

Twenty years of advances in prediction of nucleic acid binding residues in protein sequences

Sushmita Basu^{1*}, Jing Yu¹, Daisuke Kihara^{2,3}, and Lukasz Kurgan^{1*}

¹Department of Computer Science, Virginia Commonwealth University, Richmond, VA 23284, USA

²Department of Biological Sciences, Purdue University, West Lafayette, IN 47907, USA

³Department of Computer Science, Purdue University, West Lafayette, IN 47907, USA

*Corresponding authors: basudass@vcu.edu (SB); lkurgan@vcu.edu (LK)

Address: Department of Computer Science, Virginia Commonwealth University, 401 West Main Street, Room E4225, Richmond, Virginia 23284, USA; Phone: (804) 827-3986

Abstract

Computational prediction of nucleic acid binding residues in protein sequences is an active field of research, with over 80 methods that were released in the last two decades. We identify and discuss 87 sequence-based predictors that include dozens of recently published methods that are surveyed for the first time. We overview historical progress and examine multiple practical issues that include availability and impact of predictors, key features of their predictive models, and important aspects related to their training and assessment. We observe that the last decade has brought increased use of deep neural networks and protein language models, which contributed to substantial gains in the predictive performance. We also highlight advancements in vital and challenging issues that include cross-predictions between DNA and RNA-binding residues and targeting the two distinct sources of binding annotations, structure-based vs. intrinsic disorder-based. The methods trained on the structure-annotated interactions tend to perform poorly on the disorder-annotated binding and vice versa, with only a few methods that target and perform well across both annotation types. The cross predictions are a significant problem, with some predictors of DNA-binding or RNA-binding residues indiscriminately predicting interactions with both nucleic acid types. Moreover, we show that methods with web servers are cited substantially more than tools without implementation or with no longer working implementations, motivating the development and long-term maintenance of the web servers. We close by discussing future research directions that aim to drive further progress in this area.

Keywords: protein-DNA interaction; protein-RNA interaction; nucleic acid binding; DNA-binding residue; RNA-binding residue; intrinsic disorder; sequence-based prediction; machine learning; deep learning

Introduction

Interactions between biomolecules are major drivers of cellular processes. In particular, protein-nucleic acids interactions play crucial roles in a number of key cellular functions

including DNA replication, transcription, translation, and gene regulation [1-5]. Moreover, their mis-regulation is associated with several human diseases [6-8], providing further motivation to study these interactions. Several experimental techniques, such as pull-down assays, chromatin immunoprecipitation and CRISPR-Cas based approaches are used to study the protein-nucleic acids interactions [3, 9]. Additionally, atomic-level details that include information about the interacting residues and nucleotides can be obtained from experimentally solved structures of protein-nucleic acids complexes [10, 11]. However, obtaining structures of these complexes is relatively cost-intensive and challenging, especially when considering that nearly 250 million protein sequences are available as of September of 2024, with 90 million that were collected over the last five years [12]. To this point, computational methods that predict protein-nucleic acids interactions from sequences can be beneficial in bridging this large and growing function annotation gap. These computational tools are typically trained/generated from the limited amounts of the experimentally annotated data. The trained models can be used to produce predictions in a high-throughput manner for a large number of uncharacterized proteins, if their predictive performance is sufficiently good. The functional importance of protein-nucleic acids interactions combined with the availability of a sufficient amount of the corresponding experimental data for the model training and testing has motivated the development of several dozens of computational predictors of nucleic acid binding residues in proteins.

These predictors can be divided into two groups, those that use the protein structure vs. protein sequence as the input. The structure-based methods exploit the structural features, such as secondary structure, solvent accessibility, characteristics of spatial neighborhoods in the structure and shape complementarity, to derive predictive models [13-20]. Some of the popular and recent structure-based methods include (chronologically) aaRNA [17], NucleicNet [18], Graphbind [19], Geobind [20] and PNAbind [21]. With the availability of a large number of structures predicted with AlphaFold2 [22], some of the recent methods, such as EquipNAS [23], GraphSite [24] and GraphNABP [25], utilize putative structures to train their models. In this survey, we focus the sequence-based predictors that use only the protein sequences as inputs to predict the DNA-binding residues (DBRs) and the RNA-binding residues (RBRs). Since the release of the first sequence-based predictors in 2004 [26, 27], many more methods were published [28-34]. We identified over two dozen sequence-based predictors that were released in the last five years, demonstrating that this is an active field of research. Availability of a large number of methods prompted the release of several surveys that we summarize in Table 1. These reviews typically enumerate the available predictors, list the corresponding references, summarize their predictive models, and provide guidance on how to select an appropriate method from a pool of multiple choices. They usually cover dozens of methods that include tools that predict DBRs [29, 31, 35, 36], RBRs [28, 34, 37-39] and both DBRs and RBRs [30, 33, 40, 41].

Most of the surveyed predictors are trained on training datasets using machine learning algorithms. Based on the source of the binding annotations in the training datasets, they are divided into two categories. The first includes data derived from the structures of protein-nucleic acids complexes, which we refer to as structure-annotated training data. The second category involves binding interfaces that are positioned in intrinsically disordered regions (IDRs), resulting in the disorder-annotated training data. IDRs lack stable structure under

physiological conditions [42-45]. Protein sequences may have one or multiple IDRs, which in some cases may cover an entire sequence

Several studies show that IDRs are functionally important and abundant in the nucleic acid binding proteins [46-51]. Importantly, protein-nucleic acids interactions in IDRs differ in multiple ways from the interacting structured regions. The former typically have polymorphic conformations that are induced by interacting with different ligands [52, 53]. Moreover, they are enriched in disorder promoting amino acids and form larger interfaces upon binding with partner molecules when compared with the interfaces in the structured regions [54-56]. These differences may impact ability of the corresponding methods to make accurate predictions, especially when they are applied to make predictions in IDRs while the underlying model was trained on the structure-annotated dataset, and vice versa. With that in mind, Table 1 reveals that virtually all past surveys focus on the methods that were trained on the structure-annotated proteins [28, 29, 31, 33-39, 41], with just two articles that consider tools that were trained from the disorder-annotated proteins [30, 40]. Moreover, one of these two articles covers just two predictors trained from the disorder-annotated data and discusses methods that targets interactions with other ligands, such as proteins and lipids [40], and the other mentions one disorder-trained predictor without discussing this aspect of the model training [30].

Another important aspect is cross-prediction between DBRs and RBRs. Research show that some of the current methods heavily cross-predict (mis-predict) DBRs as RBRs and vice versa [33, 41]. This issue was sporadically discussed in some of the surveys [28, 33, 34, 41]. The first independent survey covering the cross-prediction was authored by Zhao *et. al.* [28] (Table 1), while the first prediction tool that assessed the cross-predictions was SPOT-Seq [57]. However, the cross-prediction analysis in both studies was limited only to the predictors of the RNA-binding proteins, which produce protein-level results that do not include prediction of RBRs. The most recent assessment, conducted in 2020, performed a residue-level analysis on the cross-prediction but it covered only predictors of RBRs [34]. We find that since the study by Yan *et. al.* [33], none of the surveys in the last eight years discussed the cross-prediction across predictors of DBRs and RBRs.

Table 1. Summary of surveys that discuss sequence-based predictors of the nucleic acid binding residues. The articles are sorted chronologically with the most recent study at the top of the table. The ‘Type of methods covered’ column identifies whether a given survey covers “Str” methods that were trained on the structure-annotated proteins and/or “Dis” methods that were trained on the disorder-annotated proteins. Bold font identifies data for this review.

Reference (Year published)	Target of assessment	Type of methods covered	No. of sequence- based predictors of DBRs/RBRs covered	No. of methods that were published after 2018 (year of the most recent method)	Types of training datasets discussed	Cross- prediction discussed
This study (NA)	DNA & RNA	Str+Dis	87	34 (2024)	yes	yes
[40] (2023)	DNA & RNA	Dis	3	2 (2022)	yes	no
[30] (2022)	DNA & RNA	Str+Dis	13	3 (2021)	no	no
[34] (2020)	RNA	Str	28	2 (2019)	no	yes
[31] (2019)	DNA	Str	8	0	no	no
[33] (2016)	DNA & RNA	Str	30	0	no	yes
[37] (2016)	RNA	Str	12	0	no	no
[41] (2015)	DNA & RNA	Str	24	0	no	yes
[29] (2015)	DNA	Str	7	0	no	no
[38] (2015)	RNA	Str	19	0	no	no
[35] (2013)	DNA	Str	13	0	no	no
[36] (2013)	DNA	Str	11	0	no	no
[28] (2013)	RNA	Str	10	0	no	yes
[39] (2012)	RNA	Str	13	0	no	no

We provide a thorough and practical overview of this active field of research, substantially improving over the past surveys that are listed in Table 1. We cover the largest number of sequence-based nucleic acid binding residue predictors, which include over two dozen new methods that were published in the last five years and which were not included in the previous review articles (Table 1). We examine how the underlying predictive models changed over time, investigate their availability and impact, and provide an expanded discussion of cross-predictions for the predictors of both DBRs and RBRs. Moreover, we categorize predictors on whether they are trained on the structure-annotated vs. disorder-annotated proteins and discuss the corresponding implications. Altogether, we comprehensively review the 20-year long journey of the development of predictors of DBRs and RBRs.

Materials and methods

Selection of methods

We performed an extensive literature search to obtain the list of sequence-based predictors of nucleic acid binding residues. We considered three main sources: i) the past surveys that covered sequence-based nucleic acid binding residue predictors [28-31, 33-39, 41]; ii) studies that cite these surveys; and iii) a manual search in PubMed using the following search keywords: ((Nucleic acid binding residue) OR (RNA binding residue) OR (DNA binding residue) OR (nucleotide binding residue)) AND ((prediction) OR (identification)), ((Nucleic acid binding site) OR (RNA binding site) OR (DNA binding site) OR (nucleotide binding site)) AND ((prediction) OR (identification)). We further filtered and selected only those methods which are published in peer-reviewed Q1 journals. Using the above protocol, we identified 87 sequence-based nucleic acid-binding residue prediction methods, which more than doubles the number of such methods covered in each of the previous surveys (Table 1).

Predictive performance assessment

Predictive performance is an important aspect of surveying tools in this area. While every published method was evaluated and compared against a selection of other predictors, the list of the other methods is typically rather short. Summarizing results from across multiple sources is challenging since several factors must be considered to ensure that results of different methods can be directly compared. Specifically, the corresponding assessments must be performed on the same benchmark dataset, with same type of annotations of binding residues, and at least one common metric of predictive performance. Additionally, these studies often covered overlapping methods, making it difficult to boost coverage, particularly for more recently published tools. We considered these factors and were able to collect, report and discuss the predictive performance of 20 methods, with 13 of them published within the last five years.

We discuss the assessments of the DBR and the RBR predictors separately. The DNA-binding test dataset that we used to assess the DBR predictors was first reported by Patiyl *et. al.* [58], and subsequently used in two recent studies [59, 60]. This dataset combines the DNA-binding datasets that were developed to assess two earlier predictors, hybridNAP [32] and ProNA2020 [61]. This test dataset contains 46 DNA-binding proteins with 965 DBRs and 9911 non-DBRs. The RNA-binding test dataset that we used to assess the RBR predictors was compiled to evaluate the Pprint2 method [62], and was also recently used in the article that introduces MuLiPred [60]. Similar to the DNA-binding dataset, it combines the datasets sourced from refs. [32] (hybridNAP) and [61] (ProNA2020). This test dataset contains 161 RNA-binding proteins with 6966 RBRs and 44349 non-RBRs. The annotations of binding residues for the test proteins sourced from the hybridNAP article were obtained from the BioLip database [63, 64], which in turn was derived from PDB [65, 66]. The binding residue annotations for the test proteins from the ProNA2020 article were collected from the Protein-DNA Interface database [67] and the Protein-RNA Interface database [68], both of which also rely on the PDB-derived data.

The 20 predictors of RBRs and DBRs output numeric propensity scores for the corresponding type of binding for each amino acid in the input protein sequence. These propensities are used to generate binary state predictions (binding vs. non-binding) using a threshold, where residues with propensities > threshold are assumed to bind, and the remaining residues are assumed not to bind. Based on their frequent use in the source evaluation articles [32, 58-62] and related comparative surveys [34, 38, 41, 69-73], we used the area under the receiver operating characteristic curve (AUC) to evaluate the predicted propensity scores and the Matthews Correlation Coefficient (MCC) to evaluate the predictions of binary states:

$$MCC = (TN * TP - FN * FP) / \sqrt{[(TP + FP)(TP + FN)(TN + FP)(TN + FN)]}$$

where TP, TN, FN and FP are the numbers of true positives (correctly predicted binding residues), true negatives (correctly predicted non-binding residues), false negatives (binding residues incorrectly predicted as non-binding), and false positives (non-binding residues incorrectly predicted as binding), respectively. MCC values range between -1 and 1, where 0 denotes predictions at random levels and a higher positive score indicates stronger correlation between predictions and the native annotations of binding. The AUC is the area under the

curve defined by the true positive rates (TPR) vs false positive rates (FPR) computed over the thresholds equal to all unique predicted propensities:

$$\text{TPR} = \text{TP}/(\text{TP}+\text{FN})$$

$$\text{FPR} = \text{FP}/(\text{FP}+\text{TN})$$

While AUC theoretically ranges from 0 to 1, the AUC of a random predictor is around 0.5 and higher values that are above 0.5 suggest stronger predictive performance. We collected the AUC and MCC scores for the predictors of DBRs from refs. [58-60, 74, 75] and for the predictors of RBRs from refs. [60, 62].

Results

We summarized details concerning the comprehensive list of the 87 predictors, such as their name, prediction target (DBRs, RBRs, or both), type of predictive models and training datasets, outputs and consideration given to cross-predictions, in Table 2. We covered 36 predictors of DBRs, 29 predictors of RBRs, and 22 predictors that target both DBRs and RBRs. Some methods are designed for specific types of nucleic acids and proteins. Specifically, SDCpred [76] predicts residues that interact with the mono- and dinucleotide-specific DNAs; SRCpred [77] targets predictions of the dinucleotide-specific RNA binding; DNAgenie [78] predicts A-DNA, B-DNA and single-stranded DNA binding residues; and TSNAAPred [79] predicts residues that interact with the A-DNA, B-DNA, single stranded DNA, mRNA, tRNA, and rRNA. Moreover, EPDRNA [80] generates predictions of nucleic acid binding residues in proteins associated with human diseases including cancer and cardiovascular and neurodegenerative diseases.

We analyzed the 87 methods from multiple complementary perspectives including a historical overview, their availability and impact, predictive performance, the consideration of the structure-based vs. disorder-based annotations of binding in their training datasets, and cross-predictions between DBRs and RBRs.

Table 2. Sequence-based predictors of DBRs and RBRs. The methods are arranged in the chronological order. The ‘Target’ column shows types of predicted binding residues: DNA-binding (D), RNA-binding (R), and DNA- and RNA-binding (DR). The ‘Predictive architecture’ column covers neural networks (NN), support vector machine (SVM), naïve Bayes (NB), logistic regression (LR), random forest (RF), template-based prediction (TB), decision tree (DT), linear regression (LR), long short-term memory (LSTM) NN; deep-learning models are marked with [D]; we also name the protein language model that a given method uses, if any, inside the round brackets. The ‘training dataset’ column differentiates between methods trained from the structure-annotated interactions (S), disorder-annotated interactions (D), and both types of annotations (S+D). The ‘Output’ column includes propensity scores (P), binary state (B), and both (B+P). The ‘Cross-prediction’ column shows whether the cross prediction was not mentioned (NM), discussed but not considered (Dis), or corrected/considered (Corr) when designing a given method.

Ref	Year published	Name	Target (D, R, DR)	Predictive architecture (PLM types used for feature extraction)	Training dataset (S, D, S+D)	Output (B, P, B+P)	Cross-prediction (NM/Dis/Corr)
[26]	2004	DBS-pred	D	NN	S	P	NM
[27]	2004	Jeong et al	R	NN	S	B	NM
[81]	2005	DBS-PSSM	D	NN	S	B+P	NM
[82]	2006	BindN	DR	SVM	S	B+P	NM
[83]	2006	DNABindR	D	NB	S	P	NM
[84]	2006	Jeong et al	R	NN	S	B	NM
[85, 86]	2006	DP-Bind	D	LR, SVM	S	B+P	NM
[87]	2007	DISIS	D	NN, SVM	S	B	NM
[88]	2007	Ho et al	D	SVM	S	B	NM
[89, 90]	2006	RNABindR	R	NB	S	B	NM
[91]	2008	Pprint	R	SVM	S	B+P	NM
[92]	2008	PRINTR	R	SVM	S	B	NM
[93]	2008	RISP	R	SVM	S	B+P	NM
[94]	2008	RNAProB	R	SVM	S	B	NM
[95]	2009	BindN-RF	D	RF	S	B+P	NM
[96]	2009	DBD-Threader	D	TB	S	B+P	NM
[97]	2009	DBindR	D	RF	S	B+P	NM
[98]	2009	ProteDNA	D	SVM	S	B	NM
[76]	2009	SDCpred	D	NN	S	P	NM
[99, 100]	2009	PiRaNhA	R	SVM	S	B+P	NM
[101]	2010	BindN+	DR	SVM	S	B+P	NM
[102]	2010	NAPS	DR	DT	S	B+P	Dis
[103]	2010	PRNA	R	RF	S	B+P	NM
[104]	2010	ProteRNA	R	SVM	S	B	NM
[105]	2010	RBRpred	R	SVM	S	B	NM
[106]	2010	RNA	R	LR	S	B	NM
[107]	2011	MetaDBSite	D	NB, NN, LR, RF, SVM	S	B	NM
[108]	2011	PRBR	R	RF	S	B+P	NM
[77]	2011	SRCpred	R	NN	S	P	NM
[109]	2011	Choi and Han	R	SVM	S	B	NM
[110]	2011	Wang et al	R	SVM	S	B	NM
[57]	2011	SPOT-Seq	R	TB	S	B	NM
[111]	2012	DNABR	D	RF	S	B+P	NM
[112]	2012	meta2	R	SVM	S	P	NM
[113]	2013	TargetS	D	SVM	S	P	NM
[114]	2014	Pan et al	R	RF	S	B	NM
[115]	2014	SPOT-Seq-DNA	D	TB	S	B	NM

[116]	2014	SPOT-Seq-RNA	R	TB	S	B	NM
[117]	2014	RNABindRplus	R	SVM, LR	S	B+P	NM
[17]	2014	aaRNA	R	NN	S	B+P	NM
[118]	2015	RBRIdent	R	RF	S	B+P	NM
[119]	2015	SNBRFinder	DR	SVM, TB	S	B+P	NM
[120, 121]	2015	DisoRDPbind	DR	LR	D	B+P	NM
[122]	2015	RBScore-SVM	R	SVM	S	B	NM
[123]	2016	Dang et al	D	RF	S	P	NM
[124]	2016	DQPred-DBR	D	SVM	S	B+P	NM
[125]	2016	FastRNABindR	R	RF, SVM	S	P	NM
[126]	2016	TargetDNA	D	SVM	S	B+P	NM
[127]	2017	DRNAPred	DR	LR	S	B+P	Corr
[128]	2017	PRODNA	D	Sparse Representation	S	P	NM
[129]	2017	EL PSSM-RT	D	RF, SVM	S	B	NM
[130]	2018	PDRLGB	D	Light Gradient Boosted DT	S	B	NM
[131]	2018	funDNAPred	D	Fuzzy Cognitive Map	S	P	NM
[132]	2019	DNAPred	D	SVM	S	B+P	NM
[32]	2019	hybridNAP	DR	LR	S	B+P	NM
[133]	2019	NucBind	DR	SVM, TB	S	B+P	Dis
[134]	2019	PSPrint-seq	R	RF	S	B	NM
[135]	2019	iProDNA-CapsNet	D	[D] Convolutional NN, Feed forward NN	S	B+P	NM
[61]	2020	ProNA2020	DR	NN, SVM (ProtVec)	S	B+P	NM
[136]	2020	EL LSTM	D	[D] LSTM NN	S	B	NM
[137]	2021	SPDH	D	SVM	S	P	NM
[78]	2021	DNAGenie	D	LR, k-nearest neighbor, NB, RF, SVM	S	B+P	Corr
[138]	2021	NCBRPred	DR	[D] Recurrent NN	S	B+P	Corr
[139]	2021	bindEmbed21	DR	[D] Convolutional NN (ProtT5)	S	B+P	NM
[140]	2022	MTDsite	DR	[D] Bidirectional-LSTM NN	S	B	Dis
[141]	2022	DeepDISOBind	DR	[D] Convolutional NN	D	B+P	Corr
[142]	2022	PredDBR	D	[D] Convolutional NN, NN	S	B	NM
[58]	2022	DBPred	D	[D] Convolutional NN	S	B	NM
[79]	2022	TSNAPred	DR	[D] CapsNet, Light Gradient Boosted DT, NN	S	B+P	Dis
[143]	2022	iDRNA-ITF	DR	[D] NN, Bidirectional gated recurrent NN (CAN-NER)	S	B+P	Corr
[62]	2023	Pprint2	R	[D] Convolutional NN	S	B	NM
[144]	2023	Guan et al	D	[D] Transformer, Convolutional NN	S	P	NM
[145]	2023	proRBR	R	RF	S	B+P	NM
[146]	2023	HybridRNAbind	R	[D] Convolutional NN, Recurrent NN, RF	S+D	B+P	Corr
[147]	2023	GLMSite	DR	[D] Graph NN (ProtTrans and ESMFold)	S	B+P	NM
[59]	2024	CLAPE-DB	D	[D] Convolutional NN (ProtBERT)	S	B+P	NM
[148]	2024	HybridDBRpred	D	[D] Transformer, NN	S+D	B+P	Corr
[80]	2024	EPDRNA	DR	RF, k-nearest neighbor, LR, XGBoost	S	B+P	NM
[149]	2024	DRBpred	DR	Light Gradient Boosted DT	S	B+P	NM
[60]	2024	MucLiPred	DR	[D] NN (BERT)	S	B+P	NM
[150]	2024	ULDNA	D	[D] LSTM NN (ESM2 and ProtTrans)	S	B+P	NM
[151]	2024	SOFB	DR	[D] Convolutional NN, Bidirectional-LSTM NN (NA Bert, ProtT5)	S	B	NM
[152]	2024	GPSFun	DR	[D] Graph NN (ProtT5-XL-U50)	S	B+P	NM
[153]	2024	GPSite	DR	[D] Graph NN (ProtTrans and ESMFold)	S	B+P	NM
[75]	2024	PDNAPred	D	[D] Convolutional NN, Bidirectional Gated Recurrent Unit NN (ESM2 and ProtT5)	S	B+P	NM
[74]	2024	DIRP	D	[D] Convolutional NN (ESM2 and ProtTrans)	S	B+P	NM
[154]	2024	DeepDBS	D	[D] LSTM NN, RF	S	B+P	NM

Historical overview

Figure 1 illustrates the timeline of the release of the 87 predictors, where we highlighted eight major milestones. The first methods that predict exclusively DBRs (DBS-Pred) or RBRs (predictor by Jeong *et. al.*) were published in 2004 [26, 27]. Since then, at least one method was released each year. DBS-pred [26] is the first predictor of DBRs that was released as a web server. This contribution introduced a simple predictive architecture in the form of a shallow feed forward neural network and defined how to collect datasets with annotations of DBRs, which were used to train and test this predictive model. At the same time, Jeong *et. al.* published two methods, one in 2004 [27] and another in 2006 [84], but neither predictor was released for public use (no code and no web server). RNABindR that was published in 2007 [89] is the first predictor of RBRs that is available as a web server.

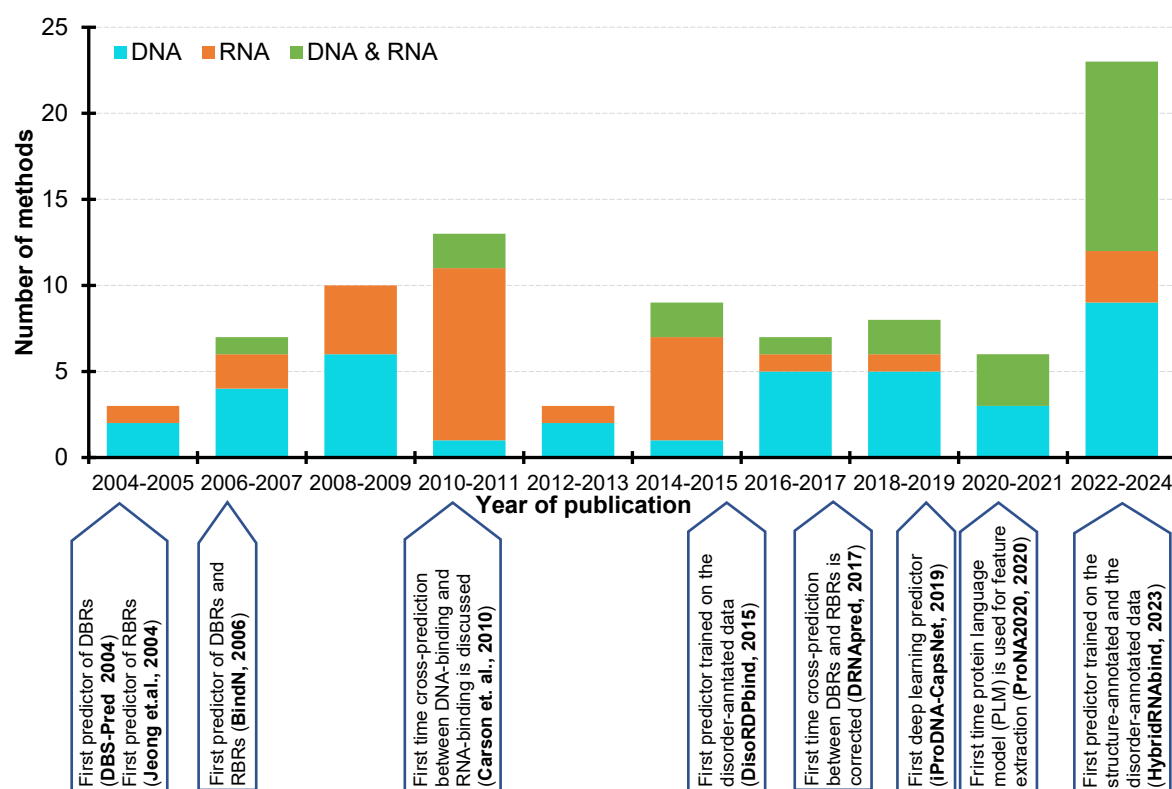


Figure 1. Timeline of the release of the 87 sequence-based nucleic acid binding residue predictors. The color-coded bars represent methods that target prediction of DBRs (blue), RBRs (orange) and both DBRs and RBRs (green). The major milestones are shown at the bottom in the blue-bordered boxes.

As another milestone (Figure 1), BindN that was published in 2006 [82] is the first method that predicts both DBRs and RBRs. Except for this aspect, BindN arguably did not contribute to moving other aspects of this field forward as it utilizes a relatively simple design, which consists of a classical support vector machine (SVM) model that uses just three biochemical features of amino acids as the input (pK_a , hydrophobicity, and molecular mass). This simplicity motivated the subsequent release of an improved BindN+ version in 2010 [101], which was designed to utilize more sophisticated inputs that rely on evolutionary information. By 2009, 20 methods were released and yet the issue of the cross-predictions between DBRs and RBRs did not

surface. Authors of NAPS [102], a tool that predicts DBRs and RBRs, were the first to briefly discuss this concern but they did not address it in their design. It took several more years until 2016 when the cross-prediction was quantified and compared across predictors in the comparative study by Yan *et. al.* [33]. This study motivated the development and release in 2017 of DRNAPred [127], which was designed to accurately predict and discriminate between DBRs and RBRs (i.e., minimize the cross-predictions). The main innovation that was introduced in DRNAPred is a two-layered architecture, where predictions of DBRs and RBRs produced by the first layer are input together into the second predictive layer that refines them to minimize cross-predictions. Importantly, until 2015 all the methods were developed using the structure-annotated training data. As a major milestone (Figure 1), DisoRDPbind that was published in 2015 [120], is the first tool that was designed using the disorder-annotated training and test datasets, extending the DBR/RBR annotation protocols that were introduced in mid 2000s.

The predictive architectures of the methods, see Table 2, are predominantly based on machine learning (ML) algorithms, with a just few exceptions where template-based approaches are used [57, 96, 115, 116]. For example, SNBRFinder [119] and NucBind [133] utilize a template-based approach to predict protein structure from the input sequence, which is followed by the application of ML-generated models, particularly SVM, that identify putative DBRs and RBRs from the predicted structures. Significant majority of the predictors of DBRs and RBRs rely on shallow ML algorithms. Many different shallow ML algorithms were tried including the most widely used SVM, which was applied by itself in 19 predictors [82, 88, 91-94, 98, 99, 101, 104, 105, 109, 110, 113, 122, 124, 126, 132, 137] and in combination with some other algorithms in additional 9 predictors [61, 78, 85, 87, 107, 117, 119, 129, 133]. The next two popular shallow ML algorithms are random forest, which was applied in 14 methods [78, 80, 95, 103, 108, 111, 114, 118, 123, 125, 129, 145, 146, 155], and shallow neural network (NN) that were used in 10 methods [17, 26, 27, 61, 76, 77, 81, 84, 87, 107]. Some of the other shallow ML algorithms include logistic regression [78, 80, 85, 107, 117, 120, 121, 127], Naïve Bayes [78, 83, 89, 107], linear regression [32, 106] and decision trees [61, 102, 130, 149]. Interestingly, we observed a substantial increase in the use of deep neural networks (DNNs) over the last five years, where 24 of the 34 predictors rely on the DNN models. Released in 2019, iProDNA-CapsNet [135] is the first method that applied the DNN model, marking another major milestone (Figure 1). The main innovation behind this predictor was the formulation and training of the predictive model that involves two two-dimensional convolutional layers connected to a fully connected feed forward layer. However, iProDNA-CapsNet relies on rather generic inputs, in the form of evolutionary information that was previously utilized to design several past predictors. Analysis of Table 2 reveals that the convolutional NNs are the most widely used architecture of DNNs among the predictors of RBRs and DBRs [58, 59, 62, 74, 75, 135, 139, 141, 142, 144, 146, 151], although several other architectures that include unidirectional and bidirectional recurrent networks, transformers, and graph networks were also used. In addition to the development of the DNN models, we identified a recent trend of using pre-trained protein language models (PLMs) as feature/input extraction tools [59-61, 74, 75, 139, 147, 150-153]. Simply put, PLMs process an input protein sequence in a way similar to processing a sentence in human language, where functional motifs and domains of the protein act as words in the sentence [156], producing a vector of numerical features for each amino acid. PLMs have been used to solve

several bioinformatics problems and literature shows that their use tends to lead to improvements in the predictive performance of models [157, 158]. In the context of the prediction of nucleic acid binding residues, the ProNA2020 method in 2020 [61] was the first to use PLM called ProtVec [159] for the feature/input extraction, denoting another milestone (Figure 1). Henceforth, eleven other predictors including bindEmbed21 (PLM: ProtTrans-ProtT5 [160]) [139], iDRNA-ITF (PLM: CAN-NER [161]) [143], CLAPE-DB (PLM: ProtTrans-ProtBERT [160]) [59], MuLiPred (PLM: ProtTrans-ProtBERT [160]) [60], GLMSite (PLM: ProtTrans [160], ESMFold[162]) [147], ULDNA (PLM: ESM2 [162], ESM-MSA [163] and ProtTrans [160])[150], SOFB (PLM: ProtTrans-ProtT5 [160]) [151], GPSFun (PLM: ProtTrans- ProtT5-XL-U50 [160]) [152], GPSite (PLM: ProtTrans [160], ESMFold [162]) [153], PDNApred (PLM: ProtTrans-ProtT5 [160] and ESM2 [162]) [75] and DIRP (PLM: ProtTrans [160] and ESM2 [162]) [74] use PLMs for the feature extraction. The latest milestone was the development and release of HybridRNAbind in 2023 [146], which is the first tool that was trained using both structure- and disorder-annotated training data, bridging the two annotations types. Additionally, this RBR predictor also minimizes cross-prediction between RBRs and DBRs [146]. Soon after, HybridDBRpred, which targets prediction of DBRs and similarly combines the structure- and disorder-annotated training data, was published [148].

The above historical overview (Figure 1) suggests that the timeline can be divided into two distinct decades, each defined by a number of unique characteristics. The first decade spans the period between 2004 and 2014, and the second decade includes years from 2014 onwards. The major focus during the first decade was on developing predictors that target either DBRs or RBRs, resulting in 16 DBR predictors, 17 RBR predictors and just three tools that predict both DBRs and RBRs. Moreover, these methods rely on relatively simple shallow ML algorithms and they were trained exclusively on the structure-annotated proteins. The second decade is a more dynamic period where five major milestones took place (Figure 1). Although many predictors of either DBRs or RBRs were still developed, 19 methods that target both types of binding residues were released in this period including 6 out of the 12 methods published in 2024 (Table 1). Furthermore, we observed a big shift in the choice of predictive architectures that increasingly included deep ML models and modern PLMs for the feature/input extraction, with the underlying objective to improve predictive performance. The second decade also brought consideration to the cross-predictions between DBRs and RBRs and predicting both structure- and disorder-annotated interactions. The former likely stems from the increased focus on predicting both types of interactions, which inevitably brings the matter of evaluating whether these predictions overlap. Altogether, the second decade featured development of more sophisticated predictive models that attempted to address the challenging issues affecting quality (cross-predictions) and scope (disorder- and structure-annotated) of DBRs and RBRs predictions.

Availability and impact

An important aspect of computational predictors is their accessibility to users. Supplementary Table S1 provides details concerning the mode of availability (a web server (WS), a standalone code (SC), both or neither, as declared in the corresponding publication) and their current availability (whether or not the declared mode is currently available). We verified the

availability of these implementations in November 2024 when we collected these data by checking the web links from the reference articles. This led to three outcomes: available (we list the corresponding links in Supplementary Table S1), not working (links in the original reference did not work as of November 2024), and never available (authors did not make their tools available in the original reference).

The two modes of availability, WS and SC, differ in multiple aspects. Users can access WSs via a web browser and these predictions are computed on the server side, typically without installing any software on the user's side. This makes a WS an arguably easier to use option, however, each prediction request is usually limited to a single protein or a small batch of proteins and the runtime might be affected by the ongoing server load. On the other hand, SC has to be downloaded and installed by the user and the computations are performed on the user's hardware. SC can be challenging to install, requiring users to install one or more third-party applications and have specific hardware and/or software infrastructure. However, SC offers certain advantages when compared to WS, including ability to be embed into other bioinformatics pipelines and to generate predictions at a large scale. We found that 75 out of the 87 listed methods provided at least one mode of availability at the time of their publication (Supplementary Table S1). Out of these 75 methods, 56 (77%) were originally released as WSs, either solely or along with SC, making WSs the most common mode of availability.

Though SC was provided for the first time in 2010 by the authors of PRNA [103], this option was rather uncommon until recently. Out of the 28 methods that were published with SC, 21 were released in the last 5 years. Moreover, 9 of these 21 methods were originally published with both WS and SC, which arguably broadens the utility of the methods when compared to tools that offer one mode of availability. However, as of November 2024 only 35 out of the 75 originally available methods have a working WS or SC; the links for the other 40 no longer work. Among the 35 currently available methods, 12 are only WS, 13 are only SC, and 10 are both WS and SC. The overall availability rate for the sequence-based predictors of RBRs and DBRs is at 40% (35 out of 87), which is same as the 40% rate for the predictors of protein-binding residues [164] and lower than the recently reported 71% rate for the predictors of the disordered binding residues [40].

We also analyzed the impact or popularity of the sequence-based predictors of RBRs and DBRs based on their citations, which we collected from the Google Scholar in November 2024 (Supplementary Table S1). While we collected the total number of citations for each method, we relied on the corresponding annual citation rates (i.e., total citations divided by the number of years since publication) which are more appropriate to compare between methods. We also excluded the methods which are published from 2023 onwards, since they are too new to reliably measure their citations.

We found that methods that did not offer either mode of availability (i.e., were not made available) are cited substantially less (median annual citations = 3) compared to the tools which originally had at least one mode of availability (median annual citations = 8). These median annual citation counts are much larger for the methods that have currently working WS and/or SC (median annual citations = 13). Among these working tools, the tools with only working WS receive the most citations (median annual citations = 15), closely followed by the methods with

both WS and SC (median annual citations = 14). Moreover, methods that are available solely as SC are comparatively less cited (median annual citations = 5). We hypothesize that the higher citations for the methods with working WSs are because this mode of availability is accessible to a much broader group of users, including those who have limited technical expertise and computational resources. On the other hand, methods that were not made available secure relatively poor citation numbers, which suggests that availability strongly affects the rate of use (citations). Moreover, we also observed that methods which were originally available but which currently do not work obtain much fewer citations when contrasted against tools with working WS and/or SC (median annual citations = 8). This implies that availability of methods should be maintained after their release, or otherwise their impact is much diminished.

Lastly, we briefly comment on a few most impactful/cited methods. BindN [82], the first tool that predicts both DNA- and RNA-binding residues, has the highest total citations of 495 (Supplementary Table S1). Since then, this tool has undergone two upgrades [95, 101], with the most recent version, BindN+ [101], released in 2010, which received over 200 citations to date. DBS-pred [26], predictor of DBRs, is the only other tool that has total citations at over 400. Among predictors of RBRs, Pprint [91] with an overall citations of 322 is the most cited. In fact, Pprint along with DP-Bind [85] (overall citations of 270) are the only two methods that have maintained their implementations for over 10 years, which likely contributed to their high citation counts. There are 21 methods which have been cited more than 100 times and hybridNAP [32], which was published in 2019, is the most recent method to collect such high number of citations.

We also highlight a few recently published applications of these highly cited tools in biological contexts to further substantiate their impact. Multiple predictors including BindN [82], BindN+ [101] and metaDBsite [107] were used in tandem to characterize the DSrC protein from sulfur oxidizing bacterium *A. vinosum* [165]. Similarly, DRNAPred [127], Pprint [91] and RNAbindPlus [117] were applied to study the GRP20 protein in the context of flower development [166], while Pprint and hybridNAP [32] were applied to investigate roles of the SIX1 protein in the iron metabolism associated with progression of endometrial cancer [167]. These predictors were also used individually, with an example of DRNAPred [127] that was recently applied to characterize proteins encoded by a viral mycobacteriophage gene [168] and CRESS DNA viruses [169]. Methods that have low runtimes were used to analyze entire proteomes. For instance, DisoRDPbind [120] was used to investigate length variation of short tandem repeats in *A. thaliana* [170], abundance and function of intrinsic disorder in the polyomavirus [171], and functions of the RNA-binding proteins in human [172]. DisoRDPbind together with DRNAPred and Pprint were applied to analyze the COVID-19 proteome [173]. Moreover, these tools were used to generate predictions at the scale of multiple proteomes and these data are conveniently available to the users via specialized databases. For instance, GPSite's predictions for the entire Swiss-Prot database are available in the GPSiteDB database [153], while DisoRDPbind's predictions for 273 reference proteomes, which cover the popular and model organisms, can be conveniently obtained from the DescribePROT database [174, 175].

To summarize, our analysis revealed a significant amount of interest in the area of the nucleic acid binding residue prediction (i.e., the corresponding predictors were collectively cited over 6600 times and 21 of them were cited over 100 times), where methods with working WSs are

the most cited/popular. These observations are in good agreement with a recent analysis for a broader collection of predictors of protein structure and function [176].

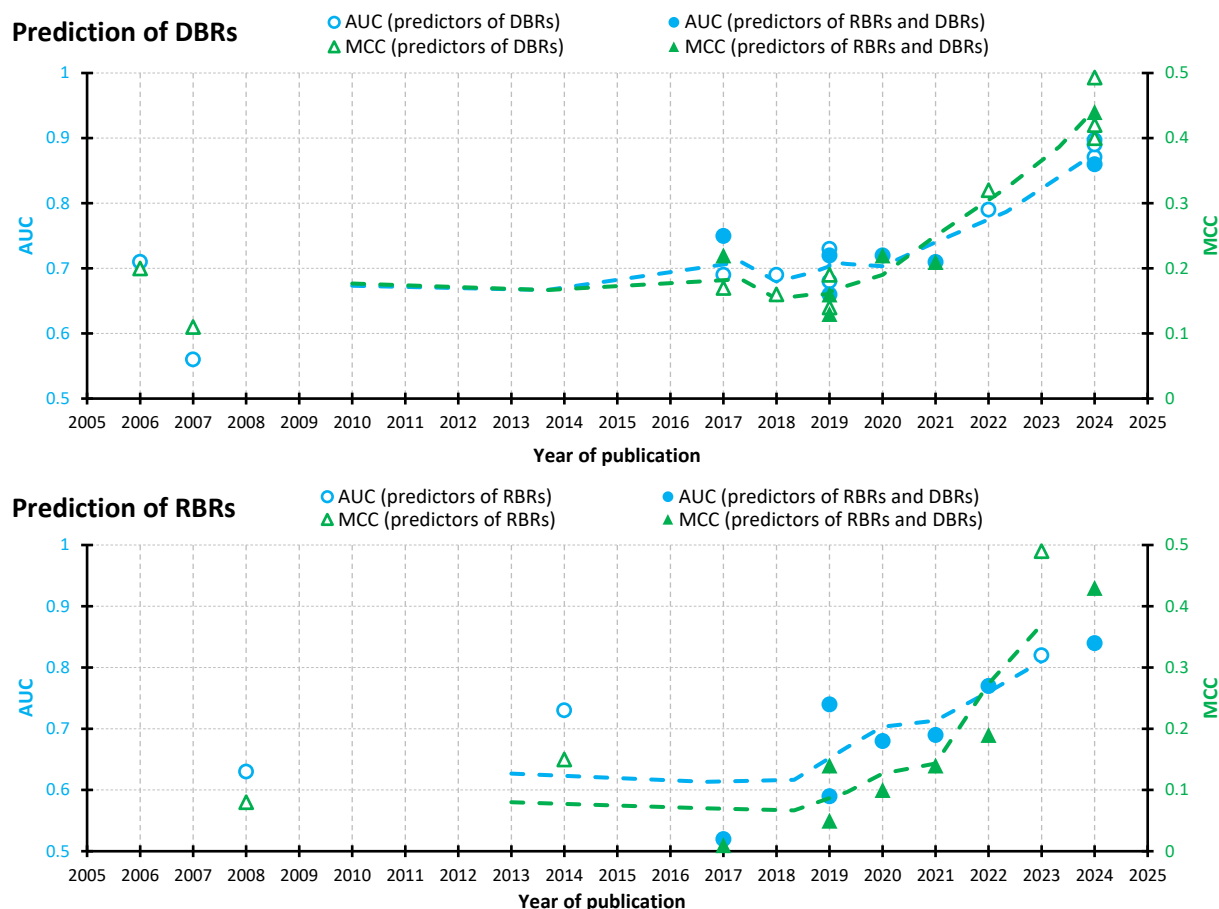


Figure 2. Relation between the publication year and predictive performance for the corresponding methods that was measured on the same benchmark dataset of 46 DNA-binding proteins that was introduced and applied in refs. [58-60, 74, 75] (top panel) and the same benchmark dataset of the 161 RNA-binding proteins that was used in refs. [60, 62] (bottom panel). Hollow markers denote methods that predict DBRs (top panel) or RBRs (bottom panel) while solid markers are for predictors of DBRs and RBRs (both panels). The primary/left y-axis quantifies the AUC values (blue markers) and the secondary/right y-axis gives the MCC values (green markers). The color-coded dashed lines are the moving averages of the corresponding metrics calculated over three consecutive methods based on the publication years. The numerical values of AUCs and MCCs are given in the Supplementary Tables S2 (results for the top panel) and S3 (results for the bottom panel).

Predictive performance

We summarized predictive performance of the current predictors and analyzed how it evolved over time in Figure 2. We compared performance of 16 predictors of DBRs and 10 predictors of RBRs. We used the same test datasets and performance metric for each of the two collections of methods to ensure that results can be directly compared across the corresponding tools. We relied on popular and recently developed test datasets that were introduced in refs. [58, 62]. The dataset for the assessment of the DBR predictors includes 46 DNA-binding proteins and was used in refs. [58-60, 74, 75], while the other dataset covers 161 RNA-binding proteins and was used in refs. [60, 62]. Figure 2 reports AUC that evaluates quality of the putative propensity

scores and MCC that quantifies quality of the binary state predictions. Further details can be found in the Materials and Methods section.

Figure 2 visualizes relations between predictive performance of the 20 predictors and their corresponding publication time. The performance of both DBR and RBR predictors varies widely between modest levels (AUC of about 0.6; MCC of about 0.2) and high levels (AUC > 0.80; MCC > 0.4), with an overall trend of improving with passing years. We computed moving average-based trends by averaging AUC and MCC scores over a window of three chronologically consecutive methods (dashed lines in Figure 2). From these trends, we found that earlier methods secure similar and relatively modest performance with an average AUC below 0.7 and MCC below 0.2. The performance trends upwards starting around 2020, with the recent methods having AUC well above 0.8 and MCC of about 0.5. The best-performing predictors of DBRs are DIPR (2024) with AUC of 0.89 and MCC of 0.49 and PDNAPred (2024) with AUC of 0.90 and MCC of 0.49; Supplementary Table S2. The best predictors of RBRs are MuLiPred (2024) with AUC of 0.84 and MCC of 0.43 and Pprint2 (2023) which secures AUC of 0.82 and MCC of 0.49; Supplementary Table S3. This demonstrates that the most accurate predictions of RBRs and DBRs are produced with similar levels of performance, with predictors of DBRs performing slightly better.

We also analyzed relation between the AUC and MCC scores and found that these scores are inconsistent for some methods, i.e., they should roughly follow a linear relation while some results deviate from this trend (Supplementary Figure S1). MCC is a threshold-dependent measure that is derived from the predicted propensities for binding while AUC directly evaluates the propensities without the use of a threshold. The articles that presented these results used a “default” threshold of 0.5 to generate the binary predictions for the calculation of MCC. Setting the same threshold value fails to account for the differences in the ranges and distributions of the propensities produced by different methods, leading to MCC values that are computed at different rates of positive (binding) to negative (non-binding) predictions. This could be the reason for the observed inconsistencies. A better strategy is to apply thresholds that standardize the predictions of different methods to a consistent prediction rate (say 5% or 10% FPR). However, the overall correlations between MCC and AUC metrics are high, with the Pearson’s correlation coefficient of 0.95 and 0.87 for the DBR and the RBR predictors, respectively.

As we discuss in the Historical Overview section, recent methods often rely on the DL-based models when compared with older methods that primarily use shallow ML-based algorithms (Table 2). Correspondingly, we investigated whether the recent improvements in the predictive importance could be attributed to the use of the more sophisticated predictive models. We limited this analysis to the predictors that were published since 2019 when the first DL-based predictor was released. Using the corresponding results from Supplementary Tables S2 and S3, we found that the DL-based DBR predictors have substantially higher predictive performance than the shallow ML-based predictors that were published in the same period of time, with median AUC = 0.86 vs 0.72 and median MCC = 0.40 vs 0.18. Similar observations are true for the RBR predictors, where the DL-based methods secure median AUC = 0.80 vs 0.68 for the other group of predictors, and MCC = 0.31 vs 0.10, respectively. Recent studies in related areas including protein function prediction [177, 178] and intrinsic disorder prediction [179] similarly

showed that DL-based predictive models substantially outperform the shallow ML-based models. Our analysis reveals that the same is true for the prediction of the nucleic acid binding residues.

We also investigated whether similar patterns of differences can be observed when comparing recently released (since 2019) predictors that use PLMs for the feature/input extraction vs those that do not utilize PLMs. Based on the data in Supplementary Tables S2 and S3, we found that predictors of DBRs that use PLMs secure median AUC = 0.87 and median MCC = 0.42 vs median AUC = 0.72 and median MCC = 0.18 for the methods that do not use PLMs. For the predictors of RBRs, we similarly observed that the PLM-utilizing methods generate on average more accurate predictions than the other methods, with median AUC = 0.77 vs 0.72 and median MCC = 0.19 vs 0.14, respectively. These results suggest that the application of PLMs produces improvements in the predictive performance for the predictors that cover both types of binding residues. Our observation is also supported by an empirical analysis that demonstrated that the use of the ProtT5 PLM produces higher levels of accuracy than the use of the popular multiple sequence alignment-based inputs for equivalent models that predict DBRs and RBRs [139].

However, we note a limitation of our analysis. The predictive performance can be impacted by sequence similarity between the training and test proteins. In principle, predictors should be tested on the test proteins that share low similarity (typically < 30%) with the training proteins. Several of the recent predictors in the above analysis were tested under this low similarity regime [58-60, 62, 74, 75], including DBpred, MuLiPred, CLAPE-DB, Pprint2, PDNApred and DIRP. The training datasets of the other methods may share higher levels of similarity, and thus their reported predictive performance might be inflated. However, this does not affect our observations since the methods for which the performance was tested on the low similarity test proteins are the most accurate. Additionally, the two test datasets are structure-annotated, which means that this assessment does not reflect the performance on the disorder-annotated protein-nucleic acids interactions. We discuss this issue in a following section.

Structure-annotated and disorder-annotated training datasets

We divided the 87 predictors into two groups based on their training data: trained from the structure-annotated interactions (structure-trained) vs. trained from the disorder-annotated interactions (disorder-trained). The 'Training dataset' column in Table 2 shows the type of interaction annotations used for training each method, where 'S' represents structure-annotated datasets and 'D' represents disorder-annotated datasets. The binding annotations for these two types of training datasets were obtained from two distinct sources. The structure-annotated binding residues were collected from the structures of the protein-nucleic acids complexes which are extracted directly from the PDB database [65, 66] or indirectly from the PDB-derived BioLip database [63, 64]. These annotations rely on certain criteria, such as distance between interacting atoms and number of atoms interacting per residue, to identify binding vs. non-binding residues. On the other hand, the disorder-annotated interactions were derived from the DisProt database [180, 181], the largest repository of experimentally validated intrinsically disordered proteins. These annotations concern binding sequence regions rather than specific binding residues, as is the case for the structure-based annotations. In other

words, disorder-annotated training datasets include disordered regions that are involved in binding, assuming (imprecisely) that all residues in the corresponding region are binding. The lack of the more precise annotation of the corresponding binding residues in the disordered regions is a result of an inherent difficulty in capturing these details without the structure. This limitation was discussed in the context of the recent community assessments of predictions of disordered binding regions [182, 183].

Table 2 shows that large majority of methods, 83 out of 87, was developed using solely the structure-annotated training datasets. On the other hand, DisoRDPbind [120, 121] and DeepDISOBind [141] are the two methods that are exclusively trained on the disorder-annotated data (Table 2). Both methods provide prediction of DBRs and RBRs alongside prediction of protein-binding residues. The rather low number of tools trained from the disorder-annotated interactions can be attributed to the fact that the corresponding data was released relatively recently, in early 2010s [184-186]. Interestingly, two recent studies reported that the structure-trained methods perform well on the structure-annotated proteins, whereas they secure poor to modest performance on the disorder-annotated proteins and vice versa [146, 148]. For example, when considering the predictors of RBRs [146], the structure-trained MTDsite performs the best on the structure-annotated interactions with AUC = 0.76, whereas, its AUC drops to 0.60 for the disorder-annotated interactions [146]. Similarly, the disorder-trained DeepDISOBind performs best on disorder-annotated interactions with AUC = 0.72, whereas, it secures a much lower AUC of 0.64 for the structure-annotated interactions [146]. Similar observations were published for the predictors of DBRs [148]. The fact that the ground truth annotations of binding for the disorder-annotated vs. structure-annotated datasets are different (residues vs. regions) and come from different source databases may explain the dichotomy in the method development efforts, i.e., methods are typically designed using either structure-annotated or disorder-annotated training sets and consequently they do not work equally well across the two annotation types. Similar observation of the dichotomy of structure-trained vs. disorder-trained predictions was reported for the sequence-based predictors of the protein-binding residues [187].

To this end, two recently published methods, HybridRNAbind [146] and HybridDBRpred [148], were designed to address this dichotomy by targeting both types of annotations, i.e., they are trained and tested on datasets composed of both types of annotations. HybridRNAbind that predicts RBRs performs relatively well with the AUC of 0.76 on the structure-annotated interactions and AUC of 0.72 for the disorder-annotated interactions [146]. Similarly, HybridDBRpred that targets prediction of DBRs secures AUCs of 0.83 and 0.77 for the structure-annotated and the disorder-annotated interactions, respectively [148]. The development of these two methods shows that it is possible to build predictors that work well across the two annotation types, and suggests that these efforts should continue.

Cross-prediction between DBRs and RBRs

DNA and RNA share relatively high levels of similarity in their physicochemical nature, as both are made up of a monomeric unit having a nitrogenous base and a sugar-phosphate group. Given their resemblance at the molecular level, it is reasonable to expect that predictors of nucleic acid binding residues may face difficulties to accurately discriminate between DBRs and

RBRs. Besides accurately predicting putative binding residues, these methods also should be free from the cross-predictions where DBRs are confused for RBRs and vice versa. High levels of cross predictions would mean that the corresponding methods predict residues that interact with nucleic acids in the type agnostic manner.

Two early surveys conducted empirical assessments of cross-predictions for several predictors of DBRs and RBRs, covering methods that were published before 2014 [33, 41]. These studies reported similar findings suggesting that none of these older methods accurately discriminates between DBRs and RBRs. Miao *et.al.* found that several accurate predictors of RBRs also obtain high AUC scores when tested on predicting DBRs in the DNA-binding proteins, which implies that these methods predict binding residues irrespective of whether they bind DNA or RNA [41]. While they also show that some methods, such as PRNA, RNABindRPlus, RBScore-SVM, discriminate between DBRs and RBRs, the predictive performance of these tools on the RNA-binding datasets is low, with AUCs around 0.5 [41]. The study by Yan *et. al.* quantified cross-predictions by measuring the fraction of DBRs (or RBRs) that are mis-predicted as RBRs (or DBRs) [33]. They found that among the RBR predictors, RNABindR has the highest cross-prediction rate, incorrectly classifying more than 60% of DBRs as RBRs while BindN+ has the lowest rate, at around 45%. For the DBR predictors, DBS-PSSM generates the most cross-predictions, with 44% RBRs predicted as DBRs. A positive exception is ProteDNA, predictor of DBRs, which accurately distinguishes DBRs from RBRs but has low sensitivity for DBRs, since it was designed to specifically predict DBRs in transcription factors [33]. These assessments emphasize the need to develop new predictors that minimize the cross-predictions between DNA and RNA. They motivated the release of DRNAPred [127], which is the first method that was specifically designed to restrict cross-predictions. Consequently, a subsequent survey of the RBR predictors shows that DRNAPred performs the best in terms of the cross-predictions [34]. However, these comparative studies are relatively old and do not cover recently published predictors [33, 34, 41]. We note that authors of several newer methods, such as NCBRpred [138], iDRNA-ITF [143], DNAGENIE [78], DeepDISOBind [141], MTDsite [140], HybridDBRpred [148] and HybridRNAbind [146], evaluated cross-predictions and designed their models to minimize them. However, these studies considered/evaluated relatively few methods and used different datasets and metrics, constraining comparative analysis to a rather limited number of recent tools. Moreover, many recently published methods that include seven methods from 2024 and several that predict both DBRs and RBRs [60-62, 74, 75, 139, 147, 149, 151-154] overlooked this crucial aspect. To this end, we note the reports of predictive performance that do not account for the cross-predictions should be interpreted with caution.

Summary and outlook

Sequence-based prediction of nucleic acid binding residues is a mature and active research area. We identified nearly 90 predictors that were published over the last two decades, including 29 that were published over the last five years and 12 that were released in 2024. We discussed multiple practical characteristics of these methods including their availability and impact, key features of their predictive models, and major aspects related to their training and assessment. We observed that the last decade produced a noteworthy progress in terms of improvements in the predictive quality, use of sophisticated predictive models and PLMs, and

advancements on multiple vital and challenging issues, such as targeting of the two distinct annotation types (structure-based vs disorder-based) and cross-predictions.

Our analysis of the availability reveals that predictors that have a web server or a standalone code are cited substantially more than tools without implementations. Methods with the arguably easier to use web servers attract more citations when compared to the tools with the code. Furthermore, we found that methods that have implementations that are working and are maintained over the long term (in particular working and maintained web server) secure much higher citation counts compared to tools with originally available implementations that currently do not work. In the context of predictive models, the recently released predictors increasingly rely on modern deep network architectures. Our analysis revealed that these models outperform the previously-dominant shallow machine-learning algorithms, which is in line with results of similar analyses in related area of protein structure and function prediction [177-179]. We also stressed the impact of training on two distinct types of datasets: structure-vs disorder-annotated. Majority of predictors were trained on the structure-annotated interactions and they perform poorly when tested on the disorder-annotated interactions. Similarly, the predictors trained on the disorder-annotated datasets perform rather poorly on the structure-annotated interactions. The underlying dichotomy of the predictive models and their inability to crossover to the other type of annotations mirrors the prediction of protein-binding residues [187]. This motivated the development of the HybridDBRpred and HybridRNAbind methods that perform relatively well for both structure and disorder-annotated interactions. However, some protein-nucleic acid interactions are driven by an assembly of protein chains, such as the one found in the ribosomal complex [188, 189]. The sequence-based methods that are trained on monomeric protein chains may underperform for these proteins, particularly when compared with the structure-based methods that use the corresponding complexes for the training. Finally, we indicated that cross-predictions between DBRs and RBRs is a crucial aspect of empirical assessments of predictive performance, and yet this aspect was not evaluated for many of the recently published tools. Consequently, while some of the recently released predictors were shown to produce accurate results ($AUC > 0.8$ and $MCC > 0.4$; Supplementary Tables S2 and S3), users should also quantify and analyze their cross-prediction rates to formulate a more holistic picture of their performance. In general, users should avoid predictors with high cross-prediction rates since this means that their predictions are nucleic acid type agnostic. The binding residues that are predicted by the accurate sequence-based tools can be used to support subsequent modelling of the protein-nucleic acid interactions. In particular, tools that specialize in modelling and predicting binding specificity for the protein-nucleic acids complexes [18, 190-194] would likely benefit from the knowledge where a given DNA or RNA binds on the protein surface. While some of these tools target specific types of proteins, such as transcription factors [190, 192, 194], other tools can be applied to a more generic class of nucleic acid binding proteins [18, 191, 193].

Our analysis motivates several considerations for future work. Many of the recently released predictors of RBRs and DBRs utilize PLMs to derive inputs to the predictive models, with ProtT5 and ProtBERT from the ProtTrans project [160] being the common selections. Our empirical analysis suggests that the use of PLMs leads to substantial improvements in predictive performance for predictors of both DBRs and RBRs, which is based on their overall higher AUC

and MCC values when compared to the predictors that do not utilize PLMs. However, PLMs used by the current predictors were produced using generic collections of protein sequences while several PLMs that were designed for specific types/classes of proteins were released in recent years. For example, IDP-BERT was designed to capture characteristics of intrinsically disordered proteins [195]; ProGen was built from sequences of five families of lysozymes [196]; and IgLM was trained using of antibody sequences [197]. We believe that similar efforts geared towards developing and using PLMs that target nucleic acid-binding proteins should drive further improvements in accuracy for the predictors of DBRs and RBRs. As a first step in this direction, the authors of the SOFB predictor adopted the generic ProtT5 PLM to make it more suitable for the recognition of nucleic acid binding residues, and named this model NABert [151]. Moreover, structural and functional aspects of proteins are typically conserved in their amino acid sequences. Consequently, evolutionary profiles generated using protein sequence databases are commonly used to predict nucleic acids binding residues [62, 74, 78, 81, 86, 88, 91, 92, 94, 101, 129, 133, 153] and in related areas, such as the prediction of secondary structures [198-200], and intrinsic disorder [201-204]. A few studies have pointed to the impact of the quality of the evolutionary profile on the predictive performance, which in turns stem from the size and quality of the underlying sequence alignment databases [81, 205]. These considerations offer additional opportunities to improve predictive performance of future RBRs and DBRs predictors.

Another important consideration concerns the ability of current and future predictors to accurately discriminate between DBRs and RBRs. The authors of future methods should measure and comparatively assess cross-predictions and design their models to minimize them. While in some aspects DNA and RNA are relatively similar, they are distinct in the structures of their binding interfaces [206, 207]. The π -stacking interactions between the aromatic amino acids and the nucleobases or sugar moieties of DNA and RNA play vital role in protein-nucleic acid recognition [208]. Previous studies highlighted differences between stacking interactions with DNA and RNA in terms of their rate of occurrences and preferences of amino acids with respective nucleobases [209, 210]. These fine details, which are typically extracted from 3D structures of protein-nucleic acids complexes, could be perhaps approximated from protein sequences or sequence-predicted protein structures, providing a way to improve the ability to distinguish between DBRs and RBRs. Moreover, future comparative assessments should cover the cross-prediction aspect to reveal which current methods accurately identify DNA vs RNA binding residues and which are nucleic acid type agnostic. The last such study was published in 2020 [34] and is relative outdated given the large number of methods that were released subsequently. Furthermore, these assessments should cover cross-predictions between DBR and RBRs and also between DBRs/RBRs and residues that interact with other types of ligands, such as proteins, peptides and small molecules. Similar evaluations were done recently in the context of developing predictors of the protein-binding residues [211].

We also emphasize the substantial impact of ensuring sustained/long-term availability of web servers for the predictors of RBRs and DBRs. The current 40% availability rate should be improved to match levels in other areas, such as the 71% rate for the intrinsic disorder predictors [40]. As shown in a recent study [176] and our analysis, this is likely to increase their scientific impact that is indirectly measured by their citation rates. This can be accomplished by

requiring the commitment to support web servers for an extended period of time at the point of publication, which would benefit both the developers and users.

Lastly, the structures of the protein-nucleic acids complexes are useful to investigate atomic level details these interactions. They can be predicted when the native structures are unavailable, which is relatively common. Many predictors are available including several docking-based tools [212-215]. The recently released Flex-LZerD that considers flexibility of the protein upon docking to nucleic acids [216], which at least partly addresses predictions for the disordered binding regions. The release of AlphaFold3 that predicts structures of protein-ligand complexes, where ligands include proteins, nucleic acids, small molecules and ions [217], is also notable. However, AlphaFold3 authors note that their model generates “*spurious structural order (hallucinations) in disordered regions*” [218], which is a major drawback in the context of prediction of nucleic acid binding residues that frequently reside in the disordered regions [46, 47, 49-51]. Moreover, these methods can be applied only when the structure of the nucleic acid is known, in contrast to the tools that we review which make predictions solely from the protein sequence.

Key points

- 87 predictors of nucleic acid-binding residues in protein sequences were developed in the last two decades
- Machine learning is the primary approach to develop these predictors
- Recent use of deep learning and protein language models resulted in substantial gains in predictive performance
- Cross-predictions between RNA-binding and DNA-binding residues are a significant challenge
- Predictors with working web servers enjoy high citation rates, motivating development and long-term maintenance of web servers

Competing Interests

The authors declare no competing interests.

Acknowledgements

This work was supported in part by the National Science Foundation (grants DBI 2146027 and IIS 2125218) and the Robert J. Matlack Endowment funds to L.K. D.K. acknowledges support from National Science Foundation (DBI 2146026, IIS 2211598, DMS 2151678) and National Institutes of Health (R01GM133840).

References

1. Djebali, S., et al., *Landscape of transcription in human cells*. Nature, 2012. **489**(7414): p. 101-108.
2. Evande, R., et al., *Protein-DNA Interactions Regulate Human Papillomavirus DNA Replication, Transcription, and Oncogenesis*. International Journal of Molecular Sciences, 2023. **24**(10).
3. Cozzolino, F., et al., *Protein-DNA/RNA Interactions: An Overview of Investigation Methods in the -Omics Era*. Journal of Proteome Research, 2021. **20**(6): p. 3018-3030.

4. Oyejobi, G.K., et al., *Regulating Protein-RNA Interactions: Advances in Targeting the LIN28/Let-7 Pathway*. International Journal of Molecular Sciences, 2024. **25**(7).
5. Peng, Z., et al., *A creature with a hundred waggly tails: intrinsically disordered proteins in the ribosome*. Cell Mol Life Sci, 2014. **71**(8): p. 1477-504.
6. Cookson, M.R., *RNA-binding proteins implicated in neurodegenerative diseases*. Wiley Interdisciplinary Reviews-Rna, 2017. **8**(1).
7. Bonczek, O., et al., *DNA and RNA Binding Proteins: From Motifs to Roles in Cancer*. International Journal of Molecular Sciences, 2022. **23**(16).
8. Tao, Y.N., et al., *Alternative splicing and related RNA binding proteins in human health and disease*. Signal Transduction and Targeted Therapy, 2024. **9**(1).
9. Valuchova, S., et al., *A rapid method for detecting protein-nucleic acid interactions by protein induced fluorescence enhancement*. Scientific Reports, 2016. **6**.
10. Chen, S., K. Yan, and B. Liu, *PDB-BRE: A ligand-protein interaction binding residue extractor based on Protein Data Bank*. Proteins, 2024. **92**(1): p. 145-153.
11. Burley, S.K., et al., *RCSB Protein Data Bank: powerful new tools for exploring 3D structures of biological macromolecules for basic and applied research and education in fundamental biology, biomedicine, biotechnology, bioengineering and energy sciences*. Nucleic Acids Res, 2021. **49**(D1): p. D437-D451.
12. UniProt, C., *UniProt: the Universal Protein Knowledgebase in 2023*. Nucleic Acids Res, 2023. **51**(D1): p. D523-D531.
13. Zhu, X.L., S.S. Ericksen, and J.C. Mitchell, *DBSI: DNA-binding site identifier*. Nucleic Acids Research, 2013. **41**(16).
14. Tsuchiya, Y., K. Kinoshita, and H. Nakamura, *Structure-based prediction of DNA-binding sites on proteins using the empirical preference of electrostatic potential and the shape of molecular surfaces*. Proteins-Structure Function and Bioinformatics, 2004. **55**(4): p. 885-894.
15. Tjong, H. and H.X. Zhou, *DISPLAR: an accurate method for predicting DNA-binding sites on protein surfaces*. Nucleic Acids Research, 2007. **35**(5): p. 1465-1477.
16. Zhou, W.Q. and H. Yan, *A discriminatory function for prediction of protein-DNA interactions based on alpha shape modeling*. Bioinformatics, 2010. **26**(20): p. 2541-2548.
17. Li, S., et al., *Quantifying sequence and structural features of protein-RNA interactions*. Nucleic Acids Res, 2014. **42**(15): p. 10086-98.
18. Lam, J.H., et al., *A deep learning framework to predict binding preference of RNA constituents on protein surface*. Nature Communications, 2019. **10**.
19. Xia, Y., et al., *GraphBind: protein structural context embedded rules learned by hierarchical graph neural networks for recognizing nucleic-acid-binding residues*. Nucleic Acids Research, 2021. **49**(9).
20. Li, P.P. and Z.P. Liu, *GeoBind: segmentation of nucleic acid binding interface on protein surface with geometric deep learning*. Nucleic Acids Research, 2023. **51**(10).
21. Sagendorf, J.M., et al., *Structure-based prediction of protein-nucleic acid binding using graph neural networks*. Biophys Rev, 2024. **16**(3): p. 297-314.
22. Varadi, M., et al., *AlphaFold Protein Structure Database: massively expanding the structural coverage of protein-sequence space with high-accuracy models*. Nucleic Acids Res, 2022. **50**(D1): p. D439-D444.
23. Roche, R., et al., *EquiPNAS: improved protein-nucleic acid binding site prediction using protein-language-model-informed equivariant deep graph neural networks*. Nucleic Acids Research, 2024. **52**(5).
24. Shi, W.T., et al., *GraphSite: Ligand Binding Site Classification with Deep Graph Learning*. Biomolecules, 2022. **12**(8).

25. Li, X., et al., *GraphNABP: Identifying nucleic acid-binding proteins with protein graphs and protein language models*. International Journal of Biological Macromolecules, 2024. **280**.
26. Ahmad, S., M.M. Gromiha, and A. Sarai, *Analysis and prediction of DNA-binding proteins and their binding residues based on composition, sequence and structural information*. Bioinformatics, 2004. **20**(4): p. 477-86.
27. Jeong, E., I.F. Chung, and S. Miyano, *A neural network method for identification of RNA-interacting residues in protein*. Genome Inform, 2004. **15**(1): p. 105-16.
28. Zhao, H.Y., Y.D. Yang, and Y.Q. Zhou, *Prediction of RNA binding proteins comes of age from low resolution to high resolution*. Molecular Biosystems, 2013. **9**(10): p. 2417-2425.
29. Si, J.N., R. Zhao, and R.L. Wu, *An Overview of the Prediction of Protein DNA-Binding Sites*. International Journal of Molecular Sciences, 2015. **16**(3): p. 5194-5215.
30. Cui, F.F., et al., *Protein-DNA/RNA interactions: Machine intelligence tools and approaches in the era of artificial intelligence and big data*. Proteomics, 2022. **22**(8).
31. Emamjomeh, A., et al., *DNA-protein interaction: identification, prediction and data analysis*. Molecular Biology Reports, 2019. **46**(3): p. 3571-3596.
32. Zhang, J., Z. Ma, and L. Kurgan, *Comprehensive review and empirical analysis of hallmarks of DNA-, RNA- and protein-binding residues in protein chains*. Brief Bioinform, 2019. **20**(4): p. 1250-1268.
33. Yan, J., S. Friedrich, and L. Kurgan, *A comprehensive comparative review of sequence-based predictors of DNA- and RNA-binding residues*. Briefings in Bioinformatics, 2016. **17**(1): p. 88-105.
34. Wang, K., et al., *Comprehensive Survey and Comparative Assessment of RNA-Binding Residue Predictions with Analysis by RNA Type*. International Journal of Molecular Sciences, 2020. **21**(18).
35. Gromiha, M.M. and R. Nagarajan, *Computational Approaches for Predicting the Binding Sites and Understanding the Recognition Mechanism of Protein-DNA Complexes*. Protein-Nucleic Acids Interactions, 2013. **91**: p. 65-99.
36. Nagarajan, R., S. Ahmad, and M.M. Gromiha, *Novel approach for selecting the best predictor for identifying the binding sites in DNA binding proteins*. Nucleic Acids Research, 2013. **41**(16): p. 7606-7614.
37. Liu, Z.P. and L.N. Chen, *Prediction and Dissection of Protein-RNA Interactions by Molecular Descriptors*. Current Topics in Medicinal Chemistry, 2016. **16**(6): p. 604-615.
38. Si, J.N., et al., *Computational Prediction of RNA-Binding Proteins and Binding Sites*. International Journal of Molecular Sciences, 2015. **16**(11): p. 26303-26317.
39. Walia, R.R., et al., *Protein-RNA interface residue prediction using machine learning: an assessment of the state of the art*. BMC Bioinformatics, 2012. **13**.
40. Basu, S., D. Kihara, and L. Kurgan, *Computational prediction of disordered binding regions*. Computational and Structural Biotechnology Journal, 2023. **21**: p. 1487-1497.
41. Miao, Z. and E. Westhof, *A Large-Scale Assessment of Nucleic Acids Binding Site Prediction Programs*. Plos Computational Biology, 2015. **11**(12).
42. Dunker, A.K., et al., *Intrinsically disordered protein*. Journal of Molecular Graphics and Modelling, 2001. **19**(1): p. 26-59.
43. Tompa, P., et al., *Close encounters of the third kind: disordered domains and the interactions of proteins*. Bioessays, 2009. **31**(3): p. 328-35.
44. Lieutaud, P., et al., *How disordered is my protein and what is its disorder for? A guide through the "dark side" of the protein universe*. Intrinsically Disord Proteins, 2016. **4**(1): p. e1259708.
45. Habchi, J., et al., *Introducing protein intrinsic disorder*. Chem Rev, 2014. **114**(13): p. 6561-88.
46. Basu, S. and R.P. Bahadur, *A structural perspective of RNA recognition by intrinsically disordered proteins*. Cellular and Molecular Life Sciences, 2016. **73**(21): p. 4075-4084.

47. Wang, C., V.N. Uversky, and L. Kurgan, *Disordered nucleome: Abundance of intrinsic disorder in the DNA- and RNA-binding proteins in 1121 species from Eukaryota, Bacteria and Archaea*. Proteomics, 2016. **16**(10): p. 1486-98.
48. Zhao, B., et al., *Intrinsic Disorder in Human RNA-Binding Proteins*. Journal of Molecular Biology, 2021. **433**(21).
49. Varadi, M., et al., *Functional Advantages of Conserved Intrinsic Disorder in RNA-Binding Proteins*. Plos One, 2015. **10**(10).
50. Dyson, J., *Role of intrinsic disorder in protein-protein and protein-nucleic acid interactions*. Febs Journal, 2012. **279**: p. 9-9.
51. Munshi, S., et al., *Tunable order-disorder continuum in protein-DNA interactions*. Nucleic Acids Research, 2018. **46**(17): p. 8700-8709.
52. Hsu, W.L., et al., *Exploring the binding diversity of intrinsically disordered proteins involved in one-to-many binding*. Protein Science, 2013. **22**(3): p. 258-273.
53. Fung, H.Y.J., M. Birol, and E. Rhoades, *IDPs in macromolecular complexes: the roles of multivalent interactions in diverse assemblies*. Current Opinion in Structural Biology, 2018. **49**: p. 36-43.
54. Williams, R.M., et al., *The protein non-folding problem: amino acid determinants of intrinsic order and disorder*. Pac Symp Biocomput, 2001: p. 89-100.
55. Wu, Z.H., et al., *In various protein complexes, disordered protomers have large per-residue surface areas and area of protein-, DNA- and RNA-binding interfaces*. Febs Letters, 2015. **589**(19): p. 2561-2569.
56. Zhao, B. and L. Kurgan, *Compositional Bias of Intrinsically Disordered Proteins and Regions and Their Predictions*. Biomolecules, 2022. **12**(7).
57. Zhao, H.Y., Y.D. Yang, and Y.Q. Zhou, *Highly accurate and high-resolution function prediction of RNA binding proteins by fold recognition and binding affinity prediction*. Rna Biology, 2011. **8**(6): p. 988-996.
58. Patiyl, S., A. Dhall, and G.P.S. Raghava, *A deep learning-based method for the prediction of DNA interacting residues in a protein*. Briefings in Bioinformatics, 2022. **23**(5).
59. Liu, Y.F. and B.X. Tian, *Protein-DNA binding sites prediction based on pre-trained protein language model and contrastive learning*. Briefings in Bioinformatics, 2024. **25**(1).
60. Zhang, J.S., R.H. Wang, and L.Y. Wei, *MuLiPred: Multi-Level Contrastive Learning for Predicting Nucleic Acid Binding Residues of Proteins*. Journal of Chemical Information and Modeling, 2024. **64**(3): p. 1050-1065.
61. Qiu, J., et al., *ProNA2020 predicts protein-DNA, protein-RNA, and protein-protein binding proteins and residues from sequence*. J Mol Biol, 2020. **432**(7): p. 2428-2443.
62. Patiyl, S., et al., *Prediction of RNA-interacting residues in a protein using CNN and evolutionary profile*. Briefings in Bioinformatics, 2023. **24**(1).
63. Yang, J.Y., A. Roy, and Y. Zhang, *BioLiP: a semi-manually curated database for biologically relevant ligand-protein interactions*. Nucleic Acids Research, 2013. **41**(D1): p. D1096-D1103.
64. Zhang, C.X., et al., *BioLiP2: an updated structure database for biologically relevant ligand-protein interactions*. Nucleic Acids Research, 2024. **52**(D1): p. D404-D412.
65. Berman, H.M., et al., *The Protein Data Bank*. Nucleic Acids Research, 2000. **28**(1): p. 235-242.
66. Burley, S.K., et al., *RCSB Protein Data Bank (RCSB.org): delivery of experimentally-determined PDB structures alongside one million computed structure models of proteins from artificial intelligence/machine learning*. Nucleic Acids Research, 2023. **51**(D1): p. D488-D508.
67. Norambuena, T. and F. Melo, *The Protein-DNA Interface database*. BMC Bioinformatics, 2010. **11**.

68. Lewis, B.A., et al., *PRIDB: a protein-RNA interface database*. Nucleic Acids Research, 2011. **39**: p. D277-D282.
69. Katuwawala, A., et al., *Computational Prediction of MoRFs, Short Disorder-to-order Transitioning Protein Binding Regions*. Computational and Structural Biotechnology Journal, 2019. **17**: p. 454-462.
70. Del Conte, A., et al., *Critical assessment of protein intrinsic disorder prediction (CAID) - Results of round 2*. Proteins-Structure Function and Bioinformatics, 2023.
71. Zhang, J. and L. Kurgan, *Review and comparative assessment of sequence-based predictors of protein-binding residues*. Brief Bioinform, 2018. **19**(5): p. 821-837.
72. Wang, C. and L. Kurgan, *Review and comparative assessment of similarity-based methods for prediction of drug-protein interactions in the druggable human proteome*. Brief Bioinform, 2019. **20**(6): p. 2066-2087.
73. Katuwawala, A., C.J. Oldfield, and L. Kurgan, *Accuracy of protein-level disorder predictions*. Brief Bioinform, 2020. **21**(5): p. 1509-1522.
74. Liu, Y.C., et al., *Deciphering the Language of Protein-DNA Interactions: A Deep Learning Approach Combining Contextual Embeddings and Multi-Scale Sequence Modeling*. Journal of Molecular Biology, 2024. **436**(22).
75. Zhang, L.R. and T.G. Liu, *PDNAPred: Interpretable prediction of protein-DNA binding sites based on pre-trained protein language models*. International Journal of Biological Macromolecules, 2024. **281**.
76. Andrabi, M., et al., *Prediction of mono- and di-nucleotide-specific DNA-binding sites in proteins using neural networks*. BMC Struct Biol, 2009. **9**: p. 30.
77. Fernandez, M., et al., *Prediction of dinucleotide-specific RNA-binding sites in proteins*. BMC Bioinformatics, 2011. **12 Suppl 13**(Suppl 13): p. S5.
78. Zhang, J., et al., *DNAgenie: accurate prediction of DNA-type-specific binding residues in protein sequences*. Briefings in Bioinformatics, 2021. **22**(6).
79. Nie, W.J. and L. Deng, *TSNAPred: predicting type-specific nucleic acid binding residues via an ensemble approach*. Briefings in Bioinformatics, 2022. **23**(4).
80. Sun, C.Z. and Y.E. Feng, *EPDRNA: A Model for Identifying DNA-RNA Binding Sites in Disease-Related Proteins*. Protein Journal, 2024. **43**(3): p. 513-521.
81. Ahmad, S. and A. Sarai, *PSSM-based prediction of DNA binding sites in proteins*. BMC Bioinformatics, 2005. **6**.
82. Wang, L. and S.J. Brown, *BindN: a web-based tool for efficient prediction of DNA and RNA binding sites in amino acid sequences*. Nucleic Acids Res, 2006. **34**(Web Server issue): p. W243-8.
83. Yan, C.H., et al., *Predicting DNA-binding sites of proteins from amino acid sequence*. BMC Bioinformatics, 2006. **7**.
84. Jeong, E.N. and S. Miyano, *A weighted profile based method for protein-RNA interacting residue prediction*. Transactions on Computational Systems Biology IV, 2006. **3939**: p. 123-139.
85. Hwang, S., Z.K. Gou, and I.B. Kuznetsov, *DP-Bind: a Web server for sequence-based prediction of DNA-binding residues in DNA-binding proteins*. Bioinformatics, 2007. **23**(5): p. 634-636.
86. Kuznetsov, I.B., et al., *Using evolutionary and structural information to predict DNA-binding sites on DNA-binding proteins*. Proteins-Structure Function and Bioinformatics, 2006. **64**(1): p. 19-27.
87. Ofra, Y., V. Mysore, and B. Rost, *Prediction of DNA-binding residues from sequence*. Bioinformatics, 2007. **23**(13): p. 1347-1353.
88. Ho, S.Y., et al., *Design of accurate predictors for DNA-binding sites in proteins using hybrid SVM-PSSM method*. Biosystems, 2007. **90**(1): p. 234-241.
89. Terribilini, M., et al., *RNABindR: a server for analyzing and predicting RNA-binding sites in proteins*. Nucleic Acids Res, 2007. **35**(Web Server issue): p. W578-84.

90. Terribilini, M., et al., *Prediction of RNA binding sites in proteins from amino acid sequence*. RNA, 2006. **12**(8): p. 1450-62.
91. Kumar, M., A.M. Gromiha, and G.P.S. Raghava, *Prediction of RNA binding sites in a protein using SVM and PSSM profile*. Proteins-Structure Function and Bioinformatics, 2008. **71**(1): p. 189-194.
92. Wang, Y., et al., *PRINTR: Prediction of RNA binding sites in proteins using SVM and profiles*. Amino Acids, 2008. **35**(2): p. 295-302.
93. Tong, J., P. Jiang, and Z.H. Lu, *RISP: A web-based server for prediction of RNA-binding sites in proteins*. Computer Methods and Programs in Biomedicine, 2008. **90**(2): p. 148-153.
94. Cheng, C.W., et al., *Predicting RNA-binding sites of proteins using support vector machines and evolutionary information*. BMC Bioinformatics, 2008. **9**.
95. Wang, L., M.Q. Yang, and J.Y. Yang, *Prediction of DNA-binding residues from protein sequence information using random forests*. BMC Genomics, 2009. **10 Suppl 1**(Suppl 1): p. S1.
96. Gao, M. and J. Skolnick, *A Threading-Based Method for the Prediction of DNA-Binding Proteins with Application to the Human Genome*. Plos Computational Biology, 2009. **5**(11).
97. Wu, J., et al., *Prediction of DNA-binding residues in proteins from amino acid sequences using a random forest model with a hybrid feature*. Bioinformatics, 2009. **25**(1): p. 30-5.
98. Chu, W.Y., et al., *ProteDNA: a sequence-based predictor of sequence-specific DNA-binding residues in transcription factors*. Nucleic Acids Research, 2009. **37**: p. W396-W401.
99. Murakami, Y., et al., *PiRaNha: a server for the computational prediction of RNA-binding residues in protein sequences*. Nucleic Acids Research, 2010. **38**: p. W412-W416.
100. Spriggs, R.V., et al., *Protein function annotation from sequence: prediction of residues interacting with RNA*. Bioinformatics, 2009. **25**(12): p. 1492-7.
101. Wang, L., et al., *BindN+ for accurate prediction of DNA and RNA-binding residues from protein sequence features*. BMC Syst Biol, 2010. **4 Suppl 1**(Suppl 1): p. S3.
102. Carson, M.B., R. Langlois, and H. Lu, *NAPS: a residue-level nucleic acid-binding prediction server*. Nucleic Acids Res, 2010. **38**(Web Server issue): p. W431-5.
103. Liu, Z.P., et al., *Prediction of protein-RNA binding sites by a random forest method with combined features*. Bioinformatics, 2010. **26**(13): p. 1616-1622.
104. Huang, Y.F., et al., *Predicting RNA-binding residues from evolutionary information and sequence conservation*. BMC Genomics, 2010. **11**.
105. Zhang, T., et al., *Analysis and prediction of RNA-binding residues using sequence, evolutionary conservation, and predicted secondary structure and solvent accessibility*. Curr Protein Pept Sci, 2010. **11**(7): p. 609-28.
106. Li, Q., Z. Cao, and H. Liu, *Improve the prediction of RNA-binding residues using structural neighbours*. Protein Pept Lett, 2010. **17**(3): p. 287-96.
107. Si, J., et al., *MetaDBSite: a meta approach to improve protein DNA-binding sites prediction*. BMC Syst Biol, 2011. **5 Suppl 1**(Suppl 1): p. S7.
108. Ma, X., et al., *Prediction of RNA-binding residues in proteins from primary sequence using an enriched random forest model with a novel hybrid feature*. Proteins, 2011. **79**(4): p. 1230-9.
109. Choi, S. and K. Han, *Prediction of RNA-binding amino acids from protein and RNA sequences*. BMC Bioinformatics, 2011. **12 Suppl 13**(Suppl 13): p. S7.
110. Wang, C.C., et al., *Identification of RNA-binding sites in proteins by integrating various sequence information*. Amino Acids, 2011. **40**(1): p. 239-48.
111. Ma, X., et al., *Sequence-Based Prediction of DNA-Binding Residues in Proteins with Conservation and Correlation Information*. Ieee-Acm Transactions on Computational Biology and Bioinformatics, 2012. **9**(6): p. 1766-1775.
112. Puton, T., et al., *Computational methods for prediction of protein-RNA interactions*. Journal of Structural Biology, 2012. **179**(3): p. 261-268.

113. Yu, D.J., et al., *Designing Template-Free Predictor for Targeting Protein-Ligand Binding Sites with Classifier Ensemble and Spatial Clustering*. Ieee-Acm Transactions on Computational Biology and Bioinformatics, 2013. **10**(4): p. 994-1008.
114. Pan, X.Y., et al., *Predicting protein-RNA interaction amino acids using random forest based on submodularity subset selection*. Computational Biology and Chemistry, 2014. **53**: p. 324-330.
115. Zhao, H.Y., et al., *Predicting DNA-Binding Proteins and Binding Residues by Complex Structure Prediction and Application to Human Proteome*. Plos One, 2014. **9**(5).
116. Yang, Y.D., et al., *SPOT-Seq-RNA: Predicting Protein-RNA Complex Structure and RNA-Binding Function by Fold Recognition and Binding Affinity Prediction*. Protein Structure Prediction, 3rd Edition, 2014. **1137**: p. 119-130.
117. Walia, R.R., et al., *RNABindRPlus: A Predictor that Combines Machine Learning and Sequence Homology-Based Methods to Improve the Reliability of Predicted RNA-Binding Residues in Proteins*. Plos One, 2014. **9**(5).
118. Xiong, D.P., J.Y. Zeng, and H.P. Gong, *RBRIdent: An algorithm for improved identification of RNA-binding residues in proteins from primary sequences*. Proteins-Structure Function and Bioinformatics, 2015. **83**(6): p. 1068-1077.
119. Yang, X.X., et al., *SNBRFinder: A Sequence-Based Hybrid Algorithm for Enhanced Prediction of Nucleic Acid-Binding Residues*. Plos One, 2015. **10**(7).
120. Peng, Z.L. and L. Kurgan, *High-throughput prediction of RNA, DNA and protein binding regions mediated by intrinsic disorder*. Nucleic Acids Research, 2015. **43**(18).
121. Peng, Z.L., et al., *Prediction of Disordered RNA, DNA, and Protein Binding Regions Using DisoRDPbind*. Prediction of Protein Secondary Structure, 2017. **1484**: p. 187-203.
122. Miao, Z.C. and E. Westhof, *Prediction of nucleic acid binding probability in proteins: a neighboring residue network based score*. Nucleic Acids Research, 2015. **43**(11): p. 5340-5351.
123. Dang, T.K.L., et al., *A Novel Sequence-Based Feature for the Identification of DNA-Binding Sites in Proteins Using Jensen-Shannon Divergence*. Entropy, 2016. **18**(10).
124. Chai, H., et al., *An evolution-based DNA-binding residue predictor using a dynamic query-driven learning scheme*. Molecular Biosystems, 2016. **12**(12): p. 3643-3650.
125. EL-Manzalawy, Y., et al., *FastRNABindR: Fast and Accurate Prediction of Protein-RNA Interface Residues*. Plos One, 2016. **11**(7).
126. Hu, J., et al., *Predicting Protein-DNA Binding Residues by Weightedly Combining Sequence-Based Features and Boosting Multiple SVMs*. Ieee-Acm Transactions on Computational Biology and Bioinformatics, 2017. **14**(6): p. 1389-1398.
127. Yan, J. and L. Kurgan, *DRNAPred, fast sequence-based method that accurately predicts and discriminates DNA- and RNA-binding residues*. Nucleic Acids Research, 2017. **45**(10).
128. Shen, C., et al., *Identification of DNA-protein Binding Sites through Multi-Scale Local Average Blocks on Sequence Information*. Molecules, 2017. **22**(12).
129. Zhou, J.Y., et al., *EL_PSSM-RT: DNA-binding residue prediction by integrating ensemble learning with PSSM Relation Transformation*. BMC Bioinformatics, 2017. **18**.
130. Deng, L., et al., *PDRLGB: precise DNA-binding residue prediction using a light gradient boosting machine*. BMC Bioinformatics, 2018. **19**.
131. Amirkhani, A., et al., *Prediction of DNA-Binding Residues in Local Segments of Protein Sequences with Fuzzy Cognitive Maps*. Ieee-Acm Transactions on Computational Biology and Bioinformatics, 2020. **17**(4): p. 1372-1382.
132. Zhu, Y.H., et al., *DNAPred: Accurate Identification of DNA-Binding Sites from Protein Sequence by Ensembled Hyperplane-Distance-Based Support Vector Machines*. Journal of Chemical Information and Modeling, 2019. **59**(6): p. 3057-3071.

133. Su, H., et al., *Improving the prediction of protein-nucleic acids binding residues via multiple sequence profiles and the consensus of complementary methods*. Bioinformatics, 2019. **35**(6): p. 930-936.
134. Jung, Y., et al., *Partner-specific prediction of RNA-binding residues in proteins: A critical assessment*. Proteins-Structure Function and Bioinformatics, 2019. **87**(3): p. 198-211.
135. Nguyen, B.P., et al., *iProDNA-CapsNet: identifying protein-DNA binding residues using capsule neural networks*. BMC Bioinformatics, 2019. **20**(1).
136. Zhou, J.Y., et al., *EL_LSTM: Prediction of DNA-Binding Residue from Protein Sequence by Combining Long Short-Term Memory and Ensemble Learning*. IEEE/ACM Transactions on Computational Biology and Bioinformatics, 2020. **17**(1): p. 124-135.
137. Yao, L.S., H.D. Wang, and Y.N. Bin, *Predicting Hot Spot Residues at Protein-DNA Binding Interfaces Based on Sequence Information*. Interdisciplinary Sciences-Computational Life Sciences, 2021. **13**(1): p. 1-11.
138. Zhang, J., Q.C. Chen, and B. Liu, *NCBRPred: predicting nucleic acid binding residues in proteins based on multilabel learning*. Briefings in Bioinformatics, 2021. **22**(5).
139. Littmann, M., et al., *Protein embeddings and deep learning predict binding residues for various ligand classes*. Scientific Reports, 2021. **11**(1).
140. Sun, Z., et al., *To Improve Prediction of Binding Residues With DNA, RNA, Carbohydrate, and Peptide Via Multi-Task Deep Neural Networks*. IEEE/ACM Trans Comput Biol Bioinform, 2022. **19**(6): p. 3735-3743.
141. Zhang, F.H., et al., *DeepDISOBind: accurate prediction of RNA-, DNA- and protein-binding intrinsically disordered residues with deep multi-task learning*. Briefings in Bioinformatics, 2022. **23**(1).
142. Hu, J., et al., *Protein-DNA Binding Residue Prediction via Bagging Strategy and Sequence-Based Cube-Format Feature*. IEEE/ACM Trans Comput Biol Bioinform, 2022. **19**(6): p. 3635-3645.
143. Wang, N., et al., *iDRNA-ITF: identifying DNA- and RNA-binding residues in proteins based on induction and transfer framework*. Brief Bioinform, 2022. **23**(4).
144. Guan, S., et al., *Protein-DNA Binding Residues Prediction Using a Deep Learning Model With Hierarchical Feature Extraction*. IEEE/ACM Trans Comput Biol Bioinform, 2023. **20**(5): p. 2619-2628.
145. Agarwal, A., S. Kant, and R.P. Bahadur, *Efficient mapping of RNA-binding residues in RNA-binding proteins using local sequence features of binding site residues in protein-RNA complexes*. Proteins, 2023. **91**(9): p. 1361-1379.
146. Zhang, F., et al., *HybridRNAbind: prediction of RNA interacting residues across structure-annotated and disorder-annotated proteins*. Nucleic Acids Res, 2023. **51**(5): p. e25.
147. Song, Y.D., et al., *Accurately identifying nucleic-acid-binding sites through geometric graph learning on language model predicted structures*. Briefings in Bioinformatics, 2023. **24**(6).
148. Zhang, J., S. Basu, and L. Kurgan, *HybridDBRpred: improved sequence-based prediction of DNA-binding amino acids using annotations from structured complexes and disordered proteins*. Nucleic Acids Res, 2024. **52**(2): p. e10.
149. Kabir, M.W.U., et al., *DRBpred: A sequence-based machine learning method to effectively predict DNA- and RNA-binding residues*. Comput Biol Med, 2024. **170**: p. 108081.
150. Zhu, Y.H., et al., *ULDNA: integrating unsupervised multi-source language models with LSTM-attention network for high-accuracy protein-DNA binding site prediction*. Brief Bioinform, 2024. **25**(2).
151. Zhang, B., et al., *SOFB is a comprehensive ensemble deep learning approach for elucidating and characterizing protein-nucleic-acid-binding residues*. Commun Biol, 2024. **7**(1): p. 679.

152. Yuan, Q.M., et al., *GPSFun: geometry-aware protein sequence function predictions with language models*. Nucleic Acids Research, 2024. **52**(W1): p. W248-W255.
153. Yuan, Q.M., C. Tian, and Y.D. Yang, *Genome-scale annotation of protein binding sites via language model and geometric deep learning*. Elife, 2024. **13**.
154. Khan, Y.D., et al., *DeepDBS: Identification of DNA-binding sites in protein sequences by using deep representations and random forest*. Methods, 2024. **231**: p. 26-36.
155. Wu, J.S., et al., *Prediction of DNA-binding residues in proteins from amino acid sequences using a random forest model with a hybrid feature*. Bioinformatics, 2009. **25**(1): p. 30-35.
156. Yu, L.J., et al., *Grammar of protein domain architectures*. Proceedings of the National Academy of Sciences of the United States of America, 2019. **116**(9): p. 3636-3645.
157. Ofer, D., N. Brandes, and M. Linial, *The language of proteins: NLP, machine learning & protein sequences*. Computational and Structural Biotechnology Journal, 2021. **19**: p. 1750-1758.
158. Zhang, S., et al., *Applications of transformer-based language models in bioinformatics: a survey*. Neuro-Oncology Advances, 2023. **5**(1).
159. Asgari, E. and M.R.K. Mofrad, *Continuous Distributed Representation of Biological Sequences for Deep Proteomics and Genomics*. Plos One, 2015. **10**(11).
160. Elnaggar, A., et al., *ProtTrans: Toward Understanding the Language of Life Through Self-Supervised Learning*. IEEE Trans Pattern Anal Mach Intell, 2022. **44**(10): p. 7112-7127.
161. Zhu, Y. and G. Wang, *CAN-NER: Convolutional Attention Network for Chinese Named Entity Recognition*. 2019 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies (Naacl Hlt 2019), Vol. 1, 2019: p. 3384-3393.
162. Lin, Z.M., et al., *Evolutionary-scale prediction of atomic-level protein structure with a language model*. Science, 2023. **379**(6637): p. 1123-1130.
163. Rao, R.M., et al., *MSA Transformer*, in *Proceedings of the 38th International Conference on Machine Learning*, M. Marina and Z. Tong, Editors. 2021, PMLR: Proceedings of Machine Learning Research. p. 8844--8856.
164. Zhang, J. and L. Kurgan, *Review and comparative assessment of sequence-based predictors of protein-binding residues*. Briefings in Bioinformatics, 2018. **19**(5): p. 821-837.
165. Ghosh, S. and A. Bagchi, *Structural study to analyze the DNA-binding properties of DsrC protein from the dsr operon of sulfur-oxidizing bacterium Allochromatium vinosum*. Journal of Molecular Modeling, 2019. **25**(3).
166. Wang, J., et al., *Regulation of micro- and small-exon retention and other splicing processes by GRP20 for flower development*. Nat Plants, 2024. **10**(1): p. 66-85.
167. Zhang, Z., et al., *Novel LncRNA LINC02936 Suppresses Ferroptosis and Promotes Tumor Progression by Interacting with SIX1/CP Axis in Endometrial Cancer*. Int J Biol Sci, 2024. **20**(4): p. 1356-1374.
168. Chong Qui, E., et al., *Mycobacteriophage Alexphander Gene 94 Encodes an Essential dsDNA-Binding Protein during Lytic Infection*. Int J Mol Sci, 2024. **25**(13).
169. Yadhav, Y., et al., *Computational studies on rep and capsid proteins of CRESS DNA viruses*. Virusdisease, 2024. **35**(1): p. 17-26.
170. Reinart, W.B., et al., *Adaptive protein evolution through length variation of short tandem repeats in*. Science Advances, 2023. **9**(12).
171. Lanclos, N., et al., *Implications of intrinsic disorder and functional proteomics in the merkel cell polyomavirus life cycle*. J Cell Biochem, 2023.
172. Zhao, B., et al., *Intrinsic Disorder in Human RNA-Binding Proteins*. J Mol Biol, 2021. **433**(21): p. 167229.

173. Giri, R., et al., *Understanding COVID-19 via comparative analysis of dark proteomes of SARS-CoV-2, human SARS and bat SARS-like coronaviruses*. Cell Mol Life Sci, 2020.
174. Basu, S., et al., *DescribePROT in 2023: more, higher-quality and experimental annotations and improved data download options*. Nucleic Acids Res, 2024. **52**(D1): p. D426-D433.
175. Zhao, B., et al., *DescribePROT: database of amino acid-level protein structure and function predictions*. Nucleic Acids Res, 2021. **49**(D1): p. D298-D308.
176. Song, J. and L. Kurgan, *Availability of web servers significantly boosts citations rates of bioinformatics methods for protein function and disorder prediction*. Bioinform Adv, 2023. **3**(1): p. vbad184.
177. Bileschi, M.L., et al., *Using deep learning to annotate the protein universe*. Nature Biotechnology, 2022. **40**(6): p. 932-+.
178. Boadu, F., A.H.Y. Lee, and J.L. Cheng, *Deep learning methods for protein function prediction*. Proteomics, 2024.
179. Zhao, B. and L. Kurgan, *Deep learning in prediction of intrinsic disorder in proteins*. Computational and Structural Biotechnology Journal, 2022. **20**: p. 1286-1294.
180. Vucetic, S., et al., *DisProt: a database of protein disorder*. Bioinformatics, 2005. **21**(1): p. 137-40.
181. Aspromonte, M.C., et al., *DisProt in 2024: improving function annotation of intrinsically disordered proteins*. Nucleic Acids Research, 2023.
182. Necci, M., et al., *Critical assessment of protein intrinsic disorder prediction*. Nat Methods, 2021. **18**(5): p. 472-481.
183. Conte, A.D., et al., *Critical assessment of protein intrinsic disorder prediction (CAID) - Results of round 2*. Proteins, 2023.
184. Piovesan, D., et al., *DisProt 7.0: a major update of the database of disordered proteins*. Nucleic Acids Res, 2016. **D1**: p. D219-D227.
185. Katuwawala, A., S. Ghadermarzi, and L. Kurgan, *Computational prediction of functions of intrinsically disordered regions*. Prog Mol Biol Transl Sci, 2019. **166**: p. 341-369.
186. Aspromonte, M.C., et al., *DisProt in 2024: improving function annotation of intrinsically disordered proteins*. Nucleic Acids Res, 2024. **52**(D1): p. D434-D441.
187. Zhang, J., S. Ghadermarzi, and L. Kurgan, *Prediction of protein-binding residues: dichotomy of sequence-based methods developed using structured complexes versus disordered proteins*. Bioinformatics, 2020. **36**(18): p. 4729-4738.
188. Hsieh, J., A.J. Andrews, and C.A. Fierke, *Roles of protein subunits in RNA-protein complexes: Lessons from ribonuclease P*. Biopolymers, 2004. **73**(1): p. 79-89.
189. Madru, C., et al., *Chaperoning 5S RNA assembly*. Genes & Development, 2015. **29**(13): p. 1432-1446.
190. Wetzel, J.L., K.Q. Zhang, and M. Singh, *Learning probabilistic protein-DNA recognition codes from DNA-binding specificities using structural mappings*. Genome Research, 2022. **32**(9): p. 1776-1786.
191. Mitra, R., et al., *Geometric deep learning of protein-DNA binding specificity*. Nature Methods, 2024. **21**(9).
192. Christensen, R.G., et al., *Recognition models to predict DNA-binding specificities of homeodomain proteins*. Bioinformatics, 2012. **28**(12): p. 184-189.
193. Alipanahi, B., et al., *Predicting the sequence specificities of DNA- and RNA-binding proteins by deep learning*. Nature Biotechnology, 2015. **33**(8): p. 831-+.
194. Yanover, C. and P. Bradley, *Extensive protein and DNA backbone sampling improves structure-based specificity prediction for C*

zinc fingers. Nucleic Acids Research, 2011. **39**(11): p. 4564-4576.

195. Pang, Y.H. and B. Liu, *IDP-LM: Prediction of protein intrinsic disorder and disorder functions based on language models*. Plos Computational Biology, 2023. **19**(11).
196. Madani, A., et al., *Large language models generate functional protein sequences across diverse families*. Nature Biotechnology, 2023. **41**(8): p. 1099-+.
197. Shuai, R.W., J.A. Ruffolo, and J.J. Gray, *IgLM: Infilling language modeling for antibody sequence design*. Cell Systems, 2023. **14**(11): p. 979-+.
198. Dong, B., et al., *Impact of Multi-Factor Features on Protein Secondary Structure Prediction*. Biomolecules, 2024. **14**(9).
199. Zhang, J., et al., *Recent Advances in Computational Prediction of Secondary and Supersecondary Structures from Protein Sequences*, in *Protein Supersecondary Structures: Methods and Protocols*, A.E. Kister, Editor. 2025, Springer US: New York, NY. p. 1-19.
200. Jiang, Q., et al., *Protein secondary structure prediction: A survey of the state of the art*. Journal of Molecular Graphics & Modelling, 2017. **76**: p. 379-402.
201. Mirabello, C. and B. Wallner, *rawMSA: End-to-end Deep Learning using raw Multiple Sequence Alignments*. Plos One, 2019. **14**(8).
202. Basu, S., D. Kihara, and L. Kurgan, *Computational prediction of disordered binding regions*. Comput Struct Biotechnol J, 2023. **21**: p. 1487-1497.
203. Zhao, B. and L. Kurgan, *Surveying over 100 predictors of intrinsic disorder in proteins*. Expert Review of Proteomics, 2021. **18**(12): p. 1019-1029.
204. Liu, Y., X. Wang, and B. Liu, *A comprehensive review and comparison of existing computational methods for intrinsically disordered protein and region prediction*. Brief Bioinform, 2019. **20**(1): p. 330-346.
205. Przybylski, D. and B. Rost, *Alignments grow, secondary structure prediction improves*. Proteins-Structure Function and Bioinformatics, 2002. **46**(2): p. 197-205.
206. Kulandaisamy, A., et al., *Dissecting and analyzing key residues in protein-DNA complexes*. J Mol Recognit, 2018. **31**(4).
207. Barik, A., et al., *Molecular architecture of protein-RNA recognition sites*. J Biomol Struct Dyn, 2015. **33**(12): p. 2738-51.
208. Baker, C.M. and G.H. Grant, *Role of aromatic amino acids in protein-nucleic acid recognition*. Biopolymers, 2007. **85**(5-6): p. 456-470.
209. Wilson, K.A., D.J. Holland, and S.D. Wetmore, *Topology of RNA-protein nucleobase-amino acid π - π interactions and comparison to analogous DNA-protein π - π contacts*. Rna, 2016. **22**(5): p. 696-708.
210. Wilson, K.A., et al., *Anatomy of noncovalent interactions between the nucleobases or ribose and π -containing amino acids in RNA-protein complexes*. Nucleic Acids Research, 2021. **49**(4): p. 2213-2225.
211. Zhang, F., et al., *PROBselect: accurate prediction of protein-binding residues from proteins sequences via dynamic predictor selection*. Bioinformatics, 2020. **36**(Supplement_2): p. i735-i744.
212. Tuszynska, I., et al., *NPDock: a web server for protein-nucleic acid docking*. Nucleic Acids Res, 2015. **43**(W1): p. W425-30.
213. Tuszynska, I. and J.M. Bujnicki, *DARS-RNP and QUASI-RNP: new statistical potentials for protein-RNA docking*. BMC Bioinformatics, 2011. **12**: p. 348.
214. Yan, Y., et al., *HDOCK: a web server for protein-protein and protein-DNA/RNA docking based on a hybrid strategy*. Nucleic Acids Res, 2017. **45**(W1): p. W365-W373.
215. Zheng, J., et al., *P3DOCK: a protein-RNA docking webserver based on template-based and template-free docking*. Bioinformatics, 2020. **36**(1): p. 96-103.

- 216. Christoffer, C. and D. Kihara, *Modeling protein-nucleic acid complexes with extremely large conformational changes using Flex-LZerD*. Proteomics, 2023. **23**(17): p. e2200322.
- 217. Abramson, J., et al., *Accurate structure prediction of biomolecular interactions with AlphaFold 3*. Nature, 2024. **630**(8016).
- 218. Abramson, J., et al., *Accurate structure prediction of biomolecular interactions with AlphaFold 3*. Nature, 2024.