

Connecting protein structure with predictions of regulatory sites

Alexandre V. Morozov, and Eric D. Siggia

PNAS 2007;104:7068-7073; originally published online Apr 16, 2007;
doi:10.1073/pnas.0701356104

This information is current as of May 2007.

Online Information & Services	High-resolution figures, a citation map, links to PubMed and Google Scholar, etc., can be found at: www.pnas.org/cgi/content/full/104/17/7068
Supplementary Material	Supplementary material can be found at: www.pnas.org/cgi/content/full/0701356104/DC1
References	This article cites 48 articles, 21 of which you can access for free at: www.pnas.org/cgi/content/full/104/17/7068#BIBL This article has been cited by other articles: www.pnas.org/cgi/content/full/104/17/7068#otherarticles
E-mail Alerts	Receive free email alerts when new articles cite this article - sign up in the box at the top right corner of the article or click here .
Rights & Permissions	To reproduce this article in part (figures, tables) or in entirety, see: www.pnas.org/misc/rightperm.shtml
Reprints	To order reprints, see: www.pnas.org/misc/reprints.shtml

Notes:

Connecting protein structure with predictions of regulatory sites

Alexandre V. Morozov[†] and Eric D. Siggia

Center for Studies in Physics and Biology, The Rockefeller University, 1230 York Avenue, New York, NY 10021

Communicated by Barry H. Honig, Columbia University, New York, NY, February 21, 2007 (received for review December 6, 2006)

A common task posed by microarray experiments is to infer the binding site preferences for a known transcription factor from a collection of genes that it regulates and to ascertain whether the factor acts alone or in a complex. The converse problem can also be posed: Given a collection of binding sites, can the regulatory factor or complex of factors be inferred? Both tasks are substantially facilitated by using relatively simple homology models for protein–DNA interactions, as well as the rapidly expanding protein structure database. For budding yeast, we are able to construct reliable structural models for 67 transcription factors and with them redetermine factor binding sites by using a Bayesian Gibbs sampling algorithm and an extensive protein localization data set. For 49 factors in common with a prior analysis of this data set (based largely on phylogenetic conservation), we find that half of the previously predicted binding motifs are in need of some revision. We also solve the inverse problem of ascertaining the factors from the binding sites by assigning a correct protein fold to 25 of the 49 cases from a previous study. Our approach is easily extended to other organisms, including higher eukaryotes. Our study highlights the utility of enlarging current structural genomics projects that exhaustively sample fold structure space to include all factors with significantly different DNA-binding specificities.

protein–DNA interactions | homology models of transcription factors | weight matrix predictions

Transcription factors (TFs) are regulatory proteins used by the cell to activate or repress gene transcription. They interact with short nucleotide sequences, typically located upstream of a gene, by means of the DNA-binding domains that recognize their cognate binding sites. As a rule, regulation of gene transcription is analyzed by the bioinformatics methods designed to detect statistically overrepresented motifs in promoter sequences. Intergenic sequences bound by the TF can be identified by using DNA microarray technology, including chromatin immunoprecipitation (ChIP-chip) (1, 2), protein binding (3), and DNA immunoprecipitation (DIP-chip) arrays (4). Of special note is a recent genome-wide study that used ChIP-chip analysis to profile *in vivo* genomic occupancies for 203 DNA-binding transcriptional regulators in *Saccharomyces cerevisiae* (2). Using these data, the authors predicted binding specificities for 65 TFs by using the genomes of related species; a number that was later increased to 98 by MacIsaac *et al.* (5).

The DNA-binding domains of TFs can be classified into a limited number of structural families (6, 7). Structural studies of the protein–DNA complexes reveal that, within each family, the overall fold of the DNA-binding domain and its mode of interaction with the cognate binding site are remarkably conserved, resulting in a characteristic pattern of amino acid contacts with DNA bases. These interactions form the basis of the sequence-specific direct readout of nucleotide sequences by amino acids in the DNA-binding domain. Another contribution to the specificity of protein–DNA interactions is indirect and comes from the curvature imposed on the DNA by the contacts with the DNA phosphate backbone. Those DNA sequences that most readily adapt to the imposed shape will bind most favorably.

In some cases, DNA-binding proteins from the same family recognize sites with similar length, symmetry, and specificity.

Alignment of such sites yields a binding profile that can be used to significantly enhance the signal-to-noise ratio in bioinformatics algorithms, allowing for the *ab initio* motif discovery in long metazoan promoter sequences (8–11). However, familial binding profiles are inappropriate when factors form dimers with varying distances between monomer-binding domains, associate with DNA in various orientations, or participate in different multimeric complexes. An averaged site in such cases has little meaning. Furthermore, sufficient binding site data are lacking in many cases. As a result of these complications, only 11 familial profiles are available in the JASPAR database (8), in contrast with 142 protein families that represent TF DNA-binding domains in the Pfam database (12).

Here we present an approach that uses structure-based biochemical constraints to guide motif discovery algorithms. The key observations are (i) that the binding site specificity is largely imparted by the contacts between DNA bases and amino acid side chains and (ii) that the number of atomic contacts with a base pair is strongly correlated with the degree of conservation of that base pair in the binding site (13). Therefore, a structure of the TF–DNA complex and a single consensus (highest affinity) sequence can be used to predict the position-specific weight matrix (PWM) for the TF (14, 15). We make PWM predictions for 67 *S. cerevisiae* TFs by using homology models (we use “homology” to imply similarity rather than evolutionary relationship). We use these predictions as the informative priors in the Bayesian Gibbs sampling algorithm (16, 17) applied to the intergenic sequences identified with the ChIP-chip experiments (2). We find genomic binding sites for TFs and TF complexes and consider whether the inferred binding specificities provide significant refinements of the PWM, based solely on the structural model. Our study extends the best current yeast regulatory map (2, 5), built around phylogenetic conservation (but employing no structural constraints), by predicting 18 additional TF specificities, correcting 16 previous predictions, and amending 10 others, mostly with respect to the length or the composition of the regulatory complex [see [supporting information \(SI\) Fig. 4](#)].

Our approach also makes it possible to solve the inverse problem of associating binding motifs of unknown origin (e.g., inferred from gene expression arrays by using standard bioinformatics methods) with TFs from the structural database, generating experimentally testable hypotheses about the identity of regulatory inputs. To demonstrate the utility of the inverse approach, we attempted to associate correct TFs with 49 PWMs from MacIsaac *et al.* (5) for which the homology models were available in the structural database. Using a database of structure-based PWMs, we were able to assign a correct fold in 25 cases (10 of which yielded our original structural template) and identified several clear instances of mis-

Author contributions: A.V.M. designed research; A.V.M. and E.D.S. performed research; and A.V.M. and E.D.S. wrote the paper.

The authors declare no conflict of interest.

Abbreviations: DIP, DNA immunoprecipitation; PWM, position-specific weight matrix; TF, transcription factor.

[†]To whom correspondence should be addressed. E-mail: morozov@edsb.rockefeller.edu.

This article contains supporting information online at www.pnas.org/cgi/content/full/0701356104/DC1.

© 2007 by The National Academy of Sciences of the USA

association (e.g., the GCN4 leucine zipper motif assigned to the ARG81 experiment, $P = 4.1 \times 10^{-6}$).

Finally, homology modeling of TFs on the genomic scale allows us to identify future targets for structural studies on the basis of the low similarity of their DNA-binding interfaces and those found in solved protein–DNA complexes. Thus, the goal of the structural genomics projects to sample all protein folds should be extended to include all factors with significantly different specificities from each DNA-binding family. One can then hope to find binding sites for the remaining ≈ 100 factors in yeast for which the protein localization data are available (2). Our method is computationally efficient, can be applied to any organism, including higher eukaryotes, and may be used independently of or in conjunction with phylogenetic footprinting. We make the structure-based modeling of DNA-binding proteins available through an interactive web site, Protein–DNA Explorer (<http://protein-dna.rockefeller.edu>).

Results

Overview. We have built a global map of transcriptional regulation in *S. cerevisiae* by using the constraints imposed by the structures of protein–DNA complexes as the informative priors in the Bayesian Gibbs sampling algorithm (see *Materials and Methods* for details). We have supplemented 10 native structures of yeast protein–DNA complexes with homology models for 92 TFs from 14 families (see *SI Table 3*). Using this data set, we were able to make reliable binding specificity predictions for 67 TFs, 57 of which are modeled by homology (see *SI Table 4*). Some models were discarded because of the low interface scores, others because it was impossible to infer the dimeric binding mode from the information available in the literature. In addition, we excluded all C2H2 zinc fingers because of the high protein–DNA interface variability in this family (7), which results in few good matches to multidomain zinc fingers in the structural database. Alternative knowledge-based methods for predicting C2H2 zinc finger binding specificities (18, 19) may be more suitable as input to the Gibbs sampling. The structure-based informative priors for the 67 TFs we chose to keep can already explain the ChIP-chip data reasonably well without further refinement (see *SI Fig. 5*). Knowledge of the structural features at the protein–DNA interface is clearly superior to knowledge of the consensus sequence alone and leads to more accurate PWM predictions (*SI Fig. 5*). Nonetheless, using the yeast genome to test and refine the initial models enables us to infer additional binding specificity caused by indirect readout and to correct the inaccuracies that are likely to occur as a result of our deliberately simple approach to structural modeling. Sometimes, we test several structural models and keep the one for which the most evidence is found in the genome, effectively using the genomic sequence to infer the structural features of the protein–DNA-binding interface. Furthermore, because we are able to reliably associate TFs with DNA sequence motifs, we can systematically determine the identity of various proteins that form multimeric regulatory complexes. In the remainder of this section, we analyze several representative cases in detail, with the idea of demonstrating the power and flexibility of our method and its utility for understanding transcriptional regulation in yeast and other organisms.

Dimeric Structural Variability and Binding Specificities in the Zn2-Cys6 Family. The Zn2-Cys6 binuclear cluster family is the most common in yeast and other fungi. TFs in this family bind DNA as homodimers and recognize sites of various lengths because of a flexible linker region that joins DNA-binding and dimerization domains (20). This structural flexibility results in the variable spacing between the monomeric 5′-CGG-3′ half-sites and in several possible orientations of the monomeric half-sites with respect to one another (Fig. 1). The observed variability of the binding sites makes it impossible to create a familial binding profile (8) for Zn2-Cys6 TFs. Besides the linker flexibility, additional discrimination between target sites is attributed to the protein–DNA inter-

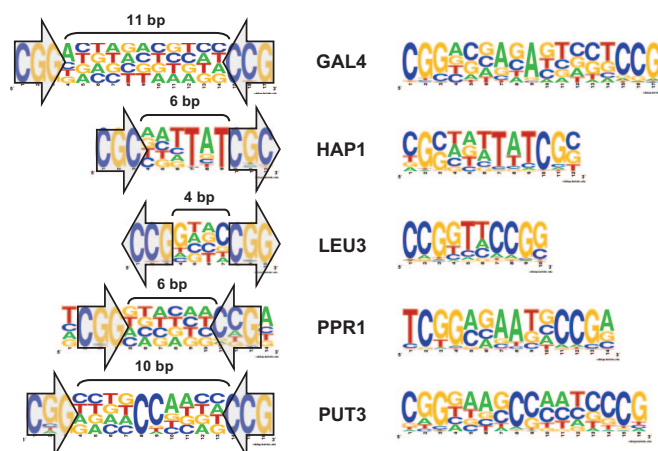


Fig. 1. PWM predictions for five TFs in the Zn2-Cys6 binuclear cluster family, with co-crystal structures showing extensive spacing and orientation variability. (Left) Structure-based priors. (Right) PWMs refined with Gibbs sampling (see *Materials and Methods*). Arrows show the relative orientation of two monomeric half-sites in the dimeric site from the crystal structure. The monomeric half-sites can be arranged in direct (tail-to-head; HAP1), inverted (head-to-head; GAL4, PPR1, PUT3), and everted (tail-to-tail; LEU3) orientations.

actions outside the canonical CGG half-sites, often caused by the asymmetric binding of the monomeric subunits (e.g., HAP1 and PUT3 in Fig. 1) (20–22). Gibbs sampling enables us to refine the initial structure-based PWM predictions with sequence data: for example, the HAP1–DNA structure was solved by using the CGC half-sites (21). However, in the Gibbs sampling PWM the half-sites become a mixture of the CGG and CGC triplets, in accordance with prior knowledge about HAP1 binding sites (21).

The base in the middle of the 17-bp GAL4 PWM is strongly conserved in the Gibbs prediction, despite the total absence of direct amino acid contacts with DNA bases (Fig. 1) but in accordance with the observed DNA conformational change in the center of the binding site (23). The PUT3–DNA crystal structure has extensive contacts between a β strand from one asymmetrically bound protein subunit and DNA minor groove, resulting in a DNA kink in the middle of the binding site (22). The additional specificity resulting from these contacts is evident from the Gibbs prediction (Fig. 1) and was previously discovered by Siddharthan *et al.* (17) using a phylogenetic approach. PPR1 activates the transcription of genes involved in regulation of pyrimidine levels; it binds extended TCGGN₆CCGA sites: van der Waals contacts are made to the bases flanking the canonical CGG triplets (24). Surprisingly, we could not find many PPR1 sites in the ChIP-chip intergenic regions for which the binding was reported. There is no overlap between ChIP-chip promoter sequences and those of nine genes likely to be involved in the pyrimidine pathway (*URA1–URA8*, *URA10*). Using the latter set, we found canonical sites in seven of these sequences ($S_{\text{track}} = 0.68$; *SI Table 4*). Finally, LEU3 is the only Zn2-Cys6 TF with a native structure to bind an everted repeat: CCGN₄CGG. The Gibbs search refines the structure-based model to reveal preference for TT in the middle of the binding site (Fig. 1).

Homology Modeling of the ARG80–ARG81–MCM1 Regulatory Complex. We predicted 14 additional Zn2-Cys6 PWMs by homology (see *SI Table 3*). For most of these proteins, the homolog has a different spacing between its monomeric half-sites, and the structure-based prior has to be modified accordingly (see *SI Table 5*). We either obtain information about the binding site length from the literature or explore a range of half-site arrangements. We illustrate our procedure with ARG81, which coordinates the expression of arginine metabolic genes and binds DNA as an ARG80–ARG81–MCM1 complex with the MADS box proteins ARG80 and MCM1

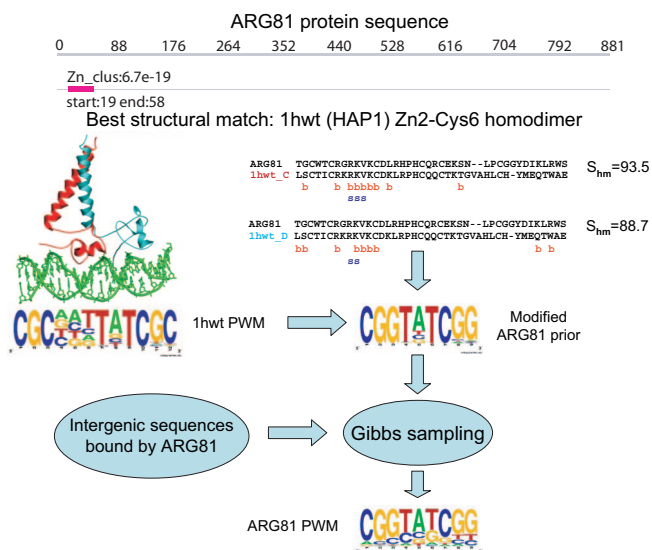


Fig. 2. Illustration of how structural and sequence data are mined in the case of ARG81. A DNA-binding domain of the Zn2-Cys6 binuclear cluster type is found in the ARG81 protein sequence. The HAP1 homodimer (PDB code 1hw1) is identified as the homolog with the highest interface scores S_{hm} (93.5 for chain C, 88.7 for chain D). The interface scores reflect the similarity of the HAP1 and ARG81 DNA-binding interfaces on the basis of their protein sequence alignments. Interface amino acids are labeled “b” for the DNA phosphate backbone contacts and “s” for the DNA base contacts. Observed amino acid mutations at the interface are sufficiently conservative and thus are assumed not to change the binding specificity significantly. However, to approximate previously characterized ARG81 binding sites (26), columns 4–6 are removed from the HAP1 PWM, and the CGC half-sites are replaced by the more common CGG half-sites. The 1hw1-based PWM modified in this way is used as the informative prior for the Gibbs sampling algorithm, which is run on the intergenic sequences known to be bound by ARG81 from the ChIP-chip experiment (2). After the ARG81 sites are identified, their alignment is used to compile the ARG81 PWM. Each site in the alignment is weighted by its posterior probability $p(s, c)$ (>0.05).

(25). Our structure-based approach allows us to differentiate among regulatory inputs from the different TFs involved in the complex.

We modeled ARG81 by using the HAP1 structure but had to make the dimer spacing consistent with the known literature sites (26) (Fig. 2). Because the reported binding sites were not delineated clearly enough to deduce the spacing unambiguously, we used alternative models with 3, 4, and 5 bp between the half-sites. The prior with 3 bp was chosen because it yielded most sites with Gibbs sampling, although in principle multiple binding modes are possible. The best model for ARG80 is in fact MCM1, which was crystallized in the MAT α 2–MCM1–MAT α 2 complex with DNA [Protein Data Bank (PDB; www.rcsb.org/pdb/home/home.do) code 1mmj]. Although the binding interface is almost completely conserved between MCM1 and ARG80, ARG80 has weaker *in vitro* binding affinity for the canonical MCM1 P-box site, CC(A/T)₆GG (27). This is attributed to the I21Q and K40R mutations in ARG80, which we classify as phosphate backbone contacts. Because of this interface similarity, we cannot differentiate between ARG80 and MCM1 binding sites. Indicative of the ARG80–ARG81–MCM1 complex formation, ARG80 and ARG81 bind some of the same intergenic regions (Table 1). Consequently, we are able to discover ARG80/MCM1 and ARG81 sites in both sets of intergenic sequences (Table 2). Furthermore, several composite sites are spaced similarly to known ARG80–ARG81–MCM1 sites (two P-boxes with the ARG81 site in between) (25).

For the set of intergenic sequences bound by ARG81, MacIsaac *et al.* (5) predict GCN4 binding specificity (Table 2). This is reasonable biologically, given that GCN4 is a master regulator of

Table 1. Partial overlaps between two sets of intergenic regions bound by the members of multiprotein complexes

Experiment 1	Experiment 2	N_1^+	N_2^+	N_{12}^+	$\langle N_{1+} \rangle^S$	$\langle N_{+2} \rangle^{\dagger 1}$
ARG80 (YPD)	ARG81 (SM)	29	27	6	0.81	0.48
	GCN4 (SM)	29	143	5	0.81	1.69
INO2 (YPD)	INO4 (YPD)	35	32	14	0.40	0.48
CBF1 (SM)	MET28 (YPD)	195	16	0	1.88	0.14
	MET4 (SM)	195	37	8	1.88	0.70
YOX1 (YPD)	YHP1 (YPD)	61	18	2	2.64	0.15

All intergenic regions are classified by Harbison *et al.* (2) as bound at the $P < 0.001$ confidence level.

[†]Number of intergenic regions bound by factor 1 (factor 2) in the ChIP-chip experiment.

[‡]Number of intergenic regions bound by both factors.

[§]Average number of intergenic regions shared between factor 1 and all other factors.

^{††}Average number of intergenic regions shared between factor 2 and all other factors.

amino acid biosynthesis (25, 26), but unreasonable structurally because GCN4 is a leucine zipper homodimer that binds symmetric AP-1 sites [ATGA(C/G)TCAT] *in vivo* (28). GCN4 binding specificity is inconsistent with the prior for ARG81, and the length of the AP-1 site is very different from what is expected for a protein complex. We find GCN4 sites in both ARG80- and ARG81-bound promoter sequences by using a prior based on the structure of GCN4 bound to the AP-1 site (PDB code 1ysa) (Table 2).

Identical Binding Specificities of TF Heterodimers. TFs known to function as heterodimers should bind the same intergenic regions and exhibit identical specificities. For example, the helix–loop–helix proteins INO2 and INO4 form a heterodimer involved in derepression of phospholipid biosynthesis genes in response to inositol deprivation (29). Consistent with the heterodimer formation, we find that 14 of the 35 intergenic regions bound by INO2 are also bound by INO4, many more than expected by chance (Table 1). The INO2–INO4 complex is homologous to the Myc–Max complex bound to the E-box (CACGTG) site (30). Using the Myc–Max-based homology model, we find many E-box sites in the sequences bound by either INO2 or INO4, in agreement with previous studies (2, 5) (see Table 2 and [SI Table 4](#)).

Table 2. Summary of DNA-binding specificity predictions for the multiprotein complexes in Table 1

Prior [†]	Experiment [‡]	$S_{\text{track}}^{\S}$	$\rho_{\text{MF}}^{\parallel}$
ARG80	ARG80 (YPD)	0.41	0.90
	ARG81 (SM)	0.22	0.80
ARG81	ARG80 (YPD)	0.65	0.54
	ARG81 (SM)	0.66	0.63
GCN4	ARG80 (YPD)	0.36	0.028
	ARG81 (SM)	0.30	9.1×10^{-6}
INO2-INO4	INO2 (YPD)	0.71	7.7×10^{-8}
	INO4 (YPD)	0.80	1.7×10^{-7}
CBF1	CBF1 (SM)	0.98	1.1×10^{-13}
	MET28 (YPD)	0.29	0.29
	MET4 (SM)	0.91	0.65
YOX1-MCM1-YOX1	YOX1 (YPD)	0.74	2.5×10^{-4}
YHP1-MCM1-YHP1	YHP1 (YPD)	0.45	0.39

[†]TF for which the structure-based informative prior was constructed.

‡Set of intergenic regions (identified by the bound factor and the environmental condition) for which Gibbs sampling with the informative prior was carried out.

⁵Motif quality as measured by the PhyloGibbs posterior probability $p(s, c)$ averaged over the top 10 sites (all sites if <10 sites were tracked; S_{track} close to 1.0 indicates well defined motifs).

[†]Probability that the optimal overlap between our PWM and that of MacIsaac *et al.* (5) is due to chance (see *Materials and Methods*).

Dual Mechanism of Gene Regulation by a Helix–Loop–Helix TF. The helix–loop–helix protein TYE7 is highly homologous to the human sterol regulatory element (StRE; ATCACCCCAC) binding protein (SREBP-1a; PDB code 1am9). The SREBP-1a binding specificity is attributed to the asymmetric homodimer conformation and to the ARG→TYR mutation at the binding interface (31). However, because SREBPs also bind E-boxes *in vitro* with comparable affinity (32), it is conceivable that their *in vivo* function is mediated by both types of sites. The E-box motif is dominant in TYE7 sequences: it is found by using either the vertebrate Max homodimer bound to the E-box (PDB code 1an2) or the SREBP-1a bound to StRE (PDB code 1am9) as the informative prior ($P = 1.2 \times 10^{-5}$). This is not surprising because 28 of the 56 intergenic sequences bound by TYE7 are also bound by another helix–loop–helix factor, CBF1. The CBF1 homodimer is a homolog of Max (1an2) and thus can be expected to bind E-boxes. However, using both 1an2 and 1am9 informative priors in a single Gibbs sampling run reveals a secondary StRE motif ($S_{\text{track}} = 0.80$). Thus, TYE7 may bind E-box sites in complex with CBF1 and may bind both E-box and StRE sites as a homodimer.

Related Binding Specificities in the Leucine Zipper Family. We made PWM predictions for 14 TFs in the leucine zipper family. Our analysis shows that many of the binding specificities in this family are similar, and thus only four structural templates are required to model all leucine zippers in yeast (see SI Table 4). As for the Zn2-Cys6 binuclear cluster family, the spacing between the monomeric half-sites is variable. Inasmuch as GCN4 is known to have comparable *in vitro* binding affinity for AP-1 (ATGA(C/G)TCAT) and ATF/CREB (ATGACGTCAT) sites (33), it is likely that the other leucine zippers are also capable of binding both types of sites. For example, all transcriptional regulators in the YAP family (CAD1, CIN5, YAP1, YAP3, YAP5, YAP6, YAP7, and ARR1) (34) are homologous to the PAP1 TF from *Schizosaccharomyces pombe* complexed with the GTTACGTAAC PAP1 site (PDB code 1gd2) (35). Similarly to GCN4, PAP1 homodimers are known to recognize shorter [GTTA(C/G)TAAC] and longer (GTTACGTAAC) sites *in vitro* (34, 35). Using both types of priors, we find more shorter sites for CAD1, YAP1, YAP3, YAP5, and ARR1 and more longer sites for CIN5 and YAP6. For YAP7, there is comparable evidence for both types of sites (SI Table 4), of which MacIsaac *et al.* (5) find only the shorter. Interestingly, only half of the palindromic site is strongly conserved in YAP5, suggesting a contribution from monomeric binding or cofactors. This notion is supported by the significant overlap of the YAP5-bound intergenic regions and those bound by PDR1 and GAT3 (data not shown).

Indirect Recruitment of MET28 and MET4. Two leucine zipper proteins, MET28 and MET4, form regulatory complexes with either the helix–loop–helix protein CBF1 or the highly related zinc finger proteins MET31 and MET32 (36, 37). In particular, the CBF1–MET28–MET4 complex acts as a transcriptional activator in the sulfur–amino acid metabolism and biosynthesis pathway. The MET4 sequence is very diverged (with an e-value of 0.0017 for a match to the bZIP₂ family and no homologous structures), whereas the MET28 DNA-binding interface is reasonably similar to GCN4. Nonetheless, it is believed that neither MET28 nor MET4 interacts with DNA directly (37, 38), being recruited instead through association with CBF1, MET31, or MET32. MacIsaac *et al.* (5) assigned MET31/32 binding specificity to MET4 sequences (even though MET4 is a leucine zipper); besides the MET31/32 motif, we also find CBF1 E-box sites in MET4 sequences (Tables 1 and 2).

Surprisingly (and contrary to previous studies), we could not find strong evidence supporting CBF1 and MET31/32 binding in MET28 sequences (data not shown and Tables 1 and 2). MacIsaac *et al.* (5) report a MET31/32 site from the literature for MET28. It is possible that MET28 participates in pathways other than sulfur metabolism by forming complexes with other factors.

Cooperative Binding of the Homeodomain TFs. Another example of synergistic TF action involves proteins in the homeodomain family, which often increase their specificity by forming homo- and heterodimers and interacting with cofactors (39). Our results suggest that all yeast homeodomains (MAT α 2, CUP9, PHO2, YHP1, and YOX1) employ this strategy to some extent. In particular, CUP9 is homologous to the extradenticle (Exd) homeodomain from the Ubx–Exd–DNA complex in *Drosophila melanogaster* (40). Even though CUP9 shares lower interface similarity with Ubx (see SI Table 4), using the whole complex yields a PWM corresponding to the CUP9 homodimeric binding mode. Similarly, PHO2 is homologous to the homeodomain homodimer from the *Drosophila* protein paired (Pax class) (41). Consistent with this observation and previous footprinting studies (42), we find PHO2 dimeric sites in the intergenic sequences (SI Table 4). Neither complex was found in the previous study (5).

YOX1 and YHP1 are the transcriptional repressors that, together with MCM1, bind early cell cycle boxes found in the promoters of genes expressed during the M/G₁ phase of the cell cycle (43). We used the MAT α 2–MCM1–MAT α 2 complex as a prior for both YOX1 and YHP1 bound sequences and found matches in more than half the genes. A previously found YOX1 site (5) corresponds to just the YOX1–MCM1 part of the larger complex, whereas for YHP1, a standard TAAT monomeric homeodomain site from the literature is reported (5).

Inferring the Identity of Regulatory Inputs with the Structure Database. Our method is also useful for associating TFs with independently discovered sequence motifs. By correlating 49 PWMs from MacIsaac *et al.* (5) for which we could build homology models, and 252 structure-based TF binding specificity predictions, we were able to assign the correct protein fold in 25 cases [we considered a fold to be predicted correctly if the P value for the correlation between the MacIsaac *et al.* (5) PWM and at least one PWM from the correct fold was among the lowest three]. The correct fold had the lowest P value in 16 of 25 instances. In 10 cases, the actual structural template used in homology modeling was among the top three (this measure is less objective because in many cases several homologs are equally plausible). Some of the failures are clearly due to the misassociation of TFs and sets of bound intergenic regions, such as the GCN4 motif reported for ARG81 or the MCM1 motif reported for YOX1 (5). Other motifs are missed because dimers with correct spacings are absent in the structural database. Indeed, by adding 24 PWMs with adjusted spacings (see SI Table 5) we were able to place 28 of 49 motifs into the correct fold and to identify 14 structural templates correctly. Thus, we have developed a computational equivalent of the yeast one-hybrid assay, a useful tool for postprocessing motifs found in microarray data analyses.

Discussion

We have developed a computational approach that uses structure-based TF binding specificity predictions as priors in the Bayesian Gibbs sampling algorithm (a related method was recently described in refs. 44 and 45). Structure-based PWM predictions are based on the observations (*i*) that for most TFs, amino acid conservation at the DNA-binding interface leads to similar specificities and (*ii*) that the number of atomic contacts with the consensus base is a good predictor of the degree of its conservation in the binding site (13). Thus, a simple probabilistic model based on the number of protein–DNA atomic contacts can be used to make PWM predictions by homology. The structural predictions are subsequently refined with sequence data from the ChIP-chip experiments. The sequence-structure approach is computationally inexpensive and thus should be used as a standard bioinformatics tool. If sequences from several species are available, the structural priors can be easily combined with phylogenetic footprinting. Our approach is not limited to the relatively simple structure-based models used in this work: any prior specificity information, including more sophisti-

cated structural predictions and experimentally inferred PWMs, can be used as input to the Gibbs sampling algorithm.

We have observed limited diversity of TF binding specificities: DNA-binding domains with 30–50% overall sequence identity routinely bind similar sites and share high interface homology. Thus, to obtain comprehensive structural coverage for the analysis of transcriptional regulation in any organism, it is sufficient to have just one representative protein–DNA structure for each distinct binding specificity subclass in every TF family. This notion extends to protein–DNA complexes the basic principle of structural genomics projects: at least one structure should be solved for every protein fold.

SI Fig. 4 summarizes our predictions for 67 *S. cerevisiae* TFs and contrasts them with prior work employing phylogenetic footprinting (2, 5). Eighteen of our predictions were not found in the previous study (5), and 26 of the 49 remaining predictions disagree with it either partially (e.g., significantly differing in length) or completely (SI Fig. 4 and SI Table 4). Phylogenetic conservation alone cannot provide a direct link between the physicochemical properties of the TF–DNA complex and the sequence-derived motifs, resulting in cases of “mistaken identity” (e.g., GCN4 PWM assigned to ARG81 and ARG80; Table 2) and incorrectly assigned specificity (YOX1–MCM1 PWM reported as YOX1; ARO80 PWM given as three monomeric sites rather than a dimer).

It is interesting to note that sometimes we do not find sites with expected specificity in the ChIP-chip intergenic sequences, even if our confidence in the informative prior is high (e.g., for PPR1 and NDT80, for which co-crystal structures are available). This can happen if a protein associates with DNA through cofactors (e.g., MET4 in the CBF1–MET28–MET4 complex) or if the binding is dominated by relatively weak, nonspecific interactions. Finally, we note that in all likelihood only a fraction of sites reported here are functional *in vivo*. Additional filters based on phylogenetic conservation (2, 5) and the proximity to the coding regions and cofactor sites can easily be used in specific situations.

Materials and Methods

Structure-Based PWM Predictions. We have built a database of 515 protein structures bound to double-stranded DNA, including 252 transcription factors from 40 families. We predicted PWMs for all protein–DNA complexes in this database by using a simple model that exploits the structure of the protein–DNA complex but does not require any detailed predictions of the protein–DNA energetics (14). The performance of this model was previously found to be only slightly inferior to that of more sophisticated but less computationally efficient all-atom models (14, 15). We construct a PWM by using the consensus DNA sequence and the number of atom–atom contacts, N_i , between all protein side chains and the DNA bases in the base pair i (we use a 4.5 Å distance cutoff; hydrogen atoms are excluded from the counts). We assume that the three nonconsensus bases occur with equal probabilities and that the consensus base is favored over a nonconsensus base by $N_i/N_{\max}$ ($0 \leq N_i \leq N_{\max}$):

$$w_i^\alpha(N_i) = \begin{cases} \frac{1}{4} (1 - N_i/N_{\max}) & (\alpha \neq \text{consensus}) \\ \frac{1}{4} (1 + 3N_i/N_{\max}) & (\alpha = \text{consensus}) \end{cases}, \quad [1]$$

where $w_i^\alpha(N_i)$ is the probability of base $\alpha = \{A, C, G, T\}$ in the PWM column i , N_i is the number of protein atoms in contact with the base pair i , and $N_{\max}$ is the number of contacts at which the native base pair becomes absolutely conserved: $w_i^{\text{consensus}}(N_i) = 1$, $N_i \geq N_{\max}$. Note that if $N_i = 0$, all four bases are equally likely. $N_{\max}$ is a free parameter of the model; its optimum value was found to be 20 by fits to experimental data (14). Fig. 3 illustrates our method, using PHO4 as an example.

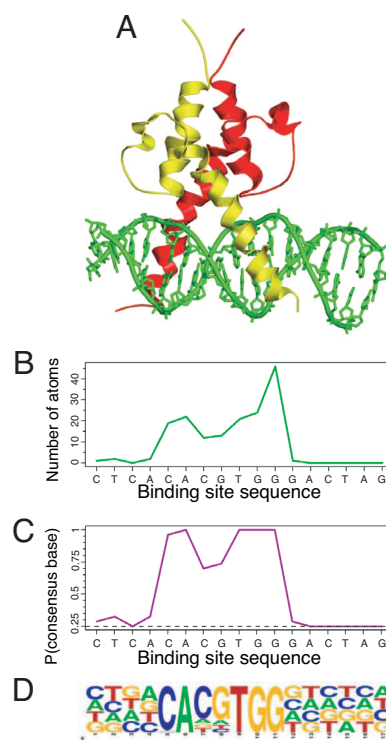


Fig. 3. Prediction of the informative prior for the phosphatase system regulator PHO4. (A) Crystal structure of the PHO4 helix-loop-helix dimer bound to its consensus site (PDB code 1a0a). (B) Atomic profile: the number of heavy atoms, N_i , within 4.5 Å of base pair i in the binding site. (C) Consensus base probability profile: the probability $w_i^\alpha(N_i)$ of the consensus base α at position i in the binding site (cf. Eq. 1). (D) Structure-based PWM prediction.

TF Homology Modeling. For each structure, we found matches to the protein families in the Pfam database (see SI Materials and Methods) and identified amino acids at the DNA-binding interface by using a distance cutoff of 4.5 Å to a DNA base pair. We classified all contacts as DNA base and/or DNA phosphate backbone/sugar ring (a given amino acid can make both types of contacts with one or several base pairs). We then searched all *S. cerevisiae* proteins for matches to those Pfam families with at least one hit in the structural database. This procedure yielded both a Pfam classification of yeast protein factors and the alignments of their sequences with the putative structural homologs, providing information about the amino acids in contact with DNA.

For each sequence-structure protein alignment, we computed the amino acid substitution score S_{hm} at the DNA-binding interface. Given an alignment of the query protein sequence with the target sequence from the protein–DNA structural database, the protein–DNA interface score is given by $S_{\text{hm}} = 0.5 S_{\text{base}} + 0.5 S_{\text{bb}}$, where S_{base} and S_{bb} are the DNA base and DNA backbone substitution scores, respectively. These scores are defined as:

$$S_{\text{base/bb}} = \frac{1}{N_{\text{cont}}} \sum_{\text{contacts}} s(\text{aa}_1, \text{aa}_2), \quad [2]$$

where the sum is over N_{cont} amino acid–DNA contacts (base contacts for S_{base} , backbone contacts for S_{bb}), aa_1 is the amino acid in the query sequence, aa_2 is the amino acid from the target sequence aligned with aa_1 and in contact with DNA, and $s(\text{aa}_1, \text{aa}_2)$ is the PET91 amino acid substitution score (46). Note that the amino acids contacting multiple bases will make a proportionately greater contribution to Eq. 2. The protein–DNA interface score defined in this way has a range 0–100, with 100 assigned when both

DNA base- and DNA backbone-contacting amino acids are fully conserved in the two-sequence alignment.

In most cases, the structure with the highest interface score was chosen as the structural template. If several structures had comparable homology scores, we chose either the most accurate one (using measures such as resolution of x-ray diffraction) or the one most relevant in the biological context (using information about cofactors and the dimerization state). Computing the score only from amino acids that contact DNA, rather than from entire aligned sequences, assumes that amino acid–DNA interactions are local: if the amino acids at the DNA-binding interface are conserved between two protein–DNA complexes, they will adopt similar geometric arrangements with respect to DNA, regardless of the rest of the protein (7, 47, 48). For example, a comparison of the engrailed and $\alpha 2$ homeodomain–DNA complexes revealed an extensive set of conserved contacts with DNA, even though the amino acid sequences were only 27% identical (7). A more recent study (48) identified a number of cases in which the local interface geometry was conserved, even if DNA conformational change was required in order to accommodate it.

Informative Priors. In a majority of cases, we modeled only those yeast factors for which a protein–DNA complex with an interface homology score S_{hm} exceeding the empirical cutoff of 80 could be found (in multimeric complexes, S_{hm} was averaged over all protein chains). Typically, the corresponding structure-based PWM was used to create the informative prior without further modification. We discarded all cases in which nonconservative amino acid mutations lowered our confidence in the homology template. More structures could be modeled if we were able to predict the new specificity with accurate descriptions of protein–DNA energetics. We also avoided updating the structure-based PWMs by using the information about interface mutations, because predicting the number of atoms in contact with DNA bases would require explicit modeling of mutated side chains. In some cases, for TFs that bind as dimers the spacing and relative orientation of the monomeric subsites were adjusted based on the information about previously characterized factor binding sites. Finally, the structure-based PWMs were multiplied by the total number of pseudocounts $\bar{n}$ and used as the informative priors in the Bayesian Gibbs sampling algorithm (17). We found empirically that setting $\bar{n} = n/2$ (where n is the number of input intergenic sequences that approximates the expected number of binding sites) biases the search fairly strongly toward the expected binding sites but at the same time is weak enough so that the final PWM prediction can be completely different if the genomic sites disagree with the prior. Fig. 2

demonstrates our method for predicting TF binding specificities with sequence and structural data in the case of ARG81.

Gibbs Sampling. We use the PhyloGibbs implementation of the Gibbs sampling algorithm because its Bayesian formulation allows us to take the prior information into account in a straightforward way (17). PhyloGibbs assigns configurations C to the input sequence S ; each configuration consists of the nonoverlapping TF binding sites and the background (modeled with a first-order Markov model) and can have multiple site instances for each TF. PhyloGibbs uses simulated annealing to find the configuration C^* with the highest posterior probability, $P(C^*|S)$ (see *SI Materials and Methods*). Once the optimal configuration C^* is identified, the algorithm samples the distribution $P(C|S)$ without restrictions to estimate the probability $p(s, c)$ that a site s belongs to a TF c in C^* (17). After this so-called tracking phase, all sites s for which $p(s, c) \geq 0.05$ are reported for each TF in the configuration C^* . Tracking is used as a convenient means of summarizing the information contained in $P(C|S)$ and is able to both discover additional sites and remove spurious sites from the reference configuration C^* .

Sequence Data. We ran PhyloGibbs in the tracking mode on the intergenic sequences from the ChIP-chip experiments in yeast (factors bound with $P < 0.001$) (2) and from the literature (see *SI Materials and Methods*). Each run of the algorithm is used to infer a single PWM, using only *S. cerevisiae* intergenic sequences bound by a specific TF and the informative prior as inputs. In several cases, when none of the sites tracked sufficiently well, we reported sites from the simulated annealing configuration. In this case, the number of sites must be guessed in advance (we assume one site per promoter sequence). Gibbs sampling PWMs are inferred from the alignments of binding sites weighted by $p(s, c)$ (if available).

PWM Correlations. For each forward and reverse complement alignment between the two PWMs, we computed Pearson's correlation coefficients for the PWM log-probabilities (15). To avoid spurious short alignments, we set the minimum allowed PWM overlap to the length of the shorter PWM, minus a constant offset (2 for GATA factors, 4 for other PWMs). We reported the overlap with the lowest P value (Bonferroni-corrected for the number of tested alignments).

We thank Erik van Nimwegen, Harmen Bussemaker, Jonathan Widom, and Amos Tanay for comments on the manuscript and Michael Mwangi and Rahul Siddharthan for useful discussions. A.V.M. was funded by The Lehman Brothers Foundation through the Leukemia and Lymphoma Society; E.D.S. was supported by National Science Foundation Grant DMR-0129848.

- Ren B, Robert F, Wyrick JJ, Aparicio O, Jennings EG, Simon I, Zeitlinger J, Schreider J, Hannett N, Kanin E, et al. (2000) *Science* 290:2306–2309.
- Harbison CT, Gordon DB, Lee TI, Rinaldi NJ, MacIsaac KD, Danford TW, Hannett NM, Tagne JB, Reynolds DB, Yoo J, et al. (2004) *Nature* 431:99–104.
- Mukherjee S, Berger MF, Jona G, Wang XS, Muzzey D, Snyder M, Young RA, Bulky ML (2004) *Nat Genet* 36:1331–1339.
- Liu X, Noll DM, Lieb JD, Clarke ND (2005) *Genome Res* 15:421–427.
- MacIsaac KD, Wang T, Gordon DB, Gifford DK, Stormo GD, Fraenkel E (2006) *BMC Bioinformatics* 7:113.
- Luscombe NM, Austin SE, Berman HM, Thornton JM (2000) *Genome Biol* 1:R001.
- Pabo CO, Sauer RT (1992) *Annu Rev Biochem* 61:1053–1095.
- Sandelin A, Wasserman WW (2004) *J Mol Biol* 338:207–215.
- Xing EP, Karp RM (2004) *Proc Natl Acad Sci USA* 101:10523–10528.
- Mahony S, Golden A, Smith TJ, Benos PV (2005) *Bioinformatics* 21(Suppl 1):i283–i291.
- MacIsaac KD, Gordon DB, Nekudova L, Odum DT, Schreiber J, Gifford DK, Young RA, Fraenkel E (2006) *Bioinformatics* 22:423–429.
- Kummerfeld SK, Teichmann SA (2006) *Nucleic Acids Res* 34:D74–D81.
- Mirny LA, Gelfand MS (2002) *Nucleic Acids Res* 30:1704–1711.
- Morozov AV, Havranek JJ, Baker D, Siggia ED (2005) *Nucleic Acids Res* 33:5781–5798.
- Foat BC, Morozov AV, Bussemaker HJ (2006) *Bioinformatics* 22:e141–e149.
- Lawrence CE, Altschul SF, Boguski MS, Liu JS, Neuwald AF, Wootton JC (1993) *Science* 262:208–214.
- Siddharthan R, Siggia ED, van Nimwegen E (2005) *PLoS Comput Biol* 1:e67.
- Benos PV, Lapedes AS, Stormo GD (2002) *J Mol Biol* 323:701–727.
- Kaplan T, Friedman N, Margalit H (2005) *PLoS Comput Biol* 1:e1.
- Schwabe JWR, Rhodes D (1997) *Nat Struct Biol* 4:680–683.
- King DA, Zhang L, Guarente L, Marmorestein R (1999) *Nat Struct Biol* 6:64–71.
- Swaminathan K, Flynn P, Reece RJ, Marmorestein R (1997) *Nat Struct Biol* 4:751–759.
- Marmorestein R, Carey M, Ptashne M, Harrison SC (1992) *Nature* 356:408–414.
- Marmorestein R, Harrison SC (1994) *Genes Dev* 8:2504–2512.
- Messenguy F, Dubois E (2000) *Food Tech Biotechnol* 38:277–285.
- De Rijcke M, Seneca S, Punyamalee B, Glansdorff N, Crabeel M (1992) *Mol Cell Biol* 12:68–81.
- Jamai A, Dubois E, Vershon AK, Messenguy F (2002) *Mol Cell Biol* 22:5741–5752.
- Garcia-Gimeno MA, Struhl K (2000) *Mol Cell Biol* 20:4340–4349.
- Robinson KA, Lopes JM (2000) *Nucleic Acids Res* 28:1499–1505.
- Nair SK, Burley SK (2003) *Cell* 112:193–205.
- Párraga A, Bellolell L, Ferré-D'Amaré AR, Burley SK (1998) *Structure (London)* 6:661–672.
- Kim JB, Spotts GD, Shih H, Halvorsen Y, Ellenberger T, Towle HC, Spiegelman BM (1995) *Mol Cell Biol* 15:2582–2588.
- Keller W, König P, Richmond TJ (1995) *J Mol Biol* 254:657–667.
- Fernandes L, Rodrigues-Pousada C, Struhl K (1997) *Mol Cell Biol* 17:6982–6993.
- Fujii Y, Shimizu T, Toda T, Yanagida M, Hakoshima T (2000) *Nat Struct Biol* 7:889–893.
- Kuras L, Barbey R, Thomas D (1997) *EMBO J* 16:2441–2451.
- Blaiseau P-L, Thomas D (1998) *EMBO J* 17:6327–6336.
- Kuras L, Cherest H, Surdin-Kerjan Y, Thomas D (1996) *EMBO J* 15:2519–2529.
- Wilson DS, Desplan C (1999) *Nat Struct Biol* 6:297–300.
- Passner JM, Ryoo HD, Shen L, Mann RS, Aggarwal AK (1999) *Nature* 397:714–719.
- Wilson DS, Guenther B, Desplan C, Kurivan J (1995) *Cell* 82:709–719.
- Vogel K, Hörz W, Hinnen A (1989) *Mol Cell Biol* 9:2050–2057.
- Pramila T, Miles S, GuhaThakurta D, Jemilo D, Breeden LL (2002) *Genes Dev* 16:3034–3045.
- Narlikar L, Hartemink AJ (2006) *Bioinformatics* 22:157–163.
- Narlikar L, Gordan R, Ohler U, Hartemink AJ (2006) *Bioinformatics* 22:e384–e392.
- Jones S, van Heyningen P, Berman HM, Thornton JM (1999) *J Mol Biol* 287:877–896.
- Pabo CO, Nekudova L (2000) *J Mol Biol* 301:597–624.
- Siggers TW, Silkov A, Honig B (2005) *J Mol Biol* 345:1027–1045.