

NVIDIA Just Helped Map 31 Million Protein Complexes and the Health Tech Investment Implications Are Enormous

APR 10, 2026



2



1

Share

Abstract

- NVIDIA, Google DeepMind, EMBL-EBI, and Seoul National University expanded the AlphaFold Protein Structure Database (AFDB) from monomeric protein structures to proteome-scale quaternary (complex) structures, predicting over 31 million homo- and heteromeric protein complexes across 4,777 proteomes
- 1.8 million high-confidence homodimer structures are now publicly available through AFDB, with the full 31M set coming for bulk download
- GPU-accelerated infrastructure running on H100 DGX Superpod clusters, using MMseqs2-GPU for multiple sequence alignment and TensorRT plus cuEquivari for deep learning inference, enabled this scale of computation
- The work used STRING database physical interaction annotations to define biologically relevant heterodimer candidates, yielding ~8M heterodimer predictions with 57K tentatively high-confidence results
- Clustering of high-confidence complexes showed extreme concentration: the top 100 of structural representatives account for ~25% of all complexes, and ~9% of clusters are conserved across superkingdoms
- Downstream applications include drug target validation, variant interpretation, protein interfaces, generative protein design benchmarking, and systems-level structural biology

- This represents a foundational shift in the computational drug discovery stack
significant implications for health tech founders and investors evaluating compa
in structural biology, protein engineering, and AI-driven therapeutics

Table of Contents

Why Protein Complexes Matter More Than Monomers

What Actually Got Built Here

The GPU Infrastructure Story

Confidence Calibration and the Heterodimer Problem

What the Clustering Reveals About Biology

The Drug Discovery and Health Tech Investment Angle

What This Means for Founders Building in This Space

Where This Goes Next

Why Protein Complexes Matter More Than Monomers

So AlphaFold2 was a massive deal. No dispute there. The Nobel Prize, the datab
200M+ predicted monomeric protein structures, the complete transformation of
computational structural biology. But here is the thing that has been nagging at
people in this space for years now: proteins almost never work alone. They form
complexes. Dimers, trimers, big gnarly multi-subunit assemblies. The biological
action happens at the interfaces between proteins, not just within the isolated 3D
of a single chain. And for most of those interfaces, structural information has be
basically nonexistent at any kind of useful scale.

The Protein Data Bank, which houses experimentally determined structures, contains a tiny fraction of known protein-protein interactions. For most organisms, the number of experimentally resolved multimeric structures is one to three orders of magnitude below what you would need to do serious systems biology or structure-based drug design against interaction surfaces. This is not a minor gap. This is the gap. If you're trying to interpret variants at protein interfaces, or validate drug targets that depend on complex formation, or benchmark generative protein design models, you have been operating with one hand tied behind your back.

What the NVIDIA-DeepMind-EMBL-EBI-Seoul National University collaboration just shipped is a direct assault on that bottleneck. Over 31 million predicted homodimer and heteromeric protein complexes. 1.8 million of them classified as high-confidence and now surfaced through the AlphaFold Database at alphafold.com. For health investors and founders working anywhere near the structural biology stack, this is worth understanding in detail because it changes the ground truth assumptions underpinning a lot of computational approaches in drug discovery and protein engineering.

What Actually Got Built Here

The team predicted 23.4 million homodimers derived from 4,777 proteomes in UniProt, including 16 model organisms and 30 WHO global health proteomes. Then they added approximately 7.6 million heterodimer candidates extracted from the STRING database using physical protein-protein interaction annotations. That is a staggering combinatorial space. The heterodimer problem in particular is nasty because the number of possible pairwise interactions grows quadratically with proteome size. You cannot just do all-against-all predictions for large proteomes; you expect to finish before the heat death of the universe.

Their approach to scoping the heterodimer set was pragmatic. They used STRING interaction evidence to filter down to physically interacting pairs, restricted to the same proteome (no inter-proteome complexes), and focused on dimers with a maximum combined sequence length of 3,000 amino acids. Critically, they did not

filter by STRING score threshold for their initial computation, choosing coverage precision. The literature suggests that filtering for STRING scores above 700 further reduces inputs while increasing prediction quality, but the team wanted maximum coverage for these priority proteomes and left that tighter filter as an option for downstream users.

For MSA generation, they used ColabFold's search tool with the MMseqs2-GPU backend, keeping only the best hit per taxon based on alignment score. This is an orthology filter that prevents paralogous sequences from diluting the evolutionary signals that AlphaFold-Multimer needs to predict complex formation accurately. For heterodimers, they just concatenated the homodimer MSAs without pairing, which sounds lazy but actually held up well in their validation. They compared taxon- or taxon-based pairing against simple concatenation and found that additional pairing did clearly yield better predictions, especially at higher confidence thresholds.

Structure prediction ran through either ColabFold or an accelerated OpenFold implementation. Both used the same parameters: one set of weights from AlphaFold-Multimer (model_1_multimer_v3), four recycles with early stopping, and no relaxation. The choice to skip relaxation saves compute without meaningfully hurting accuracy for the purposes of database-scale prediction. On a benchmark set of 1,000 ray-resolved PDB homodimers released after AlphaFold2 was introduced (minimizing training data leakage), OpenFold accelerated with TensorRT and cuEquivariance matched ColabFold interface accuracy. The accelerated pipeline hit 75.4% usable predictions (DockQ above 0.3) compared to ColabFold at 73%, with mean DockQ scores of 0.647 versus 0.637. Not a massive difference, but the throughput gains in the accelerated stack are where the real story is.

The GPU Infrastructure Story

This is where things get genuinely interesting from an infrastructure perspective. The team ran on H100 DGX Superpod clusters and faced the classic HPC problem of maximizing GPU utilization across two workloads (MSA generation and structure prediction) that scale very differently.

For MSA generation with MMseqs2-GPU, the GPUs are only used during the ungapped filter stages. The subsequent alignment stages are multithreaded CPU processes. So you end up with a lot of GPU idle time if you just run one job at a time. Their solution was to stagger multiple colabfold_search processes per GPU, monitoring output to kick off the next one as soon as the GPU was free from the previous run. On a DGX H100 node, they found that three staggered processes could increase overall throughput by up to 25%, though individual chunks process more slowly due to CPU oversubscription. Not a perfect solution but a pragmatic one.

Chunk sizing matters here too. Smaller chunks mean more per-process overhead (database loading takes a couple minutes even on fast storage), while larger chunks take longer to finish and risk hitting SLURM wall time limits. For their setup with a 4-hour wall time limit, chunks of 300 sequences worked well. They also found that pre-staging databases on node-local SSDs helped throughput.

For structure prediction with ColabFold, they got higher throughput by packing homodimers of equal length into batches sorted by MSA depth in descending order. This reduces JAX recompilations, which is a surprisingly big deal for throughput at scale. This trick does not work for heterodimers where chain lengths differ, which is annoying. For OpenFold, the recompilation problem does not exist, but sequence length still drives execution time, so they reserved longer sequences for individual jobs and overlapped CPU-bound featurization of the next query with GPU-bound inference of the current one.

The broader SLURM orchestration story involved packing multiple predictions per node, matching GPU memory to sequence length, separating short versus long sequence queues, and monitoring GPU memory fragmentation. Asynchronous I/O helped avoid disk bottlenecks. None of this is glamorous work but it is the kind of systems engineering that determines whether a project like this takes three months or three years.

Confidence Calibration and the Heterodimer Problem

This section is arguably the most important for anyone who wants to actually use these structures, because confidence calibration is where the monomeric AlphaFold experience breaks down for complexes.

For monomers, pLDDT (predicted Local Distance Difference Test) gives you a pretty good sense of per-residue confidence. Above 70 is generally good, above 90 is great. But for complexes, the problem is fundamentally harder. You need to assess not whether each chain is folded correctly but whether the interface between chains is plausible and positioned in the right pocket. That requires evaluating global and chain confidence metrics alongside local confidence metrics at the interface. With more dimensions, way less training data to calibrate against.

The team evaluated four scoring metrics against a curated ground truth set of 1,000 PDB homodimers and 2,211 PDB monomers (as negative controls), all released at AlphaFold2's training cutoff. They looked at ipTM (interface predicted TM-score), ipSAEmin (the minimum of the bidirectional interaction prediction Score from Aligned Errors), LISmin (Local Interaction Score), and pDockQ2. Of these, ipSAEmin showed the cleanest distributional separation between true homodimers and monomers, and the most stable F1 plateau across cutoffs.

They settled on a high-confidence threshold of ipSAEmin at or above 0.6, pLDDT average at or above 70, and backbone clashes at or below 10. This yielded precision of 0.859, recall of 0.655, and F1 of 0.744. Roughly 7% of homodimer predictions passed this filter, giving 1.8 million high-confidence homodimers. The AFDB website further categorizes these into “very high confidence” (ipSAEmin at or above 0.8, about 90K entries), “confident” (0.7 to 0.8, about 439K), and “low confidence” (0.6 to 0.7, about 343K).

Here is where it gets tricky. When they applied the same homodimer-derived thresholds to the 7.6 million heterodimer predictions, only about 57,000 passed. That is a tiny fraction, and the heterodimers that did pass showed a strong bias toward homodimer-like properties: smaller length differences between chains, higher inter-chain sequence identity. This is a real caveat. The current filtering criteria may be systematically excluding biologically real heterodimeric complexes that just happen

look less like homodimers. The team explicitly flags these 57K as “tentatively high confidence” and says further calibration is needed before releasing a more representative heterodimer set.

For health tech investors, this matters because a lot of the most therapeutically interesting protein-protein interactions are heteromeric. Drug targets at heteromeric interfaces, signaling pathway complexes, antibody-antigen interactions. The homodimer expansion is valuable and immediately useful, but the heterodimer set is where the bigger drug discovery value lives, and it is not fully baked yet.

What the Clustering Reveals About Biology

The team clustered all 1.8 million high-confidence structures using Foldseek Multimercluster, which compressed the dataset roughly 8-fold into about 225,000 clusters. Of these, about 87,000 were non-singletons (had at least one other member). The distribution of cluster sizes is telling.

The top 1% of non-singleton cluster representatives cover approximately 25% of entries, and the top 20% cover approximately 82%. This is a power law distribution that means predicted complex space is concentrated around a relatively small number of recurrent structural solutions. Nature keeps reusing the same interfaces. For protein engineering and generative design, this is useful information because it tells you where the structural density is and where genuinely novel folds might be hiding.

Clusters without any detectable PDB multimer match were more frequent among smaller clusters. The biggest clusters tend to overlap with known multimeric structures, which makes sense since the most common biological solutions are also the most experimentally characterized. The rare clusters, the ones with fewer members and no PDB match, are potentially the most interesting from a basic science perspective. These are predicted complex structures that nobody has crystallized or cryo-EM'd yet.

The taxonomic analysis is fascinating. About 9% of non-singleton clusters contain members from at least two different superkingdoms (bacteria, archaea, eukaryotes). These complexes likely originated in a common ancestor and have been maintained as universal building blocks of cellular life for billions of years. That is remarkable evolutionary conservation. Archaea and bacteria showed higher prediction success rates than eukaryotes, likely because prokaryotic proteins tend to be shorter, more compact, and richer in homo-oligomeric assemblies. Eukaryotic proteins are longer, more multi-domain, and more often participate in heteromeric complexes that are harder to predict.

The Drug Discovery and Health Tech Investment Angle

There are several concrete downstream applications that flow from having 1.8 million (and eventually 31M+) predicted complex structures publicly available. The first most obvious is variant interpretation at protein interfaces. When you find a variant of uncertain significance through genomic sequencing, the question is always whether it affects protein function. If the variant sits at a protein-protein interface in a pre-formed complex, that is immediately informative in a way that monomeric structure alone cannot be. This matters for clinical genomics companies, rare disease diagnostic platforms, and anyone building tools for variant classification.

Drug target validation gets a boost too. Lots of drug targets depend on protein complex formation for their biological function. Having structural hypotheses for those complexes, even at moderate confidence, gives computational chemists and medicinal chemists a starting point for structure-based drug design at interfaces. Interface-directed drug design is harder than targeting a well-defined binding pocket on a monomer, but it is also where some of the most compelling therapeutic opportunities live, particularly in oncology and immunology.

Generative protein design benchmarking is another big one. Companies building protein design platforms (de novo binders, engineered enzymes, designed protein therapeutics) need benchmark datasets to validate their models. This dataset provides

1.8 million complex structures with calibrated confidence metrics. That is a serious training and benchmarking resource for anyone in the generative bio space.

Systems-level structural biology is the broader scientific play. Being able to overlay structural information onto interaction networks from resources like STRING creates a new kind of structural systems biology that was previously impossible at proteomic scale. For health tech companies building knowledge graphs or multi-omic analysis platforms, this is another data layer to integrate.

The infrastructure itself is also investable. The fact that NVIDIA is shipping MMseqs2-GPU, cuEquivariance, and TensorRT as freely available libraries (Apache 2.0 licensing), and offering inference microservices through NIMs for MSA search, protein folding, means the barrier to running these kinds of analyses is dropping. A startup that would have needed six months and a million dollars in compute to run a large-scale complex prediction campaign can now potentially do it in weeks for much less. That changes the economics of computational structural biology startups.

What This Means for Founders Building in This Space

If you are founding or building a company anywhere in the structural biology or computational drug discovery stack, this release changes a few things worth thinking about.

First, the data moat argument for structural prediction companies just got weaker. NVIDIA and DeepMind are going to keep expanding AFDB with complex structures at this pace, and the inference tools are freely available, then simply having predicted structures is not a defensible position. The value has to come from what you do with the structures: interpretation, design, integration into clinical or drug development workflows. The raw prediction layer is being commoditized in real time.

Second, the confidence calibration problem for heterodimers is an open research and commercial opportunity. The team explicitly acknowledged that their homodimer-derived thresholds do not work well for heterodimers. If someone builds better

confidence metrics or better models for heteromeric complex prediction, that is genuinely differentiable right now. Companies like Protai (which NVIDIA has highlighted as using AlphaFold with proteomics and NVIDIA NIM for complex prediction in drug discovery) are already operating in this space.

Third, the integration opportunity is enormous. Most drug discovery and clinical genomics platforms have been built on monomeric structure assumptions. Retrofitting them to incorporate complex structure information, especially with calibrated confidence, is nontrivial engineering and science. There is real value in being the integration layer that makes complex structure predictions actionable for therapeutic development or clinical interpretation.

Fourth, compute economics continue to favor GPU-native approaches. The stage MSA generation, the sequence packing tricks, the decoupled pipeline architecture described in this work represent significant systems engineering knowledge. Startups that understand how to run these workloads efficiently on modern GPU clusters have meaningful cost advantages over those that treat compute as a black box.

The case studies in the paper are worth reading in full because they illustrate the kinds of biological insights that only emerge from complex prediction. There is a transcription elongation factor from *Dictyostelium* that has a completely fragmented low-confidence monomeric prediction (pLDDT of 50.56) but forms a clean, high confidence domain-swapped homodimer (pLDDT of 86.06). The fold literally does not exist without the partner chain. There is a membrane protein from a fungal pathogen where the monomeric prediction is mediocre but the dimeric model properly defines the membrane boundaries. There is a *Mycoplasma* transcriptional regulator where the monomer prediction is garbage (pLDDT of 56) but the dimer rescues it to high confidence (pLDDT of 85). These are not edge cases. For some proteins, monomeric prediction provides an incomplete or actively misleading structural picture. The real implications for anyone relying on AlphaFold monomer predictions as the ground truth for their analysis.

Where This Goes Next

The team has been explicit that this is not the final state. The full 31M prediction (including the 21M+ homodimers below the high-confidence threshold and the 7 heterodimers) will be released for bulk download. Better heterodimer confidence calibration is coming. The prediction tools themselves continue to improve: OpenFold3, Boltz-2, and NVIDIA's own ProteinA are all advancing the frontier of complex structure prediction accuracy.

The convergence of GPU-accelerated inference, improved prediction models, and large-scale public databases is creating a new baseline for computational structural biology. For health tech investors, the question is no longer whether accurate protein complex structures will be widely available. They will. The question is who builds the most valuable applications on top of that infrastructure. Drug discovery platforms that can exploit interface-level structural information. Clinical genomics tools that interpret variants in the context of complex formation. Protein engineering companies that design novel interactions using these structures as templates. Biosecurity and pandemic preparedness applications that leverage pathogen-host interaction predictions from WHO priority proteomes.

The AlphaFold Database expansion from monomers to complexes is not just an incremental database update. It is a shift in what kind of structural biology is computationally accessible at population scale. For anyone investing in or building companies at the intersection of AI, structural biology, and therapeutics, ignoring this would be a mistake.



2 Likes • 1 Restack

← Previous

Discussion about this post

Comments

Restacks



Write a comment...