

Towards the development of standardized methods for comparison, ranking and evaluation of structure alignments

Alex W. Slater^{1,2}, Javier I. Castellanos², Manfred J. Sippl³ and Francisco Melo^{1,2,*}

¹Molecular Bioinformatics Laboratory, Millennium Institute on Immunology and Immunotherapy, Portugal 49, Santiago, CP 8330025, ²Departamento de Genética Molecular y Microbiología, Facultad de Ciencias Biológicas, Pontificia Universidad Católica de Chile, Alameda 340, Santiago, Chile and ³Center of Applied Molecular Engineering, Division of Bioinformatics, Department of Molecular Biology, University of Salzburg, Hellbrunnerstrasse 34, 5020 Salzburg, Austria
Associate Editor: Anna Tramontano

ABSTRACT

Motivation: Pairwise alignment of protein structures is a fundamental task in structural bioinformatics. There are numerous computer programs in the public domain that produce alignments for a given pair of protein structures, but the results obtained by the various programs generally differ substantially. Hence, in the application of such programs the question arises which of the alignment programs are the most trustworthy in the sense of overall performance, and which programs provide the best result for a given pair of proteins. The major problem in comparing, evaluating and judging alignment results is that there is no clear notion of the optimality of an alignment. As a consequence, the numeric criteria and scores reported by the individual structure alignment programs are largely incomparable.

Results: Here we report on the development and application of a new approach for the evaluation of structure alignment results. The method uses the translation vector and rotation matrix to generate the superposition of two structures but discards the alignment reported by the individual programs. The optimal alignment is then generated in standardized form based on a suitably implemented dynamic programming algorithm where the length of the alignment is the single most informative parameter. We demonstrate that some of the most popular programs in protein structure research differ considerably in their overall performance. In particular, each of the programs investigated here produced in at least in one case the best and the worst alignment compared with all others. Hence, at the current state of development of structure comparison techniques, it is advisable to use several programs in parallel and to choose the optimal alignment in the way reported here.

Availability and implementation: The computer software that implements the method described here is freely available at <http://melolab.org/stovca>.

Contact: fmelo@bio.puc.cl

Received on July 27, 2012; revised on September 20, 2012; accepted on September 30, 2012

1 INTRODUCTION

An alignment of two protein sequences is completely defined by a set of pairs of amino acids, one from each protein, that are considered to be equivalent or related. On the other hand, an alignment of two protein structures can be considered as a

sequence alignment together with the geometric transformation that simultaneously superimposes the C α atoms of the equivalent amino acid pairs. In fact, for any sequence alignment, the associated optimal transformation is found by minimizing the root mean square (RMS) error of the C α distances between the equivalent pairs of residues. This unique optimal transformation can be computed efficiently (Kabsch, 1976; Kabsch, 1978; Sippl and Stegbuchner, 1991).

The remaining question is how to get the most suitable or optimal alignment. At the current state of affairs the question is unanswered. The reason is that in the construction of structure alignments, two conflicting goals need to be satisfied. One goal is to maximize the number of equivalent residues, i.e. the length of the alignment, the second is to minimize the associated RMS error. Obviously, one can lower the RMS error at the expense of shortening the alignment, or one can maximize the alignment length on the expense of increasing the RMS error. But it remains unclear what combination of alignment length and RMS value provides the most suitable result in any particular case.

Lacking a clear definition of what constitutes an optimal alignment, structural bioinformatics has created an ever-growing arsenal of computer programs (Hasegawa and Holm, 2009), all dedicated to solve one and the same problem, the pairwise superposition of two protein structures, where each program provides its own numerical criteria and scores to describe the properties of the alignments obtained. From a user's point of view, the situation is most confusing. Research on protein structure requires structure alignment tools, but it is unclear which programs produce the most valuable results in a given situation and how the results obtained have to be interpreted (Feng and Sippl, 1996; Sippl and Wiederstein, 2008).

Obviously some effort of standardization is required. In what follows, we describe an approach that can be used to compare structure alignments obtained from any program that reports the geometric transformation required to superimpose the two structures. In fact, the transformation, consisting of a translation vector and a rotation matrix, is the essential result of any structure alignment technique. Using the transformation, the structures can be superimposed and, from the superimposed structures, alignments can be constructed. Transformation and superposition on the one hand and alignment construction on the other can be executed independently. It is therefore possible to use the transformations reported by the individual programs

*To whom correspondence should be addressed.

to compute standardized structure alignments and the associated alignment length. As we demonstrate below, this standardized alignment length, which is optimal when obtained with a dynamic programming algorithm on some defined parameters, is a most convenient measure for the comparison and ranking of structure alignment results (Sippl, 2008; Sippl and Wiederstein, 2012).

2 METHODS

2.1 Benchmark dataset

A total of 215 protein structure pairs from HOMSTRAD database (Mizuguchi, *et al.*, 1998) were used in this work. These were obtained after applying the following filters to the 9538 pairs available at this database (September 2010 release): (i) individual chains must have a length in the range of 100–150 residues; (ii) the percentage of sequence identity must be $\leq 25\%$ and (iii) the structural overlap is $\leq 75\%$. The percentage sequence identity and the structural overlap were calculated from the optimal structural alignments extracted from the HOMSTRAD superpositions exactly as described in Section 2.3. The final list with the selected 215 protein structure pairs is available as supplementary data at: <http://melolab.org/sup-mat.html>.

2.2 Structure superposition methods

Seven protein structure superposition tools were tested and compared in this work: DALI (Holm, 2000), CE (Shindyalov and Bourne, 1998), MUSTANG (Konagurthu *et al.*, 2006), SALIGN (Madhusudhan *et al.*, 2009) as implemented in MODELLER (Sali and Blundell, 1993), TMALIGN (Zhang and Skolnick, 2005), MAMMOTH (Ortiz *et al.*, 2002) and TOPMATCH (Sippl and Wiederstein, 2008). All these tools were used with their default parameters. In the case of TOPMATCH, the possibility to generate the superpositions of structures with permuted sequences was turned off, and only the first superposition reported by this software was considered (i.e. TOPMATCH reports the 10 best structure superpositions found, ranked by alignment length).

2.3 Calculation of optimal structure alignments

A computer software, called STOVCA, was implemented to calculate the optimal structure alignment according to defined parameter values for a given structure superposition pair previously generated by any structure superposition software. The optimal structure alignment is the one with the largest number of equivalent residues between the two superposed structures for a fixed maximum RMS value (i.e. distance threshold to define the equivalent residue pairs). STOVCA does not alter the atomic coordinates of the input structure superposition pair and reports as an output the optimal structure alignment, along with some similarity measures calculated from it. The software relies on a Smith–Waterman dynamic programming algorithm with an affine gap penalty model (Durbin *et al.*, 1998; Ibarra and Melo, 2010; Smith and Waterman, 1981) that optimizes the structure alignment length (i.e. total number of equivalent residues) between two previously superposed structures.

In the description of structure alignments, we call the first structure the query (q) and the second structure the target (t). The structure alignments can be characterized by a small set of parameters. The most significant of these is the alignment length, which we call absolute similarity or $S(q, t)$ (Sippl, 2008). The value of $S(q, t)$ depends on the following parameters that affect the outcome of the dynamic programming algorithm: o (gap-opening penalty), e (gap-extension penalty), m (matched pair score), u (unmatched pair score), r (distance threshold) and the set $\{A\}$ (the list of atom types) used to define the equivalent residues. An equivalent position between two residues in the superposed structures is assigned when all corresponding atom pairs (i.e. with the same atom name) from the set of

atom names defined in $\{A\}$ are found at a distance equal or smaller than r Angstroms in three-dimensional space.

In addition to $S(q, t)$, other structure alignment similarity measures are calculated by STOVCA. One of these measures is the relative similarity or $s(q, t)$, which constitutes a symmetric and global similarity measure (Sippl, 2008):

$$s(q, t) = 100 \times \frac{2 \times S(q, t)}{L_q + L_t}, \quad (1)$$

where L_q and L_t are the sequence lengths of the query and target structures, respectively, measured in units of residues (i.e. amino acids, nucleotides). Another similarity measure is the structural overlap or $SO(q, t)$, which represents a measure of local similarity (i.e. the relative cover of the smallest structure with respect to the largest structure in the superposed pair) and is defined by:

$$SO(q, t) = 100 \times \frac{S(q, t)}{\min(L_q, L_t)}. \quad (2)$$

It is noteworthy to mention that $SO(q, t)$ approximates the relative query cover, $c_q = 100 \times S(q, t) / L_q$, and the relative target cover, $c_t = 100 \times S(q, t) / L_t$ (Sippl, 2008), when the two structures are similar in length, as is the case for most protein structure pairs used in this study. Alternatively, when the two structures are very different in size, $SO(q, t)$ is redundant with either c_q or c_t . However, for completeness, these three similarity measures, $SO(q, t)$, c_q and c_t , are reported in the output of STOVCA software. Finally, the percentage of sequence identity and the total RMS error (of the atom types defined) of the equivalent residues are also reported by STOVCA from the optimal structure alignment.

2.4 Parameterization of the method

The gap-opening penalty corresponds to the inverse of the smallest fragment length allowed in the alignment, and it is the most important parameter whose value should be modified by the user. Default parameter values of the software, which correspond to those used to calculate the optimal alignments in this work, are the following: $o = -3$, $e = 0$, $r = 3.5 \text{ \AA}$, $m = 1$, $\{A\} = \text{'C}\alpha'$ atoms and u = the largest negative integer representable on a given computer and therefore this value is determined at compile time.

The alignment length and the values of the similarity measures described above depend on the optimal structure alignment that is calculated, which in turns depends on the parameter values defined in the dynamic programming algorithm. After carrying out several tests with different combinations of parameter values (i.e. a sensitivity analysis), we found that the most critical parameter was the gap-opening penalty. All calculations of structural alignments reported in this work were obtained with a gap-opening penalty value of -3 . Changing this value to -2 or to -4 did not significantly modify the resulting alignments. If this penalty value is further increased to more negative values, the resulting alignments are shorter and more stringent. The user of STOVCA can adapt all parameter values of the algorithm to satisfy specific requirements. However, to preserve the expected behavior of the method, the user should change only the gap-opening penalty (o), distance threshold (r) and the set of atom types ($\{A\}$).

2.5 Software features and availability

STOVCA generates as an output the corresponding optimal structure alignment and the similarity measures derived from it. Computer scripts for rapid visualization of the calculated structural alignment can be generated for PyMOL and RASMOL molecular graphic software. The current version of STOVCA can be used to calculate the maximum structural overlap from superposed pairs of proteins, RNA molecules and DNA duplexes by adjusting the atom types. STOVCA computer software was written in C++, it is distributed under a LGPL license

and its binaries for Windows, LINUX and Mac OS X are freely available for download to everyone from our laboratory website located at: <http://melolab.org/stovca>. For advanced users, STOVCA source code is freely available upon request.

3 RESULTS

3.1 Comparison of program performance

Each of the 215 protein structure pairs in the benchmark set was superposed with DALI, CE, MUSTANG, SALIGN, TMALIGN, MAMMOTH and TOPMATCH. From the resulting structure superpositions, the optimal structural alignments were calculated with STOVCA and their corresponding similarity values computed. Several statistics derived from this data are presented in Table 1. It is noteworthy to mention that all tested software was capable of generating either the best or the worst solution in terms of similarity values for some particular cases. This result is important because the use of a single software tool does not guarantee that the best result will always be achieved. Similarly, a given software tool with a good overall performance can also produce the worst solution in some cases. Therefore, the use of several software tools along with a standardized method to calculate the optimal protein structure alignment like the one described here can be helpful to select the most appropriate solution among a set of alignments produced by various programs. From these data, it is clear that DALI, TMALIGN, TOPMATCH and SALIGN are the more robust among the tested software tools.

These four software tools produce a low number of worst solutions while achieving a high number of best solutions. They also obtain a small deviation from the best solution when they produce the worst solution, and they achieve the highest average over all tested cases. In terms of performance, a second group of software tools consists of CE, MAMMOTH and HOMSTRAD. At the bottom line of performance, MUSTANG generally produces the shortest structural alignments, the lowest average of structural overlap values, the largest

difference with the best solution when it produces the worst solution, and the largest count of worst solutions.

3.2 Statistical analysis

Plotting the cumulative distributions of SO values produced by the various software tools reveals this overall trend of performance (Fig. 1). These qualitative trends are confirmed by the statistical analysis of the data (Table 2). The observed performance difference of TOPMATCH, DALI, TMALIGN and SALIGN with other software tools is statistically significant at the 95% confidence level. The analysis of SO correlation factors between different software tools reveals that correlations among the various programs can be small (Table 2).

It is noteworthy to mention that SO values from HOMSTRAD alignments correlate poorly with those obtained from the other software tools tested here (i.e. computed *R* factors are below 0.61).

3.3 Example cases

To illustrate the large variety of solutions generated when using two independent software tools, we have selected a couple of example cases (Fig. 2). In these examples, the best and worst performing tools for particular protein structure pairs from the benchmark were selected. It is clear that totally different structure alignments are generated in these cases, with one solution being clearly better than the other one. The graphical views of all structure superpositions from our benchmark and their corresponding optimal structure alignments generated in this work are available for detailed 3D inspection as supplementary data at <http://melolab.org/supmat/stovca> (see Section 2).

A moderate example that illustrates the usefulness of the methodology described here is provided by the comparison of the catalytic domains (i.e. palm domains) of two DNA polymerases: the DNA polymerase I palm domain from *Bacillus Stearothermophilus* (PDB code 1xwl) and the palm domain of DNA polymerase Dpo4 from *Sulfolobus solfataricus* (PDB

Table 1. Statistics of structure alignments

Method or software	Number of cases		Alignment length loss from best alignment		Average structural overlap
	Best	Worst	Cumulative	Average	
CE	15/32	35	1554	7	58.0
DALI	15/46	14	820	4	60.8
MAMMOTH	8/16	61	1738	8	57.1
MUSTANG	10/18	77	3013	14	52.2
TMALIGN	13/41	14	733	3	61.2
TOPMATCH	47/84	5	529	2	62.0
SALIGN	35/72	8	603	3	61.7
HOMSTRAD	7/19	37	1821	8	57.0
Best Value	215	0	0	0	64.1

Columns with the number of best and worst cases contain the counts where a particular software tool generates the best (numbers after the slash represent the counts when the best result is shared with other software tools) and worst solutions. Cumulative and average alignment length loss represents the sum and average over the 215 superpositions in the benchmark. Last column contains the average SO value for the 215 superpositions in the benchmark. Best Value represents the maximum (structural overlap) or minimum (alignment length loss from best alignment) values obtained with any of the tested software tools and HOMSTRAD for each particular case in the benchmark.

code 2imw). The resulting optimal structure alignments from STOVCA for the superpositions of these two structures with TopMatch and Mustang show SO values of 74 and 66, respectively (Fig. 3). The percentage sequence identity of the structure alignments is 14% for TopMatch and 13% for Mustang. When the structural alignments generated by STOVCA are analysed in detail, it is observed that in the case of TopMatch, the aspartic residues involved in catalysis (Asp-829 and Asp-652 from DNA Pol I; Asp-104 and Asp-7 from DNA Pol Dpo4) are correctly aligned. This is not the case for the structural alignment obtained

from the superposition generated by Mustang, where Asp-829 from DNA Pol I is not equivalent with Asp-104 from DNA Pol Dpo4. This specific example of misalignment would have a negative impact in downstream analysis of protein function or protein evolution of DNA polymerases. In this case, the structure alignments generated by STOVCA are able to rank and identify the best structure superposition produced with these two software tools, even though the differences in SO and sequence identity values are not as large as in the previous examples (i.e. shown in Fig. 2). Thus, small differences in the structural alignments can

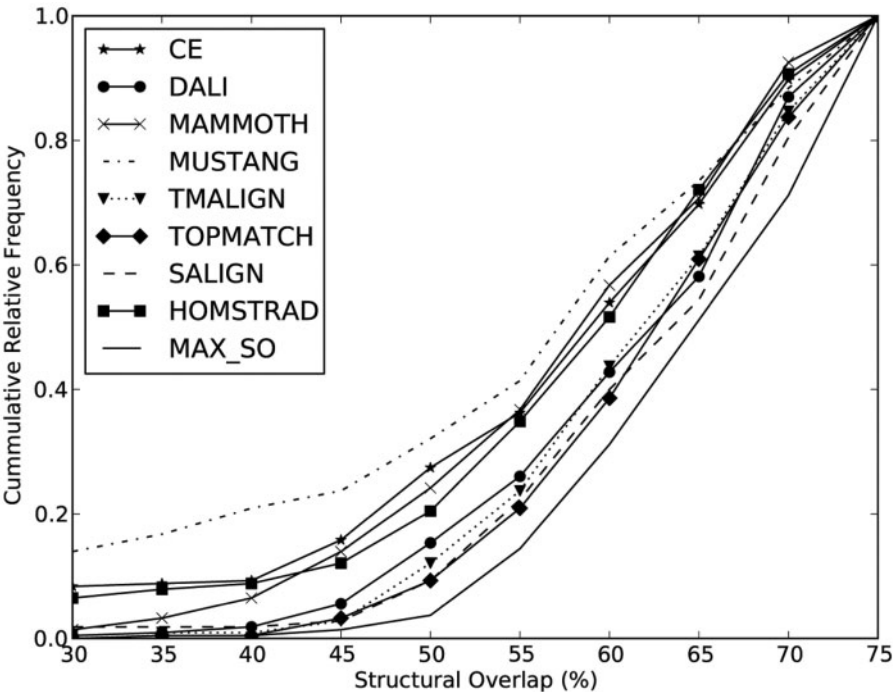


Fig. 1. Performance comparison. The cumulative distributions of SO values obtained with the seven software tools and HOMSTRAD for the 215 protein structure pairs in the benchmark are plotted. The cumulative distribution of the maximum SO value (MaxSO) among the tested tools is also included. The distribution was calculated by defining homogeneous bins of 5% width in the range 27.5–77.5% SO values. First and last bins also include data below 27.5% and above 77.5%, respectively. Middle bin values are plotted in the x-axis of this graph. Colour version of this figure is available at: <http://melolab.org/stovca>

Table 2. Statistical analysis of performance difference

	CE	DALI	MAMMOTH	MUSTANG	TMALIGN	TOPMATCH	SALIGN	HOMSTRAD	MaxSO
CE		0.75	0.67	0.62	0.78	0.76	0.73	0.53	0.80
DALI	0.021		0.77	0.73	0.89	0.89	0.82	0.59	0.92
MAMMOTH	0.653	0.015		0.62	0.81	0.76	0.74	0.44	0.79
MUSTANG	0.027	$<5 \times 10^{-4}$	0.015		0.68	0.66	0.63	0.61	0.70
TMALIGN	0.001	0.935	0.003	$<5 \times 10^{-4}$		0.88	0.89	0.58	0.93
TOPMATCH	$<5 \times 10^{-4}$	0.575	$<5 \times 10^{-4}$	$<5 \times 10^{-4}$	0.575		0.85	0.56	0.93
SALIGN	$<5 \times 10^{-4}$	0.422	$<5 \times 10^{-4}$	$<5 \times 10^{-4}$	0.422	0.738		0.55	0.88
HOMSTRAD	0.356	0.027	0.815	0.048	0.048	0.003	0.001		0.59
MaxSO	$<5 \times 10^{-4}$	0.001	$<5 \times 10^{-4}$	$<5 \times 10^{-4}$	0.006	0.008	0.102	$<5 \times 10^{-4}$	

Upper-right triangle: Pearson correlation factors of SO values from all protein structure superposition pairs in the benchmark ($n=215$). Lower-left triangle: P -values of Kolmogorov–Smirnov (KS) nonparametric statistical test for the SO values from the 215 protein structure superposition pairs in the benchmark. Bold type font indicates that the observed differences are not statistically significant ($\alpha=0.05$).

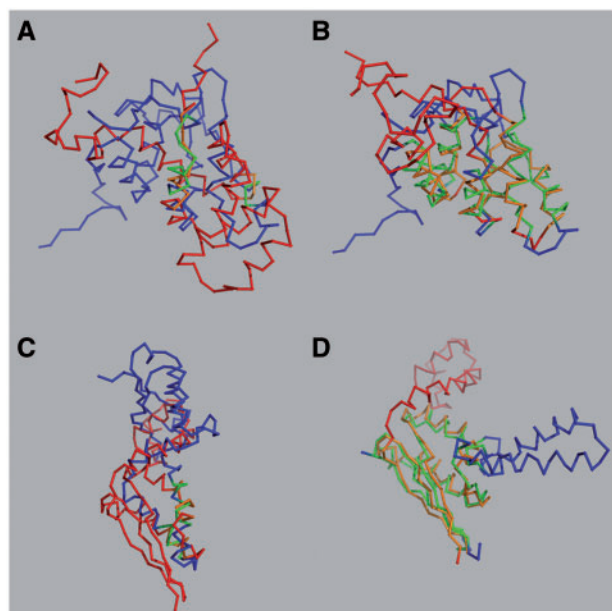


Fig. 2. Example protein structure superpositions with extreme performance difference. Resulting superposition of proteins 1d2zb and 1e41a with MAMMOTH (A) and SALIGN (B) gives SO values of 9.45 and 55.91, respectively. Resulting superposition of proteins 1efud1 and 1tfe with CE (C) and MUSTANG (D) gives SO values of 10.79 and 58.99, respectively. Structure superposition was generated by the specific software tools mentioned above, and the subsequent calculation of the optimal structure alignment with STOVCA software was carried out. Segments of protein chains that do not contain equivalent positions are shown in blue (query protein chain) and green (target protein chain) colours. Structurally aligned regions of blue and green protein chains are shown in red and orange colours, respectively

have a large impact on the conclusions drawn from the analysis of structure/function relationships in divergent protein families where the sequence conservation is low, as it is the case of DNA polymerases illustrated here. Computer tools like STOVCA can help to rank the structure alignments produced by distinct software tools to select the most appropriate solution, thus minimizing errors in downstream analysis.

4 DISCUSSION

As we have demonstrated here, the construction of structural alignments of proteins can be split in two independent parts: (i) the geometric transformation required for the superposition of two structures and (ii) the read-out of the associated sequence alignment from the superimposed coordinates. The transformation captures the performance of a particular protein structure alignment program whereas the read-out of the alignment can be standardized. In this way the geometric transformation is mapped to a single characteristic number, the alignment length, which can then be used to judge the performance of individual methods or to choose the longest alignment from a set of alternatives. Using this approach, we have shown that the various programs investigated here perform quite differently on the benchmark dataset.

On the one hand the programs split in two groups where one group clearly outperforms the other, but at the same time there is no single program consistently producing the optimal alignment. Therefore, depending on the application, it may be advantageous to use several programs in parallel to choose the most appropriate alignment from the ensemble of alternatives reported by the individual programs in each case.

It is important to highlight that the software developed here, called STOVCA, is used only to calculate sequence-based alignments of maximum length from an input pair-wise protein structure superposition generated with another software. The alignment produced by STOVCA is optimal according to the user-provided parameters. Given a fixed set of parameters, the resulting sequence alignments are standardized and thus they can be compared directly.

In benchmarking the various programs, we used single domains of comparable size corresponding to the most basic scenario encountered in protein structure research. In many applications the situation is more complex. Proteins often consist of several domains, contain repeats, the alignment may require permutations in the sequences, and the complete description of the extent of structural similarity of two proteins may require several alternative alignments. We have not addressed these issues here, as only a few programs report a full set of alignments, can handle permutations or multi-domain proteins (e.g. Sippl and Wiederstein, 2012). Therefore, besides the criteria of optimality we have discussed here, the power of a structure alignment program and the quality of the results obtained critically depend on these additional criteria.

Additionally, we have not incorporated in this benchmark any software that carries out flexible protein structure superposition, but only rigid structure superposition. Given that maximum alignment length was the criteria adopted to assess the performance of the structure superposition methods, it would be unfair to use this measure to assess the performance of flexible aligners, which normally produces many short alignments with a low RMS error. Along the same line of discussion, it is also fair to mention that not all the structure superposition software tested here have been optimized to maximize the alignment length. This may explain why on average some software tools perform better than others. The focus of the work presented here is not to compare and rank different structure superposition software but rather to emphasize that in difficult cases the results obtained from the various programs may differ considerably. A reasonable strategy to compare various alignments and to select the best one according to a given measure (like alignment length) is to generate all the alignments in standardized form. This is exactly what the method reported here does. Similar strategies may eventually be developed in the future for alternative measures, thus allowing the comparison of software performance from distinct point of views.

We expect that the use of the method presented here improves the accuracy of structural alignments of distantly related and single-domain proteins, which in turn may have a widespread and positive impact in the study of: (i) sequence-structure relationships in proteins and RNA molecules, (ii) evolutionary relationships in protein families that exhibit a high divergence at the sequence level, (iii) identification of functionally important residues and (iv) classification of protein and RNA structures.

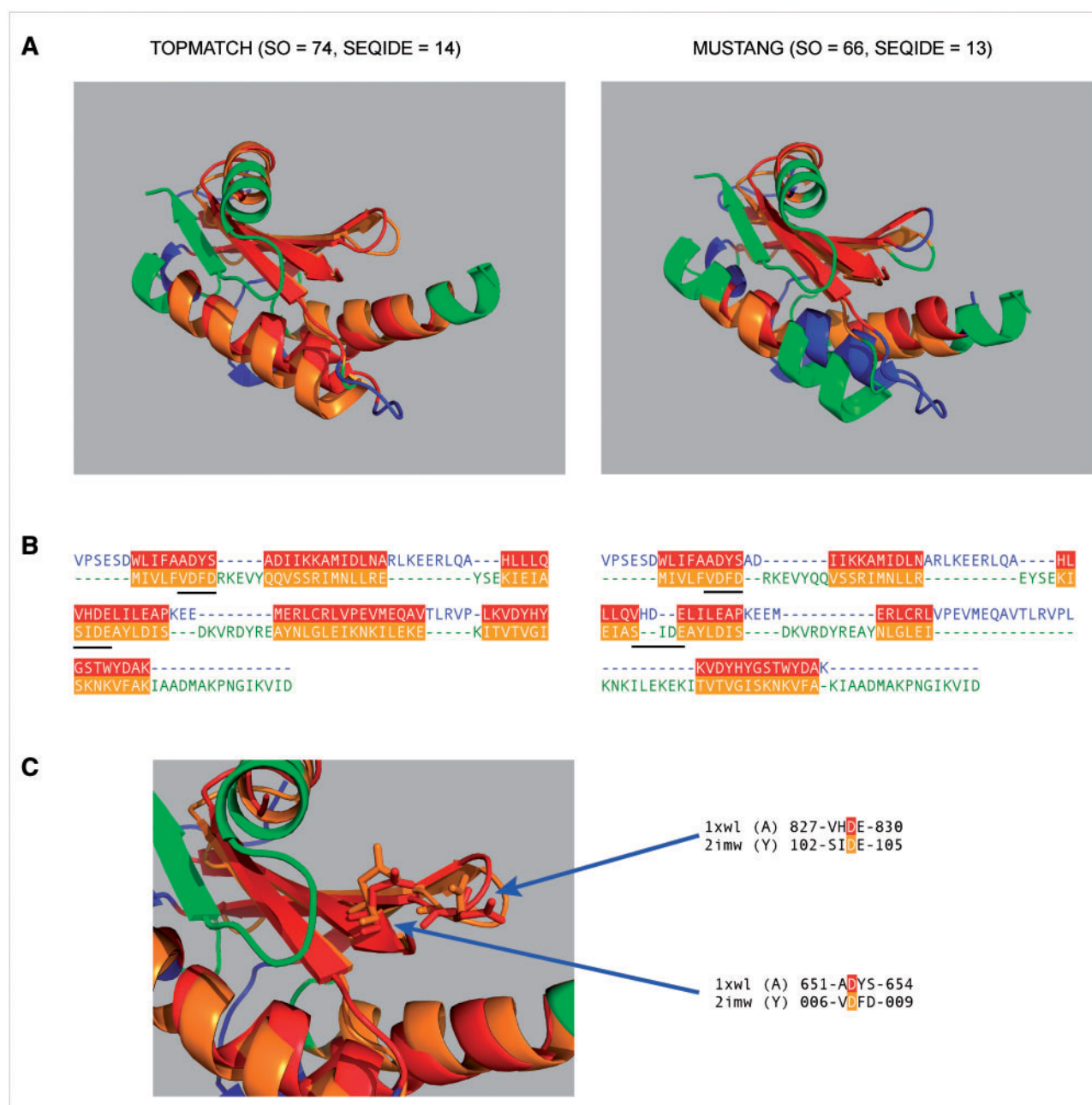


Fig. 3. Example protein structure superpositions of palm domains from DNA polymerases. (A) Structure superpositions of DNA polymerase I from *B. stearothermophilus* (PDB code 1xwl) and DNA polymerase Dpo4 from *S. solfataricus* (PDB code 2imw) generated with TOPMATCH (left) and MUSTANG (right). (B) Optimal structure alignments generated with STOVCA for the superposition reported by TOPMATCH (left) and MUSTANG (right). The two regions containing the catalytic aspartic residues are underlined in the alignments. (C) Detailed view of the catalytic aspartic residues in the superposition generated with TOPMATCH and the corresponding optimal structure alignment obtained with STOVCA. Colouring scheme used here is the same as that described in Figure 2

ACKNOWLEDGEMENT

We are very grateful to Dr Lisa Holm for sending us the Dali software, as well as for her support to install it in our computer systems.

Funding: This work was funded by CONICYT (Comisión Nacional de Investigación Científica y Tecnológica, Chile)

graduate scholarship (to A.W.S.); FWF Austria, grant number P21294-B12 (to M.J.S.); FONDECYT Chile, grant number 1110400 and ICM (Iniciativa Científica Milenio, Chile) grant number P09-016-F (to F.M.).

Conflict of Interest: none declared.

REFERENCES

- Durbin,R. *et al.* (1998) *Biological Sequence Analysis*. Vol. 2, 1st edn. Cambridge University Press, Cambridge, pp. 12–32.
- Feng,Z.K. and Sippl,M.J. (1996) Optimum superimposition of protein structures: ambiguities and implications. *Fold. Des.*, **1**, 123–132.
- Hasegawa,H. and Holm,L. (2009) Advances and pitfalls of protein structural alignment. *Curr. Opin. Struct. Biol.*, **19**, 341–348.
- Holm,L. (2000) DaliLite workbench for protein structure comparison. *Bioinformatics*, **16**, 566–567.
- Ibarra,I.L. and Melo,F. (2010) Interactive software tool to comprehend the calculation of optimal sequence alignments with dynamic programming. *Bioinformatics*, **26**, 1664–1665.
- Kabsch,W. (1976) A solution of the best rotation to relate two sets of vectors. *Acta Crystallogr. B*, **32**, 922.
- Kabsch,W. (1978) A discussion of the solution for the best rotation to relate two sets of vectors. *Acta Crystallogr. A*, **34**, 827–828.
- Konagurthu,A.S. *et al.* (2006) MUSTANG: a multiple structural alignment algorithm. *Proteins*, **64**, 559–574.
- Madhusudhan,M.S. *et al.* (2009) Alignment of multiple protein structures based on sequence and structure features. *Protein Eng. Des. Sel.*, **22**, 569–574.
- Mizuguchi,K. *et al.* (1998) HOMSTRAD: a database of protein structure alignments for homologous families. *Protein Sci.*, **7**, 2469–2471.
- Ortiz,A.R. *et al.* (2002) MAMMOTH (matching molecular models obtained from theory): an automated method for model comparison. *Protein Sci.*, **11**, 2606–2621.
- Sali,A. and Blundell,T.L. (1993) Comparative protein modelling by satisfaction of spatial restraints. *J. Mol. Biol.*, **234**, 779–815.
- Shindyalov,I.N. and Bourne,P.E. (1998) Protein structure alignment by incremental combinatorial extension (CE) of the optimal path. *Protein Eng.*, **11**, 739–747.
- Sippl,M.J. (2008) On distance and similarity in fold space. *Bioinformatics*, **24**, 872–873.
- Sippl,M.J. and Stegbuchner,H. (1991) Superposition of three-dimensional objects: a fast and numerically stable algorithm for the calculation of the matrix of optimal rotation. *Comput. Chem.*, **15**, 73–78.
- Sippl,M.J. and Wiederstein,M. (2008) Structural bioinformatics A note on difficult structure alignment problems. *Bioinformatics*, **24**, 426–427.
- Smith,T.F. and Waterman,M.S. (1981) Identification of common molecular subsequences. *J. Mol. Biol.*, **147**, 195–197.
- Sippl,M.J. and Wiederstein,M. (2012) Detection of spatial correlations in protein structures and molecular complexes. *Structure*, **20**, 718–728.
- Zhang,Y. and Skolnick,J. (2005) TM-align: a protein structure alignment algorithm based on the TM-score. *Nucleic Acids Res.*, **33**, 2302–2309.