ovchinnikov S, Schueler-Furman O. actipfTM: a refined confidence metric of AlphaFold2 predictions involving flexible regions. *Bioinformatics* 2025;**41**:btaf107.
- Wallner B. AFsample: improving multimer prediction with AlphaFold using massive sampling. *Bioinformatics* 2023;**39**:btad573.
- Wohlwend J, Corso G, Passaro S *et al.* Boltz-1 democratizing biomolecular interaction modeling. *bioRxiv*, 2024.11.19.624167, 2025, preprint: not peer reviewed.
- Yin R, Pierce BG. Evaluation of AlphaFold antibody–antigen modeling with implications for improving predictive accuracy. *Protein Sci* 2024;**33**:e4865.
- Zhang C, Shine M, Pyle AM *et al.* US-align: universal structure alignments of proteins, nucleic acids, and macromolecular complexes. *Nat Methods* 2022;**19**:1109–15.
- Zhang Y, Skolnick J. Scoring function for automated assessment of protein structure template quality. *Proteins* 2004;**57**:702–10.
- Zhu W, Shenoy A, Kundrotas P *et al.* Evaluation of AlphaFold-multimer prediction on multi-chain protein complexes. *Bioinformatics* 2023;**39**:btad424.

© The Author(s) 2026. Published by Oxford University Press.

This is an Open Access article distributed under the terms of the Creative Commons Attribution License (<https://creativecommons.org/licenses/by/4.0/>), which permits unrestricted reuse, distribution, and reproduction in any medium, provided the original work is properly cited.

Bioinformatics, 2026, 42, 1–9

<https://doi.org/10.1093/bioinformatics/btag136>

Original Paper