

Prediction of Protein Functions

Roy D. Sleator

Abstract

The recent explosion in the number and diversity of novel proteins identified by the large-scale “omics” technologies poses new and important questions to the blossoming field of systems biology – What are all these proteins, how did they come about, and most importantly, what do they do?

From a comparatively small number of protein structural domains a staggering array of structural variants has evolved, which has in turn facilitated an expanse of functional derivatives. This review considers the primary mechanisms that have contributed to the vastness of our existing, and expanding, protein repertoires, while also outlining the protocols available for elucidating their true biological function. The various function prediction programs available, both sequence and structure based, are discussed and their associated strengths and weaknesses outlined.

Key words: Protein function, Homology-based transfer, Ontologies, Sequence and structure motifs, Evolution, Protein domains, Gene duplication, Divergence, Combination, Circular permutation

1. Introduction

While the famous quote from American architect, Louis Sullivan, that “form follows function” holds for man-made structures, in protein science the reverse is true – *function follows form*.

Data from the most recent large-scale sequencing projects has facilitated detailed descriptions of the constituent protein repertoires of more than 600 distinct organisms (1). Taking protein domains (clusters of 50–200 conserved residues) to represent units of evolution, as well as their more usual designation as structural/functional motifs, it is possible to accurately trace the evolutionary relationships of approximately half of these proteins (2).

Until recently, in the absence of any experimental evidence, homology-based transfer remained the gold standard for ascribing

a functional role to such newly identified proteins (3). Based on this approach, if a query protein shares significant sequence similarity (suggesting a common evolutionary origin) to a protein of known function, then the function of the latter may be transferred to the former (referred to as the query protein). However, as the databases continue to expand at an exponential rate, the utility of homology based prediction methods continues to contract, with fewer query proteins registering significant hits to known proteins. Herein, I review the current knowledge on protein evolution with a specific focus on how gene duplications, sequence divergence and domain combinations have shaped protein evolution. Furthermore, the most recent advances in the field of automated function prediction (AFP) are discussed, along with the future challenges and outstanding questions which still remain unanswered.

2. What Is Shaping Protein Structure?

2.1. Duplication

Of the animal genomes sequenced to date, the proportion of matched domains which are the result of duplications is estimated at between 93 and 97% (4). Indeed, the haemoglobins, which were the first homologous proteins to have their structure determined, are perhaps the best example of how duplication (and subsequent mutational events) has given rise to subtle structural and functional variations such as oxygen binding profiles (5). Furthermore, in addition to the generation of whole protein homologues, partial gene duplications resulting in domain duplication and elongation are also common features of protein evolution (6). In many cases, such enlargements have resulted from the addition of subdomains, variability in loop length, and/or changes to the structural core, such as beta-sheet extensions (7). Examples of such protein duplication events include cutinase and bovine bile-salt activated cholesterol esterase. While cutinase is the smallest enzyme of the α/β hydrolases, with five strands in the main beta-sheet (8), bovine bile-salt activated cholesterol esterase has 11 strands, and loop structures up to 79 residues in length (9).

2.2. Divergence

There are essentially two types of protein structural divergence: changes to the protein's surface or peripheral regions (e.g., surface loops, surface helices, and strands on the edges of β -sheets) and the less common but far more detrimental modifications to the protein's interior or core (10). Indeed, it has been demonstrated that mutations in the protein surface are four times more biologically acceptable than those in the interior (1). In support of this is the observation that pairs of homologous proteins with identities of approximately 20% have been shown to exhibit up to 50% divergence in the peripheral regions alone (11).

In addition to subtle changes resulting from missense point mutations leading to single amino acid substitutions and the resulting gradual divergence in structure and function, more radical divergence of structure, mediated by domain shuffling (recombination or permutation) has also been reported (12). Circular permutations (CPs) in particular represent a specific form of recombination event that is characterized by the presence of the same protein subsequences in the same linear order but different positions of the N- and C-termini (13), in essence CP of a protein can be visualized as if its original termini were linked and new ones created elsewhere. First observed in plant lectins (14), a substantial number of natural examples of CP have been reported; indeed, some 120 protein clusters which appear to have segments of their sequences in different sequential order are reported in the Circular Permutation Database (15). In addition to natural evolutionary processes, artificial CPs have been engineered in an effort to study protein folding properties as well as the design of more efficient enzymes (16). A circularly permuted streptavidin, for example, has been designed to remove the flexible polypeptide loop that undergoes an open to closed conformational change when biotin is bound. The original termini have been joined by a tetrapeptide linker, and four loop residues have been removed, resulting in the creation of new N- and C-termini (16).

While domain shuffling may have dramatic effects on protein structure, protein homologues usually conserve their catalytic mechanisms, i.e., the relative positions of their functional active sites or catalytic residues may shift but they retain their functional activity. This usually occurs when divergence induces structural changes in the catalytic region, thus necessitating a reconfiguration of the position of the catalytic residues to maintain function (7). In several cases, while the functionally equivalent residues are located at non-homologous positions on the protein's 3D structure, the catalytic residues themselves are identical. An example of this is chloramphenicol acetyltransferase (PaXAT) and UDP-*N*-acetylglucosamine acyltransferase (LpxA); both of which contain an essential histidine residue thought to be involved in deprotonation of a hydroxyl group in their individual substrates. However, these residues are located at different points within the protein fold; in LpxA, the histidine is located in the core of the domain (17), whereas in PaXAT, it occurs in a loop extending from the solenoid structure.

Thus, two proteins may have quite divergent structures and/or sequences while retaining similar function; such proteins are said to be functional analogs. Such analogs may also arise as a result of convergent evolution; that is they do not diverge from a common ancestor but instead arise independently and converge on the same active configuration as a result of natural selection for a particular biochemical function. *L*-Aspartate aminotransferase and *D*-amino acid aminotransferase provide excellent examples of

convergently evolved functional analogs. Despite having a strikingly similar arrangement of residues in their active sites, the two proteins have completely different architectures, differing in size, amino acid sequence, and the fold of the protein domains.

Conversely, certain proteins share significant sequence and/or structure similarity but differ in terms of substrate specificity or indeed catalytic function. An example of such structural analogs, which arise by means of divergent evolution from a single ancestor, include Human IL-10 (hIL-10), a cytokine that modulates diverse immune responses and the Epstein-Barr virus (EBV) IL-10 homolog (vIL-10). Although vIL-10 suppresses inflammatory responses like hIL-10, it cannot activate many other immune-stimulatory functions performed by the cellular cytokine (18).

2.3. Combination

While the evolutionary impact of duplication and divergence on protein sequence, structure and function is obvious, multidomain proteins are for the most part the result of gene combinations (19). Such combinations can give rise to domain recruitment and enlargement and can significantly affect both protein structure/stability and function. For example, in the case of domain recruitment the addition of an accessory domain may affect protein function by modulating substrate selectivity; achieved either by the addition of a binding site, or, by playing a purely structural role, shaping the existing active site to accommodate substrates of different shapes and/or sizes (7). For example, prokaryotic methionine aminopeptidase exists as a monomeric single-domain protein while creatinase, is a two-domain protein. The additional domain of the second subunit of creatinase caps the active site allowing the binding of the small molecule creatine (20).

3. What Is Protein Function?

Before commencing any discussion on protein function prediction we must first consider what is meant by “function”. Biological function is highly contextual; different aspects of the function of a given protein may be viewed as occurring in different scales of space and time; from the almost instantaneous enzymatic reactions to the much slower overall biological process (21). Knowing which functional aspect is being investigated is thus extremely important and can only properly be achieved by the establishment of a standardized machine readable vocabulary.

Fortunately, significant progress has been made in the computer science arena in developing the theory and application of structured machine readable vocabularies, known as ontologies, which provide a formal explicit specification of a commonly used abstract model of the world (22). Ontologies not only allow formal

definition of concepts but also enable the creation of software tools capable of reasoning about the properties and relationships of a domain. Formats such as the Resource Description Framework (RDF) and the Web Ontology Language (OWL) have been devised that allow ontological concepts to be persisted and communicated. RDF, for example, allows the creation of statements about a particular domain by the use of triples in the form of subject–predicate–object expressions. The subject and object represents a concept, whereas the predicate defines the relationship between them.

Detailed ontologies can be created by composing further defining concepts and relationships that model the domain of interest. Ontologies that define different aspects of proteins could be used to annotate biological data with functional facets and provide the basis of a framework for machine based reasoning.

The Gene Ontology (GO) (23) goes some way to achieving this goal, formulating a definition of functional context and providing machine – legible functional annotation. GO has three “ontology trees” describing three aspects of gene product function: Molecular function, biological process and cellular location. By providing a standard vocabulary and defining relationships between terms, annotations can be computationally processed (24), thus providing a standard approach for programs to output their functional predictions.

Having defined biological “function” and the means of describing such functions we can now turn our attention to the various function prediction programs, and their associated strengths and weaknesses.

3.1. Protein Function Prediction Methods

Protein function prediction methods can be loosely divided into sequence and structure based approaches. Herein, we outline the current state of the art for sequence and structure based protein function prediction.

3.1.1. Sequence Based Approaches

Homology-Based Transfer

Homology-based transfer, using programs such as BLAST (25), is perhaps the most widely used form of computational function prediction method; assigning un-annotated proteins with the function of their annotated homologs. The rationale for this approach is based on the assumption that two sequences with a high degree of similarity most likely evolved from a common ancestor and thus must have similar functions.

While sequence similarity is undoubtedly correlated to functional similarity, exceptions have been observed on both ends of the similarity scale. Rost (26), for example, showed that even at high sequence similarity rates, enzymatic function may not necessarily be conserved, while Galperin et al. (27) observed that enzymes that are analogous on the basis of sequence dissimilarity are in fact homologous. While such errors are the exception rather than the rule, they may set the seed for further annotation errors; as more sequences

enter the databases, more are annotated by homology-based transfer, thus helping to propagate and amplify the original single erroneous annotation (28, 29).

Furthermore, as the databases continue to expand, the utility of the homology-based transfer approach begins to break down. The recent explosion of large-scale metagenomic sequencing projects (30) has resulted in an unprecedented amount of novel sequences being deposited in the databases. As a direct consequence of this sequence expansion, the number of clustered similar proteins for which no single annotated reference sequence exists is expanding rapidly, eroding the foundations of the homology-based transfer approach. Indeed, it has been estimated that <35% of all proteins could be annotated automatically when accepting errors of $\leq 5\%$, while even allowing for error rates of $>40\%$ there is no annotation for $>30\%$ of all proteins (31).

Sequence Motifs

Typically of the 100–300 amino acids in a functional protein domain <10% constitute the protein's active sites (32). Therefore, homology-based transfer from a complete protein is often not necessary to predict a protein's function. All that is required is a sequence (or structure) based signature which is associated with a particular function. Such signatures may occur at a single position on the sequence or as a “fingerprint” composed of several such patterns. A few databases are dedicated to motif searching; PROSITE (33), for example, is composed of manually selected biologically important motifs and has three types of signatures: patterns, rules, and profiles. Each signature represents a different automated method for searching motifs; while patterns and rules typically span only a few residues (e.g., A typical entry in PROSITE would be (ST)-x(2)-(DE), i.e., a Serine or Threonine, followed by any two residues, followed by Aspartate or Glutamate – the consensus sequence of a Casein kinase II phosphorylation site), profiles extend the similarity to the level of entire domains. Other well-known motif databases include BLOCKS (34) and PRINTS (35).

Genomic Context and Expression Based Prediction Methods

Genomic context based prediction, also referred to as phylogenomic profiling is a method for predicting protein function based on the observation that proteins with similar pedigrees (inter-genomic profiles) are believed to have evolved in tandem and as such are likely to share a common function (36). Furthermore, in prokaryote genomes the loci of functionally related proteins tend to be colocated on the chromosome. Combining coevolution and colocation (chromosomal proximity) has given rise to a new generation of function-prediction algorithms such as Phydbac2 (37).

As an extension of colocation, genes involved in similar cellular functions also tend to be cotranscribed. Following this logic, unknown genes coexpressed with known genes may be functionally annotated by virtue of association. This “guilt by association”

approach has given rise to an algorithm of the same name, developed by Walker et al. (38) for the analysis of gene expression arrays. Unlike the sequence motif based approach, which focuses on molecular function, annotation expression microarray based predictions are useful for annotation of the cellular aspect of protein function. Furthermore, given that most cellular processes are carried out by groups of physically interacting proteins, it is fair to assume that such interacting proteins have similar overall cellular functions. Thus, protein-protein interaction (PPI) data may also facilitate protein function annotation and several PPI databases are now available such as STRING – a database of known and predicted PPIs (39).

3.1.2. Structure Based Approaches

Given that protein structure is far more conserved than sequence, many proteins which exhibit little or no sequence similarities, due to evolutionary constraints still retain significant structure similarity (40). In this respect structure is a useful indicator of function; indeed most known protein folds are associated with a particular function or functional milieu (7). Programs that scan the Protein Data Bank (PDB) for structural similarity given a query sequence include, among others, FATCAT (41), PAST (42), and VAST (43). However, knowledge of 3D protein structure alone is not always sufficient to accurately infer function. Indeed, it is estimated that functional hypotheses can be made from 3D structures for only ~20–50% of hypothetical proteins (44, 45).

Rather than focusing on the protein as a whole, it is possible, and in some instances more desirable, to target 3D motifs associated with specific functions (e.g., binding sites or active sites). The rationale for analysing structure motifs (or patterns) is analogous to that of sequence patterns – to identify unique signatures indicative of a particular function. Libraries of 3D motifs with known function have begun to evolve (46), one example of which is PROCAT (47), a database of 3D enzyme active sites that can be queried for specific functional signatures. In addition, hybrid motifs incorporating information from sequence and structure, as well as from the literature, have also been used to predict protein function (48).

4. Conclusions and Future Prospects

Herein, I have discussed how mechanisms such as gene duplication, sequence divergence and domain combinations (49) have shaped protein evolution and how the retention of sequence and/or structural domains has facilitated the tracking of this evolutionary process through the millennia. I have also introduced the far more complex issue of protein function elucidation wherein, in contrast to protein structure in which the data is either known or easily predicted, the multifaceted and ambiguous nature of biological

function makes its elucidation a far more complex endeavor. The complexity of the problem is perhaps best illustrated by Jeffrey's (50) so called "moonlighting proteins" which perform several contextually different functions, ranging from the molecular to the cellular level. Thus, given the aggregate nature of protein function prediction, perhaps the best outcome will be achieved by adopting a multifaceted approach. For example, while biochemical function prediction is likely best served by focusing on sequence motifs, resolution of physiological function is better addressed at the genomic level, based for example on microarray expression data. Therefore, composite methods, employing a diversity of features to assess different functional aspects, are most likely to succeed. Examples of such aggregate functional prediction programs include InterPro, ProKnow and ProFunc, which utilize several data sources and/or algorithms to predict function.

However, despite the emergence of ever more sophisticated and versatile function prediction algorithms; the proper assessment of such programs still remains a significant limitation to the development of the field. Unlike assessment of protein structure, function prediction methods still lack a viable blind benchmark for which to assess program efficacy. This obstacle may eventually be overcome by emulating successful collaborative efforts of computational and experimental structural biologists in the form of CASP (Critical Assessment of Structure Prediction) for the benchmarking of protein structure.

References

1. Chothia, C., and Gough, J. (2009) Genomic and structural aspects of protein evolution, *Biochem J* **419**, 15–28.
2. Sleator, R. D. (2010) An overview of the processes shaping protein evolution, *Science Progress* **93**, 1–6.
3. Sleator, R. D., and Walsh, P. (2010) An overview of in silico protein function prediction, *Arch Microbiol* **192**, 151–155.
4. Wilson, D., Pethica, R., Zhou, Y., Talbot, C., Vogel, C., Madera, M., Chothia, C., and Gough, J. (2009) SUPERFAMILY – sophisticated comparative genomics, data mining, visualization and phylogeny, *Nucleic Acids Res* **37**, D380–386.
5. Blanchetot, A., Wilson, V., Wood, D., and Jeffreys, A. J. (1983) The seal myoglobin gene: an unusually long globin gene, *Nature* **301**, 732–734.
6. Moore, A. D., Bjorklund, A. K., Ekman, D., Bornberg-Bauer, E., and Elofsson, A. (2008) Arrangements in the modular evolution of proteins, *Trends Biochem Sci* **33**, 444–451.
7. Todd, A. E., Orengo, C. A., and Thornton, J. M. (2001) Evolution of function in protein superfamilies, from a structural perspective, *J Mol Biol* **307**, 1113–1143.
8. Longhi, S., Czjzek, M., Lamzin, V., Nicolas, A., and Cambillau, C. (1997) Atomic resolution (1.0 Å) crystal structure of *Fusarium solani* cutinase: stereochemical analysis, *J Mol Biol* **268**, 779–799.
9. Chen, J. C., Miercke, L. J., Krucinski, J., Starr, J. R., Saenz, G., Wang, X., Spilburg, C. A., Lange, L. G., Ellsworth, J. L., and Stroud, R. M. (1998) Structure of bovine pancreatic cholesterol esterase at 1.6 Å: novel structural features involved in lipase activation, *Biochemistry* **37**, 5107–5117.
10. Gerstein, M., Sonnhammer, E. L., and Chothia, C. (1994) Volume changes in protein evolution, *J Mol Biol* **236**, 1067–1078.
11. Chothia, C., and Lesk, A. M. (1986) The relation between the divergence of sequence and structure in proteins, *EMBO J* **5**, 823–826.

12. Kawashima, T., Kawashima, S., Tanaka, C., Murai, M., Yoneda, M., Putnam, N. H., Rokhsar, D. S., Kanehisa, M., Satoh, N., and Wada, H. (2009) Domain shuffling and the evolution of vertebrates, *Genome Res* **19**, 1393–1403.
13. Vogel, C., and Morea, V. (2006) Duplication, divergence and formation of novel protein topologies, *Bioessays* **28**, 973–978.
14. Lindqvist, Y., and Schneider, G. (1997) Circular permutations of natural protein sequences: structural evidence, *Curr Opin Struct Biol* **7**, 422–427.
15. Lo, W. C., Lee, C. C., Lee, C. Y., and Lyu, P. C. (2009) CPDB: a database of circular permutation in proteins, *Nucleic Acids Res* **37**, D328–332.
16. Heinemann, U., Ay, J., Gaiser, O., Muller, J. J., and Ponnuswamy, M. N. (1996) Enzymology and folding of natural and engineered bacterial beta-glucanases studied by X-ray crystallography, *Biol Chem* **377**, 447–454.
17. Wyckoff, T. J., and Raetz, C. R. (1999) The active site of *Escherichia coli* UDP-N-acetylglucosamine acyltransferase. Chemical modification and site-directed mutagenesis, *J Biol Chem* **274**, 27047–27055.
18. Yoon, S. I., Jones, B. C., Logsdon, N. J., and Walter, M. R. (2005) Same structure, different function crystal structure of the Epstein-Barr virus IL-10 bound to the soluble IL-10RI chain, *Structure* **13**, 551–564.
19. Apic, G., Gough, J., and Teichmann, S. A. (2001) Domain combinations in archaeal, eubacterial and eukaryotic proteomes, *J Mol Biol* **310**, 311–325.
20. Hoeffken, H. W., Knof, S. H., Bartlett, P. A., Huber, R., Moellering, H., and Schumacher, G. (1988) Crystal structure determination, refinement and molecular model of creatine amidinohydrolase from *Pseudomonas putida*, *J Mol Biol* **204**, 417–433.
21. Godzik, A., Jambon, M., and Friedberg, I. (2007) Computational protein function prediction: are we making progress? *Cell Mol Life Sci* **64**, 2505–2511.
22. Losko, S., and Heumann, K. (2009) Semantic data integration and knowledge management to represent biological network associations, *Methods Mol Biol* **563**, 241–258.
23. Ashburner, M., and Lewis, S. (2002) On ontologies for biologists: the Gene Ontology – untangling the web, *Novartis Found Symp* **247**, 66–80; discussion 80–63, 84–90, 244–252.
24. Smith, C. L., Goldsmith, C. A., and Eppig, J. T. (2005) The Mammalian Phenotype Ontology as a tool for annotating, analyzing and comparing phenotypic information, *Genome Biol* **6**, R7.
25. Altschul, S. F., Madden, T. L., Schaffer, A. A., Zhang, J., Zhang, Z., Miller, W., and Lipman, D. J. (1997) Gapped BLAST and PSI-BLAST: a new generation of protein database search programs, *Nucleic Acids Res* **25**, 3389–3402.
26. Rost, B. (2002) Enzyme function less conserved than anticipated, *J Mol Biol* **318**, 595–608.
27. Galperin, M. Y., Walker, D. R., and Koonin, E. V. (1998) Analogous enzymes: independent inventions in enzyme evolution, *Genome Res* **8**, 779–790.
28. Bork, P. (2000) Powers and pitfalls in sequence analysis: the 70% hurdle, *Genome Res* **10**, 398–400.
29. Gilks, W. R., Audit, B., de Angelis, D., Tsoka, S., and Ouzounis, C. A. (2005) Percolation of annotation errors through hierarchically structured protein sequence databases, *Math Biosci* **193**, 223–234.
30. Sleator, R. D., Shortall, C., and Hill, C. (2008) Metagenomics, *Lett Appl Microbiol* **47**, 361–366.
31. Rost, B., Liu, J., Nair, R., Wrzeszczynski, K. O., and Ofran, Y. (2003) Automatic prediction of protein function, *Cell Mol Life Sci* **60**, 2637–2650.
32. Friedberg, I. (2006) Automated protein function prediction – the genomic challenge, *Brief Bioinform* **7**, 225–242.
33. Hulo, N., Bairoch, A., Bulliard, V., Cerutti, L., Cuče, B. A., de Castro, E., Lachaize, C., Langendijk-Genevaux, P. S., and Sigrist, C. J. (2008) The 20 years of PROSITE, *Nucleic Acids Res* **36**, D245–249.
34. Henikoff, J. G., Greene, E. A., Pietrokovski, S., and Henikoff, S. (2000) Increased coverage of protein families with the blocks database servers, *Nucleic Acids Res* **28**, 228–230.
35. Attwood, T. K., Bradley, P., Flower, D. R., Gaulton, A., Maudling, N., Mitchell, A. L., Moulton, G., Nordle, A., Paine, K., Taylor, P., Uddin, A., and Zygouri, C. (2003) PRINTS and its automatic supplement, prePRINTS, *Nucleic Acids Res* **31**, 400–402.
36. Eisenberg, D., Marcotte, E. M., Xenarios, I., and Yeates, T. O. (2000) Protein function in the post-genomic era, *Nature* **405**, 823–826.
37. Enault, F., Suhre, K., and Claverie, J. M. (2005) Phylbac “Gene Function Predictor”: a gene annotation tool based on genomic context analysis, *BMC Bioinformatics* **6**, 247.
38. Walker, M. G., Volkmut, W., Sprinzak, E., Hodgson, D., and Klingler, T. (1999) Prediction of gene function by genome-scale expression analysis: prostate cancer-associated genes, *Genome Res* **9**, 1198–1203.

39. Zhao, X. M., Chen, L., and Aihara, K. (2008) Protein function prediction with high-throughput data, *Amino Acids* **35**, 517–530.
40. Watson, J. D., Laskowski, R. A., and Thornton, J. M. (2005) Predicting protein function from sequence and structural data, *Curr Opin Struct Biol* **15**, 275–284.
41. Ye, Y., and Godzik, A. (2004) FATCAT: a web server for flexible structure comparison and structure similarity searching, *Nucleic Acids Res* **32**, W582–585.
42. Taubig, H., Buchner, A., and Griebisch, J. (2006) PAST: fast structure-based searching in the PDB, *Nucleic Acids Res* **34**, W20–23.
43. Gibrat, J. F., Madej, T., and Bryant, S. H. (1996) Surprising similarities in structure comparison, *Curr Opin Struct Biol* **6**, 377–385.
44. Laskowski, R. A., Watson, J. D., and Thornton, J. M. (2003) From protein structure to biochemical function? *J Struct Funct Genomics* **4**, 167–177.
45. Goldsmith-Fischman, S., and Honig, B. (2003) Structural genomics: computational methods for structure analysis, *Protein Sci* **12**, 1813–1821.
46. Jones, S., and Thornton, J. M. (2004) Searching for functional sites in protein structures, *Curr Opin Chem Biol* **8**, 3–7.
47. Wallace, A. C., Laskowski, R. A., and Thornton, J. M. (1996) Derivation of 3D coordinate templates for searching structural databases: application to Ser-His-Asp catalytic triads in the serine proteinases and lipases, *Protein Sci* **5**, 1001–1013.
48. Di Gennaro, J. A., Siew, N., Hoffman, B. T., Zhang, L., Skolnick, J., Neilson, L. I., and Fetrow, J. S. (2001) Enhanced functional annotation of protein sequences via the use of structural descriptors, *J Struct Biol* **134**, 232–245.
49. Chothia, C., Gough, J., Vogel, C., and Teichmann, S. A. (2003) Evolution of the protein repertoire, *Science* **300**, 1701–1703.
50. Jeffery, C. J. (2003) Moonlighting proteins: old proteins learning new tricks, *Trends Genet* **19**, 415–417.



<http://www.springer.com/978-1-61779-423-0>

Functional Genomics

Methods and Protocols

Kaufmann, M.; Klinger, C. (Eds.)

2012, XV, 418 p., Hardcover

ISBN: 978-1-61779-423-0

A product of Humana Press