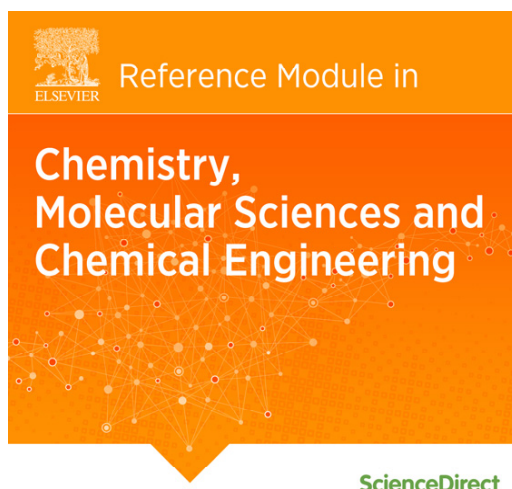


**Provided for non-commercial research and educational use.
Not for reproduction, distribution or commercial use.**

This article was published in the Elsevier Reference Module in *Chemistry, Molecular Sciences and Chemical Engineering*, and the attached copy is provided by Elsevier for the author's benefit and for the benefit of the author's institution, for non-commercial research and educational use including without limitation use in instruction at your institution, sending it to specific colleagues who you know, and providing a copy to your institution's administrator.



All other uses, reproduction and distribution, including without limitation commercial reprints, selling or licensing copies or access, or posting on open internet sites, your personal or institution's website or repository, are prohibited. For exceptions, permission may be sought for such use through Elsevier's permissions site at:

<http://www.elsevier.com/locate/permissionusematerial>

Webb B., Eswar N., Fan H., Khuri N., Pieper U., Dong G.Q. and Sali A. (2014) Comparative Modeling of Drug Target Proteins. In: Reedijk, J. (Ed.) Elsevier Reference Module in Chemistry, Molecular Sciences and Chemical Engineering. Waltham, MA: Elsevier. 29-Sep-14 doi: 10.1016/B978-0-12-409547-2.11133-3.

© 2014 Elsevier Inc. All rights reserved.

Comparative Modeling of Drug Target Proteins[☆]

B Webb, N Eswar, H Fan, N Khuri, U Pieper, GQ Dong, and A Sali, University of California at San Francisco, San Francisco, CA, USA

© 2014 Elsevier Inc. All rights reserved.

Introduction	2
Structure-Based Drug Discovery	2
The Sequence–Structure Gap	2
Structure Prediction Addresses the Sequence–Structure Gap	2
The Basis of Comparative Modeling	2
Comparative Modeling Benefits from Structural Genomics	2
Outline	2
Steps in Comparative Modeling	3
Fold Assignment and Sequence–Structure Alignment	3
Fold assignment	3
Three levels of similarity	3
Sequence–sequence methods	4
Sequence–profile methods	4
Profile–profile methods	4
Sequence–structure threading methods	5
Iterative sequence–structure alignment	5
Most Alignment Errors are Unrecoverable	5
Template Selection	5
Model Building	5
Three Approaches to Comparative Model Building	5
Modeller: Comparative Modeling by Satisfaction of Spatial Restraints	6
Relative Accuracy, Flexibility, and Automation	6
Refinement of Comparative Models	6
Loop Modeling	6
Definition of the problem	6
Two classes of methods	6
Side-Chain Modeling	7
Fixed backbone	7
Rotamers	7
Methods	7
Errors in Comparative Models	7
Selection of Incorrect Templates	7
Errors due to Misalignments	7
Errors in Regions without a Template	8
Distortions and Shifts in Correctly Aligned Regions	8
Errors in Side-Chain Packing	9
Prediction of Model Errors	9
Initial Assessment of the Fold	9
Self-Consistency	9
Final Assessment of the Fold (Model)	9
Evaluation of Comparative Modeling Methods	9
Applications of Comparative Models	10
Comparative Models versus Experimental Structures in Virtual Screening	11
Use of Comparative Models to Obtain Novel Drug Leads	11
Future Directions	12
Automation and Availability of Resources for Comparative Modeling and Ligand Docking	13
Acknowledgments	15
References	16

[☆]*Change History:* August 2014. B Webb updated the reference section and brought the entire text up to date.

H Fan updated section 'Comparative Models versus Experimental Structures in Virtual Screening' with new studies, and removed the old Figure 5.

N Khuri added new studies to section 'Use of Comparative Models to Obtain Novel Drug Leads.'

U Pieper updated section 'Automation and Availability of Resources for Comparative Modeling and Ligand Docking,' Table 1, and Figure 5.

GQ Dong added a new section 'Final Assessment of the Fold (Model selection).'

Introduction

Structure-Based Drug Discovery

Structure-based or rational drug discovery has already resulted in a number of drugs on the market and many more in the development pipeline.^{1–4} Structure-based methods are now routinely used in almost all stages of drug development, from target identification to lead optimization.^{5–8} Central to all structure-based discovery approaches is the knowledge of the three-dimensional (3D) structure of the target protein or complex because the structure and dynamics of the target determine which ligands it binds. The 3D structures of the target proteins are best determined by experimental methods that yield solutions at atomic resolution, such as X-ray crystallography and nuclear magnetic resonance (NMR) spectroscopy.⁹ While developments in the techniques of experimental structure determination have enhanced the applicability, accuracy, and speed of these structural studies,^{10,11} structural characterization of sequences remains an expensive and time-consuming task.

The Sequence–Structure Gap

The publicly available Protein Data Bank (PDB)¹² currently contains ~92 000 structures and grows at a rate of approximately 40% every 2 years. On the other hand, the various genome-sequencing projects have resulted in over 40 million sequences, including the complete genetic blueprints of humans and hundreds of other organisms.^{13,14} This achievement has resulted in a vast collection of sequence information about possible target proteins with little or no structural information. Current statistics show that the structures available in the PDB account for less than 1% of the sequences in the UniProt database.¹³ Moreover, the rate of growth of the sequence information is more than twice that of the structures, and is expected to accelerate even more with the advent of readily available next-generation sequencing technologies. Due to this wide sequence–structure gap, reliance on experimentally determined structures limits the number of proteins that can be targeted by structure-based drug discovery.

Structure Prediction Addresses the Sequence–Structure Gap

Fortunately, domains in protein sequences are gradually evolving entities that can be clustered into a relatively small number of families with similar sequences and structures.^{15,16} For instance, 75–80% of the sequences in the UniProt database have been grouped into fewer than 15 000 domain families.^{17,18} Similarly, all the structures in the PDB have been classified into about 1000 distinct folds.^{19,20} Computational protein structure prediction methods, such as threading²¹ and comparative protein structure modeling,^{22,23} strive to bridge the sequence–structure gap by utilizing these evolutionary relationships. The speed, low cost, and relative accuracy of these computational methods have led to the use of predicted 3D structures in the drug discovery process.^{24,25} The other class of prediction methods, *de novo* or *ab initio* methods, attempts to predict the structure from sequence alone, without reliance on evolutionary relationships. However, despite progress in these methods,^{26–28} especially for small proteins with fewer than 100 amino acid residues, comparative modeling remains the most reliable method of predicting the 3D structure of a protein, with an accuracy that can be comparable to a low-resolution, experimentally determined structure.⁹

The Basis of Comparative Modeling

The primary requirement for reliable comparative modeling is a detectable similarity between the sequence of interest (target sequence) and a known structure (template). As early as 1986, Chothia and Lesk²⁹ showed that there is a strong correlation between sequence and structural similarities. This correlation provides the basis of comparative modeling, allows a coarse assessment of model errors, and also highlights one of its major challenges: modeling the structural differences between the template and target structures³⁰ (Figure 1).

Comparative Modeling Benefits from Structural Genomics

Comparative modeling stands to benefit greatly from the structural genomics initiative.³¹ Structural genomics aims to achieve significant structural coverage of the sequence space with an efficient combination of experimental and prediction methods.³² This goal is pursued by careful selection of target proteins for structure determination by X-ray crystallography and NMR spectroscopy, such that most other sequences are within ‘modeling distance’ (e.g., >30% sequence identity) of a known structure.^{15,16,31,33} The expectation is that the determination of these structures combined with comparative modeling will yield useful structural information for the largest possible fraction of sequences in the shortest possible timeframe. The impact of structural genomics is illustrated by comparative modeling based on the structures determined by the New York Structural Genomics Research Consortium. For each new structure without a close homolog in the PDB, on average, 3500 protein sequences without any prior structural characterization could be modeled at least at the level of the fold.³⁴ Thus, the structures of most proteins will eventually be predicted by computation, not determined by experiment.

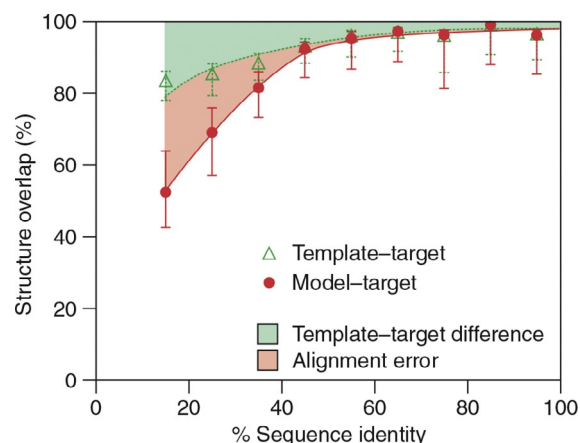


Figure 1 Average model accuracy as a function of sequence identity.³⁰ As the sequence identity between the target sequence and the template structure decreases, the average structural similarity between the template and the target also decreases (dashed line, triangles).²⁹ Structural overlap is defined as the fraction of equivalent C α atoms. For the comparison of the model with the actual structure (filled circles), two C α atoms were considered equivalent if they belonged to the same residue and were within 3.5 Å of each other after least-squares superposition. For comparisons between the template structure and the actual target structure (triangles), two C α atoms were considered equivalent if they were within 3.5 Å of each other after alignment and rigid-body superposition. The difference between the model and the actual target structure is a combination of the target–template differences (green area) and the alignment errors (red area). The figure was constructed by calculating 3993 comparative models based on a single template of varying similarity to the targets. All targets had known (experimentally determined) structures.³⁰

Outline

In this review, we begin by describing the various steps involved in comparative modeling. Next, we emphasize two aspects of model refinement, loop modeling and side-chain modeling, due to their relevance in ligand docking and rational drug discovery. We then discuss the errors in comparative models. Finally, we describe the role of comparative modeling in drug discovery, focusing on ligand docking against comparative models. We compare successes of docking against models and X-ray structures, and illustrate the computational docking against models with a number of examples. We conclude with a summary of topics that will impact on the future utility of comparative modeling in drug discovery, including an automation and integration of resources required for comparative modeling and ligand docking.

Steps in Comparative Modeling

Comparative modeling consists of four main steps²³ (Figure 2(a)): (1) fold assignment that identifies similarity between the target sequence of interest and at least one known protein structure (the template); (2) alignment of the target sequence and the template(s); (3) building a model based on the alignment with the chosen template(s); and (4) predicting model errors.

Fold Assignment and Sequence–Structure Alignment

Although fold assignment and sequence–structure alignment are logically two distinct steps in the process of comparative modeling, in practice almost all fold assignment methods also provide sequence–structure alignments. In the past, fold assignment methods were optimized for better sensitivity in detecting remotely related homologs, often at the cost of alignment accuracy. However, recent methods simultaneously optimize both the sensitivity and alignment accuracy. Therefore, in the following discussion, we will treat fold assignment and sequence–structure alignment as a single protocol, explaining the differences as needed.

Fold assignment

As mentioned earlier, the primary requirement for comparative modeling is the identification of one or more known template structures with detectable similarity to the target sequence. The identification of suitable templates is achieved by scanning structure databases, such as PDB,¹² SCOP,¹⁹ DALI,³⁶ and CATH,²⁰ with the target sequence as the query. The detected similarity is usually quantified in terms of sequence identity or statistical measures, such as *E*-value or *z*-score, depending on the method used.

Three levels of similarity

Sequence–structure relationships are coarsely classified into three different regimes in the sequence similarity spectrum: (1) the easily detected relationships characterized by >30% sequence identity; (2) the ‘twilight zone,’³⁷ corresponding to relationships

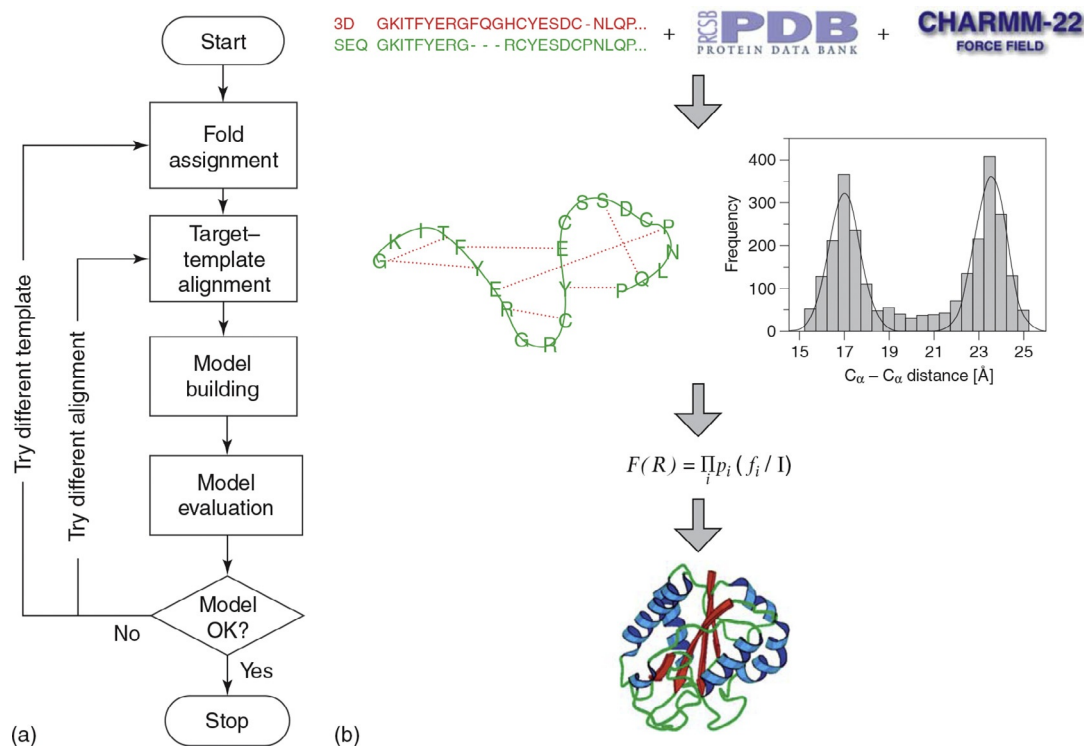


Figure 2 Comparative protein structure modeling. (a) A flowchart illustrating the steps in the construction of a comparative model.²³ (b) Description of comparative modeling by extraction of spatial restraints as implemented in Modeller.³⁵ By default, spatial restraints in Modeller include: (1) homology-derived restraints from the aligned template structures; (2) statistical restraints derived from all known protein structures; and (3) stereochemical restraints from the CHARMM-22 molecular mechanics force field. These restraints are combined into an objective function that is then optimized to calculate the final 3D model of the target sequence.

with statistically significant sequence similarity in the 10–30% range; and (3) the 'midnight zone',³⁷ corresponding to statistically insignificant sequence similarity.

Sequence–sequence methods

For closely related protein sequences with identities higher than 30–40%, the alignments produced by all methods are almost always largely correct. The quickest way to search for suitable templates in this regime is to use simple pairwise sequence alignment methods such as SSEARCH,³⁸ BLAST,³⁹ and FASTA.³⁸ Brenner et al. showed that these methods detect only ~18% of the homologous pairs at less than 40% sequence identity, while they identify more than 90% of the relationships when sequence identity is between 30% and 40%.⁴⁰ Another benchmark, based on 200 reference structural alignments with 0–40% sequence identity, indicated that BLAST is able to correctly align only 26% of the residue positions.⁴¹

Sequence–profile methods

The sensitivity of the search and accuracy of the alignment become progressively difficult as the relationships move into the twilight zone.^{37,42} A significant improvement in this area was the introduction of profile methods by Gribskov and co-workers.⁴³ The profile of a sequence is derived from a multiple sequence alignment and specifies residue-type occurrences for each alignment position. The information in a multiple sequence alignment is most often encoded as either a position-specific scoring matrix (PSSM)^{39,44,45} or as a hidden Markov model (HMM).^{46,47} To identify suitable templates for comparative modeling, the profile of the target sequence is used to search against a database of template sequences. The profile–sequence methods are more sensitive in detecting related structures in the twilight zone than the pairwise sequence-based methods; they detect approximately twice the number of homologs under 40% sequence identity.^{41,48,49} The resulting profile–sequence alignments correctly align approximately 43–48% of residues in the 0–40% sequence identity range;^{41,50} this number is almost twice as large as that of the pairwise sequence methods. Frequently used programs for profile–sequence alignment are PSI-BLAST,³⁹ SAM,⁵¹ HMMER,⁴⁶ HHsearch,⁵² HHBlits,⁵³ and BUILD_PROFILE.⁵⁴

Profile–profile methods

As a natural extension, the profile–sequence alignment methods have led to profile–profile alignment methods that search for suitable template structures by scanning the profile of the target sequence against a database of template profiles, as opposed to a database of template sequences. These methods have proven to include the most sensitive and accurate fold assignment and

alignment protocols to date.^{50,55–57} Profile–profile methods detect ~28% more relationships at the superfamily level and improve the alignment accuracy by 15–20% compared to profile–sequence methods.^{50,58} There are a number of variants of profile–profile alignment methods that differ in the scoring functions they use.^{50,55,58–64} However, several analyses have shown that the overall performances of these methods are comparable.^{50,55–57} Some of the programs that can be used to detect suitable templates are FFAS,⁶⁵ SP3,⁵⁸ SALIGN,⁵⁰ HHBlits,⁵³ HHsearch,⁵² and PPSCAN.⁵⁴

Sequence–structure threading methods

As the sequence identity drops below the threshold of the twilight zone, there is usually insufficient signal in the sequences or their profiles for the sequence-based methods discussed above to detect true relationships.⁴⁸ Sequence–structure threading methods are most useful in this regime as they can sometimes recognize common folds, even in the absence of any statistically significant sequence similarity.²¹ These methods achieve higher sensitivity by using structural information derived from the templates. The accuracy of a sequence–structure match is assessed by the score of a corresponding coarse model and not by sequence similarity, as in sequence comparison methods.²¹ The scoring scheme used to evaluate the accuracy is either based on residue substitution tables dependent on structural features such as solvent exposure, secondary structure type, and hydrogen bonding properties,^{58,66–68} or on statistical potentials for residue interactions implied by the alignment.^{69–73} The use of structural data does not have to be restricted to the structure side of the aligned sequence–structure pair. For example, SAM-T08 makes use of the predicted local structure for the target sequence to enhance homolog detection and alignment accuracy.⁷⁴ Commonly used threading programs are GenTHREADER,^{66,75} 3D-PSSM,⁷⁶ FUGUE,⁶⁸ SP3,⁵⁸ SAM-T08 multitrack HMM,^{67,74,77} and MUSTER.⁷⁸

Iterative sequence–structure alignment

Yet another strategy is to optimize the alignment by iterating over the process of calculating alignments, building models, and evaluating models. Such a protocol can sample alignments that are not statistically significant and identify the alignment that yields the best model. Although this procedure can be time-consuming, it can significantly improve the accuracy of the resulting comparative models in difficult cases.⁷⁹

Most Alignment Errors are Unrecoverable

Regardless of the method used, searching in the twilight and midnight zones of the sequence–structure relationship often results in false negatives, false positives, or alignments that contain an increasingly large number of gaps and alignment errors. Improving the performance and accuracy of methods in this regime remains one of the main tasks of comparative modeling today.⁸⁰ It is imperative to calculate an accurate alignment between the target–template pair. Although some progress has been made recently,⁸¹ comparative modeling can rarely recover from an alignment error.⁸²

Template Selection

After a list of all related protein structures and their alignments with the target sequence have been obtained, template structures are prioritized depending on the purpose of the comparative model. Template structures may be chosen purely based on the target–template sequence identity or a combination of several other criteria, such as experimental accuracy of the structures (resolution of X-ray structures, number of restraints per residue for NMR structures), conservation of active-site residues, holo-structures that have bound ligands of interest, and prior biological information that pertains to the solvent, pH, and quaternary contacts. It is not necessary to select only one template. In fact, the use of several templates approximately equidistant from the target sequence generally increases the model accuracy.^{83,84}

Model Building

Three Approaches to Comparative Model Building

Once an initial target–template alignment is built, a variety of methods can be used to construct a 3D model for the target protein.^{23,82,85–88} The original and still widely used method is modeling by rigid-body assembly.^{86,87,89} This method constructs the model from a few core regions, and from loops and side chains that are obtained by dissecting related structures. Commonly used programs that implement this method are COMPOSER,^{90–93} 3D-JIGSAW,⁹⁴ RosettaCM,⁸¹ and SWISS-MODEL.⁹⁵ Another family of methods, modeling by segment matching, relies on the approximate positions of conserved atoms from the templates to calculate the coordinates of other atoms.^{96–100} An instance of this approach is implemented in SegMod.⁹⁹ The third group of methods, modeling by satisfaction of spatial restraints, uses either distance geometry or optimization techniques to satisfy spatial restraints obtained from the alignment of the target sequences with the template structures.^{35,101–104} Specifically, Modeller,^{35,105,106} our own program for comparative modeling, belongs to this group of methods.

Modeller: Comparative Modeling by Satisfaction of Spatial Restraints

Modeller implements comparative protein structure modeling by the satisfaction of spatial restraints that include: (1) homology-derived restraints on the distances and dihedral angles in the target sequence, extracted from its alignment with the template structures;³⁵ (2) stereochemical restraints such as bond length and bond angle preferences, obtained from the CHARMM-22 molecular mechanics force field;¹⁰⁷ (3) statistical preferences for dihedral angles and nonbonded interatomic distances, obtained from a representative set of known protein structures;¹⁰⁸ and (4) optional manually curated restraints, such as those from NMR spectroscopy, rules of secondary structure packing, cross-linking experiments, fluorescence spectroscopy, image reconstruction from electron microscopy, site-directed mutagenesis, and intuition (Figure 2(b)). The spatial restraints, expressed as probability density functions, are combined into an objective function that is optimized by a combination of conjugate gradients and molecular dynamics with simulated annealing. This model-building procedure is similar to structure determination by NMR spectroscopy.

Relative Accuracy, Flexibility, and Automation

Accuracies of the various model-building methods are relatively similar when used optimally.^{109,110} Other factors, such as template selection and alignment accuracy, usually have a larger impact on the model accuracy, especially for models based on less than 30% sequence identity to the templates. However, it is important that a modeling method allows a degree of flexibility and automation to obtain better models more easily and rapidly. For example, a method should allow for an easy recalculation of a model when a change is made in the alignment; it should be straightforward to calculate models based on several templates; and the method should provide tools for incorporation of prior knowledge about the target (e.g., cross-linking restraints and predicted secondary structure).

Refinement of Comparative Models

Protein sequences evolve through a series of amino acid residue substitutions, insertions, and deletions. While substitutions can occur throughout the length of the sequence, insertions and deletions mostly occur on the surface of proteins in segments that connect regular secondary structure segments (i.e., loops). While the template structures are helpful in the modeling of the aligned target backbone segments, they are generally less valuable for the modeling of side chains and irrelevant for the modeling of insertions such as loops. The loops and side chains of comparative models are especially important for ligand docking; thus, we discuss them in the following two sections.

Loop Modeling

Definition of the problem

Loop modeling is an especially important aspect of comparative modeling in the range from 30% to 50% sequence identity. In this range of overall similarity, loops among the homologs vary while the core regions are still relatively conserved and aligned accurately. Loops often play an important role in defining the functional specificity of a given protein, forming the active and binding sites. Loop modeling can be seen as a mini protein folding problem because the correct conformation of a given segment of a polypeptide chain has to be calculated mainly from the sequence of the segment itself. However, loops are generally too short to provide sufficient information about their local fold. Even identical decapeptides in different proteins do not always have the same conformation.^{111,112} Some additional restraints are provided by the core anchor regions that span the loop and by the structure of the rest of the protein that cradles the loop. Although many loop-modeling methods have been described, it is still challenging to correctly and confidently model loops longer than approximately 10–12 residues.^{105,113,114}

Two classes of methods

There are two main classes of loop-modeling methods: (1) database search approaches that scan a database of all known protein structures to find segments fitting the anchor core regions^{98,115}; and (2) conformational search approaches that rely on optimizing a scoring function.^{116–118} There are also methods that combine these two approaches.^{119,120}

Database-based loop modeling

The database search approach to loop modeling is accurate and efficient when a database of specific loops is created to address the modeling of the same class of loops, such as β -hairpins,¹²¹ or loops on a specific fold, such as the hypervariable regions in the immunoglobulin fold.^{115,122} There are attempts to classify loop conformations into more general categories, thus extending the applicability of the database search approach.^{123–125} However, the database methods are limited because the number of possible conformations increases exponentially with the length of a loop, and until the late 1990s only loops up to 7 residues long could be modeled using the database of known protein structures.^{126,127} However, the growth of the PDB in recent years has largely eliminated this problem.¹²⁸

Optimization-based methods

There are many optimization-based methods, exploiting different protein representations, objective functions, and optimization or enumeration algorithms. The search algorithms include the minimum perturbation method,¹²⁹ dihedral angle search through a rotamer library,^{114,130} molecular dynamics simulations,^{119,131} genetic algorithms,¹³² Monte Carlo and simulated annealing,^{133–135} multiple-copy simultaneous search,¹³⁶ self-consistent field optimization,¹³⁷ robotics-inspired kinematic closure¹³⁸ and enumeration based on graph theory.¹³⁹ The accuracy of loop predictions can be further improved by clustering the sampled loop conformations and partially accounting for the entropic contribution to the free energy.¹⁴⁰ Another way of improving the accuracy of loop predictions is to consider the solvent effects. Improvements in implicit solvation models, such as the Generalized Born solvation model, motivated their use in loop modeling. The solvent contribution to the free energy can be added to the scoring function for optimization, or it can be used to rank the sampled loop conformations after they are generated with a scoring function that does not include the solvent terms.^{105,141–143}

Side-Chain Modeling

Fixed backbone

Two simplifications are frequently applied in the modeling of side-chain conformations.¹⁴⁴ First, amino acid residue replacements often leave the backbone structure almost unchanged,¹⁴⁵ allowing us to fix the backbone during the search for the best side-chain conformations. Second, most side chains in high-resolution crystallographic structures can be represented by a limited number of conformers that comply with stereochemical and energetic constraints.¹⁴⁶ This observation motivated Ponder and Richards¹⁴⁷ to develop the first library of side-chain rotamers for the 17 types of residues with dihedral angle degrees of freedom in their side chains, based on 10 high-resolution protein structures determined by X-ray crystallography. Subsequently, a number of additional libraries have been derived.^{148–155}

Rotamers

Rotamers on a fixed backbone are often used when all the side chains need to be modeled on a given backbone. This approach reduces the combinatorial explosion associated with a full conformational search of all the side chains, and is applied by some comparative modeling⁸⁶ and protein design approaches.¹⁵⁶ However, ~15% of the side chains cannot be represented well by these libraries.¹⁵⁷ In addition, it has been shown that the accuracy of side-chain modeling on a fixed backbone decreases rapidly when the backbone errors are larger than 0.5 Å.¹⁵⁸

Methods

Earlier methods for side-chain modeling often put less emphasis on the energy or scoring function. The function was usually greatly simplified, and consisted of the empirical rotamer preferences and simple repulsion terms for nonbonded contacts.¹⁵¹ Nevertheless, these approaches have been justified by their performance. For example, a method based on a rotamer library compared favorably with that based on a molecular mechanics force field,¹⁵⁹ and new methods continue to be based on the rotamer library approach.^{160–162} The various optimization approaches include a Monte Carlo simulation,¹⁶³ simulated annealing,¹⁶⁴ a combination of Monte Carlo and simulated annealing,¹⁶⁵ the dead-end elimination theorem,^{166,167} genetic algorithms,¹⁵⁵ neural network with simulated annealing,¹⁶⁸ mean field optimization,¹⁶⁹ and combinatorial searches.^{151,170,171} Several studies focused on the testing of more sophisticated potential functions for conformational search^{171,172} and development of new scoring functions for side-chain modeling,¹⁷³ reporting higher accuracy than earlier studies.

Errors in Comparative Models

The major sources of error in comparative modeling are discussed in the relevant sections above. The following is a summary of these errors, dividing them into five categories (Figure 3).

Selection of Incorrect Templates

This error is a potential problem when distantly related proteins are used as templates (i.e., less than 30% sequence identity). Distinguishing between a model based on an incorrect template and a model based on an incorrect alignment with a correct template is difficult. In both cases, the evaluation methods (below) will predict an unreliable model. The conservation of the key functional or structural residues in the target sequence increases the confidence in a given fold assignment.

Errors due to Misalignments

The single source of errors with the largest impact on comparative modeling is misalignments, especially when the target–template sequence identity decreases below 30%. Alignment errors can be minimized in two ways. Using the profile-based methods

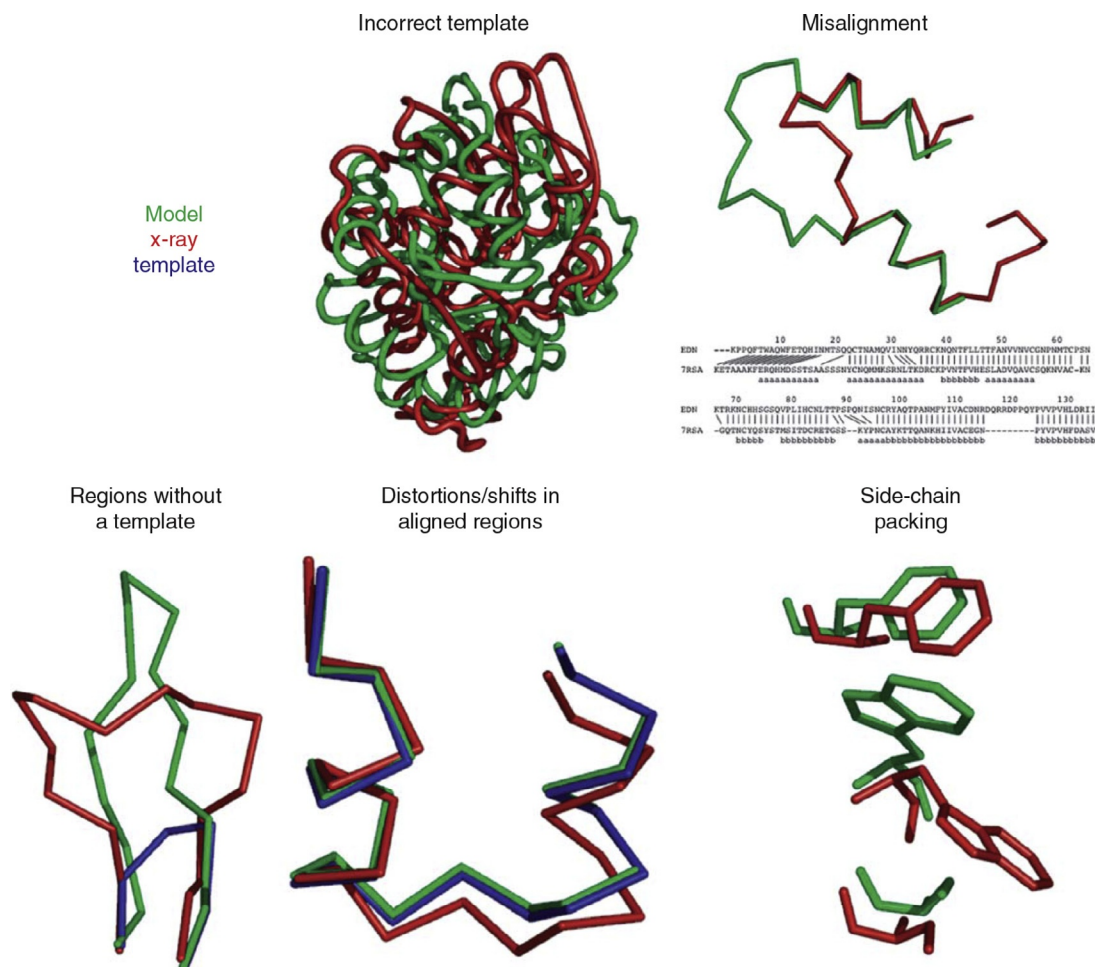


Figure 3 Typical errors in comparative modeling.²³ Shown are the typical sources of errors encountered in comparative models. Two of the major sources of errors in comparative modeling are due to incorrect templates or incorrect alignments with the correct templates. The modeling procedure can rarely recover from such errors. The next significant source of errors arises from regions in the target with no corresponding region in the template, i.e., insertions or loops. Other sources of errors, which occur even with an accurate alignment, are due to rigid-body shifts, distortions in the backbone, and errors in the packing of side chains.

discussed above usually results in more accurate alignments than those from pairwise sequence alignment methods. Another way of improving the alignment is to modify those regions in the alignment that correspond to predicted errors in the model.⁸³

Errors in Regions without a Template

Segments of the target sequence that have no equivalent region in the template structure (i.e., insertions or loops) are one of the most difficult regions to model. Again, when the target and template are distantly related, errors in the alignment can lead to incorrect positions of the insertions. Using alignment methods that incorporate structural information can often correct such errors. Once a reliable alignment is obtained, various modeling protocols can predict the loop conformation, for insertions of fewer than 8–10 residues.^{105,113,119,174}

Distortions and Shifts in Correctly Aligned Regions

As a consequence of sequence divergence, the main-chain conformation changes, even if the overall fold remains the same. Therefore, it is possible that in some correctly aligned segments of a model, the template is locally different ($<3 \text{ \AA}$) from the target, resulting in errors in that region. The structural differences are sometimes not due to differences in sequence, but are a consequence of artifacts in structure determination or structure determination in different environments (e.g., packing of subunits in a crystal). The simultaneous use of several templates can minimize this kind of error.^{83,84}

Errors in Side-Chain Packing

As the sequences diverge, the packing of the atoms in the protein core changes. Sometimes even the conformation of identical side chains is not conserved – a pitfall for many comparative modeling methods. Side-chain errors are critical if they occur in regions that are involved in protein function, such as active sites and ligand-binding sites.

Prediction of Model Errors

The accuracy of the predicted model determines the information that can be extracted from it. Thus, estimating the accuracy of a model in the absence of the known structure is essential for interpreting it.

Initial Assessment of the Fold

As discussed earlier, a model calculated using a template structure that shares more than 30% sequence identity is indicative of an overall accurate structure. However, when the sequence identity is lower, the first aspect of model evaluation is to confirm whether or not a correct template was used for modeling. It is often the case, when operating in this regime, that the fold assignment step produces only false positives. A further complication is that at such low similarities the alignment generally contains many errors, making it difficult to distinguish between an incorrect template on one hand and an incorrect alignment with a correct template on the other hand. There are several methods that use 3D profiles and statistical potentials,^{70,175,176} which assess the compatibility between the sequence and modeled structure by evaluating the environment of each residue in a model with respect to the expected environment, as found in native high-resolution experimental structures. These methods can be used to assess whether or not the correct template was used for the modeling. They include VERIFY3D,¹⁷⁵ Prosa2003,^{177,178} HARMONY,¹⁷⁹ ANOLEA,¹⁸⁰ DFIRE,¹⁸¹ DOPE,¹⁸² QMEAN local,¹⁸³ ProQ2,¹⁸⁴ and TSVMMod.¹⁸⁵

Even when the model is based on alignments that have >30% sequence identity, other factors, including the environment, can strongly influence the accuracy of a model. For instance, some calcium-binding proteins undergo large conformational changes when bound to calcium. If a calcium-free template is used to model the calcium-bound state of the target, it is likely that the model will be incorrect, irrespective of the target–template similarity or accuracy of the template structure.¹⁸⁶

Self-Consistency

The model should also be subjected to evaluations of self-consistency to ensure that it satisfies the restraints used to calculate it. Additionally, the stereochemistry of the model (e.g., bond lengths, bond angles, backbone torsion angles, and nonbonded contacts) may be evaluated using programs such as PROCHECK¹⁸⁷ and WHATCHECK.¹⁸⁸ Although errors in stereochemistry are rare and less informative than errors detected by statistical potentials, a cluster of stereochemical errors may indicate that there are larger errors (e.g., alignment errors) in that region.

Final Assessment of the Fold (Model)

When multiple models are calculated for the target based on a single template or when multiple loops are built for a single or multiple models, it is practical to select a subset of models or loops that are judged to be most suitable for subsequent docking calculations. If some known ligands or other information for the desired model is available, model selection should be guided by this known information.¹⁸⁹ If this extra information is not available, model selection should aim to select the most accurate model. While models or loops can be selected by the energy function used for guiding the building of comparative models or the sampling of loop configurations, using a separate statistical potential for selecting the most accurate models or loops is often more successful.^{181,182,190,191}

Evaluation of Comparative Modeling Methods

It is crucial for method developers and users alike to assess the accuracy of their methods. An attempt to address this problem has been made by the Critical Assessment of Techniques for Proteins Structure Prediction (CASP)¹⁹² and in the past by the Critical Assessment of Fully Automated Structure Prediction (CAFASP) experiments,¹⁹³ which is now integrated into CASP. However, CASP assesses methods only over a limited number of target protein sequences, and is conducted only every 2 years.^{109,194} To overcome this limitation, the new CAMEO web server continuously evaluates the accuracy and reliability of a number of comparative protein structure prediction servers, in a fully automated manner.¹⁹⁵ Every week, CAMEO provides each tested server with the prerelease sequences of structures that are to be shortly released by the PDB. Each server then has 4 days to build and return a 3D model of these sequences. When PDB releases the structures, CAMEO compares the models against the experimentally determined structures, and presents the results on its web site. This enables developers, non-expert users, and reviewers to determine the performance of

the tested prediction servers. CAMEO is similar in concept to two prior such continuous testing servers, LiveBench¹⁹⁴ and EVA.^{196,197}

Applications of Comparative Models

There is a wide range of applications of protein structure models (Figure 4).^{1,198–204} For example, high- and medium-accuracy comparative models are frequently helpful in refining functional predictions that have been based on a sequence match alone because ligand binding is more directly determined by the structure of the binding site than by its sequence. It is often possible to predict correctly features of the target protein that do not occur in the template structure.^{205,206} For example, the size of a ligand may be predicted from the volume of the binding site cleft and the location of a binding site for a charged ligand can be predicted from a cluster of charged residues on the protein. Fortunately, errors in the functionally important regions in comparative models are many times relatively low because the functional regions, such as active sites, tend to be more conserved in evolution than the rest of the fold. Even low-accuracy comparative models may be useful, for example, for assigning the fold of a protein. Fold

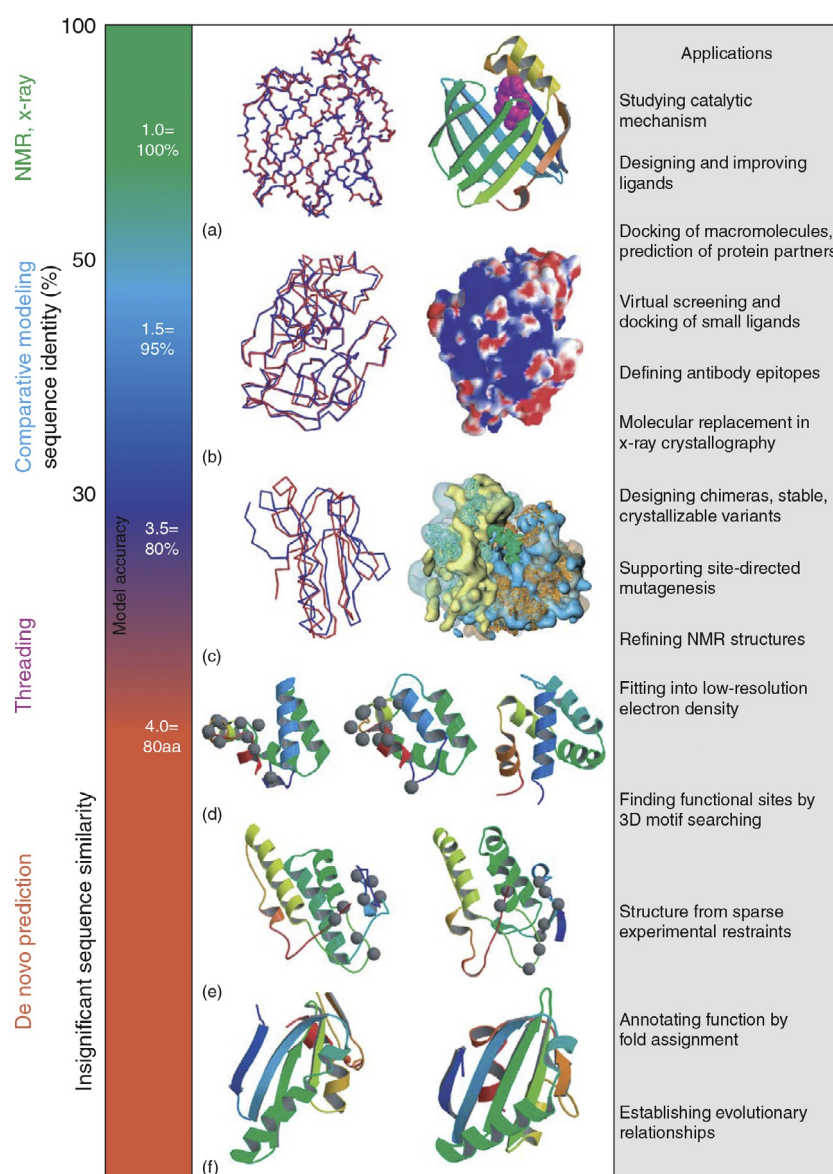


Figure 4 Accuracy and applications of protein structure models.⁹ Shown are the different ranges of applicability of comparative protein structure modeling, threading, and *de novo* structure prediction, their corresponding accuracies, and their sample applications.

assignment can be very helpful in drug discovery, because it can shortcut the search for leads by pointing to compounds that have been previously developed for other members of the same family.^{207,208}

Comparative Models versus Experimental Structures in Virtual Screening

The remainder of this review focuses on the use of comparative models for ligand docking.^{209–211} Comparative protein structure modeling extends the applicability of virtual screening beyond the atomic structures determined by X-ray crystallography or NMR spectroscopy. In fact, comparative models have been used in virtual screening to detect novel ligands for many protein targets,²⁰¹ including the G-protein coupled receptors (GPCR),^{210,212–223} protein kinases,^{224–227} nuclear hormone receptors, and many different enzymes.^{228–241} Nevertheless, the relative utility of comparative models versus experimentally determined structures has only been relatively sparsely assessed.^{212,224,225,242–244} The utility of comparative models for molecular docking screens in ligand discovery has been documented²⁴⁵ with the aid of 38 protein targets selected from the 'directory of useful decoys' (DUD).²⁴⁶ For each target sequence, templates for comparative modeling were obtained from the PDB, including at least one holo (ligand bound) and one apo (ligand free) template structure for each of the eight 10% sequence identity ranges from 20% to 100%. In total, 222 models were generated based on 222 templates for the 38 test proteins using Modeller 9v2.³⁵ DUD ligands and decoys (98 266 molecules) were screened against the holo X-ray structure, the apo X-ray structure, and the comparative models of each target using DOCK 3.5.54.²⁴⁷ The accuracy of virtual screening was evaluated by the overall ligand enrichment that was calculated by integrating the area under the enrichment plot (logAUC). A key result was that, for 63% and 79% of the targets, at least one comparative model yielded ligand enrichment better or comparable to that of the corresponding holo and apo X-ray structure.²⁴⁵ This result indicates that comparative models can be useful docking targets when multiple templates are available. However, it was not possible to predict which model, out of all those used, led to the highest enrichment. Therefore, a 'consensus' enrichment score was computed by ranking each library compound by its best docking score against all comparative models and/or templates. For 47% and 70% of the targets, the consensus enrichment for multiple models was better or comparable to that of the holo and apo X-ray structures, respectively, suggesting that multiple comparative models can be useful for virtual screening.

Use of Comparative Models to Obtain Novel Drug Leads

Despite problems with comparative modeling and ligand docking, comparative models have been successfully used in practice in conjunction with virtual screening to identify novel inhibitors. We briefly review a few of these success stories to highlight the potential of the combined comparative modeling and ligand-docking approach to drug discovery.

Comparative models have been employed to aid rational drug design against parasites for more than 20 years.^{132,231,232,240} As early as 1993, Ring et al.¹³² used comparative models for computational docking studies that identified low micromolar nonpeptidic inhibitors of proteases in malarial and schistosome parasite lifecycles. Li et al.²³¹ subsequently used similar methods to develop nanomolar inhibitors of falcipain that are active against chloroquine-resistant strains of malaria. In a study by Selzer et al.²³² comparative models were used to predict new nonpeptide inhibitors of cathepsin L-like cysteine proteases in *Leishmania major*. Sixty-nine compounds were selected by DOCK 3.5 as strong binders to a comparative model of protein cpB, and of these, 21 had experimental IC₅₀ values below 100 mmol l⁻¹. Finally, in a study by Que et al.,²⁴⁰ comparative models were used to rationalize ligand-binding affinities of cysteine proteases in *Entamoeba histolytica*. Specifically, this work provided an explanation for why proteases ACP1 and ACP2 had substrate specificity similar to that of cathepsin B, although their overall structure is more similar to that of cathepsin D.

Enyedy et al.²⁴⁸ discovered 15 new inhibitors of matriptase by docking against its comparative model. The comparative model employed thrombin as the template, sharing only 34% sequence identity with the target sequence. Moreover, some residues in the binding site are significantly different; a trio of charged Asp residues in matriptase correspond to 1 Tyr and 2 Trp residues in thrombin. Thrombin was chosen as the template, in part because it prefers substrates with positively charged residues at the P1 position, as does matriptase. The National Cancer Institute database was used for virtual screening that targeted the S1 site with the DOCK program. The 2000 best-scoring compounds were manually inspected to identify positively charged ligands (the S1 site is negatively charged), and 69 compounds were experimentally screened for inhibition, identifying the 15 inhibitors. One of them, hexamidine, was used as a lead to identify additional compounds selective for matriptase relative to thrombin. The Wang group has also used similar methods to discover seven new, low-micromolar inhibitors of Bcl-2, using a comparative model based on the NMR solution structure of Bcl-X_L.²³³

Schapira et al.²⁴⁹ discovered a novel inhibitor of a retinoic acid receptor by virtual screening using a comparative model. In this case, the target (RAR- α) and template (RAR- γ) are very closely related; only three residues in the binding site are not conserved. The ICM program was used for virtual screening of ligands from the Available Chemicals Directory (ACD). The 5364 high-scoring compounds identified in the first round were subsequently docked into a full atom representation of the receptor with flexible side chains to obtain a final set of 300 good-scoring hits. These compounds were then manually inspected to choose the final 30 for testing. Two novel agonists were identified, with 50-nanomolar activity.

Zuccotto et al.²⁵⁰ identified novel inhibitors of dihydrofolate reductase (DHFR) in *Trypanosoma cruzi* (the parasite that causes Chagas disease) by docking into a comparative model based on ~50% sequence identity to DHFR in *L. major*, a related parasite. The virtual screening procedure used DOCK for rigid docking of over 50 000 selected compounds from the Cambridge Structural

Database (CSD). Visual inspection of the top 100 hits was used to select 36 compounds for experimental testing. This work identified several novel scaffolds with micromolar IC_{50} values. The authors report attempting to use virtual screening results to identify compounds with greater affinity for *T. cruzi* DHFR than human DHFR, but it is not clear how successful they were.

Following the outbreak of the severe acute respiratory syndrome (SARS) in 2003, Anand et al.²⁵¹ used the experimentally determined structures of the main protease from human coronavirus (M^{PRO}) and an inhibitor complex of porcine coronavirus (transmissible gastroenteritis virus, TGEV) M^{PRO} to calculate a comparative model of the SARS coronavirus M^{PRO} . This model then provided a basis for the design of anti-SARS drugs. In particular, a comparison of the active site residues in these and other related structures suggested that the AG7088 inhibitor of the human rhinovirus type 2 3C protease is a good starting point for design of anticoronaviral drugs.²⁵²

Comparative models of protein kinases combined with virtual screening have also been intensely used for drug discovery.^{224,225,253–255} The >500 kinases in the human genome, the relatively small number of experimental structures available, and the high level of conservation around the important adenosine triphosphate-binding site make comparative modeling an attractive approach toward structure-based drug discovery.

G protein-coupled receptors are another interesting class of proteins that in principle allow drug discovery through comparative modeling.^{212,256–259} Approximately 40% of current drug targets belong to this class of proteins. However, these proteins have been extremely difficult to crystallize and most comparative modeling has been based on the atomic resolution structure of the bovine rhodopsin.²⁶⁰ Despite this limitation, a rather extensive test of docking methods with rhodopsin-based comparative models shows encouraging results.

The applicability of structure-based modeling and virtual screening has recently been expanded to membrane proteins that transport solutes, such as ions, metabolites, peptides, and drugs. In humans, these transporters contribute to the absorption, distribution, metabolism, and excretion of drugs, and often, mediate drug-drug interactions. Additionally, several transporters can be targeted directly by small molecules. For instance, methylphenidate (Ritalin) inhibiting the norepinephrine transporter (NET) and, consequently, inhibiting the reuptake of norepinephrine, is used in the treatment of attention-deficit hyperactivity disorder (ADHD).²⁶¹ Schlessinger et al.²⁶² predicted 18 putative ligands of human NET by docking 6436 drugs from the Kyoto Encyclopedia of Genes (KEGG DRUG) into a comparative model based on ~25% sequence identity to leucine transporter (LeuT) from *Aquifex aeolicus*. Of these 18 predicted ligands, ten were validated by *cis*-inhibition experiments; five of them were chemically novel. Close examination of the predicted primary binding site helped rationalize interactions of NET with its primary substrate, norepinephrine, as well as positive and negative pharmacological effects of other NET ligands. Subsequently, Schlessinger et al.²⁶³ modeled two different conformations of the human GABA transporter 2 (GAT-2), using the LeuT structures in occluded-outward-facing and outward-facing conformations. Enrichment calculations were used to assess the quality of the models in molecular dynamics simulations and side-chain refinements. The key residue, Glu48, interacting with the substrate was identified during the refinement of the models and validated by site-directed mutagenesis. Docking against two conformations of the transporter enriches for different physicochemical properties of ligands. For example, top-scoring ligands found by docking against the outward-facing model were bulkier and more hydrophobic than those predicted using the occluded-outward-facing model. Among twelve ligands validated in *cis*-inhibition assays, six were chemically novel (e.g., homotaurine). Based on the validation experiments, GAT-2 is likely to be a high-selectivity/low-affinity transporter. Following these two studies, a combination of comparative modeling, ligand docking, and experimental validation was used to rationalize toxicity of an anti-cancer agent, acivicin.²⁶⁴ The toxic side-effects are thought to be facilitated by the active transport of acivicin through the blood–brain-barrier (BBB) via the large-neutral amino acid Transporter 1 (LAT-1). In addition, four small-molecule ligands of LAT-1 were identified by docking against a comparative model based on two templates, the structures of the outward-occluded arginine-bound arginine/arginine transporter AdiC from *E. coli*²⁶⁵ and the inward-apo conformation of the amino acid, polyamine, and organo-cation transporter ApcT from *Methanococcus jannaschii*.²⁶⁶ Two of the four hits, acivicin and fenclonine, were confirmed as substrates by a *trans*-stimulation assay. These studies clearly illustrate the applicability of combined comparative modeling and virtual screening to ligand discovery for transporters.

Future Directions

Although reports of successful virtual screening against comparative models are encouraging, such efforts are not yet a routine part of rational drug design. Even the successful efforts appear to rely strongly on visual inspection of the docking results. Much work remains to be done to improve the accuracy, efficiency, and robustness of docking against comparative models. Despite assessments of relative successes of docking against comparative models and native X-ray structures,^{225,244} relatively little has been done to compare the accuracy achievable by different approaches to comparative modeling and to identify the specific structural reasons why comparative models generally produce less accurate virtual screening results than the holo structures. Among the many issues that deserve consideration are the following:

- The inclusion of cofactors and bound water molecules in protein receptors is often critical for success of virtual screening; however, cofactors are not routinely included in comparative models

- Most docking programs currently retain the protein receptor in a rigid conformation. While this approach is appropriate for 'lock-and-key' binding modes, it does not work when the ligand induces conformational changes in the receptor upon binding. A flexible receptor approach is necessary to address such induced-fit cases^{267,268}
- The accuracy of comparative models is frequently judged by the C α root mean square error or other similar measures of backbone accuracy. For virtual screening, however, the precise positioning of side chains in the binding site is likely to be critical; measures of accuracy for binding sites are needed to help evaluate the suitability of comparative modeling algorithms for constructing models for docking
- Knowledge of known inhibitors, either for the target protein or the template, should help to evaluate and improve virtual screening against comparative models. For example, comparative models constructed from holo template structures implicitly preserve some information about the ligand-bound receptor conformation
- Improvement in the accuracy of models produced by comparative modeling will require methods that finely sample protein conformational space using a free energy or scoring function that has sufficient accuracy to distinguish the native structure from the nonnative conformations. Despite many years of development of molecular simulation methods, attempts to refine models that are already relatively close to the native structure have met with relatively little success. This failure is likely to be due in part to inaccuracies in the scoring functions used in the simulations, particularly in the treatment of electrostatics and solvation effects. A combination of physics-based energy function with the statistical information extracted from known protein structures may provide a route to the development of improved scoring functions
- Improvements in sampling strategies are also likely to be necessary, for both comparative modeling and flexible docking

Automation and Availability of Resources for Comparative Modeling and Ligand Docking

Given the increasing number of target sequences for which no experimentally determined structures are available, drug discovery stands to gain immensely from comparative modeling and other *in silico* methods. Despite unsolved problems in virtually every step of comparative modeling and ligand docking, it is highly desirable to automate the whole process, starting with the target sequence and ending with a ranked list of its putative ligands. Automation encourages development of better methods, improves their testing, allows application on a large scale, and makes the technology more accessible to both experts and non-specialists alike. Through large-scale application, new questions, such as those about ligand-binding specificity, can in principle be addressed. Enabling a wider community to use the methods provides useful feedback and resources toward the development of the next generation of methods.

There are a number of servers for automated comparative modeling (Table 1). However, in spite of automation, the process of calculating a model for a given sequence, refining its structure, as well as visualizing and analyzing its family members in the sequence and structure spaces can involve the use of scripts, local programs, and servers scattered across the internet and not necessarily interconnected. In addition, manual intervention is generally still needed to maximize the accuracy of the models in the difficult cases. The main repository for precomputed comparative models, the Protein Model Portal,^{195,198,279} begins to address these deficiencies by serving models from several modeling groups, including the SWISS-MODEL⁹⁵ and ModBase³⁴ databases. It provides access to web-based comparative modeling tools, cross-links to other sequence and structure databases, and annotations of sequences and their models.

A number of databases containing comparative models and web servers for computing comparative models are publicly available. The Protein Model Portal (PMP)^{195,198,279} centralizes access to these models created by different methodologies. The PMP is being developed as a module of the Protein Structure Initiative Knowledgebase (PSI KB)³¹⁶ and functions as a meta server for comparative models from external databases, including SWISS-MODEL⁹⁵ and ModBase,³⁴ additionally to being a repository for comparative models that are derived from structures determined by the PSI centers. It provides quality estimations of the deposited models, access to web-based comparative modeling tools, cross-links to other sequence and structure databases, annotations of sequences and their models, and detailed tutorials on comparative modeling and the use of their tools. The PMP currently contains 19.5 million comparative models for 4.4 million UniProt sequences (August 2013).

A schematic of our own attempt at integrating several useful tools for comparative modeling is shown in Figure 5.^{34,291} ModBase is a database that currently contains ~29 million predicted models for domains in approximately 4.7 million unique sequences from UniProt, Ensembl,²⁶⁹ GenBank,¹⁴ and private sequence datasets. The models were calculated using ModPipe^{30,291} and Modeller.³⁵ The web interface to the database allows flexible querying for fold assignments, sequence–structure alignments, models, and model assessments. An integrated sequence–structure viewer, Chimera,³⁰⁴ allows inspection and analysis of the query results. Models can also be calculated using ModWeb,^{291,309} a web interface to ModPipe, and stored in ModBase, which also makes them accessible through the PMP. Other resources associated with ModBase include a comprehensive database of multiple protein structure alignments (DBALI),²⁸¹ structurally defined ligand-binding sites,³¹⁹ structurally defined binary domain interfaces (PIBASE),^{320,321} predictions of ligand-binding sites, interactions between yeast proteins, and functional consequences of human nsSNPs (LS-SNP).^{199,322,323} A number of associated web services handle modeling of loops in protein structures (ModLoop),^{324,325} evaluation of models (ModEval), fitting of models against Small Angle X-ray Scattering (SAXS) profiles (FoXS),^{326–328} modeling of ligand-induced protein dynamics such as allostery (AllosMod),^{329,330} prediction of the ensemble of conformations that best fit a given SAXS profile (AllosMod-FoXS),³³¹ prediction of cryptic binding sites,³³² scoring protein-ligand complexes based on a

Table 1 Programs, databases and web servers useful in comparative protein structure modeling

<i>Name</i>	<i>World Wide Web address</i>
<i>Databases</i>	
Protein Sequence Databases	
Ensembl ²⁶⁹	http://www.ensembl.org
GENBANK ¹⁴	http://www.ncbi.nlm.nih.gov/Genbank/
Protein Information Resource ²⁷⁰	http://pir.georgetown.edu/
UniprotKB ¹³	http://www.uniprot.org
Domains and Superfamilies	
CATH/Gene3D ²⁰	http://www.cathdb.info
InterPro ²⁷¹	http://www.ebi.ac.uk/interpro/
MEME ²⁷²	http://meme.nbcr.net/meme/
PFAM ¹⁷	http://pfam.sanger.ac.uk/
PRINTS ²⁷³	http://www.bioinf.manchester.ac.uk/dbbrowser/PRINTS/index.php
ProDom ²⁷⁴	http://prodom.prabi.fr
ProSite ²⁷⁵	http://prosite.expasy.org/
SCOP ¹⁹	http://scop.mrc-lmb.cam.ac.uk/scop/
SFLD ²⁷⁶	http://sfld.rbvi.ucsf.edu/
SMART ²⁷⁷	http://smart.embl-heidelberg.de/
SUPERFAMILY ²⁷⁸	http://supfam.cs.bris.ac.uk/SUPERFAMILY/
Protein Structures and Models	
ModBase ³⁴	http://www.salilab.org/modbase/
PDB ¹²	http://www.pdb.org/
Protein Model Portal ^{195,279}	http://www.proteinmodelportal.org/
SwissModel Repository ²⁸⁰	http://swissmodel.expasy.org/repository/
Miscellaneous	
DBALI ²⁸¹	http://www.salilab.org/dbali
GENECENSUS ²⁸²	http://bioinfo.mbb.yale.edu/genome/
<i>Alignment</i>	
Sequence and structure based sequence alignment	
AlignMe ²⁸³	http://www.bioinfo.mpg.de/AlignMe/
CLUSTALW ²⁸⁴	http://www2.ebi.ac.uk/clustalw/
COMPASS ⁶²	ftp://iole.swmed.edu/pub/compass/
EXPRESSO ²⁸⁵	http://igs-server.cnrs-mrs.fr/Tcoffee/tcoffee_cgi/index.cgi
FastA ²⁸⁶	http://www.ebi.ac.uk/Tools/sss/fasta/
FFAS03 ⁶⁵	http://ffas.burnham.org/
FUGUE ⁶⁸	http://www-cryst.bioc.cam.ac.uk/fugue
GENTHREADER ^{66,75}	http://bioinf.cs.ucl.ac.uk/psipred/
HHBlits/HHsearch ⁵³	http://toolkit.lmb.uni-muenchen.de/hhsuite
MAFFT ²⁸⁷	http://mafft.cbrc.jp/alignment/software/
MUSCLE ²⁸⁸	http://www.drive5.com/muscle
MUSTER ⁷⁸	http://zhanglab.ccmb.med.umich.edu/MUSTER
PROMALS3D ²⁸⁹	http://prodata.swmed.edu/promals3d/promals3d.php
PSI-BLAST ³⁹	http://blast.ncbi.nlm.nih.gov/Blast.cgi
PSIPRED ²⁹⁰	http://bioinf.cs.ucl.ac.uk/psipred/
SALIGN ²⁹¹	http://www.salilab.org/salign/
SAM-T08 ^{74,77}	http://compbio.soe.ucsc.edu/HMM-apps/
Staccato ²⁹²	http://bioinfo3d.cs.tau.ac.il/staccato/
T-Coffee ^{293,294}	http://www.tcoffee.org/
Structure	
CE ²⁹⁵	http://source.rcsb.org/jfatcatserver/ceHome.jsp
GANGSTA+ ²⁹⁶	http://agknapp.chemie.fu-berlin.de/gplus/index.php
HHsearch ⁵²	ftp://toolkit.lmb.uni-muenchen.de/hhsearch/
Mammoth ²⁹⁷	http://ub.cbm.uam.es/software/mammoth.php
Mammoth-mult ²⁹⁸	http://ub.cbm.uam.es/software/mammothm.php
MASS ²⁹⁹	http://bioinfo3d.cs.tau.ac.il/MASS/
MultiProt ³⁰⁰	http://bioinfo3d.cs.tau.ac.il/MultiProt
MUSTANG ³⁰¹	http://www.csse.monash.edu.au/~karun/Site/mustang.html
PDBeFold ³⁶	http://www.ebi.ac.uk/msd-srv/ssm/
SALIGN ²⁹¹	http://www.salilab.org/salign/
TM-align ³⁰²	http://zhanglab.ccmb.med.umich.edu/TM-align/

(Continued)

Table 1 (Continued)

Name	World Wide Web address
Alignment modules in molecular graphics programs	
Discovery Studio	http://www.accelrys.com
PyMol	http://www.pymol.org/
Swiss-PDB Viewer ³⁰³	http://spdbv.vital-it.ch/
UCSF Chimera ³⁰⁴	http://www.cgl.ucsf.edu/chimera
<i>Comparative Modeling, Threading and Refinement</i>	
Web servers	
3d-jigsaw ⁹⁴	http://www.bmm.icnet.uk/servers/3djigsaw/
HHPred ³⁰⁵	http://toolkit.genzentrum.lmu.de/hhpred
IntFold ³⁰⁶	http://www.reading.ac.uk/bioinf/IntFOLD/
i-TASSER ³⁰⁷	http://zhanglab.ccmb.med.umich.edu/I-TASSER/
M4T ³⁰⁸	http://manaslu.aecom.yu.edu/M4T/
ModWeb ^{291,309}	http://salilab.org/modweb/
Phyre2 ³¹⁰	http://www.sbg.bio.ic.ac.uk/phyre2
RaptorX ³¹¹	http://raptorx.uchicago.edu/
Robetta ⁸¹	http://rosetta.bakerlab.org/
SWISS-MODEL ⁹⁵	http://www.expasy.org/swissmod
Programs	
HHsuite ⁵²	ftp://toolkit.genzentrum.lmu.de/pub/HH-suite/
Modeller ³⁵	http://www.salilab.org/modeller/
MolIDE ³¹²	http://dunbrack.fccc.edu/molide/
Rosetta@home	http://boinc.bakerlab.org/rosetta/
RosettaCM ⁸¹	https://www.rosettacommons.org/home
SCWRL ¹⁶⁰	http://dunbrack.fccc.edu/scwrl4/SCWRL4.php
Quality estimation	
ANOLEA ¹⁸⁰	http://melolab.org/anolea/index.html
ERRAT ³¹³	http://nihserver.mbi.ucla.edu/ERRAT/
ModEval	http://salilab.org/modeval/
ProQ2 ¹⁸⁴	http://proq2.theophys.kth.se/
PROCHECK ¹⁸⁷	http://www.ebi.ac.uk/thornton-srv/software/PROCHECK/
Prosa2003 ^{177,178}	http://www.came.sbg.ac.at
QMEAN local ¹⁸³	http://www.openstructure.org/download/
SwissModel Workspace ³¹⁴	http://swissmodel.expasy.org/workspace/index.php?func=tools_structureassessment1
VERIFY3D ¹⁷⁵	http://www.doe-mbi.ucla.edu/Services/Verify_3D/
WHATCHECK ¹⁸⁸	http://www.cmbi.kun.nl/gv/whatcheck/
Methods evaluation	
CAMEO ¹⁹⁵	http://cameo3d.org/
CASP ³¹⁵	http://predictioncenter.llnl.gov

statistical potential (Pose & Rank),¹⁹⁰ protein-protein docking filtered by a SAXS profile (FoXSDock),^{333,334} and merging SAXS profiles (SAXS Merge).^{335,336}

Compared to protein structure prediction, the attempts at automation and integration of resources in the field of docking for virtual screening are still in their nascent stages. One of the successful efforts in this direction is ZINC,^{317,318} a publicly available database of commercially available drug-like compounds, developed in the laboratory of Brian Shoichet. ZINC contains more than 21 million 'ready-to-dock' compounds organized in several subsets and allows the user to query the compounds by molecular properties and constitution. The Shoichet group also provides a DOCKBLASTER service³³⁷ that enables end-users to dock the ZINC compounds against their target structures using DOCK.^{247,338}

In the future, we will no doubt see efforts to improve the accuracy of comparative modeling and ligand docking. But perhaps as importantly, the two techniques will be integrated into a single protocol for more accurate and automated docking of ligands against sequences without known structures. As a result, the number and variety of applications of both comparative modeling and ligand docking will continue to increase.

Acknowledgments

This article is partially based on papers by Jacobson and Sali,²⁰¹ Fiser and Sali,³³⁹ and Madhusudhan et al.³⁴⁰ We also acknowledge the funds from Sandler Family Supporting Foundation, NIH R01 GM54762, P01 GM71790, P01 A135707, and U54 GM62529, as well as Sun, IBM, and Intel for hardware gifts.

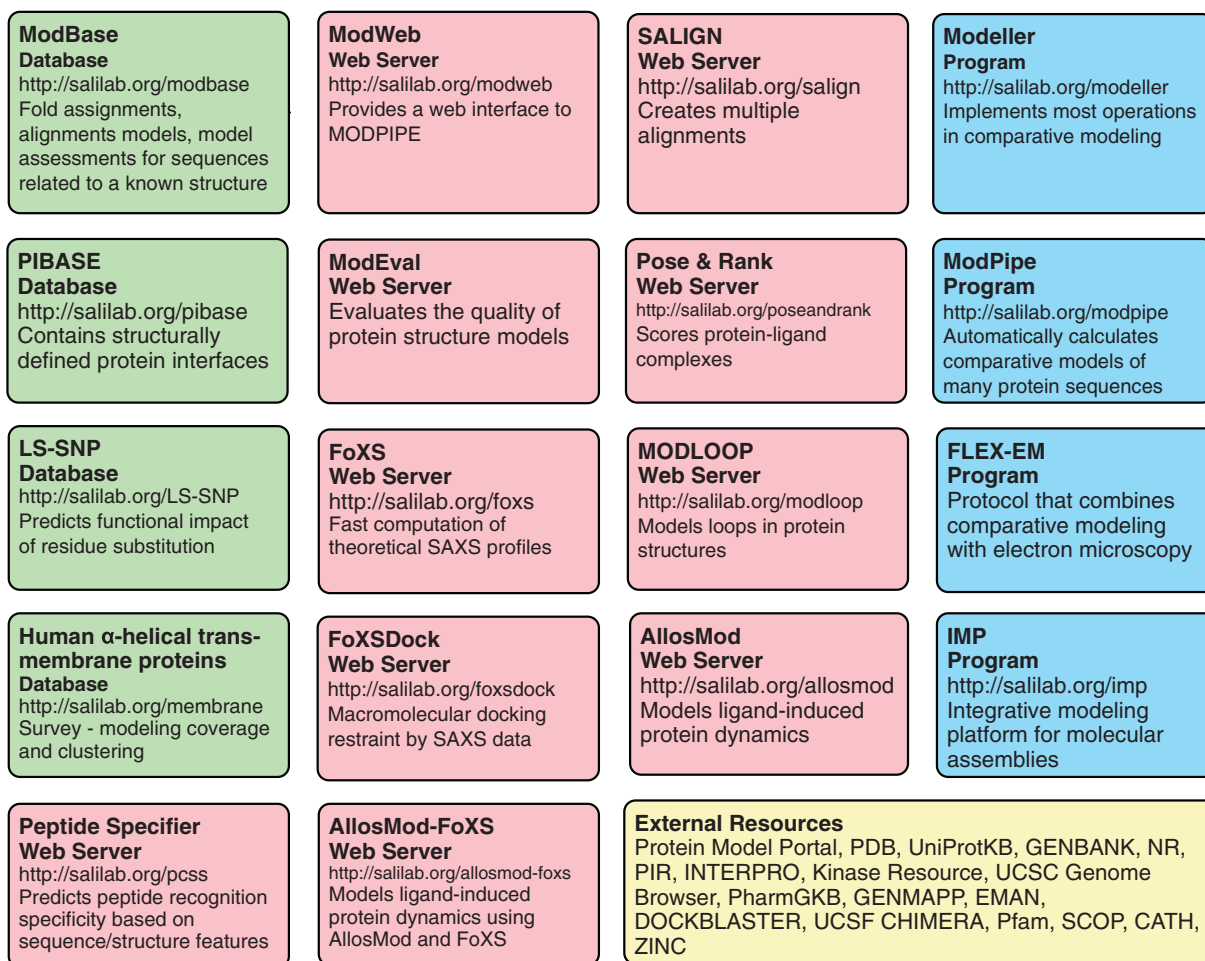


Figure 5 An integrated set of resources for comparative modeling.³⁴ Various databases and programs required for comparative modeling and docking are usually scattered over the internet, and require manual intervention or a good deal of expertise to be useful. Automation and integration of these resources are efficient ways to put these resources in the hands of experts and non-specialists alike. We have outlined a comprehensive interconnected set of resources for comparative modeling and hope to integrate it with a similar effort in the area of ligand docking made by the Shoichet group.^{317,318}

References

- Congreve, M.; Murray, C. W.; Blundell, T. L. *Drug Discov. Today* **2005**, *10*, 895–907.
- Hardy, L. W.; Malikayil, A. *Curr. Drug Disc.* **2003**, *3*, 15–20.
- Lombardino, J. G.; Lowe, J. A., 3rd. *Nat. Rev. Drug Discov.* **2004**, *3*, 853–862.
- van Dongen, M.; Weigelt, J.; Uppenberg, J.; Schultz, J.; Wikstrom, M. *Drug Discov. Today* **2002**, *7*, 471–478.
- Maryanoff, B. E. *J. Med. Chem.* **2004**, *47*, 769–787.
- Pollack, V. A.; Savage, D. M.; Baker, D. A.; Tsaparikos, K. E.; Sloan, D. E.; Moyer, J. D.; Barbacci, E. G.; Pustilnik, L. R.; Smolarek, T. A.; Davis, J. A.; Vaidya, M. P.; Arnold, L. D.; Doty, J. L.; Iwata, K. K.; Morin, M. J. *J. Pharmacol. Exp. Ther.* **1999**, *291*, 739–748.
- von Itzstein, M.; Wu, W. Y.; Kok, G. B.; Pegg, M. S.; Dyason, J. C.; Jin, B.; Van Phan, T.; Smythe, M. L.; White, H. F.; Oliver, S. W.; et al. *Nature* **1993**, *363*, 418–423.
- Zimmermann, J.; Caravatti, G.; Mett, H.; Meyer, T.; Muller, M.; Lydon, N. B.; Fabbro, D. *Arch. Pharm.* **1996**, *329*, 371–376.
- Baker, D.; Sali, A. *Science* **2001**, *294*, 93–96.
- Arzt, S.; Beteva, A.; Cipriani, F.; Delageniere, S.; Felisaz, F.; Forstner, G.; Gordon, E.; Launer, L.; Lavault, B.; Leonard, G.; Mairs, T.; McCarthy, A.; McCarthy, J.; McSweeney, S.; Meyer, J.; Mitchell, E.; Monaco, S.; Nurizzo, D.; Ravelli, R.; Rey, V.; Shepard, W.; Spruce, D.; Svensson, O.; Thevenneau, P. *Prog. Biophys. Mol. Biol.* **2005**, *89*, 124–152.
- Pusey, M. L.; Liu, Z. J.; Tempel, W.; Praissman, J.; Lin, D.; Wang, B. C.; Gavira, J. A.; Ng, J. D. *Prog. Biophys. Mol. Biol.* **2005**, *88*, 359–386.
- Berman, H. M.; Westbrook, J.; Feng, Z.; Gilliland, G.; Bhat, T. N.; Weissig, H.; Shindyalov, I. N.; Bourne, P. E. *Nucleic Acids Res.* **2000**, *28*, 235–242.
- Bairoch, A.; Apweiler, R.; Wu, C. H.; Barker, W. C.; Boeckmann, B.; Ferro, S.; Gasteiger, E.; Huang, H.; Lopez, R.; Magrane, M.; Martin, M. J.; Natale, D. A.; O'Donovan, C.; Redaschi, N.; Yeh, L. S. *Nucleic Acids Res.* **2005**, *33*, D154–D159.
- Benson, D. A.; Cavanaugh, M.; Clark, K.; Karsch-Mizrachi, I.; Lipman, D. J.; Ostell, J.; Sayers, E. W. *Nucleic Acids Res.* **2013**, *41*, D36–D42.
- Chandonia, J. M.; Brenner, S. E. *Proteins* **2005**, *58*, 166–179.
- Vitkup, D.; Melamud, E.; Moulton, J.; Sander, C. *Nat. Struct. Biol.* **2001**, *8*, 559–566.

17. Bateman, A.; Coin, L.; Durbin, R.; Finn, R. D.; Hollich, V.; Griffiths-Jones, S.; Khanna, A.; Marshall, M.; Moxon, S.; Sonnhammer, E. L.; Studholme, D. J.; Yeats, C.; Eddy, S. R. *Nucleic Acids Res.* **2004**, *32*, D138–D141.
18. Mulder, N. J.; Apweiler, R.; Attwood, T. K.; Bairoch, A.; Bateman, A.; Binns, D.; Bradley, P.; Bork, P.; Bucher, P.; Cerutti, L.; Copley, R.; Courcelle, E.; Das, U.; Durbin, R.; Fleischmann, W.; Gough, J.; Haft, D.; Harte, N.; Hulo, N.; Kahn, D.; Kanapin, A.; Krestyaninova, M.; Lonsdale, D.; Lopez, R.; Letunic, I.; Madera, M.; Maslen, J.; McDowall, J.; Mitchell, A.; Nikolskaya, A. N.; Orchard, S.; Pagni, M.; Ponting, C. P.; Quevillon, E.; Selengut, J.; Sigrist, C. J.; Silventoinen, V.; Studholme, D. J.; Vaughan, R.; Wu, C. H. *Nucleic Acids Res.* **2005**, *33*, D201–D205.
19. Andreeva, A.; Howorth, D.; Brenner, S. E.; Hubbard, T. J.; Chothia, C.; Murzin, A. G. *Nucleic Acids Res.* **2004**, *32*, D226–D229.
20. Pearl, F.; Todd, A.; Sillitoe, I.; Dibley, M.; Redfern, O.; Lewis, T.; Bennett, C.; Marsden, R.; Grant, A.; Lee, D.; Akpor, A.; Maibaum, M.; Harrison, A.; Dallman, T.; Reeves, G.; Diboun, I.; Addou, S.; Lise, S.; Johnston, C.; Sillero, A.; Thornton, J.; Orengo, C. *Nucleic Acids Res.* **2005**, *33*, D247–D251.
21. Godzik, A. *Methods Biochem. Anal.* **2003**, *44*, 525–546.
22. Fiser, A.; Sali, A. Comparative Protein Structure Modeling. In *Protein Structure*; Chasman, D., Ed.; Marcel Dekker, Inc.: New York, 2003; pp 167–206.
23. Marti-Renom, M. A.; Stuart, A. C.; Fiser, A.; Sanchez, R.; Melo, F.; Sali, A. *Annu. Rev. Biophys. Biomol. Struct.* **2000**, *29*, 291–325.
24. Hillisch, A.; Pineda, L. F.; Hilgenfeld, R. *Drug Discov. Today* **2004**, *9*, 659–669.
25. Jorgensen, W. L. *Science* **2004**, *303*, 1813–1818.
26. Das, R.; Qian, B.; Raman, S.; Vernon, R.; Thompson, J.; Bradley, P.; Khare, S.; Tyka, M. D.; Bhat, D.; Chivian, D.; Kim, D. E.; Sheffler, W. H.; Malmstrom, L.; Wollacott, A. M.; Wang, C.; Andre, I.; Baker, D. *Proteins* **2007**, *69* (Suppl 8), 118–128.
27. Raman, S.; Vernon, R.; Thompson, J.; Tyka, M.; Sadreyev, R.; Pei, J.; Kim, D.; Kellogg, E.; DiMaio, F.; Lange, O.; Kinch, L.; Sheffler, W.; Kim, B. H.; Das, R.; Grishin, N. V.; Baker, D. *Proteins* **2009**, *77* (Suppl 9), 89–99.
28. Qian, B.; Raman, S.; Das, R.; Bradley, P.; McCoy, A. J.; Read, R. J.; Baker, D. *Nature* **2007**, *450*, 259–264.
29. Chothia, C.; Lesk, A. M. *EMBO J.* **1986**, *5*, 823–826.
30. Sanchez, R.; Sali, A. *Proc. Natl. Acad. Sci. U. S. A.* **1998**, *95*, 13597–13602.
31. Sanchez, R.; Pieper, U.; Melo, F.; Eswar, N.; Marti-Renom, M. A.; Madhusudhan, M. S.; Mirkovic, N.; Sali, A. *Nat. Struct. Biol.* **2000**, *7*, 986–990.
32. Sali, A. *Nat. Struct. Biol.* **1998**, *5*, 1029–1032.
33. Sali, A. *Nat. Struct. Biol.* **2001**, *8*, 482–484.
34. Pieper, U.; Webb, B. M.; Barkan, D. T.; Schneidman-Duhovny, D.; Schlessinger, A.; Braberg, H.; Yang, Z.; Meng, E. C.; Pettersen, E. F.; Huang, C. C.; Datta, R. S.; Sampathkumar, P.; Madhusudhan, M. S.; Sjolander, K.; Ferrin, T. E.; Burley, S. K.; Sali, A. *Nucleic Acids Res.* **2011**, *39*, 465–474.
35. Sali, A.; Blundell, T. L. *J. Mol. Biol.* **1993**, *234*, 779–815.
36. Dietmann, S.; Park, J.; Notredame, C.; Heger, A.; Lappe, M.; Holm, L. *Nucleic Acids Res.* **2001**, *29*, 55–57.
37. Rost, B. *Protein Eng.* **1999**, *12*, 85–94.
38. Pearson, W. R. *Methods Mol. Biol.* **1994**, *24*, 307–331.
39. Altschul, S. F.; Madden, T. L.; Schaffer, A. A.; Zhang, J.; Zhang, W.; Lipman, D. J. *Nucleic Acids Res.* **1997**, *25*, 3389–3402.
40. Brenner, S. E.; Chothia, C.; Hubbard, T. J. *Proc. Natl. Acad. Sci. U. S. A.* **1998**, *95*, 6073–6078.
41. Sauder, J. M.; Arthur, J. W.; Dunbrack, R. L., Jr. *Proteins* **2000**, *40*, 6–22.
42. Saqi, M. A.; Russell, R. B.; Sternberg, M. J. *Protein Eng.* **1998**, *11*, 627–630.
43. Gribskov, M.; McLachlan, A. D.; Eisenberg, D. *Proc. Natl. Acad. Sci. U. S. A.* **1987**, *84*, 4355–4358.
44. Henikoff, J. G.; Henikoff, S. *Comput. Appl. Biosci.* **1996**, *12*, 135–143.
45. Henikoff, S.; Henikoff, J. G. *J. Mol. Biol.* **1994**, *243*, 574–578.
46. Eddy, S. R. *Bioinformatics* **1998**, *14*, 755–763.
47. Krogh, A.; Brown, M.; Mian, I. S.; Sjolander, K.; Haussler, D. *J. Mol. Biol.* **1994**, *235*, 1501–1531.
48. Lindahl, E.; Elofsson, A. *J. Mol. Biol.* **2000**, *295*, 613–625.
49. Park, J.; Karplus, K.; Barrett, C.; Hughey, R.; Haussler, D.; Hubbard, T.; Chothia, C. *J. Mol. Biol.* **1998**, *284*, 1201–1210.
50. Marti-Renom, M. A.; Madhusudhan, M. S.; Sali, A. *Protein Sci.* **2004**, *13*, 1071–1087.
51. Karplus, K.; Barrett, C.; Hughey, R. *Bioinformatics* **1998**, *14*, 846–856.
52. Soding, J. *Bioinformatics* **2005**, *21*, 951–960.
53. Remmert, M.; Biegert, A.; Hauser, A.; Soding, J. *Nat. Methods* **2012**, *9*, 173–175.
54. Modeller 9.12. Available from: <http://salilab.org/modeller/>.
55. Edgar, R. C.; Sjolander, K. *Bioinformatics* **2004**, *20*, 1301–1308.
56. Ohlson, T.; Wallner, B.; Elofsson, A. *Proteins* **2004**, *57*, 188–197.
57. Wang, G.; Dunbrack, R. L., Jr. *Protein Sci.* **2004**, *13*, 1612–1626.
58. Zhou, H.; Zhou, Y. *Proteins* **2005**, *58*, 321–328.
59. Panchenko, A. R. *Nucleic Acids Res.* **2003**, *31*, 683–689.
60. Pietrovski, S. *Nucleic Acids Res.* **1996**, *24*, 3836–3845.
61. Rychlewski, L.; Zhang, B.; Godzik, A. *Fold. Des.* **1998**, *3*, 229–238.
62. Sadreyev, R.; Grishin, N. *J. Mol. Biol.* **2003**, *326*, 317–336.
63. von Ohlsen, N.; Sommer, I.; Zimmer, R. *Pac. Symp. Biocomput.* **2003**, *8*, 252–263.
64. Yona, G.; Levitt, M. *J. Mol. Biol.* **2002**, *315*, 1257–1275.
65. Jaroszewski, L.; Rychlewski, L.; Li, Z.; Li, W.; Godzik, A. *Nucleic Acids Res.* **2005**, *33*, W284–W288.
66. McGuffin, L. J.; Jones, D. T. *Bioinformatics* **2003**, *19*, 874–881.
67. Karchin, R.; Cline, M.; Mandel-Gutfreund, Y.; Karplus, K. *Proteins* **2003**, *51*, 504–514.
68. Shi, J.; Blundell, T. L.; Mizuguchi, K. *J. Mol. Biol.* **2001**, *310*, 243–257.
69. Bowie, J. U.; Luthy, R.; Eisenberg, D. *Science* **1991**, *253*, 164–170.
70. Sippl, M. *J. Mol. Biol.* **1990**, *213*, 859–883.
71. Sippl, M. *J. Curr. Opin. Struct. Biol.* **1995**, *5*, 229–235.
72. Skolnick, J.; Kihara, D. *Proteins* **2001**, *42*, 319–331.
73. Xu, J.; Li, M.; Kim, D.; Xu, Y. *J. Bioinf. Comput. Biol.* **2003**, *1*, 95–117.
74. Karplus, K.; Karchin, R.; Draper, J.; Casper, J.; Mandel-Gutfreund, Y.; Diekhans, M.; Hughey, R. *Proteins* **2003**, *53* (Suppl 6), 491–496.
75. Jones, D. T. *J. Mol. Biol.* **1999**, *287*, 797–815.
76. Kelley, L. A.; MacCallum, R. M.; Sternberg, M. J. *J. Mol. Biol.* **2000**, *299*, 499–520.
77. Karplus, K. *Nucleic Acids Res.* **2009**, *37*, W492–W497.
78. Wu, S.; Zhang, Y. *Proteins* **2008**, *72*, 547–556.
79. John, B.; Sali, A. *Nucleic Acids Res.* **2003**, *31*, 3982–3992.
80. Moul, J. *Curr. Opin. Struct. Biol.* **2005**, *15*, 285–289.
81. Song, Y.; Dimaio, F.; Wang, R. Y.; Kim, D.; Miles, C.; Brunette, T.; Thompson, J.; Baker, D. *Structure* **2013**, *21*, 1735–1742.

82. Sanchez, R.; Sali, A. *Curr. Opin. Struct. Biol.* **1997**, *7*, 206–214.
83. Sanchez, R.; Sali, A. *Proteins* **1997**, (Suppl 1), 50–58.
84. Srinivasan, N.; Blundell, T. L. *Protein Eng.* **1993**, *6*, 501–512.
85. Bajorath, J.; Aruffo, A. *Bioconj. Chem.* **1994**, *5*, 173–181.
86. Blundell, T. L.; Sibanda, B. L.; Sternberg, M. J.; Thornton, J. M. *Nature* **1987**, *326*, 347–352.
87. Browne, W. J.; North, A. C.; Phillips, D. C.; Brew, K.; Vanaman, T. C.; Hill, R. L. *J. Mol. Biol.* **1969**, *42*, 65–86.
88. Johnson, M. S.; Srinivasan, N.; Sowdhamini, R.; Blundell, T. L. *Crit. Rev. Biochem. Mol. Biol.* **1994**, *29*, 1–68.
89. Greer, J. J. *J. Mol. Biol.* **1981**, *153*, 1027–1042.
90. Nagarajaram, H. A.; Reddy, B. V.; Blundell, T. L. *Protein Eng.* **1999**, *12*, 1055–1062.
91. Sutcliffe, M. J.; Haneef, I.; Carney, D.; Blundell, T. L. *Protein Eng.* **1987**, *1*, 377–384.
92. Sutcliffe, M. J.; Hayes, F. R.; Blundell, T. L. *Protein Eng.* **1987**, *1*, 385–392.
93. Topham, C. M.; McLeod, A.; Eisenmenger, F.; Overington, J. P.; Johnson, M. S.; Blundell, T. L. *J. Mol. Biol.* **1993**, *229*, 194–220.
94. Bates, P. A.; Kelley, L. A.; MacCallum, R. M.; Sternberg, M. J. *Proteins* **2001**, (Suppl 5), 39–46.
95. Schwede, T.; Kopp, J.; Guex, N.; Peitsch, M. C. *Nucleic Acids Res.* **2003**, *31*, 3381–3385.
96. Bystroff, C.; Baker, D. *J. Mol. Biol.* **1998**, *281*, 565–577.
97. Claessens, M.; Van Cutsem, E.; Lasters, I.; Wodak, S. *Protein Eng.* **1989**, *2*, 335–345.
98. Jones, T. A.; Thirup, S. *EMBO J.* **1986**, *5*, 819–822.
99. Levitt, M. J. *J. Mol. Biol.* **1992**, *226*, 507–533.
100. Unger, R.; Harel, D.; Wherland, S.; Sussman, J. L. *Proteins* **1989**, *5*, 355–373.
101. Aszodi, A.; Taylor, W. R. *Fold. Des.* **1996**, *1*, 325–334.
102. Brocklehurst, S. M.; Perham, R. N. *Protein Sci.* **1993**, *2*, 626–639.
103. Havel, T. F.; Snow, M. E. *J. Mol. Biol.* **1991**, *217*, 1–7.
104. Srinivasan, S.; March, C. J.; Sudarsanam, S. *Protein Sci.* **1993**, *2*, 277–289.
105. Fiser, A.; Do, R. K. G.; Sali, A. *Protein Sci.* **2000**, *9*, 1753–1773.
106. Fiser, A.; Feig, M.; Brooks, C. L.; Sali, A. *Acc. Chem. Res.* **2002**, *35*, 413–421.
107. MacKerell, A. D.; Bashford, D.; Bellott, M.; Dunbrack, R. L., Jr.; Evanseck, J. D.; Field, M. J.; Fischer, S.; Gao, J.; Guo, H.; Ha, S.; Joseph-McCarthy, D.; Kuchnir, L.; Kuczera, K.; Lau, F. T. K.; Mattos, C.; Michnick, S.; Ngo, T.; Nguyen, D. T.; Prodhom, B.; Reiher, W. E.; Roux, B.; Schlenkrich, M.; Smith, J. C.; Stote, R.; Straub, J.; Watanabe, M.; Wiorkiewicz-Kuczera, J.; Yin, D.; Karplus, M. *J. Phys. Chem. B* **1998**, *102*, 3586–3616.
108. Sali, A.; Overington, J. P. *Protein Sci.* **1994**, *3*, 1582–1596.
109. Marti-Renom, M. A.; Madhusudhan, M. S.; Fiser, A.; Rost, B.; Sali, A. *Structure* **2002**, *10*, 435–440.
110. Wallner, B.; Elofsson, A. *Protein Sci.* **2005**, *14*, 1315–1327.
111. Kabsch, W.; Sander, C. *Proc. Natl. Acad. Sci. U. S. A.* **1984**, *81*, 1075–1078.
112. Mezei, M. *Protein Eng.* **1998**, *11*, 411–414.
113. Jacobson, M. P.; Pincus, D. L.; Rapp, C. S.; Day, T. J.; Honig, B.; Shaw, D. E.; Friesner, R. A. *Proteins* **2004**, *55*, 351–367.
114. Zhu, K.; Pincus, D. L.; Zhao, S.; Friesner, R. A. *Proteins* **2006**, *65*, 438–452.
115. Chothia, C.; Lesk, A. M. *J. Mol. Biol.* **1987**, *196*, 901–917.
116. Moulton, J.; James, M. N. *Proteins* **1986**, *1*, 146–163.
117. Bruccoleri, R. E.; Karplus, M. *Biopolymers* **1987**, *26*, 137–168.
118. Shenkin, P. S.; Yarmush, D. L.; Fine, R. M.; Wang, H. J.; Levinthal, C. *Biopolymers* **1987**, *26*, 2053–2085.
119. van Vlijmen, H. W.; Karplus, M. *J. Mol. Biol.* **1997**, *267*, 975–1001.
120. Deane, C. M.; Blundell, T. L. *Protein Sci.* **2001**, *10*, 599–612.
121. Sibanda, B. L.; Blundell, T. L.; Thornton, J. M. *J. Mol. Biol.* **1989**, *206*, 759–777.
122. Chothia, C.; Lesk, A. M.; Tramontano, A.; Levitt, M.; Smith-Gill, S. J.; Air, G.; Sheriff, S.; Padlan, E. A.; Davies, D.; Tulip, W. R.; et al. *Nature* **1989**, *342*, 877–883.
123. Rufino, S. D.; Donate, L. E.; Canard, L. H.; Blundell, T. L. *J. Mol. Biol.* **1997**, *267*, 352–367.
124. Oliva, B.; Bates, P. A.; Querol, E.; Aviles, F. X.; Sternberg, M. J. *J. Mol. Biol.* **1997**, *266*, 814–830.
125. Ring, C. S.; Kneller, D. G.; Langridge, R.; Cohen, F. E. *J. Mol. Biol.* **1992**, *224*, 685–699.
126. Fidelis, K.; Stern, P. S.; Bacon, D.; Moulton, J. *Protein Eng.* **1994**, *7*, 953–960.
127. Lessel, U.; Schomburg, D. *Protein Eng.* **1994**, *7*, 1175–1187.
128. Fernandez-Fuentes, N.; Fiser, A. *BMC Struct. Biol.* **2006**, *6*, 15.
129. Fine, R. M.; Wang, H.; Shenkin, P. S.; Yarmush, D. L.; Levinthal, C. *Proteins* **1986**, *1*, 342–362.
130. Sellers, B. D.; Zhu, K.; Zhao, S.; Friesner, R. A.; Jacobson, M. P. *Proteins* **2008**, *72*, 959–971.
131. Bruccoleri, R. E.; Karplus, M. *Biopolymers* **1990**, *29*, 1847–1862.
132. Ring, C. S.; Sun, E.; McKerrow, J. H.; Lee, G. K.; Rosenthal, P. J.; Kuntz, I. D.; Cohen, F. E. *Proc. Natl. Acad. Sci. U. S. A.* **1993**, *90*, 3583–3587.
133. Abagyan, R.; Totrov, M. *J. Mol. Biol.* **1994**, *235*, 983–1002.
134. Collura, V.; Higo, J.; Garnier, J. *Protein Sci.* **1993**, *2*, 1502–1510.
135. Higo, J.; Collura, V.; Garnier, J. *Biopolymers* **1992**, *32*, 33–43.
136. Zheng, Q.; Rosenfeld, R.; Vajda, S.; DeLisi, C. *Protein Sci.* **1993**, *2*, 1242–1248.
137. Koehl, P.; Delarue, M. *Nat. Struct. Biol.* **1995**, *2*, 163–170.
138. Mandell, D. J.; Coutsiaris, E. A.; Kortemme, T. *Nat. Methods* **2009**, *6*, 551–552.
139. Samudrala, R.; Moulton, J. *J. Mol. Biol.* **1998**, *279*, 287–302.
140. Xiang, Z.; Soto, C. S.; Honig, B. *Proc. Natl. Acad. Sci. U. S. A.* **2002**, *99*, 7432–7437.
141. de Bakker, P. I.; DePristo, M. A.; Burke, D. F.; Blundell, T. L. *Proteins* **2003**, *51*, 21–40.
142. DePristo, M. A.; de Bakker, P. I.; Lovell, S. C.; Blundell, T. L. *Proteins* **2003**, *51*, 41–55.
143. Felts, A. K.; Gallicchio, E.; Wallqvist, A.; Levy, R. M. *Proteins* **2002**, *48*, 404–422.
144. Dunbrack, R. L., Jr. *Curr. Opin. Struct. Biol.* **2002**, *12*, 431–440.
145. Bradley, P.; Misura, K. M.; Baker, D. *Science* **2005**, *309*, 1868–1871.
146. Janin, J.; Chothia, C. *Biochemistry (Mosc.)* **1978**, *17*, 2943–2948.
147. Ponder, J. W.; Richards, F. M. *J. Mol. Biol.* **1987**, *193*, 775–791.
148. Scouras, A. D.; Daggett, V. *Protein Sci.* **2011**, *20*, 341–352.
149. De Maeyer, M.; Desmet, J.; Lasters, I. *Fold. Des.* **1997**, *2*, 53–66.
150. Dunbrack, R. L., Jr.; Cohen, F. E. *Protein Sci.* **1997**, *6*, 1661–1681.
151. Dunbrack, R. L., Jr.; Karplus, M. *J. Mol. Biol.* **1993**, *230*, 543–574.
152. Lovell, S. C.; Word, J. M.; Richardson, J. S.; Richardson, D. C. *Proteins* **2000**, *40*, 389–408.

153. McGregor, M. J.; Islam, S. A.; Sternberg, M. J. *J. Mol. Biol.* **1987**, *198*, 295–310.
154. Schrauber, H.; Eisenhaber, F.; Argos, P. *J. Mol. Biol.* **1993**, *230*, 592–612.
155. Tuffery, P.; Etchebest, C.; Hazout, S.; Lavery, R. *J. Biomol. Struct. Dyn.* **1991**, *8*, 1267–1289.
156. Desjarlais, J. R.; Handel, T. M. *J. Mol. Biol.* **1999**, *290*, 305–318.
157. De Filippis, V.; Sander, C.; Vriend, G. *Protein Eng.* **1994**, *7*, 1203–1208.
158. Chung, S. Y.; Subbiah, S. *Pac. Symp. Biocomput.* **1996**, *1*, 126–141.
159. Cregut, D.; Liautard, J. P.; Chiche, L. *Protein Eng.* **1994**, *7*, 1333–1344.
160. Krivov, G. G.; Shapovalov, M. V.; Dunbrack, R. L., Jr. *Proteins* **2009**, *77*, 778–795.
161. Canutescu, A. A.; Shelenkov, A. A.; Dunbrack, R. L., Jr. *Protein Sci.* **2003**, *12*, 2001–2014.
162. Xiang, Z.; Honig, B. *J. Mol. Biol.* **2001**, *311*, 421–430.
163. Eisenmenger, F.; Argos, P.; Abagyan, R. *J. Mol. Biol.* **1993**, *231*, 849–860.
164. Lee, G. M.; Varma, A.; Palsson, B. O. *Biotechnol. Prog.* **1991**, *7*, 72–75.
165. Holm, L.; Sander, C. *Proteins* **1992**, *14*, 213–223.
166. Lasters, I.; Desmet, J. *Protein Eng.* **1993**, *6*, 717–722.
167. Looger, L. L.; Hellinga, H. W. *J. Mol. Biol.* **2001**, *307*, 429–445.
168. Hwang, J. K.; Liao, W. F. *Protein Eng.* **1995**, *8*, 363–370.
169. Koehl, P.; Delarue, M. *J. Mol. Biol.* **1994**, *239*, 249–275.
170. Bower, M. J.; Cohen, F. E.; Dunbrack, R. L., Jr. *J. Mol. Biol.* **1997**, *267*, 1268–1282.
171. Petrella, R. J.; Lazaridis, T.; Karplus, M. *Fold. Des.* **1998**, *3*, 353–377.
172. Jacobson, M. P.; Kaminski, G. A.; Friesner, R. A. *J. Phys. Chem. B* **2002**, *106*, 11673–11680.
173. Liang, S.; Grishin, N. V. *Protein Sci.* **2002**, *11*, 322–331.
174. Coutsiias, E. A.; Seok, C.; Jacobson, M. P.; Dill, K. A. *J. Comput. Chem.* **2004**, *25*, 510–528.
175. Luthy, R.; Bowie, J. U.; Eisenberg, D. *Nature* **1992**, *356*, 83–85.
176. Melo, F.; Sanchez, R.; Sali, A. *Protein Sci.* **2002**, *11*, 430–448.
177. Sippl, M. *J. Proteins* **1993**, *17*, 355–362.
178. Wiederstein, M.; Sippl, M. *J. Nucleic Acids Res.* **2007**, *35*, W407–W410.
179. Topham, C. M.; Srinivasan, N.; Thorpe, C. J.; Overington, J. P.; Kalsheker, N. A. *Protein Eng.* **1994**, *7*, 869–894.
180. Melo, F.; Feytmans, E. *J. Mol. Biol.* **1998**, *277*, 1141–1152.
181. Zhou, H.; Zhou, Y. *Protein Sci.* **2002**, *11*, 2714–2726.
182. Shen, M. Y.; Sali, A. *Protein Sci.* **2006**, *15*, 2507–2524.
183. Benkert, P.; Biasini, M.; Schwede, T. *Bioinformatics* **2011**, *27*, 343–350.
184. Ray, A.; Lindahl, E.; Wallner, B. *BMC Bioinformatics* **2012**, *13*, 224.
185. Eramian, D.; Eswar, N.; Shen, M.; Sali, A. *Protein Sci.* **2008**, *17*, 1881–1893.
186. Pawlowski, K.; Bierzynski, A.; Godzik, A. *J. Mol. Biol.* **1996**, *258*, 349–366.
187. Laskowski, R.; MacArthur, M.; Moss, D.; Thornton, J. *J. Appl. Cryst.* **1993**, *26*, 283–291.
188. Hooft, R. W.; Vriend, G.; Sander, C.; Abola, E. E. *Nature* **1996**, *381*, 272.
189. Fan, H.; Hitchcock, D.; Seidel, R.; Hillerich, B.; Lin, H.; Almo, S.; Sali, A.; Shoichet, B.; Raushel, F. *J. Am. Chem. Soc.* **2013**, *135*, 795–803.
190. Fan, H.; Schneidman, D.; Irwin, J. J.; Dong, G.; Shoichet, B.; Sali, A. *J. Chem. Inf. Model.* **2011**, *51*, 3078–3092.
191. Dong, G. Q.; Fan, H.; Schneidman-Duhovny, D.; Webb, B.; Sali, A. *Bioinformatics* **2013**, *29*, 3158–3166. PMID3842762.
192. Zemla, A.; Venclovas, J.; Moul, J.; Fidelis, K. *Proteins* **2001**, (Suppl 5), 13–21.
193. Fischer, D.; Elofsson, A.; Rychlewski, L.; Pazos, F.; Valencia, A.; Rost, B.; Ortiz, A. R.; Dunbrack, R. L., Jr. *Proteins* **2001**, (Suppl 5), 171–183.
194. Bujnicki, J. M.; Elofsson, A.; Fischer, D.; Rychlewski, L. *Protein Sci.* **2001**, *10*, 352–361.
195. Haas, J.; Roth, S.; Arnold, K.; Kiefer, F.; Schmidt, T.; Bordoli, L.; Schwede, T. *Database (Oxford)* **2013**, *2013*, bat031.
196. Eylich, V. A.; Marti-Renom, M. A.; Przybylski, D.; Madhusudhan, M. S.; Fiser, A.; Pazos, F.; Valencia, A.; Sali, A.; Rost, B. *Bioinformatics* **2001**, *17*, 1242–1243.
197. Koh, I. Y.; Eylich, V. A.; Marti-Renom, M. A.; Przybylski, D.; Madhusudhan, M. S.; Eswar, N.; Grana, O.; Pazos, F.; Valencia, A.; Sali, A.; Rost, B. *Nucleic Acids Res.* **2003**, *31*, 3311–3315.
198. Schwede, T.; Sali, A.; Honig, B.; Levitt, M.; Berman, H. M.; Jones, D.; Brenner, S. E.; Burley, S. K.; Das, R.; Dokholyan, N. V.; Dunbrack, R. L., Jr.; Fidelis, K.; Fiser, A.; Godzik, A.; Huang, Y. J.; Humblet, C.; Jacobson, M. P.; Joachimiak, A.; Krystek, S. R., Jr.; Kortemme, T.; Kryzhanovych, A.; Montelione, G. T.; Moul, J.; Murray, D.; Sanchez, R.; Sosnick, T. R.; Standley, D. M.; Stouch, T.; Vajda, S.; Vasquez, M.; Westbrook, J. D.; Wilson, I. A. *Structure* **2009**, *17*, 151–159.
199. Karchin, R.; Diekhans, M.; Kelly, L.; Thomas, D. J.; Pieper, U.; Eswar, N.; Haussler, D.; Sali, A. *Bioinformatics* **2005**, *21*, 2814–2820.
200. Thiel, K. A. *Nat. Biotechnol.* **2004**, *22*, 513–519.
201. Jacobson, M.; Sali, A. Comparative Protein Structure Modeling and Its Applications to Drug Discovery. In *Annu Rep Med Chem*; Overington, J., Ed.; Inpharmatica Ltd.: London, 2004; pp 259–276.
202. Gao, H. X.; Sengupta, J.; Valle, M.; Korostelev, A.; Eswar, N.; Stagg, S. M.; Van Roey, P.; Agrawal, R. K.; Harvey, S. C.; Sali, A.; Chapman, M. S.; Frank, J. *Cell* **2003**, *113*, 789–801.
203. Spahn, C. M.; Beckmann, R.; Eswar, N.; Penczek, P. A.; Sali, A.; Blobel, G.; Frank, J. *Cell* **2001**, *107*, 373–386.
204. Blundell, T. L.; Johnson, M. S. *Protein Sci.* **1993**, *2*, 877–883.
205. Chakravarty, S.; Sanchez, R. *Structure* **2004**, *12*, 1461–1470.
206. Chakravarty, S.; Wang, L.; Sanchez, R. *Nucleic Acids Res.* **2005**, *33*, 244–259.
207. von Grothuss, M.; Wyrwicz, L. S.; Rychlewski, L. *Cell* **2003**, *113*, 701–702.
208. Gordon, R. K.; Ginalski, K.; Rudnicki, W. R.; Rychlewski, L.; Pankaskie, M. C.; Bujnicki, J. M.; Chiang, P. K. *Eur. J. Biochem.* **2003**, *270*, 3507–3517.
209. Evers, A.; Gohlke, H.; Klebe, G. *J. Mol. Biol.* **2003**, *334*, 327–345.
210. Evers, A.; Klebe, G. *Angew. Chem.* **2004**, *43*, 248–251.
211. Schafferhans, A.; Klebe, G. *J. Mol. Biol.* **2001**, *307*, 407–427.
212. Bissantz, C.; Bernard, P.; Hibert, M.; Rognan, D. *Proteins* **2003**, *50*, 5–25.
213. Cavasotto, C. N.; Orry, A. J.; Abagyan, R. A. *Proteins* **2003**, *51*, 423–433.
214. Evers, A.; Klabunde, T. *J. Med. Chem.* **2005**, *48*, 1088–1097.
215. Evers, A.; Klebe, G. *J. Med. Chem.* **2004**, *47*, 5381–5392.
216. Moro, S.; Defflorian, F.; Bacilieri, M.; Spalluto, G. *Curr. Med. Chem.* **2006**, *13*, 639–645.
217. Nowak, M.; Kolaczowski, M.; Pawlowski, M.; Bojarski, A. *J. Med. Chem.* **2006**, *49*, 205–214.
218. Chen, J. Z.; Wang, J.; Xie, X. Q. *J. Chem. Inf. Model.* **2007**, *47*, 1626–1637.
219. Zylberg, J.; Ecke, D.; Fischer, B.; Reiser, G. *Biochem. J.* **2007**, *405*, 277–286.
220. Radesstock, S.; Weil, T.; Renner, S. *J. Chem. Inf. Model.* **2008**, *48*, 1104–1117.

221. Singh, N.; Cheve, G.; Ferguson, D. M.; McCurdy, C. R. *J. Comput. Aided Mol. Des.* **2006**, *20*, 471–493.
222. Kiss, R.; Kiss, B.; Konczol, A.; Szalai, F.; Jelinek, I.; Laszlo, V.; Noszal, B.; Falus, A.; Keseru, G. M. *J. Med. Chem.* **2008**, *51*, 3145–3153.
223. de Graaf, C.; Foata, N.; Engkvist, O.; Rognan, D. *Proteins* **2008**, *71*, 599–620.
224. Diller, D. J.; Li, R. *J. Med. Chem.* **2003**, *46*, 4638–4647.
225. Oshiro, C.; Bradley, E. K.; Eksterowicz, J.; Evensen, E.; Lamb, M. L.; Lancot, J. K.; Putta, S.; Stanton, R.; Grootenhuys, P. D. *J. Med. Chem.* **2004**, *47*, 764–767.
226. Nguyen, T. L.; Gussio, R.; Smith, J. A.; Lannigan, D. A.; Hecht, S. M.; Scudiero, D. A.; Shoemaker, R. H.; Zaharevitz, D. W. *Biorg. Med. Chem.* **2006**, *14*, 6097–6105.
227. Rockey, W. M.; Elcock, A. H. *Curr. Protein Pept. Sci.* **2006**, *7*, 437–457.
228. Schapira, M.; Abagyan, R.; Totrov, M. *J. Med. Chem.* **2003**, *46*, 3045–3059.
229. Marhefka, C. A.; Moore, B. M., 2nd; Bishop, T. C.; Kirkovsky, L.; Mukherjee, A.; Dalton, J. T.; Miller, D. D. *J. Med. Chem.* **2001**, *44*, 1729–1740.
230. Kasuya, A.; Sawada, Y.; Tsukamoto, Y.; Tanaka, K.; Toya, T.; Yanagi, M. *J. Mol. Model.* **2003**, *9*, 58–65.
231. Li, R.; Chen, X.; Gong, B.; Selzer, P. M.; Li, Z.; Davidson, E.; Kurzbau, G.; Miller, R. E.; Nuzum, E. O.; McKerrow, J. H.; Fletterick, R. J.; Gillmor, S. A.; Craik, C. S.; Kuntz, I. D.; Cohen, F. E.; Kenyon, G. L. *Biorg. Med. Chem.* **1996**, *4*, 1421–1427.
232. Selzer, P. M.; Chen, X.; Chan, V. J.; Cheng, M.; Kenyon, G. L.; Kuntz, I. D.; Sakanari, J. A.; Cohen, F. E.; McKerrow, J. H. *Exp. Parasitol.* **1997**, *87*, 212–221.
233. Enyedy, I. J.; Ling, Y.; Nacro, K.; Tomita, Y.; Wu, X.; Cao, Y.; Guo, R.; Li, B.; Zhu, X.; Huang, Y.; Long, Y. Q.; Roller, P. P.; Yang, D.; Wang, S. *J. Med. Chem.* **2001**, *44*, 4313–4324.
234. de Graaf, C.; Oostenbrink, C.; Keizers, P. H.; van der Wijst, T.; Jongejan, A.; Vermeulen, N. P. *J. Med. Chem.* **2006**, *49*, 2417–2430.
235. Katritch, V.; Byrd, C. M.; Tseitlin, V.; Dai, D.; Raush, E.; Totrov, M.; Abagyan, R.; Jordan, R.; Hraby, D. E. *J. Comput. Aided Mol. Des.* **2007**, *21*, 549–558.
236. Mukherjee, P.; Desai, P. V.; Srivastava, A.; Tekwani, B. L.; Avery, M. A. *J. Chem. Inf. Model.* **2008**, *48*, 1026–1040.
237. Song, L.; Kalyanaraman, C.; Fedorov, A. A.; Fedorov, E. V.; Glasner, M. E.; Brown, S.; Imker, H. J.; Babbitt, P. C.; Almo, S. C.; Jacobson, M. P.; Gerlt, J. A. *Nat. Chem. Biol.* **2007**, *3*, 486–491.
238. Kalyanaraman, C.; Imker, H. J.; Fedorov, A. A.; Fedorov, E. V.; Glasner, M. E.; Babbitt, P. C.; Almo, S. C.; Gerlt, J. A.; Jacobson, M. P. *Structure* **2008**, *16*, 1668–1677.
239. Rotkiewicz, P.; Sicińska, W.; Kolinski, A.; DeLuca, H. F. *Proteins* **2001**, *44*, 188–199.
240. Que, X.; Brinen, L. S.; Perkins, P.; Herdman, S.; Hirata, K.; Torian, B. E.; Rubin, H.; McKerrow, J. H.; Reed, S. L. *Mol. Biochem. Parasitol.* **2002**, *119*, 23–32.
241. Parrill, A. L.; Echols, U.; Nguyen, T.; Pham, T. C.; Hoeglund, A.; Baker, D. L. *Biorg. Med. Chem.* **2008**, *16*, 1784–1795.
242. Fernandes, M. X.; Kairys, V.; Gilson, M. K. *J. Chem. Inf. Comput. Sci.* **2004**, *44*, 1961–1970.
243. Kairys, V.; Fernandes, M. X.; Gilson, M. K. *J. Chem. Inf. Model.* **2006**, *46*, 365–379.
244. McGovern, S. L.; Shoichet, B. K. *J. Med. Chem.* **2003**, *46*, 2895–2907.
245. Fan, H.; Irwin, J. J.; Webb, B. M.; Klebe, G.; Shoichet, B.; Sali, A. *J. Chem. Inf. Model.* **2009**, *49*, 2512–2527.
246. Huang, N.; Shoichet, B. K.; Irwin, J. J. *J. Med. Chem.* **2006**, *49*, 6789–6801.
247. Lorber, D. M.; Shoichet, B. K. *Curr. Top. Med. Chem.* **2005**, *5*, 739–749.
248. Enyedy, I. J.; Lee, S. L.; Kuo, A. H.; Dickson, R. B.; Lin, C. Y.; Wang, S. *J. Med. Chem.* **2001**, *44*, 1349–1355.
249. Schapira, M.; Raaka, B. M.; Samuels, H. H.; Abagyan, R. *BMC Struct. Biol.* **2001**, *1*, 1.
250. Zuccotto, F.; Zvebil, M.; Brun, R.; Chowdhury, S. F.; Di Lucrezia, R.; Leal, I.; Maes, L.; Ruiz-Perez, L. M.; Gonzalez Pacanowska, D.; Gilbert, I. H. *Eur. J. Med. Chem.* **2001**, *36*, 395–405.
251. Anand, K.; Ziebuhr, J.; Wadhwani, P.; Mesters, J. R.; Hilgenfeld, R. *Science* **2003**, *300*, 1763–1767.
252. Rajnarayanan, R. V.; Dakshnamurthy, S.; Pattabiraman, N. *Biochem. Biophys. Res. Commun.* **2004**, *321*, 370–378.
253. Diller, D. J.; Merz, K. M., Jr. *Proteins* **2001**, *43*, 113–124.
254. Rockey, W. M.; Elcock, A. H. *Proteins* **2002**, *48*, 664–671.
255. Vangrevelinghe, E.; Zimmermann, K.; Schoepfer, J.; Portmann, R.; Fabbro, D.; Furet, P. *J. Med. Chem.* **2003**, *46*, 2656–2662.
256. Becker, O. M.; Shacham, S.; Marantz, Y.; Noiman, S. *Curr. Opin. Drug Discov. Devel.* **2003**, *6*, 353–361.
257. Bissantz, C.; Logean, A.; Rognan, D. *J. Chem. Inf. Comput. Sci.* **2004**, *44*, 1162–1176.
258. Shacham, S.; Topf, M.; Avisar, N.; Glaser, F.; Marantz, Y.; Bar-Haim, S.; Noiman, S.; Naor, Z.; Becker, O. M. *Med. Res. Rev.* **2001**, *21*, 472–483.
259. Vaidehi, N.; Floriano, W. B.; Trabanino, R.; Hall, S. E.; Freddolino, P.; Choi, E. J.; Zamanakos, G.; Goddard, W. A., 3rd. *Proc. Natl. Acad. Sci. U. S. A.* **2002**, *99*, 12622–12627.
260. Palczewski, K.; Kumasaka, T.; Hori, T.; Behnke, C. A.; Motoshima, H.; Fox, B. A.; Le Trong, I.; Teller, D. C.; Okada, T.; Stenkamp, R. E.; Yamamoto, M.; Miyano, M. *Science* **2000**, *289*, 739–745.
261. Chen, N. H.; Reith, M. E.; Quick, M. W. *Pflugers Arch.* **2004**, *447*, 519–531.
262. Schlessinger, A.; Geier, E.; Fan, H.; Irwin, J.; Shoichet, B.; Giacomini, K.; Sali, A. *Proc. Natl. Acad. Sci. U. S. A.* **2011**, *108*, 15810–15815.
263. Schlessinger, A.; Wittwer, M. B.; Dahlin, A.; Khuri, N.; Bonomi, M.; Fan, H.; Giacomini, K.; Sali, A. *J. Biol. Chem.* **2012**, *287*, 37745–37756.
264. Geier, E.; Schlessinger, A.; Fan, H.; Gable, J.; Irwin, J.; Sali, A.; Giacomini, K. *Proc. Natl. Acad. Sci. U. S. A.* **2013**, *110*, 5480–5485.
265. Gao, X.; Zhou, L.; Jiao, X.; Lu, F.; Yan, C.; Zeng, X.; Wang, J.; Shi, Y. *Nature* **2010**, *463*, 828–832.
266. Shaffer, P. L.; Goehring, A.; Shankaranarayanan, A.; Gouaux, E. *Science* **2009**, *325*, 1010–1014.
267. Barril, X.; Morley, S. D. *J. Med. Chem.* **2005**, *48*, 4432–4443.
268. Carlson, H. A.; McCammon, J. A. *Mol. Pharmacol.* **2000**, *57*, 213–218.
269. Flicek, P.; Ahmed, I.; Amode, M. R.; Barrell, D.; Beal, K.; Brent, S.; Carvalho-Silva, D.; Clapham, P.; Coates, G.; Fairley, S.; Fitzgerald, S.; Gil, L.; Garcia-Giron, C.; Gordon, L.; Hourlier, T.; Hunt, S.; Juettemann, T.; Kahari, A. K.; Keenan, S.; Komorowska, M.; Kulesha, E.; Longden, I.; Maurel, T.; McLaren, W. M.; Muffato, M.; Nag, R.; Overduin, B.; Pignatelli, M.; Pritchard, B.; Pritchard, E.; Riat, H. S.; Ritchie, G. R.; Ruffier, M.; Schuster, M.; Sheppard, D.; Sobral, D.; Taylor, K.; Thormann, A.; Trevanion, S.; White, S.; Wilder, S. P.; Aken, B. L.; Birney, E.; Cunningham, F.; Dunham, I.; Harrow, J.; Herrero, J.; Hubbard, T. J.; Johnson, N.; Kinsella, R.; Parker, A.; Spudich, G.; Yates, A.; Zadissa, A.; Searle, S. M. *Nucleic Acids Res.* **2013**, *41*, D48–D55.
270. Huang, H.; Hu, Z. Z.; Arighi, C. N.; Wu, C. H. *Front. Biosci.* **2007**, *12*, 5071–5088.
271. Hunter, S.; Jones, P.; Mitchell, A.; Apweiler, R.; Attwood, T. K.; Bateman, A.; Bernard, T.; Binns, D.; Bork, P.; Burge, S.; de Castro, E.; Coghill, P.; Corbett, M.; Das, U.; Daugherty, L.; Duquenne, L.; Finn, R. D.; Fraser, M.; Gough, J.; Haft, D.; Hulo, N.; Kahn, D.; Kelly, E.; Letunic, I.; Lonsdale, D.; Lopez, R.; Madera, M.; Maslen, J.; McAnulla, C.; McDowall, J.; McMenamin, C.; Mi, H.; Mutowo-Mueller, P.; Mulder, N.; Natale, D.; Orengo, C.; Pesce, S.; Punta, M.; Quinn, A. F.; Rivoire, C.; Sangrador-Vegas, A.; Selengut, J. D.; Sigrist, C. J.; Scheremetjew, M.; Tate, J.; Thimmajananthan, M.; Thomas, P. D.; Wu, C. H.; Yeats, C.; Yong, S. Y. *Nucleic Acids Res.* **2012**, *40*, D306–D312.
272. Bailey, T. L.; Elkan, C. *Proc. Int. Conf. Intell. Syst. Mol. Biol.* **1994**, *2*, 28–36.
273. Attwood, T. K.; Coletta, A.; Muirhead, G.; Pavlopoulou, A.; Philippou, P. B.; Popov, I.; Roma-Mateo, C.; Theodosiou, A.; Mitchell, A. L. *Database (Oxford)* **2012**, *2012*, bas019.
274. Bru, C.; Courcelle, E.; Carrere, S.; Beausse, Y.; Dalmar, S.; Kahn, D. *Nucleic Acids Res.* **2005**, *33*, D212–D215.
275. Hulo, N.; Bairoch, A.; Bulliard, V.; Cerutti, L.; De Castro, E.; Langendijk-Genevaux, P. S.; Pagni, M.; Sigrist, C. J. *Nucleic Acids Res.* **2006**, *34*, D227–D230.
276. Brown, S. D.; Babbitt, P. C. *J. Biol. Chem.* **2012**, *287*, 35–42.
277. Letunic, I.; Doerks, T.; Bork, P. *Nucleic Acids Res.* **2012**, *40*, D302–D305.
278. Gough, J.; Karplus, K.; Hughey, R.; Chothia, C. *J. Mol. Biol.* **2001**, *313*, 903–919.
279. Arnold, K.; Kiefer, F.; Kopp, J.; Battey, J. N.; Podvinec, M.; Westbrook, J. D.; Berman, H. M.; Bordoli, L.; Schwede, T. *J. Struct. Funct. Genomics* **2009**, *10*, 1–8.
280. Kiefer, F.; Arnold, K.; Kunzli, M.; Bordoli, L.; Schwede, T. *Nucleic Acids Res.* **2009**, *37*, D387–D392.
281. Marti-Renom, M. A.; Ilyin, V. A.; Sali, A. *Bioinformatics* **2001**, *17*, 746–747.

282. Lin, J.; Qian, J.; Greenbaum, D.; Bertone, P.; Das, R.; Echols, N.; Senes, A.; Stenger, B.; Gerstein, M. *Nucleic Acids Res.* **2002**, *30*, 4574–4582.
283. Khafizov, K.; Staritzbichler, R.; Stamm, M.; Forrest, L. R. *Biochemistry (Mosc.)* **2010**, *49*, 10702–10713.
284. Thompson, J. D.; Higgins, D. G.; Gibson, T. J. *Nucleic Acids Res.* **1994**, *22*, 4673–4680.
285. Armougom, F.; Moretti, S.; Poirot, O.; Audic, S.; Dumas, P.; Schaeli, B.; Keduas, V.; Notredame, C. *Nucleic Acids Res.* **2006**, *34*, W604–W608.
286. Pearson, W. R. *Methods Mol. Biol.* **2000**, *132*, 185–219.
287. Katoh, K.; Standley, D. M. *Mol. Biol. Evol.* **2013**, *30*, 772–780.
288. Edgar, R. C. *Nucleic Acids Res.* **2004**, *32*, 1792–1797.
289. Pei, J.; Kim, B. H.; Grishin, N. V. *Nucleic Acids Res.* **2008**, *36*, 2295–2300.
290. McGuffin, L. J.; Bryson, K.; Jones, D. T. *Bioinformatics* **2000**, *16*, 404–405.
291. Eswar, N.; John, B.; Mirkovic, N.; Fiser, A.; Ilyin, V. A.; Pieper, U.; Stuart, A. C.; Marti-Renom, M. A.; Madhusudhan, M. S.; Yerkovich, B.; Sali, A. *Nucleic Acids Res.* **2003**, *31*, 3375–3380.
292. Shatsky, M.; Nussinov, R.; Wolfson, H. J. *Proteins* **2006**, *62*, 209–217.
293. Notredame, C. *Curr. Protoc. Bioinformatics* **2010**, *29*, 1–25, Chapter 3, Unit 3.8.
294. Notredame, C.; Higgins, D. G.; Heringa, J. J. *Mol. Biol.* **2000**, *302*, 205–217.
295. Prlic, A.; Bliven, S.; Rose, P. W.; Bluhm, W. F.; Bizon, C.; Godzik, A.; Bourne, P. E. *Bioinformatics* **2010**, *26*, 2983–2985.
296. Guerler, A.; Knapp, E. W. *Protein Sci.* **2008**, *17*, 1374–1382.
297. Ortiz, A. R.; Strauss, C. E.; Olmea, O. *Protein Sci.* **2002**, *11*, 2606–2621.
298. Lupyán, D.; Leo-Macias, A.; Ortiz, A. R. *Bioinformatics* **2005**, *21*, 3255–3263.
299. Dror, O.; Benyamini, H.; Nussinov, R.; Wolfson, H. *Bioinformatics* **2003**, *19*(Suppl 1), i95–i104.
300. Shatsky, M.; Nussinov, R.; Wolfson, H. J. *Proteins* **2004**, *56*, 143–156.
301. Konagurthu, A. S.; Whisstock, J. C.; Stuckey, P. J.; Lesk, A. M. *Proteins* **2006**, *64*, 559–574.
302. Zhang, Y.; Skolnick, J. *Nucleic Acids Res.* **2005**, *33*, 2302–2309.
303. Kaplan, W.; Littlejohn, T. G. *Brief. Bioinform.* **2001**, *2*, 195–197.
304. Huang, C. C.; Novak, W. R.; Babbitt, P. C.; Jewett, A. I.; Ferrin, T. E.; Klein, T. E. *Pac. Symp. Biocomput.* **2000**, *12*, 230–241.
305. Soding, J.; Biegert, A.; Lupas, A. N. *Nucleic Acids Res.* **2005**, *33*, W244–W248.
306. Roche, D. B.; Buenavista, M. T.; Tetchner, S. J.; McGuffin, L. J. *Nucleic Acids Res.* **2011**, *39*, W171–W176.
307. Roy, A.; Kucukural, A.; Zhang, Y. *Nat. Protoc.* **2010**, *5*, 725–738.
308. Fernandez-Fuentes, N.; Rai, B. K.; Madrid-Aliste, C. J.; Fajardo, J. E.; Fiser, A. *Bioinformatics* **2007**, *23*, 2558–2565.
309. ModWeb. Available from: <http://salilab.org/modweb/>.
310. Kelley, L. A.; Sternberg, M. J. *Nat. Protoc.* **2009**, *4*, 363–371.
311. Kallberg, M.; Wang, H.; Wang, S.; Peng, J.; Wang, Z.; Lu, H.; Xu, J. *Nat. Protoc.* **2012**, *7*, 1511–1522.
312. Wang, Q.; Canutescu, A. A.; Dunbrack, R. L., Jr. *Nat. Protoc.* **2008**, *3*, 1832–1847.
313. Colovos, C.; Yeates, T. O. *Protein Sci.* **1993**, *2*, 1511–1519.
314. Arnold, K.; Bordoli, L.; Kopp, J.; Schwede, T. *Bioinformatics* **2006**, *22*, 195–201.
315. Moulton, J.; Fidelis, K.; Zemla, A.; Hubbard, T. *Proteins* **2003**, *53*(Suppl 6), 334–339.
316. Gifford, L. K.; Carter, L. G.; Gabanyi, M. J.; Berman, H. M.; Adams, P. D. *J. Struct. Funct. Genomics* **2012**, *13*, 57–62.
317. Irwin, J. J.; Shoichet, B. K. *J. Chem. Inf. Model.* **2005**, *45*, 177–182.
318. Irwin, J. J.; Sterling, T.; Mysinger, M. M.; Bolstad, E. S.; Coleman, R. G. *J. Chem. Inf. Model.* **2012**, *52*, 1757–1768.
319. Stuart, A. C.; Ilyin, V. A.; Sali, A. *Bioinformatics* **2002**, *18*, 200–201.
320. Davis, F.; Sali, A. *Bioinformatics* **2005**, *21*, 1901–1907.
321. PiBASE. Available from: <http://salilab.org/pibase/>.
322. Mirkovic, N.; Marti-Renom, M. A.; Weber, B. L.; Sali, A.; Monteiro, A. N. *Cancer Res.* **2004**, *64*, 3790–3797.
323. LS-SNP. Available from: <http://salilab.org/LS-SNP>.
324. Fiser, A.; Sali, A. *Bioinformatics* **2003**, *19*, 2500–2501.
325. ModLoop. Available from: <http://salilab.org/modloop/>.
326. Schneidman-Duhovny, D.; Hammel, M.; Sali, A. *Nucleic Acids Res.* **2010**, *38*, 541–544.
327. Schneidman-Duhovny, D.; Kim, S. J.; Sali, A. *BMC Struct. Biol.* **2012**, *12*, 17.
328. FoXS. Available from: <http://salilab.org/foxs/>.
329. Weinkam, P.; Pons, J.; Sali, A. *Proc. Natl. Acad. Sci. U. S. A.* **2012**, *109*, 4875–4880.
330. AllosMod. Available from: <http://salilab.org/allosmod/>.
331. AllosMod-FoXS. Available from: <http://salilab.org/allosmod-foxs/>.
332. CryptoSite. Available from: <http://salilab.org/cryptosite/>.
333. Schneidman-Duhovny, D.; Hammel, M.; Sali, A. *J. Struct. Biol.* **2011**, *3*, 461–471.
334. FoXSDock. Available from: <http://salilab.org/foxsdock/>.
335. Spill, Y.; Kim, S. J.; Schneidman-Duhovny, D.; Russel, D.; Webb, B.; Sali, A.; Nilges, M. *J. Synchrotron Radiat.* **2014**, *21*, 203–208. PMID3874021.
336. SAXS Merge. Available from: <http://salilab.org/saxsmerge/>.
337. Irwin, J. J.; Shoichet, B. K.; Mysinger, M. M.; Huang, N.; Colizzi, F.; Wassam, P.; Cao, Y. *J. Med. Chem.* **2009**, *52*, 5712–5720.
338. Wei, B. Q.; Baase, W. A.; Weaver, L. H.; Matthews, B. W.; Shoichet, B. K. *J. Mol. Biol.* **2002**, *322*, 339–355.
339. Fiser, A.; Sali, A. *Methods Enzymol.* **2003**, *374*, 461–491.
340. Madhusudhan, M. S.; Marti-Renom, M. A.; Eswar, N.; John, B.; Pieper, U.; Karchin, R.; Shen, M.-Y.; Sali, A. Comparative Protein Structure Modeling. In *Proteomics Protocols Handbook*; Walker, J. M., Ed.; Humana Press Inc.: Totowa, NJ, 2005; pp 831–860.