

Evaluation and Improvement of Multiple Sequence Methods for Protein Secondary Structure Prediction

James A. Cuff^{1,2} and Geoffrey J. Barton^{1,2*}

¹Laboratory of Molecular Biophysics, Oxford, United Kingdom

²European Molecular Biology Laboratory Outstation, The European Bioinformatics Institute, Wellcome Trust Genome Campus, Hinxton, Cambridge, United Kingdom

ABSTRACT A new dataset of 396 protein domains is developed and used to evaluate the performance of the protein secondary structure prediction algorithms DSC, PHD, NNSSP, and PREDATOR. The maximum theoretical Q_3 accuracy for combination of these methods is shown to be 78%. A simple consensus prediction on the 396 domains, with automatically generated multiple sequence alignments gives an average Q_3 prediction accuracy of 72.9%. This is a 1% improvement over PHD, which was the best single method evaluated. Segment Overlap Accuracy (SOV) is 75.4% for the consensus method on the 396-protein set. The secondary structure definition method DSSP defines 8 states, but these are reduced by most authors to 3 for prediction. Application of the different published 8- to 3-state reduction methods shows variation of over 3% on apparent prediction accuracy. This suggests that care should be taken to compare methods by the same reduction method. Two new sequence datasets (CB513 and CB251) are derived which are suitable for cross-validation of secondary structure prediction methods without artifacts due to internal homology. A fully automatic World Wide Web service that predicts protein secondary structure by a combination of methods is available via <http://barton.ebi.ac.uk/>. *Proteins* 1999;34:508–519. © 1999 Wiley-Liss, Inc.

Key words: protein; secondary structure prediction; combination of methods; benchmarks

INTRODUCTION

The most successful techniques for prediction of the protein three dimensional structure rely on aligning the sequence of a protein of unknown structure to a homolog of known structure (e.g., see Sali for review¹). Such methods fail if there is no homolog in the structural database, or if the technique for searching the structural database is unable to identify homologs that are present. While absence of a homolog must await further X-ray or NMR structures, up to 4/5 of known homologues may be missed even by the best conventional pairwise sequence comparison methods.²

Techniques that exploit evolutionary information from protein families^{3–9} or use empirical pair-potentials^{10,11} can normally detect more homologs than pairwise sequence

comparison methods. An even greater challenge is to detect proteins that share similar folds, but are not clearly derived from a common ancestor (e.g. Rossman fold domains of lactate dehydrogenase and glycogen phosphorylase, and SH2-BirA¹²).

Techniques for the prediction of protein secondary structure provide information that is useful both in ab initio structure prediction and as an additional constraint for fold-recognition algorithms.^{13–15} Knowledge of secondary structure alone can help in the design of site-directed or deletion mutants that will not destroy the native protein structure. However, for all these applications it is essential that the secondary structure prediction be accurate, or at least that, the reliability for each residue can be assessed.

The majority of secondary structure prediction algorithms derive parameters or rules from an analysis of proteins of known three dimensional structure. The parameters are then applied by the algorithm to the sequence of unknown structure. Such approaches rely on having sufficient data to obtain reliable parameters and to avoid over-training for a specific data set.

Early algorithms to predict protein secondary structure^{16–18} claimed high accuracy for prediction, but on small datasets that were also used in training the methods. For example, Lim (1974)¹⁶ quoted 70% Q_3 ¹⁹ accuracy on a dataset of 25 proteins, Garnier et al. (1978)¹⁸ achieved 63% accuracy for a different set of 26 proteins, and Chou and Fasman (1974)¹⁷ quoted 77% for yet another different set of 19 proteins.

The use of different datasets in training and testing each algorithm makes it difficult to make an objective comparison of methods. For this reason, Kabsch and Sander (1983)²⁰ carried out a test of prediction methods by applying the algorithms to proteins that were not used in their development. In this independent test, the GOR¹⁸ accuracy reduced by 7% to 56%. The Lim¹⁶ accuracy reduced by 14% to 56%, and Chou-Fasman¹⁷ dropped by 27% to 50%. Cross-validation techniques, where test proteins are removed from the training set, have allowed more realistic evaluation of prediction accuracy to be obtained.

Grant sponsor: Medical Research Council; Grant sponsor: Royal Society.

*Correspondence to: G.J. Barton, EMBL-European Bioinformatics Institute, Wellcome Trust Genome Campus, Hinxton, Cambridge, CB10 1SD, United Kingdom.

Received 25 August 1998; Accepted 6 November 1998.

Prediction from a multiple alignment of protein sequences rather than a single sequence has long been recognized as a way to improve prediction accuracy.¹⁸ During evolution, residues with similar physico-chemical properties are conserved if they are important to the fold or function of the protein. This makes patterns of hydrophobic residues characteristic of particular secondary structures easier to identify.²¹ Analysis of conservation in protein families has been effective in many secondary structure predictions performed before knowledge of the protein structure.^{22–25} Zvelebil et al. (1987)²⁶ developed an automatic procedure that showed a 9% improvement in prediction accuracy on a small set of protein families when multiple sequence data was included. Most current secondary structure prediction algorithms exploit similar principles to gain higher accuracy than is possible from a single sequence.^{27–30} The recent CASP³¹ series of experiments in which predictions are made blind have shown that recent claims for secondary structure prediction algorithms³² are within reasonable limits.

Prediction accuracy has also been improved by combining more than one algorithm on a single sequence.^{33–37} For example, Zhang et al. (1992)³⁴ obtained 66.4% accuracy on a set of 107 proteins, an improvement of 2% over the best method they considered.

In this paper we describe datasets and procedures for the evaluation of current techniques for secondary structure prediction. We discuss the effects of homology within the training and test datasets and describe new non-redundant datasets appropriate for developing secondary structure prediction algorithms. We evaluate the accuracy of four recently published algorithms that exploit multiple sequence data NNSSP,³⁰ PHD,²⁷ DSC,²⁹ and PREDATOR²⁸ and two older methods, ZPRED²⁶ and MULPRED (Barton, unpublished). We develop an algorithm that combines the predictions of PHD, DSC, PREDATOR, and NNSSP and show that it gives a 1% improvement in average accuracy over the best single method. Finally, we investigate the effect of the quality of multiple sequence alignment used in prediction, the effect of secondary structure assignment algorithm (DSSP,³⁸ DEFINE,³⁹ and STRIDE⁴⁰) and influence of redundancy in the multiple alignments.

METHODS

The Problem of Objectively Testing Secondary Structure Prediction Methods

If a protein sequence shows clear similarity to a protein of known three dimensional structure, then the most accurate method of predicting the secondary structure is to align the sequences by standard dynamic programming algorithms,⁴¹ as homology modeling is much more accurate than secondary structure prediction for high levels of sequence identity. Secondary structure prediction methods are of most use when sequence similarity to a protein of known structure is undetectable. Accordingly, it is important that there is no detectable sequence similarity between sequences used to train and test secondary structure prediction methods.

Most secondary structure prediction methods include a set of parameters that must be estimated. Values for the parameters are obtained by statistical analysis or learning from a set of proteins for which the tertiary structure is known. This is the *training set* of proteins. Testing predictive accuracy on the training set leads to unrealistically high accuracies. An objective test of a secondary structure prediction method will predict the structures of a *test set* of proteins that are not in the training set and show no detectable sequence similarity with the training set. If the test is to be balanced, then both training and test sets should have a similar distribution of secondary structure classes and types.

Since the number of proteins of known structure is limited, it is normal to develop secondary structure prediction methods by cross-validation techniques, or jack-knife. In a full jack-knife test of N proteins, one protein is removed from the set, the parameters are developed on the remaining $N - 1$ proteins, then the structure of the removed protein is predicted and its accuracy measured. This process is repeated N times by removing each protein in turn. Since some training techniques are very time consuming, a more limited cross-validation is often performed. The set of proteins might be split into M equally balanced subsets rather than N . Parameters are developed on $(M - 1)N/M$ proteins, then tested on the remaining N/M proteins. This process is repeated M times, once for each subset. As described, the jack-knife process may also be referred to as a leave-one-out technique, although the two terminologies have become somewhat synonymous.

Cross-validation appears to remove the problem of a limited data set for training and test. However, artificially high accuracies can be obtained for some methods if the set of proteins used in the cross-validation show sequence similarity to each other. Accordingly, cross-validation sets must be pruned stringently to remove internal sequence similarities, or if this is not possible, then a completely independent test set must be used.

Selection of suitable test and training sets rests with the definition of “undetectable” sequence similarity. Appropriate measures of sequence similarity are discussed in the following section.

There are now available ≈ 500 sequence dissimilar proteins of known three-dimensional structure, suitable for developing and testing secondary structure prediction techniques. However, many of the current generation of secondary structure prediction methods were developed on a set of 126 protein chains proposed by Rost and Sander²⁷ (referred to here as RS126). In this paper we develop a new, non-redundant set of 396 protein domains (the CB396 set) that does not include proteins from the RS126 set.

Training and Test Sets of Protein Structures

Rost and Sander (1993) selected 126 proteins with which to train and test secondary structure prediction algorithms.²⁷ They defined non-redundancy to mean that no two proteins in the set share more than 25% sequence identity over a length of more than 80 residues. Unfortunately, as shown below, the RS126 set contains pairs of

TABLE I. Pairs in the RS126 Set That Have an SD Score of Greater Than Five[†]

(1)	(2)	SD score	Fold (1)	Fold (2)
1eca ⁷⁵	2lhb ⁵²	5.12	Globin-like	Globin-like
1azu ⁷⁶	2pcy ⁵³	5.40	Cupredoxins	Cupredoxins
2rspa ⁷⁷	5hvp ⁵⁶	5.81	Acid proteases	Acid proteases
1paz ⁷⁸	2pcy ⁵³	7.22	Cupredoxins	Cupredoxins
2lhb ⁵²	4sdha ⁷⁹	7.70	Globin-like	Globin-like
1cdta ⁸⁰	3ebx ⁵⁴	8.26	Snake toxin-like	Snake toxin-like
3cln ⁸¹	4cpv ⁵⁵	8.27	EF Hand-like	EF Hand-like
2gbp ⁸²	8abp ⁵⁷	8.86	Periplasmic binding	Periplasmic binding
1ovoa ⁸³	1tgsi ⁵¹	9.45	Ovomucoid/PCI-1 like	Ovomucoid/PCI-1 like
1fdlh ⁸⁴	1mcp ⁵⁰	12.66	Immunoglobulin	Immunoglobulin
4cms ⁴⁸	5er2e ⁴⁹	15.98	Acid proteases	Acid proteases

[†]Alignments were generated by the AMPS package⁶ a blosum62 matrix, and gap penalty of 10, with 100 randomizations. Fold definitions come from the current release (1.37) of the SCOP database.⁴⁷

proteins that are clearly sequence similar when compared by more sophisticated methods than percentage identity.

Percentage identity has long been known to be a poor measure of sequence similarity, particularly for values below 30%. Percentage identity is dependent upon both the length of the alignment⁴² and the composition of the sequences. Thus, two sequences of similar unusual amino acid composition may give high values of percentage identity, even when unrelated.

Recently, deficiencies in percentage identity have been quantified by Brenner et al. (1998) when scoring protein sequence database searches.² Even with the length correction suggested by Sander and Schneider (1991),⁴² percentage identity was significantly worse than measures that consider conservative substitutions as well as identities, and attempt corrections for length and composition.

Fortunately, techniques exist that overcome the deficiencies of percentage identity or other simple measures of sequence similarity. A long established method^{6,43} to measure the similarity between two protein sequences *A* and *B* is first to align the proteins by a standard dynamic programming algorithm (e.g. Needleman and Wunsch (1970)⁴⁴) and obtain the score for the alignment *V*. The order of amino acids in each protein sequence is then randomized and a dynamic programming alignment of the randomized sequences performed. This process is repeated typically 100 or more times and the mean $\bar{x}$ and standard deviation σ of the scores for comparison of the randomised sequences is calculated. The SD score, or *Z* score for comparison of the native sequences is given by: $(V - \bar{x})/\sigma$. Unlike the percentage identity, SD score corrects for bias due to the length and composition of the sequences. Accordingly, we use SD scores to derive our non-redundant test set of protein sequences.

PHD,²⁷ NNSSP,³⁰ DSC,²⁹ and PREDATOR²⁸ have been trained on the Rost and Sander set of 126 proteins. The release versions of PREDATOR and NNSSP available for this analysis were trained on larger sets, that included the 126 proteins. In principle, this should give PREDATOR and NNSSP an advantage over PHD.

The sequences in the test set developed here came from the 3Dee⁴⁵ database of structural domain definitions. In

3Dee, a non-redundant sequence set was created by the use of a sensitive sequence comparison algorithm and cluster analysis, rather than a simple percentage identity cutoff. This provided a set of 1,233 domains where no pair shared obvious sequence similarity. The new test set was derived from these domains by first removing multi-segment domains, to reduce the set size from 1,233 to 988 sequences. The sequences were then filtered only to permit X-ray crystal structures with resolutions of ≤ 2.5 Angstroms. This left a representative set of 554 domain sequences, referred to as CB554.

To ensure that the CB554 domain set had no sequence similarity to the RS126 set, the two sets were combined and all pairs of sequences compared by AMPS⁶ with a blosum62 matrix, and gap penalty of 10. Alignments with an SD score of ≥ 5 were regarded as sequence similar.^{6,46} According to this stringent definition of similarity, there were 11 sequence-similar pairs within the RS126 protein set, 119 pairs between CB554 domain set and RS126, and 21 pairs within CB554. Thus, there were 140 sequences in CB554 that matched either a sequence in CB554, or in the RS126 protein set. Of the 140, 3 sequences matched more than once, leaving 137 unique sequences. The 137 sequences were removed from CB554, leaving 417 sequences that were not sequence similar either to any sequence within the set of 417 sequences, or the RS126 sequence set. Of the 417 domain sequences remaining, 21 that did not have "full DSSP definitions," (i.e., those with more than 9 consecutive residues with incomplete backbones for which DSSP³⁸ does not define a state), were also removed, leaving a test set of 396 proteins (CB396).

The process of deriving CB396 showed up homologies in the RS126 set, with 11 proteins showing sequence similarity to at least one other within the RS126 set. These pairs are summarized in Table I. Table I shows each pair to have the same fold according to SCOP.⁴⁷ For example, 4cms⁴⁸ and 5er2e⁴⁹ are present in the RS126 set, yet have an SD score of 15.9. Both proteins are acid proteases with an all β , closed barrel structure.

Although not applied in this paper, three further non-redundant datasets suitable for cross-validated training and testing of secondary structure prediction methods

were generated. The CB396 and RS126 sequence sets were combined. One of each of the 11 pairs that had an SD score of ≥ 5 were removed from the RS126 set. Since 2pcy and 1lhb matched more than one protein in this subset, this left 9 unique homologues (1mcpl,⁵⁰ 1tgsi,⁵¹ 2lhb,⁵² 2pcy,⁵³ 3ebx,⁵⁴ 4cms,⁴⁸ 4cpv,⁵⁵ 5hvpa,⁵⁶ 8abp⁵⁷) that were removed from RS126. This set added to CB396 gave CB513. Protein chains of ≤ 30 residues often do not have well defined secondary structure. The CB497 set was constructed by removing the 16 domains from CB513 of ≤ 30 residues.

The 5SD cutoff used to derive the sets CB396, CB497 and CB513 is more stringent than scores used in previous studies of secondary structure prediction. However, although the SD score is a good measure of pairwise sequence similarity, it still will not identify all known homologs within the data set. In the SCOP⁴⁷ classification of protein structure superfamilies are defined from careful analysis of structure, evolution, and function. The SCOP superfamilies contain protein domains that have the same fold and are likely to have evolved from a common ancestor. Accordingly, we derived a further dataset from an analysis of all domains in SCOP_1.37. We took a representative domain from each superfamily, screened out multi-segment domains, NMR structures and those with a resolution ≥ 2.5 Angstroms to give the CB251 dataset.

All datasets, including secondary structure definitions and automatically generated multiple sequence alignments will be distributed via <http://barton.ebi.ac.uk/>.

Generating the Multiple Sequence Alignments

With the exception of PREDATOR^{28,58} all methods considered here, required a multiple sequence alignment as input, where as PREDATOR only required the multiple sequences in an unaligned format. In order to simplify the generation of multiple sequence alignments for large numbers of proteins, in this study we developed an automatic procedure.

We first perform a BLAST⁵⁹ database search of the OWL v29.4 database, which contains 198,742 entries.⁶⁰ The BLAST output is then screened by SCANPS, an implementation of the Smith Waterman dynamic programming algorithm,^{61,62} with length dependent statistics. Sequences are rejected if their SCANPS probability score is higher than 1×10^{-4} . Sequences are also rejected if they do not fit a length cutoff of 1.5. For example, if the query sequence is 90 residues long, the sequence length would have to range between 60 and 135 residues to be included. If sequences exceed the length criterion, they are truncated by removing end residues until the length of the sequence satisfies the cut off value. Sequences falling short of the lower length limit are discarded. The value of 1.5 for the length cutoff was reached by visual inspection of a number of multiple sequence alignments, produced with different cut-off values. The method removes both ridiculously long, short and unrelated sequences. However it does allow sequences that are longer than the query, and are related, to be included after truncation. The sequence similar proteins selected by this method, are then aligned by CLUSTALW (version 1.7),⁶³ with default parameters.

The multiple sequence alignments are modified so that they do not contain gaps in the first or "query" sequence, since with the current algorithms, gaps in the first sequence tend to reduce the accuracy of the prediction, or cause the program to fail to execute (NNSSP³⁰). A slightly different method is used for PHD,⁶⁴ whereby only gaps at the end of the target sequence are removed. Without this modification, the conversion of MSF to HSSP file format fails, as a correct insertion table is not constructed.

The reference secondary structure for each domain was defined by DSSP,³⁸ STRIDE,⁴⁰ and DEFINE.³⁹ All definitions were reduced to 3 state models, as follows:

1. DSSP: H and G to H, E and B to E, all other states to C.
2. STRIDE: H and G to H, E and b to E, all other states to C.
3. DEFINE: H and G to H, E to E, all other states to C.

Where H is α -helix, G is 3_{10} -helix, B and b are isolated β -bridge and E is β -strand.

The effect of alternative reduction methods for the DSSP algorithm is discussed in the results section.

Prediction Methods Analyzed

Six different secondary structure prediction methods were run on the alignments, each is briefly described here.

PHD⁶⁴ is a 3-level artificial neural network. The different levels consist of a sequence to secondary structure network, with a window of 13 amino acids, a structure to structure network, with a window of 17 amino acids, and finally an arithmetic average over a number of independently trained networks. The structure to structure network, improves prediction of the final length distributions of secondary structures. The arithmetic average has the effect of smoothing random noise that is seen in all artificial neural networks. The method also applies balanced training, percentage amino acid composition and conservation, sequence length, and insertions and deletions (indels) to enhance prediction accuracy.

DSC²⁹ applies GOR¹⁸ residue attributes, with the addition of hydrophobicity and amino acid position, which are combined with information from the multiple sequence alignment (conservation and indels). Optimal weights are deduced by linear discrimination, with filtering applied to remove erroneous predictions. This method has an advantage in that the prediction method is both implicit and effective.

NNSSP³⁰ is a scored nearest-neighbor method. It is based upon the environmental scoring scheme proposed by Bowie.⁶⁶ The NNSSP method extends the Bowie method by considering N and C terminal positions of α -helices and β -strands. The size of the database used for scanning is also altered to reflect similarity to the query sequence, reducing computation time, and improving the final accuracy.

PREDATOR²⁸ is slightly different to other methods discussed here, in that it uses an internal pairwise alignment method, rather than reading a global multiple sequence alignment. The SIM software⁶⁷ is applied to pro-

duce local alignments between sequence pairs. The original PREDATOR algorithm⁵⁸ is then used to predict the secondary structure segments. This algorithm also includes propensities for hydrogen bonding characteristics of β -sheets. Seven different secondary structure propensities are generated for the query sequence, with a nearest-neighbor implementation applied to calculate propensities for α -helix, β -strand and coil.

ZPRED²⁶ is also based on the GOR¹⁸ method, but with the addition of weights from calculated conservation values. The conservation value is calculated from amino acid properties as proposed by Taylor.⁶⁸ The ZPRED method improved the accuracy of the GOR method by noting that insertions and high sequence variability tend to occur in loop regions.

MULPRED (Barton, unpublished) is a combination of single sequence methods that are combined to give a prediction profile, from which a consensus is taken. The methods within MULPRED are Lim,¹⁶ GOR,¹⁸ Chou-Fasman,¹⁷ Rose⁶⁹ and Wilmot and Thornton⁷⁰ turn prediction methods.

Consensus Prediction Method

The observed Q_3 accuracy of ZPRED and MULPRED was between 3 and 8% lower than the other methods, so a consensus was calculated only from DSC, PHD, PREDATOR, and NNSSP. The standard consensus was calculated by examining the prediction for each method, at each position and taking the most popular state. For example if a residue had the following predictions:

NNSSP = Helix, PREDATOR = Helix,

DSC = Strand, PHD = Helix

the consensus prediction would be Helix. If there was no consensus for a particular residue, the result from the PHD method was used. More complex methods for combining the different predictions were investigated and tested, as discussed in the results section.

Assessment of Accuracy

Two methods were applied to assess the accuracy of the predictions. Average Q_3 ,¹⁹ and Segment Overlap.⁶⁴

Q_3 is a measure of the overall percentage of predicted residues, to observed:

$$Q_3 = \sum_{(i=H,E,C)} \frac{\text{predicted}_i}{\text{observed}_i} \times 100. \quad (1)$$

Segment overlap calculation⁶⁴ was performed for each data set. Segment overlap values attempt to capture segment prediction, and vary from an ignorance level of 37% (random protein pairs) to an average 90% level for homologous protein pairs. Segment overlap is calculated by:

$$Sov = \frac{1}{N} \sum_s \frac{\minov(s_{obs}; s_{pred}) + \delta}{\maxov(s_{obs}; s_{pred})} \times \text{len}(s_1). \quad (2)$$

TABLE II. Percentages of Secondary Structural State Per Secondary Structure Definition Method

	DSSP ³⁸	STRIDE ⁴⁰	DEFINE ³⁹
Helix	28.9	29.8	30.2
Sheet	22.9	24.4	30.0
Coil	48.1	45.8	39.7

Where N is the total number of residues, $\minov$ is the actual overlap, with $\maxov$ is the extent of the segment. δ is the accepted variation which assures a ratio of 1.0 where there are only minor deviations at the ends of segments.⁶⁴

RESULTS AND DISCUSSION

Comparison of Secondary Structure Definition Methods

All secondary structure prediction methods are trained and tested on secondary structure definitions from known structures. Defining secondary structure from coordinates is an inexact process due to differences in the concept of what is a secondary structure, as well as errors and inconsistencies in the experimental structure. This was illustrated in the comparison by Colloc'h et al. of DSSP,³⁸ DEFINE,³⁹ and P-curve⁷¹ on a non-redundant set of 154 proteins where all three methods agree at only 63% of positions.

Here we show the differences between DSSP,³⁸ DEFINE³⁹ and STRIDE⁴⁰ definitions on the RS126 protein set. When compared pairwise, DSSP and STRIDE agree to 95%, whereas DSSP and DEFINE agree at 73%, with STRIDE and DEFINE agreeing at 74%. All three methods agree at only 71% of positions.

Table II shows that DEFINE defines more sheet. 7.1% relative to DSSP and 5.6% relative to STRIDE. Helix is also defined more often by DEFINE. As a consequence of these two factors, DSSP and STRIDE define more coil than DEFINE, at 8.4% and 6.1% respectively.

The length of secondary structure elements as defined by DSSP, STRIDE, and DEFINE is summarized in Table III and Figure 1. DEFINE does not define sheet regions of less than 4 residues. The mean segment length values for DEFINE are higher than those of STRIDE and DSSP for all secondary structure states. Figure 1 shows DSSP to have a peak in the helix distribution at 4 residues. However, this is not found with the STRIDE or DEFINE definitions. With the exception of the peak at 4, the overall shape of DSSP and STRIDE length distributions are similar.

When assessing prediction methods, the average Q_3 was calculated for all the definition methods, for all runs, but because DEFINE is so dissimilar to DSSP and STRIDE, all results from DEFINE have been omitted from discussion.

Analysis of the Test and Training Alignments

Table IV summarizes an analysis of the automatic multiple sequence alignments that were generated for the RS126 and CB396 sets. Both sets have a similar average length of sequence, and average percentage identity within

TABLE III. Ranges of Length for Secondary Structural Elements as Defined by DSSP³⁸ STRIDE⁴⁰, and DEFINE³⁹ for the RS126 Set

State	Method	Min	Mean	Max	Total number of secondary structures
Helix	DSSP ³⁸	3	9	54	817
	STRIDE ⁴⁰	2	10	51	753
	DEFINE ³⁹	5	14	65	553
Strand	DSSP ³⁸	1	4	19	1302
	STRIDE ⁴⁰	1	4	19	1303
	DEFINE ³⁹	4	6	26	1030

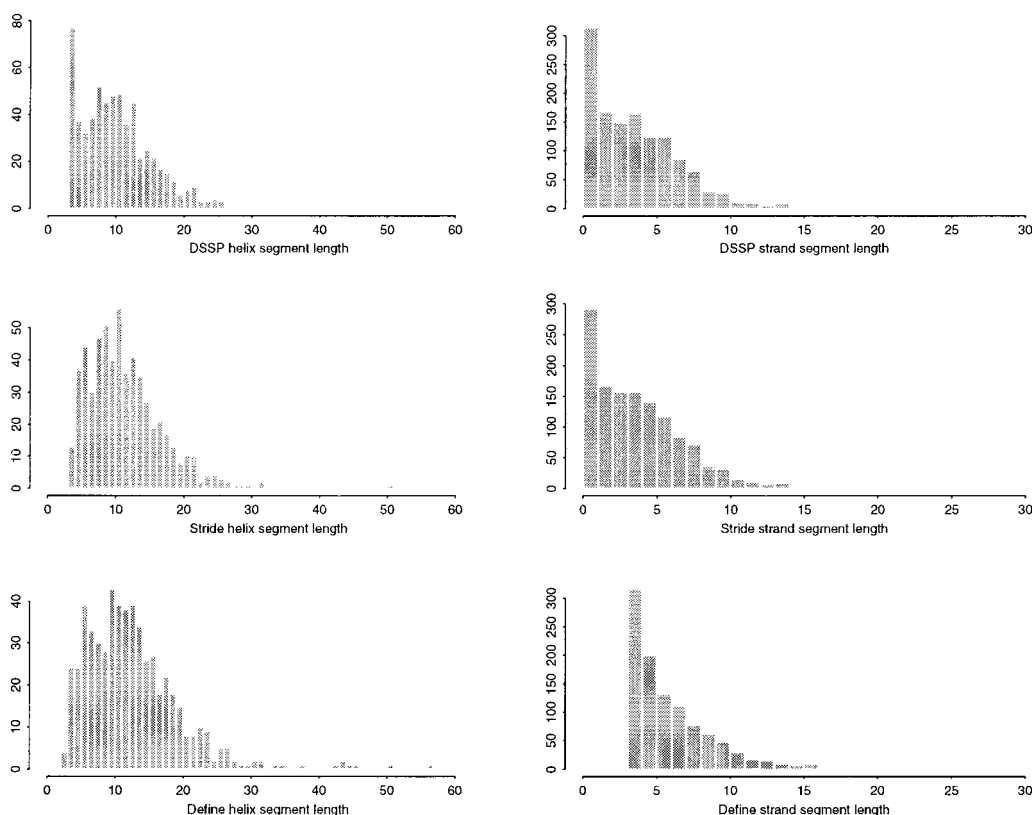


Fig. 1. Comparison of segment length distributions for each definition method.

TABLE IV. Summary Statistics of the Alignments Used in the Predictions

	Average percent identity between sequences	Average sequence length	Average number of sequences per alignment
CB396 set	34	157 residues	18
RS126 set	31	185 residues	30

the set. However, there is a significant difference between the average number of sequences per alignment between the two sets, even though both sets of alignments were generated using the same method. The older RS126 protein set has significantly (1.6 times) more sequence-similar

proteins in each alignment. The distribution of the number of sequences in the RS126 protein set was not biased by one or two large families.

A comparison between the CB396 set and the RS126 set showed the same distribution. The difference is therefore that each sequence family in the RS126 set is on average larger than any found in the CB396 set. This observation may simply reflect the fact that RS126 was derived from protein families whose first known members were characterized longer ago.

To verify that there was no bias to a particular structural class, the SCOP⁴⁷ classifications were examined for the proteins within the two sets as shown in Table V. There is a higher proportion of small proteins in the RS126 protein set (14% against 7%), while the CB396 protein set

TABLE V. Data for Class Types Used for the Predictions

Class definition	RS126 set number (%)	CB396 set number (%)
Alpha and beta (a/b)	25 (20)	107 (27)
Alpha and beta (a + b)	17 (13)	101 (26)
All alpha	27 (21)	68 (17)
All beta	38 (20)	78 (20)
Multi-domain	3 (2)	0 (0)
Small proteins	18 (14)	27 (7)
Membrane	1 (≤ 1)	3 (≤ 1)
Peptides	1 (≤ 1)	12 (3)

TABLE VI. Q_3 and Segment Overlap Results for the Set of RS126, and CB396 Proteins

Method	RS126 protein set		CB396 protein set	
	Q_3	SOV	Q_3	SOV
PHD ⁶⁴	73.5	73.5	71.9	75.3
DSC ²⁹	71.1	71.6	68.4	72.0
PREDATOR ²⁸	70.3	69.9	68.6	69.8
NNSSP ³⁰	72.7	70.6	71.4	71.3
CONSENSUS	74.8	74.5	72.9	75.4

has a higher proportion of $\alpha + \beta$ proteins (26% against 13%). However, the overall composition within each of the two sets is balanced.

Alignment Quality

The RS126 set of proteins was used to check that the alignments generated by our method could reproduce previous results for PHD,²⁷ and DSC.²⁹ PHD²⁷ (run in cross-validation mode) increased in average Q_3 accuracy by 1.9% from 71.6 to 73.5%, over the published accuracy. The accuracy obtained here for DSC improved by 1.0%. However, the published value of 70.1% for DSC was for a full jack-knife test, whereas our test utilized all the data. These results confirm that the automatic method used to build multiple alignments in this study is appropriate for secondary structure prediction.

Comparison of Prediction Accuracy for the CB396 Test and RS126 Training Sets

Table VI shows the differences between the RS126 and CB396 set of proteins. For all methods, the average accuracy drops by between 1.3% (NNSSP) and 2.7% (DSC) for the CB396 protein set. The NNSSP and PREDATOR programs used in this analysis were trained on larger numbers of proteins than RS126, and so should be less degraded by evaluation on a different test set. However, these methods still show a decrease in accuracy with the CB396 set.

Table VI illustrates that the percentages for the segment overlap correlate well with the Q_3 values. The SOV score for the consensus method is somewhat higher (74.5%) than the previous published value for PHD of 72%.⁶⁴ Table VI also shows that the PHD method does exceedingly well at

TABLE VII. Family Size for the Automatically Generated Alignments for the RS126 Protein Set, Considering 2 Levels of BLAST⁵⁹ p -Value Cutoff

	p -Value cutoff 10^{-10}	p -Value cutoff 10^{-2}
Total number of residues	1716356	2013632
Total number of sequences	7013	8974
Average number of sequences per family	55.6	71.2

predicting segments. As measured by the SOV method, it is on average 2% better than any of the other methods tested. The difference between the consensus and PHD segment overlap scores is smaller than the corresponding Q_3 value. This may be due to segment overlap being a more sensitive method to assess secondary structure predictions, or that the PHD method is the only one that has been optimized to predict segments scored by Equation 2.

Table VI summarizes the differences between SOV and Q_3 accuracies for each method on the RS126 and CB396 sets. Although Q_3 shows a consistent reduction on moving from RS126 to CB396, SOV shows a general improvement. Of the individual methods, NNSSP increases in accuracy the most (0.7%) while the consensus method increases by 0.9% to 75.4% SOV accuracy.

Effect on Q_3 of Changing the Number of Related Sequences

Prediction methods that use multiple sequence alignments gain accuracy over single-sequence methods by exploiting the patterns of residue conservation that are seen in protein families. Inclusion of more distantly related sequences in the alignment should improve the clarity of such patterns, but in an automated alignment building procedure, the risk is that unrelated protein sequences will pollute the alignment. Here, we investigated the effect of using a more permissive BLAST p -value cutoff⁵⁹ in the first phase of our alignment building procedure. The cutoff was lowered from 1×10^{-10} to 1×10^{-2} while leaving thresholds for SCANPS alone.

Table VII shows that the change in p -value cutoff increased the total number of residues after filtering with SCANPS by 297,276, and the total number of sequences by 1,961. This gives an increase in the average number of sequences per alignment of 15. Table VIII shows that increasing the number of sequences improves Q_3 by approximately 1% for all methods.

Table VIII also shows the marked difference between the prediction methods. The older methods, ZPRED and MULPRED were between 3 and 8 percent worse than the newer methods.

Effect on Q_3 of Reducing Redundancy in Multiple Alignments

While all sequences that are not 100% identical in a multiple sequence alignment will contribute to the prediction, the most informative sequences are those with the

TABLE VIII. Comparison of the Q_3 Accuracy for a Decrease in the BLAST⁵⁹ p -Value Cut-Off From 10^{-10} to 10^{-2} with the RS126 Set[†]

Method	p -Value cutoff 10^{-10}		p -Value cutoff 10^{-2}	
	DSSP	STRIDE	DSSP	STRIDE
PHD ⁶⁴	72.4	72.4	73.2	73.2
DSC ²⁹	70.2	70.0	71.0	70.7
PREDATOR ⁵⁸	69.7	69.3	70.7	70.3
NNSP ³⁰	71.8	71.2	72.4	71.7
MULPRED	66.7	65.4	67.2	66.8
ZPRED ²⁶	65.5	64.7	66.7	65.9
CONSENSUS	73.9	73.7	74.5	74.3

[†]The alignments used for these predictions did not use a percentage identity filter.

TABLE IX. Effect on Q_3 Accuracy by Removing all Sequences Similar to the Query at Different Percent Identity Thresholds, Using the RS126 Protein Set

	100%	95%	80%	75%	60%
PHD ⁶⁴	73.2	73.3	73.4	73.5	73.3
DSC ²⁹	71.0	71.0	71.0	71.1	70.9
PREDATOR ²⁸	70.7	70.4	70.1	70.3	70.4
NNSP ³⁰	72.4	72.5	72.7	72.7	72.8
CONSENSUS	74.5	74.5	74.6	74.8	74.7
Total number of sequences	8974	5907	4320	3833	2681
Ave number of sequences/alignment	71	47	34	30	20

greatest variation from the query. Here we test the effect of systematically removing sequences from the alignment that were similar to the query sequence at better than 95, 80, 75, and 60% identity. Table IX summarizes the effect of these thresholds. The average Q_3 accuracy improves slightly as the percentage identity threshold is reduced. The consensus method improves by 0.3% at the 75% level. Since the predictions do not get any worse by removing redundant sequences and prediction methods run faster with fewer sequences, the 75% cutoff was used for all predictions, other than those shown in Table VIII.

The average Q_3 for each prediction when compared to DEFINE secondary structure definitions was between 3 and 8% worse than for DSSP and STRIDE definitions. As none of the prediction methods examined here were trained on DEFINE definitions, we do not consider comparison of predictions to DEFINE definitions any further.

The Effect on Accuracy of Alternative 8- to 3-State Reductions

DSSP³⁸ provides an 8-state assignment of secondary structure denoted by single letter codes: H (α -helix), T (β -turn), S (bend), I (π helix), G (3_{10} -helix), E (β -strand), B (β -bridge) and C (not HTSIGE or B). However, prediction methods are normally trained and assessed for only 3 states (H,C,E), so the 8 states must be reduced to 3. Here we consider the effect on accuracy of applying three

TABLE X. Mean Percentages of Secondary Structure State Defined by DSSP³⁸ When Different 8- to 3-State Reduction Methods are Used

	Mean percent of helix	Mean percent of sheet	Mean percent of coil
Method A	28.9	22.9	48.1
Method B	25.3	21.2	52.6
Method C	25.6	21.2	52.3

TABLE XI. Changing 8- to 3-State Reduction, for DSSP and Resultant Q_3 Accuracy for the Consensus Method, Based on the RS126 Set of Proteins

Change	Q_3
Reduction method A ⁶⁴	74.8
B $\rightarrow$ Coil only	75.7
G $\rightarrow$ Coil only	76.6
B and G $\rightarrow$ Coil	77.5
GGGHHHH $\rightarrow$ HHHHHHH, B and G $\rightarrow$ Coil (Method C) ³⁰	77.5
B, G and HHHH EE $\rightarrow$ Coil (Method B) ⁵⁸	77.9

different published 8- to 3-state reduction methods when testing.

Method A: E and B to E, G and H to H. Rest to coil.²⁷

Method B: E as E, H as H. Rest to coil including EE and HHHH.²⁸

Method C: GGGHHHH redefined as HHHHHHH, then B, and GGG to coil, with H to H and E to E.³⁰

Method **A** treats both isolated β -bridges and residues that are part of a β -sheet as "extended." 3_{10} -helix and α -helix are treated as "helix," and all other states treated as "coil." Method **B** translates more secondary structures into coil. This includes all 3_{10} -helix, α -helix that is a single turn, isolated β -bridges and β -strands that are only two residues long. The rationale for this reduction is that short secondary structures are normally of marginal stability and also variable within protein families. Table X shows reduction Methods **B** and **C** on average to contain 4% more coil, 3% less helix and 2% less strand than Method **A**.

Table XI summarizes the difference in Q_3 accuracy obtained by varying the reduction of 8-state DSSP definition to 3 states. The original value for the consensus method was 74.8% (Method **A**). This increases to 77.9% for Method **B**. Table XI shows that reducing both single strand "B" to coil and 3_{10} -helix "G" to coil, gives the greatest stepwise increase (2.7%).

Table XII shows that *all* prediction methods appear to improve in accuracy with comparison to Method **A**, when one uses Method **B**, as the method for the 8-state reduction. The apparent improvement is between 2.2 and 4.9%. The PREDATOR²⁸ method improves by the largest extent (4.9%), as a reflection of PREDATOR being trained with

TABLE XII. Results for the RS126 Protein Set, by Reducing the Definition to 3-State by Methods A and B

Method	DSSP (A)	STRIDE (A)	DSSP (B)	STRIDE (B)
PHD ⁶⁴	73.5	73.5	76.3	76.3
DSC ²⁹	71.1	70.9	73.3	73.4
PREDATOR ²⁸	70.3	69.6	75.2	74.0
NNSSP ³⁰	72.7	72.2	77.3	76.5
CONSENSUS	74.8	74.7	77.9	77.9

TABLE XIII. Results for Single Sequence Prediction Methods Via a Full Jack-Knife Test†

Method	RS126	CB396	Author
PHD ⁶⁴	76.3	74.2	—
SIMPA ⁷²	67.3	67.6	67.7
GOR IV ⁷⁴	53.3	64.6	64.4
SOPM ⁷³	66.8	64.7	69.0

†The column “Author” is the authors jack-knife value for the method with their dataset, and definition reduction method. All results are calculated using reduction method A, and also converting G and B states to coil. For PHD⁶⁴ the authors quote 71.6% as their cross-validated accuracy. However, G and B states were considered in the accuracy calculation for PHD⁶⁴.

Method **B** as the reduction scheme. STRIDE⁴⁰ values have also been included in Table XII for comparison. Prior to this study, PHD⁶⁴ and DSC²⁹ used the same reduction method (Method **A**), but NNSSP³⁰ and PREDATOR^{28,58} used methods **B** and **C** respectively. Unless the same reduction method is used, an objective comparison can not be made.

Single Sequence Prediction Methods

The prediction methods we have carefully selected for this work represent current state-of-the-art prediction methods, that use multiple sequences. However for completeness, the SIMPA,⁷² SOPM,⁷³ and GORIV,⁷⁴ single sequence methods were also examined. These methods do not have precalculated propensity tables, and as such we could perform a full jack-knife test with the new datasets. We only compare the results for the single sequence methods to those obtained for the PHD algorithm, as PHD was the only other method for which we were able to carry out cross validation. The SIMPA,⁷² SOPM⁷³ and GORIV⁷⁴ methods have quoted accuracies based on removing helices shorter than 4 residues and strands less than 2 residues. For testing these methods, we used method **A** as the reduction method and also converted G and B states to coil. If reduction method **A** alone is used, SIMPA,⁷² SOPM,⁷³ and GORIV⁷⁴ reduce in accuracy from those shown in Table XIII by 2–4%.

The difference between the single sequence methods we examined the PHD ranges between 23.3%, and 6.6% depending upon the method and database used. Table XIII shows the GORIV method to improve remarkably (11.3%)

with an increased database size 126 proteins to 396 proteins). This is to be expected as GORIV no longer uses “dummy frequencies”⁷⁴ instead relying on a large database to calculate its propensity tables. To examine if this feature scaled, we also applied the GORIV method to the CB513 dataset. The accuracy improved by 1.1% from 64.6% to 65.7%. SOPM only achieved 66.8% on the RS126 protein set, and 64.6% for the CB396 set. The authors of the SOPM method quoted 69%.⁷³ However, the database used in their study, was non-redundant at 50% sequence identity, and so included a number of clear homologues.

Improving the Consensus Prediction

In order to establish the upper limit of accuracy possible by combining the prediction methods, we took the most accurate prediction for each residue in the RS126 data set by PHD, DSC, NNSSP, or PREDATOR. This gave the theoretical best accuracy for a combination of these methods of $Q_3 = 78\%$.

We investigated a variety of techniques for combining the prediction methods, in an attempt to raise the average Q_3 on RS126 from 74.8% towards 78%. All possible combinations of methods were tried to calculate the consensus, but no combination of methods improved upon the average Q_3 of the consensus of DSC, PREDATOR, NNSSP, and PHD, with PHD taken if there was a tie. However, the next highest combination was only 0.3% worse at 74.2% and used NNSSP, PREDATOR, and DSC, predictions relying on PREDATOR's definition if there was no consensus. Experiments with filtering single-residue helix predictions and other unlikely secondary structures did not improve the overall Q_3 .

The reliability information from the PHD and PREDATOR predictions was also investigated. When a method predicted with a reliability of greater than 7, that prediction was taken. No further increase in average Q_3 accuracy could be achieved using this approach.

The predictions for each method were weighted by adding constants. All combinations of all values from 1 to 10 were applied to all predictions for each method. The consensus was then calculated in the same manner as before, but now using the weighted predictions. The optimal weighting scheme was 2,1,2,2 where PREDATOR was down weighted by one point. The Q_3 accuracy for this approach was no higher than that of the non-weighted majority-wins method.

An artificial neural network, with 9 hidden nodes was trained with the output from the NNSSP, PHD, DSC, and PREDATOR methods. A 17-residue window was used. The inputs were coded as binary, with 001, 010, and 100 representing the helix, strand, and coil states respectively. Seven-fold cross validation was performed. This yielded 73.2% for the 126 protein set. This result was still lower than the simple consensus approach. No further improvement in accuracy was seen by changing the free parameters of the network, for example, hidden nodes or number of training epochs. The target sequence was also included in the input layer, but this also proved unsuccessful. We suggest that better accuracies may be achieved if propensi-

ties for the different states are used, rather than the binary input, and this idea forms the basis of future work.

When the lower accuracy predictions from ZPRED and MULPRED were included, the overall accuracy of the consensus method was reduced. SIMPA, SOPM, and GORIV were not included at any stage in the consensus method. Further work aims to discover if the single sequence prediction methods can be incorporated into a more accurate consensus method.

SUMMARY AND CONCLUSIONS

In this study we have developed a new, nonredundant test set of 396 protein domains (CB396). The set does not include any of the 126 proteins with which many current methods have been trained, nor does it contain homologs of those 126 proteins as measured by a stringent test of sequence similarity. We have shown that by combining four secondary structure prediction methods DSC,²⁹ PHD,²⁷ PREDATOR,²⁸ and NNSSP³⁰ by a simple majority-wins method, the average three-state Q_3 prediction accuracy can be improved by 1% from 71.9% (PHD) to 72.9% on the CB396 set. A fair comparison of the accuracy of the constituent methods is only possible for PHD²⁷ and DSC²⁹ as all other algorithms included some of our test proteins in their training set. Despite this, PHD²⁷ still gave the highest accuracy on the new test set (71.9%) of any of the methods considered.

An automatic procedure for database searching to build a multiple sequence alignment has been developed. Alignments from this procedure give a 1.9% increase in the average accuracy of prediction compared to previous published results for the PHD algorithm on the 126 protein set.⁶⁴ The increase may be attributed to better alignments and the increased size of the current sequence databases.

In the literature there are different standards for reducing DSSP³⁸ 8-state (H,C,B,E,T,S,G,I) assignments to 3 states (H,C,E). It was found that changing the reduction method can alter the apparent prediction accuracy by over 3% on average. Although we were unable to train the methods using different 8- to 3-state reductions, testing all methods with different reduction methods showed that Method B⁵⁸ consistently gave higher accuracy. This may be attributed to Method B assigning more of the protein to Coil (C).

Secondary structure definition methods DSSP,³⁸ DEFINE,³⁹ and STRIDE⁴⁰ were compared. All three agree at only 75% of positions. This is mainly due to differences between DEFINE and DSSP/STRIDE. DSSP and STRIDE agree at 95% of positions, though DSSP defines many more 4 residue helices than STRIDE.

In summary, with the alignment method presented here, the method with the highest average accuracy on the new non-redundant test set of 396 proteins was PHD⁶⁴ with 71.9%. While the new combination of NNSSP,³⁰ PHD,²⁷ DSC,²⁹ and PREDATOR²⁸ presented here improves upon this figure by 1% to 72.9%.

The non-redundant datasets constructed during this analysis will facilitate the future development and testing

of secondary structure prediction methods. The datasets, alignments and definitions are available via <http://barton.ebi.ac.uk>.

ACKNOWLEDGEMENTS

We would like to thank the developers of the algorithms we have considered for their helpful discussions and cooperation in this study. In particular we thank Drs. B. Rost, D. Frishman V. Solovyev, R. King, and M. Zvelebil, for providing their software. We also thank Matt Finlay for coding many of the routines to parse prediction output, as well as Dr. A.S. Siddiqui for his automatic alignment building algorithm. We thank Prof. L.N. Johnson for her encouragement and support. JC is an Oxford Centre for Molecular Sciences/MRC student.

REFERENCES

1. Sali A. Modelling mutations and homologous proteins. *Curr Opin Biotechnol* 1995;6:437–451.
2. Brenner SE, Chothia C, Hubbard TJP. Assessing sequence comparison methods with reliable structurally identified distant evolutionary relationships. *Proc Natl Acad Sci* 1998;95:6073–6078.
3. Taylor WR. Identification of protein sequence homology by consensus template alignment. *J Mol Biol* 1986;188:233–258.
4. Gribskov M, McLachlan AD, Eisenberg D. Profile analysis: detection of distantly related proteins. *Proc Natl Acad Sci* 1987;84:4355–4358.
5. Barton GJ. Protein multiple sequence alignment and flexible pattern matching. *Methods Enzymol* 1990;183:403–428.
6. Barton GJ, Sternberg MJE. A strategy for the rapid multiple alignment of protein sequences: confidence levels from tertiary structure comparisons. *J Mol Biol* 1987;198:327–337.
7. Eddy SR. Hidden markov models. *Curr Opin Struct Biol* 1996;6:361–365.
8. Karplus K, Sjolander K, Barrett C, et al. Predicting protein structure using hidden markov models. *Proteins* 1997; Suppl 1:134–139.
9. Park J, Teichmann SA, Hubbard T, Chothia C. Intermediate sequences increase the detection of distant sequence homologues. *J Mol Biol* 1997;273:349–354.
10. Jones DT, Taylor WR, Thornton JM. A new approach to protein fold recognition. *Nature* 1992;358:86–89.
11. Lemer C, Rooman MJ, Wodak SJ. Protein structure prediction by threading methods: evaluation of current techniques. *Proteins* 1996;23:337–355.
12. Russell RB, Barton GJ. An SH2-SH3 domain hybrid. *Nature* 1993;364:765–765.
13. Russell RB, Copley RR, Barton GJ. Protein fold recognition by mapping predicted secondary structures. *J Mol Biol* 1996;259:349–365.
14. Rost B. TOPITS: threading one-dimensional predictions into three-dimensional structures. *Proceedings from the Third International Conference on Intelligent Systems for Molecular Biology*. Menlo Park, CA: AAAI Press, 1995. p 314–321.
15. Rost B. Protein fold recognition by prediction-based threading. *J Mol Biol* 1997;270:1–10.
16. Lim VI. Algorithms for prediction of α helices and β structural regions in globular proteins. *J Mol Biol* 1974;88:873–894.
17. Chou PY, Fasman GD. Conformational parameters for amino acids in helical, β -sheet, and random coil regions calculated from proteins. *Biochemistry* 1974;13:211–222.
18. Garnier J, Osguthorpe DJ, Robson B. Analysis and implications of simple methods for predicting the secondary structure of globular proteins. *J Mol Biol* 1978;120:97–120.
19. Schulz GE, Schirmer RH. *Principles of Proteins Structure*. New York: Springer-Verlag, 1979. p 1–314.
20. Kabsch W, Sander C. How good are predictions of protein secondary structure? *FEBS Lett* 1983;155:179–182.
21. Livingstone CD, Barton GJ. Identification of functional residues and secondary structure from protein multiple sequence alignment. *Methods Enzymol* 1996;266:497–512.

22. Crawford IP, Niermann T, Kirchner K. Prediction of secondary structure by evolutionary comparison: application to the alpha subunit of tryptophan synthase. *Proteins* 1987;2:118–129.
23. Barton GJ, Newman RH, Freemont PF, Crumpton MJ. Amino acid sequence analysis of the annexin super-gene family of proteins. *Eur J Biochem* 1991;198:749–760.
24. Russell RB, Breed J, Barton GJ. Conservation analysis and structure prediction of the SH2 family of phosphotyrosine binding domains. *FEBS Lett* 1992;304:15–20.
25. Benner SA, Gerloff D. Patterns of divergence in homologous proteins as indicators of secondary and tertiary structure: a prediction of the structure of the catalytic domain of protein kinases. *Adv Enzyme Regul* 1990;31:121–181.
26. Zvelebil MJJM, Barton GJ, Taylor WR, Sternberg MJE. Prediction of protein secondary structure and active sites using the alignment of homologous sequences. *J Mol Biol* 1987;195:957–961.
27. Rost B, Sander C. Prediction of protein secondary structure at better than 70% accuracy. *J Mol Biol* 1993;232:584–599.
28. Frishman D, Argos P. Seventy-five percent accuracy in protein secondary structure prediction. *Proteins* 1997;27:329–335.
29. King RD, Sternberg MJE. Identification and application of the concepts important for accurate and reliable protein secondary structure prediction. *Protein Sci* 1996;5:2298–2310.
30. Salamov AA, Solovyev VV. Prediction of protein secondary structure by combining nearest-neighbor algorithms and multiple sequence alignments. *J Mol Biol* 1995;247:11–15.
31. *Proteins* 1997; Suppl 1:1–230.
32. Rost B. Better 1D predictions by experts with machines. *Proteins* 1997; Suppl 1:192–197.
33. Biou V, Gilbrat JF, Robson B, Garnier J. Secondary structure prediction: combination of three different methods. *Protein Eng* 1995;2:185–191.
34. Zhang X, Mesriov D, Waltz J. A hybrid system for protein secondary structure prediction. *J Mol Biol* 1992;225:1049–1063.
35. Nishikawa K, Ooi T. Amino acid sequence homology applied to the prediction of protein secondary structures, and joint prediction with existing methods. *Biochim Biophys Acta* 1986;871:45–54.
36. Nishikawa K, Noguchi T. Predicting protein secondary structure based on amino acid sequence. *Methods Enzymol* 1995;202:31–44.
37. Geourjon C, Deleage G. Sopma: significant improvements in protein secondary structure prediction by consensus prediction from multiple alignments. *Comput Appl Biosci* 1995;11:681–684.
38. Kabsch W, Sander C. A dictionary of protein secondary structure. *Biopolymers* 1983;22:2577–2637.
39. Richards FM, Kundrot CE. Identification of structural motifs from protein coordinate data: secondary structure and first-level super-secondary structure. *Proteins* 1988;3:71–84.
40. Frishman D, Argos P. Knowledge-based protein secondary structure assignment. *Proteins* 1995;23:566–579.
41. Boscott PE, Barton GJ, Richards WG. Secondary structure prediction for modelling by homology. *Protein Eng* 1993;6:261–266.
42. Sander C, Schneider R. Database of homology-derived protein structures and the structural meaning of sequence alignment. *Proteins* 1991;9:56–68.
43. Feng DF, Johnson MS, Doolittle RF. Aligning amino acid sequences: comparison of commonly used methods. *J Mol Evol* 1985;21:112–125.
44. Needleman SB, Wunsch CD. A general method applicable to the search for similarities in the amino acid sequence of two proteins. *J Mol Biol* 1970;48:443–453.
45. Siddiqui A, Barton GJ. 3DDee—database of protein domain definitions. 1998; submitted.
46. Barton GJ, Sternberg MJ. Evaluation and improvements in the automatic alignment of protein sequences. *Protein Eng* 1987;1: 89–94.
47. Murzin A, Brenner SE, Hubbard T, Chothia C. Scop: a structural classification of proteins database and the investigation of sequences and structures. *J Mol Biol* 1995;247:536–540.
48. Newman M, Frazao C, Khan G, Tickle IJ, Blundell TL, Safo M, Andreeva N, Zdanov A. X-ray analyses of aspartic proteinases. Structure and refinement at 2.2 angstroms resolution of bovine chymosin. *J Mol Biol* 1991;221:1295–1309.
49. Sali A, Veerapandian B, Cooper JB, Foundling SI, Hoover DJ, Blundell TL. High resolution x-ray diffraction study of the complex between endothiapepsin and an oligopeptide inhibitor. The analysis of the inhibitor binding and description of the rigid body shift in the enzyme. *EMBO J* 1989;8:2179–2188.
50. Satow Y, Cohen GH, Padlan EA, Davies DR. Phosphocholine binding immunoglobulin study at 2.7 angstroms. *J Mol Biol* 1987;190:593–604.
51. Bolognesi M, Gatti G, Menegatti E, et al. Three dimensional structure of the complex between pancreatic secretory inhibitor (kazal type) and trypsinogen at 1.8 angstroms resolution. *J Mol Biol* 1982;162:839–868.
52. Honzatko RB, Hendrickson WA, Love WE. Refinement of a molecular model for lamprey hemoglobin from *perromyzon marinus*. *J Mol Biol* 1985;184:147–164.
53. Garrett TPJ, Guss JM, Freeman HC. The crystal structure of poplar apoplastocyanin at 1.8 Angstroms resolution. *J Biol Chem* 1984;259:2822–2825.
54. Smith JL, Corfields PWR, Hendrickson WA, Low BW. Refinement at 1.4 angstroms resolution of a model of erabutoxin B. Treatment of ordered solvent and discrete order. *Acta Crystallogr Sect A* 1988;44:357–362.
55. Kumar VD, Lee L, Edwards BFP. Refined crystal structure of calcium liganded carp parvalbumin 4.25 at 1.5 angstroms resolution. *Biochemistry* 1990;29:1404–1412.
56. Fitzgerald PMD, Mc Keever BM, Van Middlesworth JF, Springer JP. Crystallographic analysis of a complex between human immunodeficiency virus type 1 protease and acetyl pepstatin at 2.0 angstroms resolution. *J Biol Chem* 1990;265:14209–14219.
57. Quiocho FA, Wilson DK, Vyas NK. Substrate specificity and affinity of a protein modulated by bound water molecules. *Nature* 1989;340:404–407.
58. Frishman D, Argos P. Incorporation of non-local interactions in protein secondary structure prediction from the amino acid sequence. *Protein Eng* 1996;9:133–142.
59. Altschul SF, Gish W, Miller W, Myers EW, Lipman DJ. Basic local alignment search tool. *J Mol Biol* 1990;215:403–410.
60. Bleasby AJ, Akrigg D, Attwood TK. OWL—a non-redundant, composite protein sequence database. *Nucleic Acids Res* 1994;22: 3574–3577.
61. Smith TF, Waterman MS. Identification of common molecular subsequences. *J Mol Biol* 1981;147:195–197.
62. Barton GJ. Alscript: a tool to format multiple sequence alignments. *Protein Eng* 1993;6:37–40.
63. Thompson JD, Higgins DG, Gibson TJ. CLUSTAL W: improving the sensitivity of progressive multiple sequence alignment through sequence weighting, positions-specific gap penalties and weight matrix choice. *Nucleic Acids Res* 1994;22:4673–4680.
64. Rost BR, Sander C, Schneider R. Redefining the goals of protein secondary structure prediction. *J Mol Biol* 1994;235:13–26.
65. Reference deleted in proofs.
66. Bowie JU, Luthy R, Eisenberg D. A method to identify protein sequences that fold into a known three-dimensional structure. *Science* 1991;253:164–170.
67. Huang X, Miller WA. A time efficient, linear-space local similarity algorithm. *Adv Appl Math* 1991;12:337–357.
68. Taylor WR. Classification of amino acid conservation. *J Theor Biol* 1986;119:205–218.
69. Rose GD. Prediction of chain turns in globular proteins on a hydrophobic basis. *Nature* 1978;272:586–591.
70. Wilmot ACM, Thornton JM. Analysis and prediction of the different types of beta-turn in proteins. *J Mol Biol* 1988;203:221–232.
71. Sklenar H, Etchebest C, Lavery R. Describing protein structure—a general algorithm yielding complete helicoidal parameters and unique overall axis. *Proteins* 1989;6:46–60.
72. Levin JM. Exploring the limits of nearest neighbour secondary structure prediction. *Protein Eng* 1997;10:771–776.
73. Geourjon C, Deleage G. SOPM: a self optimised method for protein secondary structure prediction. *Protein Eng* 1994;7:157–164.
74. Robson B, Garnier J, Giblat J. GOR method for predicting protein secondary structure from amino acid sequence. *Methods Enzymol* 1996;266:540–553.
75. Steigemann W, Webber E. Structure of Erythrocrucorin in different ligand states refined at 1.4 Angstroms resolution. *J Mol Biol* 1979;127:309–338.

76. Adman ET, Sieker LC, Jensen LH. Structural features of Azurin at 2.7 Angstroms. *Isr J Chem* 1981;21:8–11.
77. Wlodawer A, Miller M, Jaskolski M. Crystal structure of a retroviral protease proves relationship to aspartic protease family. *Nature* 1989;337:576–579.
78. Petratos K, Dauter Z, Wilson KS. Refinement of the structure of Pseudoazurin from *Alcaligenes Faecalis* S-6 at 1.55 Angstroms. *Acta Crystallogr Sect B*. 1988;44:628–636.
79. Royer WE. Jr. High resolution crystallographic analysis of a cooperative dimeric Haemoglobin. *J Mol Biol* 1994;657:657–681.
80. Rees B, Bilwes A, Samama JP, Moras D. Cardiotoxin from *Naja Mossanmica*: the refined crystal structure. *J Mol Biol* 1990;214: 281–297.
81. Babu YS, Bugg CE, Cook WJ. Structure of Calmodulin refined at 2.2 Angstroms resolution. *J Mol Biol* 1988;204:191–204.
82. Vyas NK, Vyas MN, Quijcho FA. Sugar and signal transducer binding sites of the *Escherichia coli* galactose chemoreceptor protein. *Science* 1988;242:1290–1295.
83. Weber E, Papamokos E, Bode W, Huber R, Kato I, Laskowski M. Ovomucoid, a Kazal-type inhibitor, and model building studies of complexes with serine proteases. *J Mol Biol* 1982;158: 515–537.
84. Fischmann TO, Poljak RJ. Crystallographic refinement of the three-dimensional structure of FAB D1.2 Lysozyme complex at 2.5 Angstroms. *J Biol Chem* 1991;266:12915–12920.