

Structural bioinformatics

Text mining for modeling of protein complexes enhanced by machine learning

Varsha D. Badal¹, Petras J. Kundrotas^{1,*} and Ilya A. Vakser^{1,2,*}

¹Computational Biology Program and ²Department of Molecular Biosciences, The University of Kansas, Lawrence, Kansas 66045, USA

*To whom correspondence should be addressed.

Associate Editor: XXXXXXXX

Received on XXXXX; revised on XXXXX; accepted on XXXXX

Abstract

Motivation: Procedures for structural modeling of protein-protein complexes (protein docking) produce a number of models which need to be further analyzed and scored. Scoring can be based on independently determined constraints on the structure of the complex, such as knowledge of amino acids essential for the protein interaction. Previously, we showed that text mining of residues in freely available PubMed abstracts of papers on studies of protein-protein interactions may generate such constraints. However, absence of post-processing of the spotted residues reduced usability of the constraints, as a significant number of the residues were not relevant for the binding of the specific proteins.

Results: We explored filtering of the irrelevant residues by two machine learning approaches, Deep Recursive Neural Network (DRNN) and Support Vector Machine (SVM) models with different training/testing schemes. The results showed that the DRNN model is superior to the SVM model when training is performed on the PMC-OA full-text articles and applied to classification (interface or non-interface) of the residues spotted in the PubMed abstracts. When both training and testing is performed on full-text articles or on abstracts, the performance of these models is similar. Thus, in such cases, there is no need to utilize computationally demanding DRNN approach, which is computationally expensive especially at the training stage. The reason is that SVM success is often determined by the similarity in data/text patterns in the training and the testing sets, whereas the sentence structures in the abstracts are, in general, different from those in the full text articles.

Availability: The code and the datasets generated in this study are available at <https://gitlab.ku.edu/vakser-lab-public/text-mining/-/tree/2020-09-04>.

Contact: vakser@ku.edu or pkundro@ku.edu

Supplementary information: Supplementary data are available at *Bioinformatics* online.

1 Introduction

Protein-protein interactions (PPI) play a key role in cellular mechanisms. Computational approaches, such as protein docking, are important for the structural characterization of PPI. Protein docking determines the structure of a protein-protein complex, given the structure of the interacting proteins (Vakser, 2014). A typical docking pipeline involves three major steps: (i) global scan generating multiple tentative protein-protein matches (docking poses), (ii) evaluation of these poses by physics-based or knowledge-based scoring functions, and (iii) structural refinement of the top-scoring matches. The ability of the modeling protocol to differentiate between correct (near native) and incorrect docking poses determines the overall docking success. Knowledge of even a single residue at the

protein-protein interface is a powerful constraint for the docking search, dramatically reducing the number of docking poses to be evaluated, thus significantly increasing reliability of the resulting docking models.

Such knowledge can be acquired from the PPI-related scientific publications. However, rapidly growing number of biomedical publications in public repositories, such as PubMed, renders manual extraction of relevant information nearly impossible. This necessitates utilization of automated text mining (TM) procedures for generating docking constraints (extracting interface residues from the text of publications). The textual content of the publications is easily understandable by human experts, but processing of that information by computers requires TM algorithms, specific for each particular field of study. So far, TM applications have been mainly focused on prediction of interactions between biological macromolecules (Caufield and Ping, 2019; Li, et al., 2016; Papanikolaou, et al.,

2015; Raja, et al., 2020; Tagore, et al., 2019; Yu, et al., 2018). However, the TM algorithms applicable to protein-protein docking currently are underdeveloped.

The TM techniques extract usable bits of information from the body of text. A scientific text has varying information concentration and coverage depending on a section. Abstracts of scientific publications typically are readily and freely available, have high information density, but have limited content coverage compared to the full-text papers (Lan and Su, 2010; Martin, et al., 2004; Schuemie, et al., 2004). The full texts use longer sentences and parenthesized material (Cohen, et al., 2010) and have heterogeneous distribution of information (as measured by density of keywords in various sections) (Shah, et al., 2003). Access to the full-text papers creates a more comprehensive source (corpus) for the text mining and increases the recall compared to the abstracts (Caporaso, et al., 2008; Westergaard, et al., 2018). However, copyright restrictions generally limit the use of full text articles in the automated TM protocols (Cohen and Hersh, 2005; Rodriguez-Esteban, 2009). The number of PMC-OA (the repository of freely available full-text papers) articles is not increasing at the same rate as the number of PubMed abstracts. The full-text articles have statistical properties (such as a term frequency in the document) that are more robust, but have more noise compared to the abstracts (Lin, 2009). Text mining of the full-text papers has helped in extraction of various biological information (Corney, et al., 2004; Fink, et al., 2008; Friedman, et al., 2001; Gerner, et al., 2010; Gerner, et al., 2012; Mallory, et al., 2015; McIntosh and Curran, 2009), including one on non-structural aspects of PPI (Dogan, et al., 2017; Hakenberg, et al., 2010; Huang, et al., 2004; Krallinger, et al., 2008; Peng, et al., 2016).

Understanding the textual context of publications requires recognition of specific patterns in texts that can be identified by machine learning (ML) techniques, especially, deep learning (DL) approaches implemented using neural networks (NN) with several hidden layers (Ching, et al., 2018; Habibi, et al., 2017). Each successive layer learns higher level of abstraction (Bengio, et al., 2013; LeCun, et al., 2015). NN are trained using the back-propagation algorithm (LeCun, et al., 2015) where the error (the difference between actual and desired output) is projected backwards layer-by-layer, with the connection weights adjusted in proportion (Rumelhart, et al., 1986). The NN applications include, but not limited to automatic speech recognition (Schwenk, 2007), machine translation (Mikolov, 2012), paraphrasing (Turney, 2013), image and scene annotation (Socher, et al., 2011; Weston, et al., 2011), as well as prediction of protein-protein interactions (Yao, et al., 2019). However, DL is computationally demanding, especially at the training stage (which is necessary to repeat, e.g. when the model changes), and thus such approaches may be employed when simpler ML algorithms, e.g. Support Vector Machine (SVM), do not suffice.

For computational procedures, it is desirable to represent words using numbers. Simplistic approaches may assign a unique single number (scalar) to each word of a language (e.g., in lexicographical order). The next step is to represent a word as a series of numbers (word vector or word embedding), so that vector operations can be meaningfully applied (Mikolov, et al., 2013a; Mikolov, et al., 2013b). Then, the inner product of the two vectors would be a measure of similarity of the two words, the sum of the two vectors would reflect the combined meaning of the two words, and the subtraction of the two vectors (offset) would capture the relations (e.g. plural relations, like "molecules vs. molecule" and "residues vs. residue" would have similar offsets). Word vectors, e.g. implemented in the word2vec software, can be efficiently estimated on a large scale (Mikolov, et al., 2013a). They are widely used as a first generic step in a united architecture for solving a specific Natural Language Processing (NLP) task using deep neural networks (Collobert and Weston, 2008;

Irsoy and Cardie, 2014a; Irsoy and Cardie, 2014b; Mikolov, et al., 2013a; Mikolov, et al., 2013b; Mikolov, et al., 2013c; Socher, et al., 2011), e.g. for machine translation requiring large vocabulary across multiple languages (Brants, et al., 2007). The word vectors are used in analysis of sentence-level sentiment, a quantitative score of a subjective information (e.g. "tone of a speaker" or "attitude of a customer") (Socher, et al., 2011; Socher, et al., 2013).

Earlier, we implemented an algorithm that searches for the patterns of letters and digits typically used by authors referring to a specific residue in a protein (referred to as basic TM in this paper). We showed that such information, although mined in a simplistic manner, efficiently excludes incorrect docking models from consideration and thus significantly improves docking success rate (Badal, et al., 2015). However, without interpretation of the context in which the residue appears in the text, the initial pool of the extracted data inevitably contains residues that are not relevant for the binding of specific proteins. Thus, the initially extracted set of residues needs further post-processing. Recently, we investigated filtering of the non-relevant residues by several Natural Language Processing (NLP) techniques, such as keywords semantic similarities, dictionaries look-up and analysis of sentence parse trees with and without SVM model (Badal, et al., 2018). However, the amount of non-relevant residues still remained high. In this paper, for the analysis of context in which the residue is mentioned, we use a deep recursive neural network (DRNN) model based on the concept of the word vectors. We compared the performance of the DRNN and SVM models in various training/testing schemes and showed that the DRNN model is superior, albeit slightly, to the SVM model when training is performed on the PMC-OA full-text articles and applied to classification (interface vs. non-interface) of the residues spotted in the PubMed abstracts. When both training and testing is performed on the full-texts articles, the performance of the models is similar. Thus, the use of DRNN, which is computationally expensive especially at the training stage, in such cases may be unnecessary.

2 Methods

2.1 Basic text mining protocol

Our basic TM tool consists of information retrieval (IR; retrieval of publications relevant to a particular pair of interacting proteins), and information extraction (IE; extraction of residues in the abstracts of identified publications) (Badal, et al., 2015). In this study, in addition to using PubMed resources from <https://www.ncbi.nlm.nih.gov/> (NCBI, 2013), we also downloaded and stored locally the PMC-OA full text articles from <http://www.ncbi.nlm.nih.gov/pmc/tools/ftp/>. Thus, the IR stage was modified to incorporate the local availability of the full-text articles, as opposed to E-fetch from the E-utilities for the PubMed abstracts. PMID (a unique ID for a PubMed abstract) and PMCID (a unique ID for a PMC-OA full-text articles) are different for the same article. For computational efficiency, mapping between them, allowing fetching of a full-text article given a PMID of its abstract, was implemented as a PostgreSQL table. As in our previous study (Badal, et al., 2015), articles, relevant to a protein pair, were retrieved by AND-queries (requiring that both proteins in the complex are mentioned in the text) and OR-queries (either of the proteins is mentioned) using NCBI E-utilities (<http://www.ncbi.nlm.nih.gov/books/NBK25501>). Then using PMID-PMCID mapping, the available full-text articles for that protein pair were identified (for 2,640,816 PMID only 196,912 PMCID were mapped). These full-text articles and abstracts were subjected to the IE stage of the protocol (Fig. 1), which spots different variations of the residue name and its number (Badal, et al., 2015). A simple residue filtering was performed

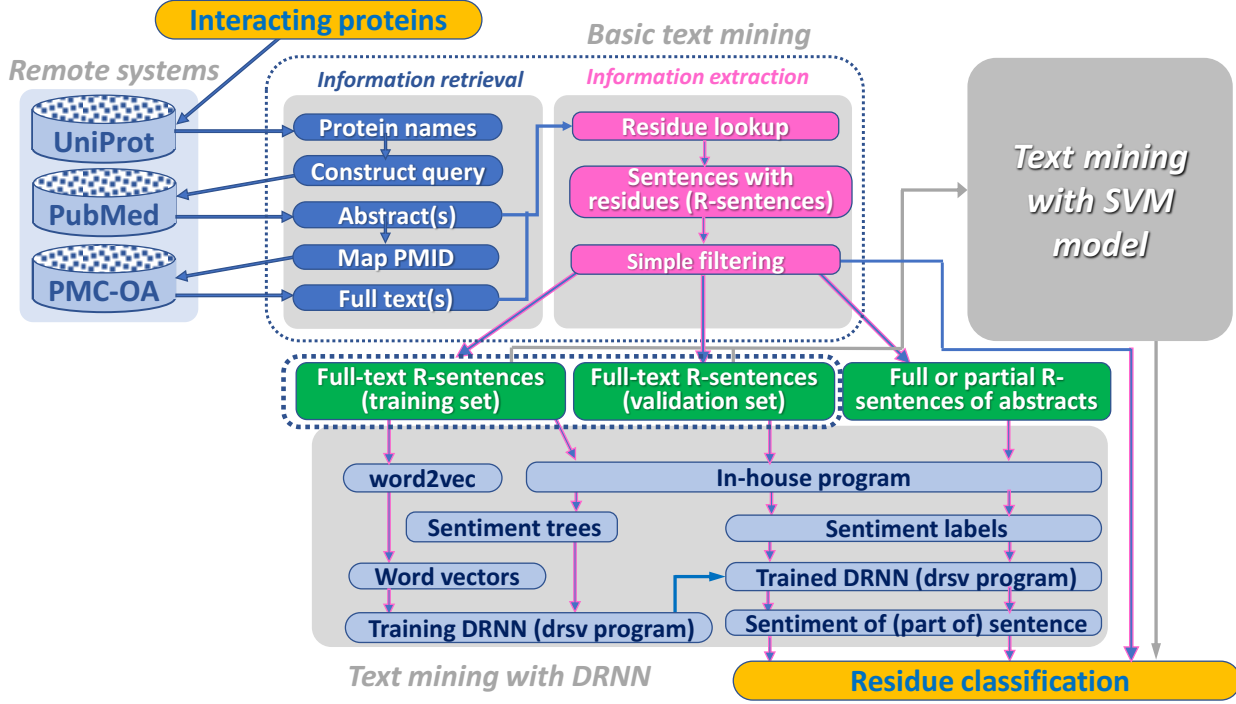


Figure 1. Flowchart of the text-mining system. Algorithm of the text mining with the SVM model is the same as in (Badal, et al., 2018) and thus is not shown in detail. Both full text sets are unified in one training set (shown by the dashed line) when the trained neural network is tested on the PubMed abstracts.

by checking that the residues are on the protein surfaces. All or part of residue-containing sentences (hereafter termed as R-sentences) from the full-text articles were used to estimate the effectiveness of the basic TM on the full-texts and to train the DRNN and the SVM models. The trained models were used to classify residues found in the PubMed abstracts and full-text articles as interface or non-interface. In the same training/testing scenarios, both SVM and DRNN models used the same list of raw-mined residues. Thus, the difference in performance of the machine learning models generated by the two approaches originates only from the difference in the text manipulation steps (see below).

2.2 Evaluating performance of the text mining protocol

Whereas residues essential for protein recognition could be outside the protein-protein interface, the goal of our text mining approach is to generate constraints for docking, which implies residues at the protein-protein interface. Thus, the performance of the TM protocol for a particular PPI, for which N residue-containing articles (abstract-only or full-text) were retrieved, was evaluated as (Badal, et al., 2015)

$$P_{TM} = \frac{\sum_{i=1}^N N_i^{int}}{\sum_{i=1}^N (N_i^{int} + N_i^{non})}, \quad (1)$$

where N_i^{int} and N_i^{non} are the numbers of interface and non-interface residues, correspondingly, mentioned in article i for this PPI, which were not filtered out by an algorithm. If all residues in an article were purged, this article was excluded from the P_{TM} calculations. Distribution of P_{TM} for all

complexes in the dataset provides detailed description of efficiency of the TM algorithm for that dataset. We also compared the performance of two algorithms for residue filtering (Badal, et al., 2018) by

$$\Delta N(P_{TM}) = N_{tar}^{X_1}(P_{TM}) - N_{tar}^{X_2}(P_{TM}), \quad (2)$$

where $N_{tar}^{X_1}(P_{TM})$ and $N_{tar}^{X_2}(P_{TM})$ are the number of targets with P_{TM} value yielded by algorithms X_1 and X_2 , respectively. Since the major contributions to P_{TM} distribution are from all-false-positive ($P_{TM} = 0$) and all-true-positive ($P_{TM} = 1$) cases, algorithm efficiency can be roughly assessed just by the two extreme values of $\Delta N(P_{TM})$ at $P_{TM} = 0$ and $P_{TM} = 1$. The values $\Delta N(0) < 0$ and $\Delta N(1) > 0$ indicate better performance of model X_1 with respect to model X_2 in purging of the PPI-irrelevant residues from the mined articles.

2.3 Datasets

The approaches were benchmarked on the set of 579 non-redundant (at 30% sequence identity) binary protein-protein complexes from the DOCKGROUND resource (<https://dockground.compbio.ku.edu>) (Kundrotas, et al., 2018). The dataset for training ML models (DRNN and SVM) consisted of 4,982 residue-containing sentences (hereafter referred to as R-sentences), which passed the initial screening of residue-containing sentences automatically extracted by the OR-queries from the full-text PMC-OA articles (full training set). By querying the native PDB structures, these sentences were classified into 1,605 positive (interface residue) and 3,377 negative (non-interface residue) R-sentences. The interface residues were defined by 6 Å distance between any atoms across the chains interface. The dataset for testing the ML models comprised 5,786 R-sentences

(or their parts around identified residues, see Results), extracted from the PubMed abstracts by the OR-queries (abstract testing set). Since only a small fraction of the training text came from PMC-OA abstracts of the PMC-OA articles, we did not exclude PubMed abstracts that have PMC-OA full-text articles available.

For testing the ML models on the full-text articles, the training set consisted of 803 positive and 1689 negative sentences extracted from the articles describing studies of the top 290 complexes from the full list of complexes (sorted in alphabetical order of corresponding PDB codes). The remaining sentences from the articles describing studies of the rest of the complexes comprised the test set.

2.4 Generation of keywords

To identify keywords relevant to the protein binding, we computed the differences (bias) in word frequencies (percent of the sentences with that word) calculated separately for the positive and negative R-sentences in the full training set. Words with the biases between 1 and -1, stop words and names of a protein, an amino acid or species were omitted from consideration. This resulted in 47 PPI+ive (PPI-relevant, positive bias) and 37 PPI-ive (PPI-irrelevant, negative bias) keywords (Table S1).

2.5 Support Vector Machine model

The features for the SVM model consisted of the scores, calculated from the parse trees of the R- and the context (immediately preceding and following the R-sentence) sentences. This score reflects an effective count of the edges on the parse tree between a residue (or root for the context sentences) and all the keywords from the Table S1 (detailed description is published elsewhere (Badal, et al., 2018)). Sentence parse trees were built by the Perl module of the Stanford parser (De Marneffe and Manning, 2008a; De Marneffe and Manning, 2008b) <http://nlp.stanford.edu/software/index.shtml> downloaded from <http://search.cpan.org> site. The SVM model was trained and validated using program SVMlight with linear, polynomial and RBF kernels (Joachims, 1998; Joachims, 1999; Morik, et al., 1999). Analysis of the results indicated (data not shown) that the best SVM performance was achieved using RBF kernel with gamma 0.25 for both training-testing schemes (training on full-texts, testing on abstracts and training and testing on full-texts only).

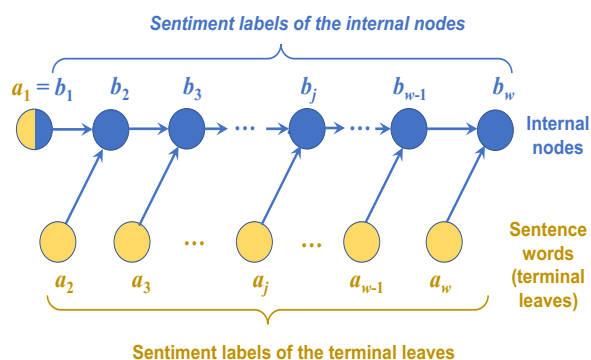


Figure 2. Schematic representation of a sentence binary tree and associated sentiment labels.

2.6 Docking protocol

Docking by our GRAMM procedure (Vakser, 1996) was evaluated on the unbound protein structures from the DOCKGROUND X-ray benchmark set 4 (Kundrotas, et al., 2018). The set originally consisted of 389 protein complexes, out of which 57 were also present in the set used in development of the TM protocols and thus were excluded from docking. The residues identified by the TM protocols were used to generate a confidence score (for details, see Badal, et al., 2018) to increase the weight of protein-protein matches that had these residues at the predicted interface (Badal, et al., 2018). The quality of a match was assessed by C^α ligand (smaller protein in the complex) interface root-mean-square deviation (i-RMSD) between the interface of the docked unbound ligand and corresponding atoms of the unbound ligand superimposed on the bound one in the co-crystallized complex. Docking was evaluated by the percentage of the successfully predicted complexes. The low-resolution protein-protein match was considered correct if it was inside the binding funnel (i-RMSD ≤ 8 Å; Hunjan, et al., 2008), making it a practical starting point for the refinement trajectories (Hunjan, et al., 2008). Protein complex was considered predicted successfully at low resolution if such match was among top 100. Docking was also assessed by the enrichment of the prediction pool by the correct matches across all complexes (overall increase in the number of the top-100 predictions). Text mining provides constraints for docking regardless of the complexity of their modeling (e.g. for protein interactions involving large conformational changes upon binding), which should be instrumental for the refinement. In practical docking, the number of models for the refinement is limited by the computational efficiency of the refinement protocol. Currently available protocols, including those with the GPU implementation, make 100 refinement trajectories practical (Dauzhenka, et al., 2018).

3 Results and Discussion

3.1 Architecture of neural network

For the DRNN training, we generated PPI-specific sentiment tree bank, which is a set of binary trees (Fig. 2) of the R-sentences from the training set with each leaf and internal node tagged by a sentiment labels a_j and b_j , respectively (j counts words in the sentence). According to the Stanford Sentiment Treebank (Socher, et al., 2013), we utilized five standard sentiment classes: very +ive (labeled 4), +ive (3), neutral (2), -ive (1), very -ive (0). In our case, a sentiment (or more precisely, its label), quantifies the degree to which the information in a sentence is relevant to protein-protein docking, i.e. how likely residues mentioned in the sentence are at the protein-protein interface (label 4 is most probable and label 0 is least probable). In a sentence of N words, the b_j is calculated as

$$b_j = \max(b_{j-1}, a_j) \quad \text{or} \quad b_j = \min(b_{j-1}, a_j), \quad j = 2, \dots, N \quad (3)$$

for the positive and negative sentences, respectively. The sentiment label a_j is

$$a_j = \begin{cases} F_j & , \text{ if word } j \text{ is a keyword} \\ \text{round}\left(2 \pm \frac{j}{w}\right) & , \text{ for the rest of the words} \end{cases} \quad (4)$$

where F_j is the fixed sentiment label for PPI+ive and PPI-ive keywords (Table S1).

Such labeling scheme ensures that the final sentiment label of a sentence mentioning interface residues is 3 or 4 (0 or 1 for sentences with non-interface residues) and captures the baseline trend of a sentiment,

steadily increasing for the positive and decreasing for the negative sentences. The sets of a_j and b_j were the first part of the input, necessary for the DRNN training.

The second part of the training input was a set of initial word vectors (numeric weights associated with the word), $\{\vec{v}_i\}$, for each of the 74,438 unique words in the sentences of the training set (out of ~20M total words). The vectors were generated by the word2vec program with skip-gram model (a predictive language model that works well for even rarely used words) and the default training window size of 10 (the number of considered words in the context) (Mikolov, et al., 2013a). The

dimensionality of the word vectors was set to 300, considered sufficient for complex NLP tasks (Jurafsky and Martin, 2017; Mikolov, et al., 2013b; Pennington, et al., 2014). The word vectors corresponding to similar words were distributed close to each other (Fig. 3). The amino acids were in one region of the vector space, as were the words associated with shapes. Similarly, co-localized were words such as "interaction" and "complex." Antonyms, such as hydrophobic, hydrophilic, are also in the proximity of each other, indicating that these terms are linguistically interchangeable.

Table 1. Overall performance of basic TM on abstracts (PubMed and PMC-OA) and full texts (PMC-OA)

Dataset	Query Type ^a	L_{tot} ^b	L_{int} ^c	Coverage (%) ^d	Success (%) ^e	Accuracy (%) ^f
PubMed abstracts	AND	128	108	22.1	18.7	84.4
PubMed abstracts	OR	328	273	56.6	47.2	83.2
PMC-OA abstracts	AND	37	21	6.3	3.6	56.7
PMC-OA abstracts	OR	164	89	28.3	15.3	54.2
PMC-OA full-text	AND	103	70	17.7	12.0	67.9
PMC-OA full-text	OR	313	238	54.0	41.1	76.0

^a AND and OR query requires that the name of both proteins (AND) or either protein (OR) is mentioned in the returned document

^b Number of complexes, for which TM retrieved at least one article with residues

^c Number of complexes with at least one interface residue found in the retrieved articles

^d Ratio of L_{tot} and total number of complexes (579)

^e Ratio of L_{int} and total number of complexes (579)

^f Ratio of L_{int} and L_{tot}

Table 2. Overall TM performance on test set of PMC-OA full-text articles retrieved by OR-queries with simplified residue filtering (basic TM) and with residue filtering by SVM and DRNN models trained on reduced full-text training set.

Model	L_{tot} ^a	L_{int} ^b	Coverage (%) ^c	Success (%) ^d	Accuracy (%) ^e	$\Delta N(0)$ ^f	$\Delta N(1)$ ^f
Basic TM	157	115	60.6	44.4	73.2	—	—
SVM	87	58	33.6	22.4	66.7	−22	+15
DRNN	75	46	28.9	17.8	61.3	−24	+11

^a Number of complexes for which TM found at least one article with residues

^b Number of complexes with at least one interface residue found in articles

^c Ratio of L_{tot} and total number of complexes (259)

^d Ratio of L_{int} and total number of complexes (259)

^e Ratio of L_{int} and L_{tot}

^f From Eq. 6 with values from basic TM (first row) as X_2

Table 3. Overall TM performance on PubMed abstracts retrieved by OR-queries with simplified residue filtering (basic TM) and with residue filtering by the SVM and DRNN models. Trained DRNN model was used for classifying residues in the entire sentence, as well as using 7-words window around mined residues. Both SVM and DRNN models were trained on the complete full-text training set.

Model	L_{tot} ^a	L_{int} ^b	Coverage (%) ^c	Success (%) ^d	Accuracy (%) ^e	$\Delta N(0)$ ^f	$\Delta N(1)$ ^f
Basic TM	328	273	56.6	47.2	83.2	–	–
SVM	182	135	31.4	23.3	74.1	-15	-3
DRNN (Whole sentence)	179	120	30.9	20.7	67.0	-3	+6
DRNN (7-words window)	150	104	25.9	18.0	69.3	-16	+13

^a Number of complexes for which TM protocol found at least one abstract with residues

^b Number of complexes with at least one interface residue found in abstracts

^c Ratio of L_{tot} and total number of complexes (579)

^d Ratio of L_{int} and total number of complexes (579)

^e Ratio of L_{int} and L_{tot}

^f From Eq. 6 with values from basic TM (first row) as X_2

Both input components were submitted to the program *drsv* (<https://github.com/oir/deep-recursive>) (Irsoy and Cardie, 2014a) to train 3-layers DRNN model. The DRNN learned over ~10 epochs (epoch is defined as a sweep through the entire training set). Beyond 10 epochs, DRNN was getting over-trained (Fig. S1). The same program was used to evaluate the sentiment for the entire or partial sentences using trained DRNN model. In this case, the input consisted of the sentiment labels a_i (Eq. 2) assigned to the words of a sentence or its parts. Such DRNN architecture with corresponding sentiment treebanks (domain knowledge specific or generic) is widely used (e.g. in the analysis of Netflix movie reviews (Irsoy and Cardie, 2014a; Socher, et al., 2013)).

3.2 Mining of full-text articles

The full text of a paper provides much more information than its abstract. But due to copyright restrictions only just over one million articles are freely available in the PMC-OA database, compared to ~26 million entries in the PubMed database of freely available abstracts. This causes significantly better TM performance on the PubMed abstracts than on the abstracts of the PMC-OA articles (Table 1).

The limited access to the full texts is counterweighted by the abundant information in them, as the overall TM performance on the PMC-OA full-text articles is comparable to that on the PubMed abstracts (Table 1). Significantly better TM performance on PMC-OA full-texts than on the PMC-OA abstracts (Table 1) points to more frequent mentioning of residues in the full texts (for 149 complexes, all mined residues were in the full texts only). However, due to lesser space constraints in the full texts, residues there are mentioned in a variety of contexts. This leads to a significantly larger number of PPI-irrelevant residues in the full texts than in the abstracts (corresponding bars at $P_{\text{TM}} = 0$ and $P_{\text{TM}} = 1$ in Fig. 4).

Research on a specific protein interaction could be published only in journals with limited access to their full texts. Our results indicated that for a significant part of the complexes in our set (75 out of 579, or ~13%) this is indeed the case (one such example is shown in Figure S2 with the detailed description in Text S1). Thus, we argue that, at least presently,

PMC-OA full-text articles are more suitable for thorough analysis of residue-mentioning context (with consequent application to the residue purging in the PubMed abstracts) rather than for the extraction of the raw information.

Both SVM and DRNN models similarly affected the TM of the full-text articles (Table 2 and Fig. 5) and their abstracts (Table S2 and Fig. S3). SVM model purged all initially mined residues (i.e. all mined residues were considered non-interface ones) for a smaller number of complexes (second column in Table 2). At the same time, it was slightly better in removing non-interface residues from the full-text articles (last column in Table 2). Compared to the basic TM, both SVM and DRNN models for the full-text articles significantly increased the fraction of complexes, for which all mined residues are located at the complex interfaces. In the abstracts of the PMC-OA articles, DRNN and SVM removed all retrieved abstracts with the residues for ~70% of the complexes. Thus, the difference between the two models was not statistically significant. The performance of the SVM model only slightly depends on the keywords used for the SVM training and testing (Table S3 and Figure S4 show the results for the SVM model with the manually selected keywords from our previous study (Badal, et al., 2018)). In our simplified scheme, we assigned a definite sentiment only to frequently appearing words designated as PPI keywords. Thus, we could miss infrequently occurring words or word groups that carry a strong sentiment. This is illustrated in Figure 6 where residue Asp53 at the interface of heat shock HSP82 and AHA1 proteins was incorrectly filtered by both SVM and DRNN models. However, human reader can easily identify this residue as the interface one (additional examples in Text S3). In the future, performance of a DRNN model can be potentially improved by the use of PPI-specific hand-curated sentiment tree bank.

3.3 Text mining of abstracts

When both SVM and DRNN models are trained on full-text articles, and applied to classification of residues in the abstracts, the results are similar, with somewhat better performance of the DRNN model (reflected in the last column of Table 3 and the rightmost columns of Figure 7). For this model, however, the larger fraction of complexes with only interface

residues in the final list is counterweighted by the largest fraction of complexes, for which only non-interface residues were mined (left- and right-most columns in Figure 7). The training of the DRNN model was done on the entire set of full-text articles, but its performance only weakly depends on the size of the training set (Table S4 and Fig. S5). The SVM model was better in removing complexes with only non-interface residues mined, but failed in increasing the number of complexes, for which all mined residues are PPI-relevant (Table 3).

We argue that performance of the SVM model suffered from the different structure of sentences in the full-texts and the abstracts. The DRNN model learns data/text patterns at a higher level of generality, and thus easily adapts to different domains, as diverse as, for example, protein docking and Netflix movie reviews. This suggests the use of DL algorithms for analysis of TM results when, for example, a particular PPI is widely studied by a variety of authors using different lexical semantic styles. On the other hand, SVM models may be better in finer analysis of articles of the same group of authors with similar writing styles.

Besides better adaptation of DRNN to the different sentence structures, it also has an advantage of an easy implementation of independent classification of multiple residues in a sentence by limiting the context to a few words around the residue (contextual window) and estimating a sentiment for that part of the sentence only. Obviously, a smaller contextual window allows independent classification of a larger number of residues in the sentence. However, due to the loss of broader contextual information embedded in the trained DRNN model, the sentiment accuracy may decrease. Our results indicate that the optimal TM performance is achieved when the sentiment is calculated for sentence fragments of seven words around the residue (Fig. 8). Overall, the DRNN model with the contextual window significantly improves filtering of the non-PPI residues, while only slightly reducing the coverage of the dataset (Table 3 and Fig. 7). Figure S6 illustrates the advantage of sentiment calculation using context window for cationic trypsin - trypsin inhibitor complex.

3.4 Application to modeling of protein complexes

We tested the applicability of the above protocols to protein docking. The TM protocols used were basic TM, SVM and DRNN (trained on whole sentences) - all applied to full-text papers. The results showed the number of successfully predicted complexes increase over the basic TM by 5% and 10% using SVM and DRNN correspondingly. The total number of the top-100 correct matches across all complexes increased by 12% and 19%, respectively (Fig. 9). The current added value of the text mining in docking is limited in part by the lack of uniform standard of residue numbering. In different studies, the numbering often depends on the sequence fragment used in the experiments. The numbering of residues in the PDB structures is even more complicated as it may depend not only on what fragment was crystallized, but also on other factors (e.g., for comparison with homologous proteins, numbers like 20A and 20B may be inserted between 20 and 21). Fortunately, in the course of time, the numbering of residues is starting to follow that of the sequences (or their isoforms) in the UNIPROT database. However, currently, improving the mapping of the mined residues, especially those from the older literature, to structures needed for docking would require complicated analysis of the textual context, which is outside the scope of this study. Also, the number of the full-text articles in the open access, from which the residues could be mined, is still relatively small. With the rapid growth of popularity of the open access publishing (Piwowar, et al., 2018), the increase in the success rate should grow as well. Thus, the utility of the machine learning approaches in application to protein docking, the basic principles of which we explored in this study, will expand accordingly, leading to more accurate modeling of protein interactions.

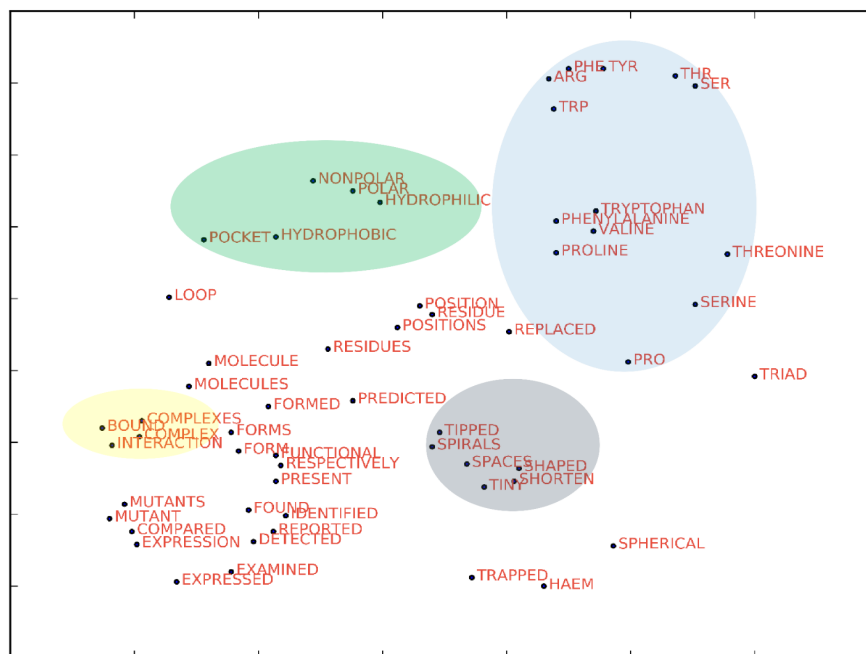


Figure 3. Example of the initial word vectors distribution. Highlighted areas show similar words in the same region of the vector space. The distribution was generated by arbitrarily choosing a set of 55 words that are typically found in PPI publications, not meeting any scientific criteria, but representing a rich mix of domain vocabulary. The words were put in a list, along with the file containing word-vector lookup table (output of word2vec). The t-SNE software (van der Maaten and Hinton, 2008) extracted relevant 55 word-vectors and performed dimensionality reduction, until the required criteria of given perplexity (40) is met and the final dimensionality is reduced to 2. The points were plotted and labelled using pyplot in a python script. The highlighted areas were overlaid on the graph.

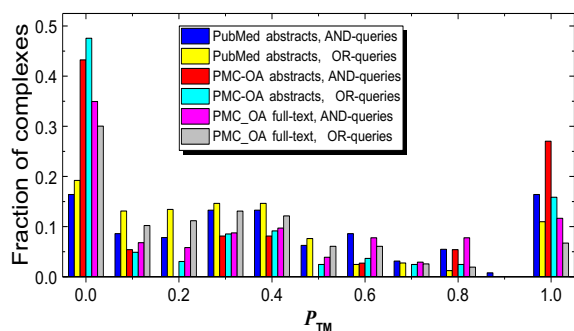


Figure 4. Basic text mining on abstracts and full-texts. The performance P_{TM} for individual complexes was calculated by Eq. 1. The distribution is normalized to the total number of complexes for which residues were extracted (Table 1).

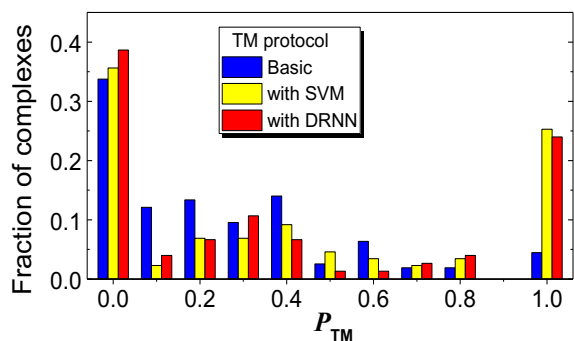


Figure 5. Comparison of text-mining protocols on full texts. Both SVM and DRNN were trained on reduced full-text training set. The performance P_{TM} was calculated by Eq. 1. The distribution is normalized to the total number of complexes for which residues were extracted (Table 2).

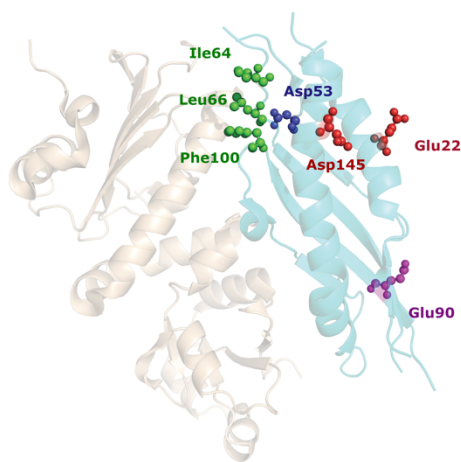


Figure 6. Example of residues mined from full texts. The structure is 1usu, chain A (gray) and B (cyan). Out of six residues identified by the basic TM (four at the interface, in green and blue, and two not at the interface, in red) NLP SVM and DRNN models correctly classified three interface residues (green) and failed to identify Asp53 (which could be easily identified by a human reader, see Text S2). Only one non-interface residue (magenta) was mined from the PubMed abstracts. The abstracts from PMC-OA did not predict any residues in basic TM. Details are in Text S2.

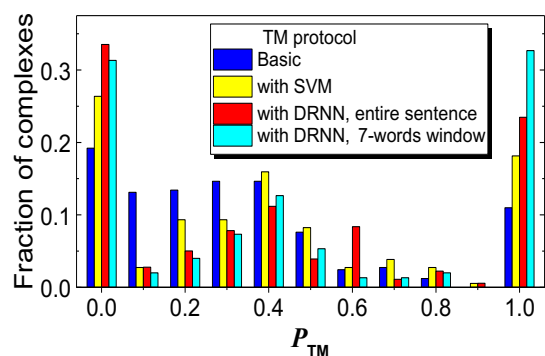


Figure 7. Comparison of text-mining protocols on abstracts. The PubMed abstracts were retrieved by the OR-queries with simplified residue filtering (basic TM) and with the residue filtering by SVM model and DRNN. Both SVM and DRNN were trained on entire full-text training set. DRNN was applied for classifying residues in the entire sentence and with 7-words window around the mined residues. The performance P_{TM} for individual complexes was calculated by Eq. 1. The distribution is normalized to the total number of complexes for which residues were extracted (Table 3).

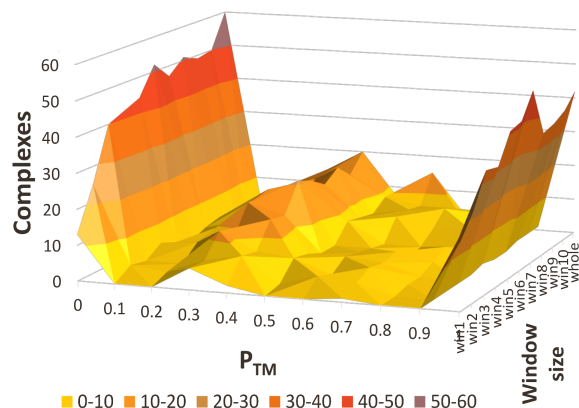


Figure 8. TM performance with residue filtering by DRNN using different window sizes around mined residues. DRNN was trained on the entire training set of PMC-OA full-text articles. The performance P_{TM} was calculated by Eq. 1.

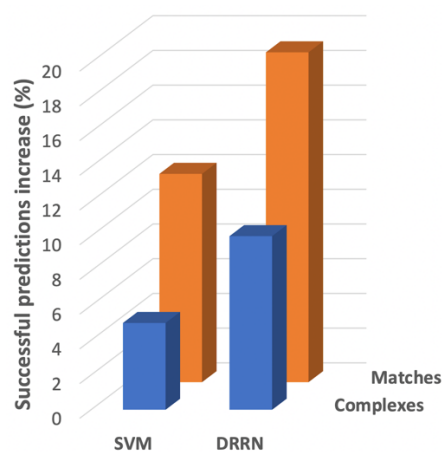


Figure 9. The increase in successful docking predictions over the basic TM protocol. The change of the number of correctly predicted complexes is in blue and the change of the number of correct matches in top 100 is in orange.

4 Conclusion

We continued development of the methodology for generating constraints from publicly available literature for application to structural modeling of protein complexes. Capitalizing on our earlier results on generating such constraints from the basic text mining of PubMed abstracts (Badal, et al., 2015), improved by natural language processing techniques (Badal, et al., 2018), in this study, we focused on comparing performances of machine learning models generated by two methods (Deep Recursive Neural Network and Support Vector Machine) in filtering non-interface residues from the same list of initially mined residues either in PubMed abstracts or PMC-OA subset of freely available full text articles.

The PMC-OA full text articles, despite representing a small subset of scientific publications, provide a useful source for training of DRNN model. The networks can be applied to classification of residues found in the abstracts, where the sentence structures are, in general, different from those in the full-text articles. In such case, the DRNN model is superior to SVM model, because the success of the latter is often determined by the similarity in data/text patterns in the training and the testing sets. Our study provides an insight into the optimal context size for TM applications, based on the significant improvement of the DRNN model's performance when the sentiment was calculated for a part of the sentence around the mined residue rather than for the entire sentence. The results indicate that the bank of sentiment trees, specific for protein-protein interactions and curated by the experts in the field, is essential for further performance improvement of the ML-enhanced text-mining. Overall, following our previous results on NLP application to abstracts (Badal, et al., 2018), we showed that DRNN model similarly outperform the basic TM on the abstracts, and both SVM and DRNN models outperform the basic TM when applied to the full-text papers. Greater availability of the full-text papers should increase usefulness of this source of information for structural modeling of protein complexes. A simpler SVM approach is often sufficient for filtering residues mined from the texts with patterns similar to those used for training (abstracts - abstracts or full texts - full texts). Thus, the use of DL approaches (which are computationally demanding, especially, at the training stage) in such cases may be unnecessary.

By its nature, the approach based on full-text articles depends on the pool of published open access articles about the target protein complex. Thus, naturally its application to the recent challenging targets is still limited. However, with the growth of popularity of the open access publishing, the utility of the approach will grow. Our study focused on applicability of the text mining to modeling of protein complexes, in anticipation of the inevitable future growth of the open access publishing.

Funding

This study was supported by NIH grant R01GM074255 and NSF grants DBI1565107 and DBI1917263.

Conflict of Interest: none declared.

References

- Badal, V.D., Kundrotas, P.J. and Vakser, I.A. Text mining for protein docking. *PLoS Comp. Biol.* 2015;11:e1004630.
- Badal, V.D., Kundrotas, P.J. and Vakser, I.A. Natural language processing in text mining for structural modeling of protein complexes. *BMC Bioinformatics* 2018;19:84.
- Bengio, Y., Courville, A. and Vincent, P. Representation learning: A review and new perspectives. *IEEE transactions on pattern analysis and machine intelligence* 2013;35:1798-1828.
- Brants, T., et al. Large language models in machine translation. In, *Proc. Joint Conf. Empirical Methods in Natural Language Processing and Computational Natural Language Learning*. Citeseer; 2007.
- Caporaso, J.G., et al. Intrinsic evaluation of text mining tools may not predict performance on realistic tasks. In, *Pacific Sympos. Biocomputing*. NIH Public Access; 2008. p. 640.
- Caufield, J.H. and Ping, P. New advances in extracting and learning from protein-protein interactions within unstructured biomedical text data. *Emerging Topics in Life Sciences* 2019;3(4):357-369.
- Ching, T., et al. Opportunities and obstacles for deep learning in biology and medicine. *Journal of the Royal Society Interface* 2018;15(141).
- Cohen, A.M. and Hersh, W.R. A survey of current work in biomedical text mining. *Brief. Bioinf.* 2005;6:57-71.
- Cohen, K.B., et al. The structural and content aspects of abstracts versus bodies of full text journal articles are different. *BMC Bioinformatics* 2010;11:492.
- Collobert, R. and Weston, J. A unified architecture for natural language processing: Deep neural networks with multitask learning. In, *Proc. 25th Int. Conf. Machine Learning*. ACM; 2008. p. 160-167.
- Corney, D.P., et al. BioRAT: Extracting biological information from full-length papers. *Bioinformatics* 2004;20:3206-3213.
- Dauzhenka, T., Kundrotas, P.J. and Vakser, I.A. Computational feasibility of an exhaustive search of side-chain conformations in protein-protein docking. *J. Comput. Chem.* 2018;39:2012-2021.
- De Marneffe, M.-C. and Manning, C.D. Stanford typed dependencies manual. In: Technical report, Stanford University; 2008a. p. 338-345.
- De Marneffe, M.C. and Manning, C.D. The Stanford typed dependencies representation. In, *Proc. Workshop Cross-Framework Cross-Domain Parser Evaluation*. Manchester, UK: Association for Computational Linguistics; 2008b. p. 1-8.
- Dogan, R.I., et al. The BioC-BioGRID corpus: Full text articles annotated for curation of protein-protein and genetic interactions. *Database* 2017;2017.
- Fink, J.L., et al. BioLit: Integrating biological literature with databases. *Nucl. Acids Res.* 2008;36(suppl 2):W385-W389.
- Friedman, C., et al. GENIES: A natural-language processing system for the extraction of molecular pathways from journal articles. *Bioinformatics* 2001;17(suppl 1):S74-S82.
- Gerner, M., Nenadic, G. and Bergman, C.M. An exploration of mining gene expression mentions and their anatomical locations from biomedical text. In, *Proc. 2010 Workshop Biomed. Natural Language Processing*. Association for Computational Linguistics; 2010. p. 72-80.
- Gerner, M., et al. BioContext: An integrated text mining system for large-scale extraction and contextualization of biomolecular events. *Bioinformatics* 2012;28:2154-2161.
- Habibi, M., et al. Deep learning with word embeddings improves biomedical named entity recognition. *Bioinformatics* 2017;33(14):I37-I48.
- Hakenberg, J., et al. Efficient extraction of protein-protein interactions from full-text articles. *IEEE-ACM T Comput Biol Bioinf* 2010;7:481-494.
- Huang, M., et al. Discovering patterns to extract protein-protein interactions from full texts. *Bioinformatics* 2004;20:3604-3612.
- Hunjan, J., et al. The size of the intermolecular energy funnel in protein-protein interactions. *Proteins* 2008;72:344-352.

- Irsoy, O. and Cardie, C. Deep recursive neural networks for compositionality in language. In, *Adv. Neural Information Processing Systems*. 2014a. p. 2096-2104.
- Irsoy, O. and Cardie, C. Modeling compositionality with multiplicative recurrent neural networks. *arXiv:1412.6577* 2014b.
- Joachims, T. Text categorization with Support Vector Machines: Learning with many relevant features. In: Nédellec, C. and Rouveilol, C., editors, *Machine Learning: ECML-98*. Springer Berlin Heidelberg; 1998. p. 137-142.
- Joachims, T. Making large-scale support vector machine learning practical. In, *Advances in Kernel Methods*. MIT Press; 1999. p. 169-184.
- Jurafsky, D. and Martin, J.H. Semantics with Dense Vectors In, *Speech and Language Processing: An Introduction to Natural Language Processing, Computational Linguistics, and Speech Recognition*. 2017.
- Krallinger, M., et al. Overview of the protein-protein interaction annotation extraction task of BioCreative II. *Genome Biol.* 2008;9(Suppl 2):S4.
- Kundrotas, P.J., et al. Dockground: A comprehensive data resource for modeling of protein complexes. *Protein Sci.* 2018;27:172-181.
- Lan, M. and Su, J. Empirical investigations into full-text protein interaction Article Categorization Task (ACT) in the BioCreative II. 5 Challenge. *IEEE/ACM Transactions on Computational Biology and Bioinformatics (TCBB)* 2010;7:421-427.
- LeCun, Y., Bengio, Y. and Hinton, G. Deep learning. *Nature* 2015;521:436-444.
- Li, A., et al. A text feature-based approach for literature mining of lncRNA-protein interactions. *Neurocomputing* 2016;206:73-80.
- Lin, J. Is searching full text more effective than searching abstracts? *BMC Bioinformatics* 2009;10:46.
- Mallory, E.K., et al. Large-scale extraction of gene interactions from full text literature using deepdive. *Bioinformatics* 2015:btv476.
- Martin, E.P., et al. Analysis of protein/protein interactions through biomedical literature: Text mining of abstracts vs. text mining of full text articles. In, *Knowledge Exploration in Life Science Informatics*. Springer; 2004. p. 96-108.
- McIntosh, T. and Curran, J.R. Challenges for automatically extracting molecular interactions from full-text articles. *BMC Bioinformatics* 2009;10:311.
- Mikolov, T. Statistical language models based on neural networks. *PhD Thesis, Brno University of Technology* 2012.
- Mikolov, T., et al. Efficient estimation of word representations in vector space. *arXiv:1301.3781* 2013a.
- Mikolov, T., et al. Distributed representations of words and phrases and their compositionality. In, *Adv. Neural Information Processing Systems*. 2013b. p. 3111-3119.
- Mikolov, T., Yih, W.-t. and Zweig, G. Linguistic Regularities in Continuous Space Word Representations. In, *Hlt-naacl*. 2013c. p. 746-751.
- Morik, K., Brockhausen, P. and Joachims, T. Combining statistical learning with a knowledge-based approach: A case study in intensive care monitoring (No. 1999, 24). In.: Technical Report, SFB 475: Komplexitätsreduktion in Multivariaten Datenstrukturen, Universität Dortmund; 1999.
- NCBI, R.C. Database resources of the National Center for Biotechnology Information. *Nucl. Acids Res.* 2013;41:D8.
- Papanikolaou, N.L., et al. Protein-protein interaction predictions using text mining methods. *Methods* 2015;74:47-53.
- Peng, Y., et al. BioC-compatible full-text passage detection for protein-protein interactions using extended dependency graph. *Database* 2016;2016:baw072.
- Pennington, J., Socher, R. and Manning, C. Glove: Global vectors for word representation. In, *Proc. 2014 Conf. Empirical Methods in Natural Language Processing (EMNLP)*. 2014. p. 1532-1543.
- Piwowar, H., et al. The state of OA: A large-scale analysis of the prevalence and impact of Open Access articles. *Peer J.* 2018;6:e4375.
- Raja, K., et al. Automated Extraction and Visualization of Protein-Protein Interaction Networks and Beyond: A Text-Mining Protocol. *Methods in molecular biology (Clifton, N.J.)* 2020;2074:13-34.
- Rodriguez-Esteban, R. Biomedical text mining and its applications. *PLoS Comput. Biol.* 2009;5:e1000597.
- Rumelhart, D., Hinton, G. and Williams, R. Learning representations by back-propagating errors. *Nature* 1986;323:533-538.
- Schuemie, M.J., et al. Distribution of information in biomedical abstracts and full-text publications. *Bioinformatics* 2004;20:2597-2604.
- Schwenk, H. Continuous space language models. *Comp. Speech & Language* 2007;21:492-518.
- Shah, P.K., et al. Information extraction from full text scientific articles: Where are the keywords? *BMC Bioinformatics* 2003;4:20.
- Socher, R., et al. Parsing natural scenes and natural language with recursive neural networks. In, *Proc. 28th Int. Conf. Machine Learning (ICML-11)*. 2011. p. 129-136.
- Socher, R., et al. Semi-supervised recursive autoencoders for predicting sentiment distributions. In, *Proc. Conf. Empirical Methods in Natural Language Processing*. Association for Computational Linguistics; 2011. p. 151-161.
- Socher, R., et al. Recursive deep models for semantic compositionality over a sentiment treebank. In, *Proc. Conf. Empirical Methods in Natural Language Processing (EMNLP)*. Citeseer; 2013. p. 1642.
- Tagore, S., et al. ProtFus: A Comprehensive Method Characterizing Protein-Protein Interactions of Fusion Proteins. *PLoS Computational Biology* 2019;15(8).
- Turney, P.D. Distributional semantics beyond words: Supervised learning of analogy and paraphrase. *Trans. Assoc. Comput. Linguistics (TACL)* 2013;1:353-366.
- Vakser, I.A. Low-resolution docking: Prediction of complexes for under-determined structures. *Biopolymers* 1996;39:455-464.
- Vakser, I.A. Protein-protein docking: From interaction to interactome. *Biophys. J.* 2014;107:1785-1793.
- van der Maaten, L. and Hinton, G. Visualizing data using t-SNE. *J. Machine Learning Res.* 2008;9:2579-2605.
- Westergaard, D., et al. A comprehensive and quantitative comparison of text-mining in 15 million full-text articles versus their corresponding abstracts. *PLoS Computational Biology* 2018;14(2).
- Weston, J., Bengio, S. and Usunier, N. Wsabie: Scaling up to large vocabulary image annotation. In, *IJCAI*. 2011. p. 2764-2770.
- Yao, Y., et al. An integration of deep learning with feature embedding for protein-protein interaction prediction. *PeerJ* 2019;7.
- Yu, K., et al. Automatic extraction of protein-protein interactions using grammatical relationship graph. *BMC Medical Informatics and Decision Making* 2018;18.

SUPPLEMENTARY MATERIAL

Table S1. List of automatically generated keywords and associated sentiment labels. Positive and negative scores (biases) indicate keywords relevant (PPI+ive) and irrelevant (PPI-ive), respectively, to protein-protein binding sites. Details are in Methods (main text).

Score	Keyword	Sentiment	Score	Keyword	Sentiment
-6.2086	Mutations	0	1.0936	Spots	3
-5.9771	Mutants	0	1.1334	Compounds	3
-4.6471	Domain	0	1.1391	Active	3
-3.9662	Mutant	0	1.1903	intermolecular	3
-2.9279	Mutation	0	1.2049	Interacting	3
-2.6234	Anti	0	1.2172	Bond	3
-2.5567	Cells	0	1.2236	Extensive	3
-2.1993	Sites	1	1.2634	Defined	3
-2.1408	Observed	1	1.3243	Ubiquitin	3
-2.0669	Position	1	1.3600	Docking	3
-1.9548	Gene	1	1.3602	Structure	3
-1.8865	Additional	1	1.3658	Form	3
-1.6241	Phosphorylation	1	1.3942	Expected	3
-1.4819	Effects	1	1.4608	Stable	3
-1.4340	Substitutions	1	1.5494	Pocket	3
-1.4193	Using	1	1.6241	Complexes	3
-1.4072	Effect	1	1.7257	Salt	3
-1.3690	Previously	1	1.7428	Loops	3
-1.3186	Reduced	1	1.7436	Terminal	3
-1.2951	Values	1	1.7469	Surface	3
-1.2812	Results	1	1.8078	Buried	3
-1.2804	Motif	1	1.8793	Aromatic	3
-1.2780	Substitution	1	2.0572	Sequence	3
-1.2552	Kinase	1	2.0645	Cluster	3
-1.1301	Levels	1	2.2051	Main	3
-1.1065	Reported	1	2.3774	Interface	3
-1.1025	Linker	1	2.4139	Interact	3
-1.0708	Catalytic	1	2.4456	Affinity	3
-1.0642	Amino	1	2.4764	Catenin	3
-1.0562	Specific	1	2.5009	Helix	4
-1.0432	Identified	1	2.5481	Site	4
-1.0416	Direct	1	2.5764	Chains	4
-1.0399	Chemical	1	2.6105	Contacts	4
-1.0107	Factor	1	2.7096	Patch	4
-1.0050	Described	1	2.7569	Interaction	4
-1.0049	Helices	1	2.7836	Complex	4
-1.0042	Core	1	3.2955	Hydrogen	4
1.0270	Burring	3	3.9495	Formed	4
1.0278	Contributing	3	4.1811	Chain	4
1.0603	Protonation	3	4.4273	Interactions	4
1.0855	Forms	3	5.2863	Binding	4
1.0900	Interacts	3	9.1639	Hydrophobic	4

Table S2. Overall TM performance on the abstracts of PMC-OA test set of full-text articles with simplified residue filtering (basic TM) and with the residue filtering by the SVM model and DRNN. Both SVM model and DRNN were trained on the reduced full-text training set.

Method	^a	^b	Coverage (%) ^c	Success (%) ^d	Accuracy (%) ^e	$\Delta N(0)$ ^f	$\Delta N(1)$ ^f
Basic TM			30.5	15.4	50.6	—	—
SVM			8.9	6.2	69.6		
DRNN			7.7	5.4	70.0		

^a Number of complexes for which TM found at least one abstract with residues

^b Number of complexes with at least one interface residue found in abstracts

^c Ratio of L_{tot} and total number of complexes (259)

^d Ratio of L_{int} and total number of complexes (259)

^e Ratio of L_{int} and L_{tot}

^f From Eq. 2 in the main text with values from basic TM (first row) as X_2

Table S3. Overall TM performance on full texts with residue filtering using SVM model with automatically (Table S0) and manually (Badal, et al., 2018) selected keywords.

Keywords	L_{tot} ^a	L_{int} ^b	Coverage (%) ^c	Success (%) ^d	Accuracy (%) ^e
Auto		58	33.6	22.4	66.7
Manual		62	36.7	23.9	65.3

^a Number of complexes for which TM protocol found at least one article with residues

^b Number of complexes with at least one interface residue found in articles

^c Ratio of L_{tot} and total number of complexes (259)

^d Ratio of L_{int} and total number of complexes (259)

^e Ratio of L_{int} and L_{tot}

Table S4. Influence of the training set on TM performance with residue filtering by DRNN.

Training dataset	^a	^b	Coverage (%) ^c	Success (%) ^d	Accuracy (%) ^e	$\Delta N(0)$ ^f	$\Delta N(1)$ ^f
Full PMC-OA set, (4,982 sentences)			30.9	20.7	67.0	−3	+6
Reduced PMC-OA set (2,492 sentences)			20.0	13.3	66.4		

^a Number of complexes for which TM protocol found at least one abstract with residues

^b Number of complexes with at least one interface residue found in abstracts

^c Ratio of L_{tot} and total number of complexes (579)

^d Ratio of L_{int} and total number of complexes (579)

^e Ratio of L_{int} and L_{tot}

^f From Eq. 6 in the main text with values from the second row in Table 1 as X_2

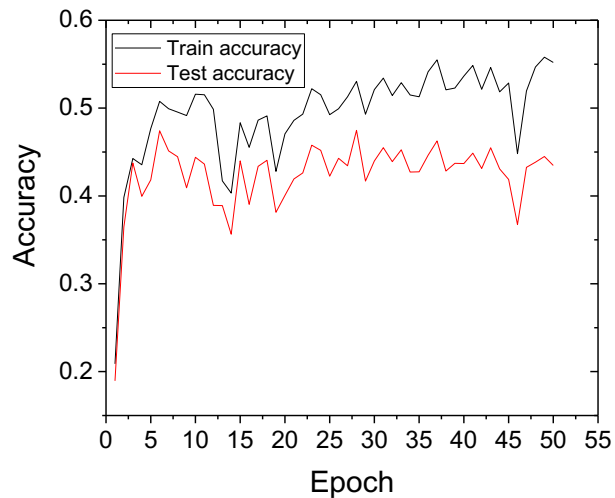


Figure S1. *Dependence of Deep Recursive Neural Network accuracy on the training length.* Training and testing were performed on the 4,982 sentences of the PMC-OA full-text articles and 5,786 sentences of the PubMed abstracts, respectively. The accuracy was defined as the fraction of sentences in the dataset, for which the correct sentiment was assigned. The DRNN learned over ~10 epochs. Beyond that it appeared to be over-trained (accuracy on the training set increases without the corresponding increase in the accuracy on the test set). Sharp falls and rises in the DRNN accuracy can be attributed to the drop-out method used to avoid the exploding gradient issue (Irsoy and Cardie, 2014). Drop-out nodes were chosen at random and at times could correspond to a weight representing important learning, thus causing a sudden drop in test and training accuracies learned in subsequent epochs.

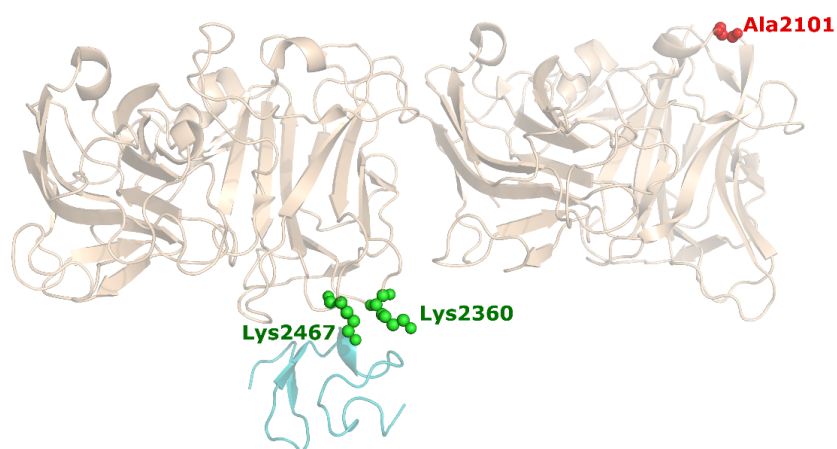


Figure S2. *Example of residues mined from abstracts but not from full texts.* The abstracts are from PubMed and the full texts are from PMC-OA. The structure is 3a7q chain A (gray) and B (cyan). Interface (green) and non-interface (red) residues are mined by the basic TM protocol. Details are in Text S1.

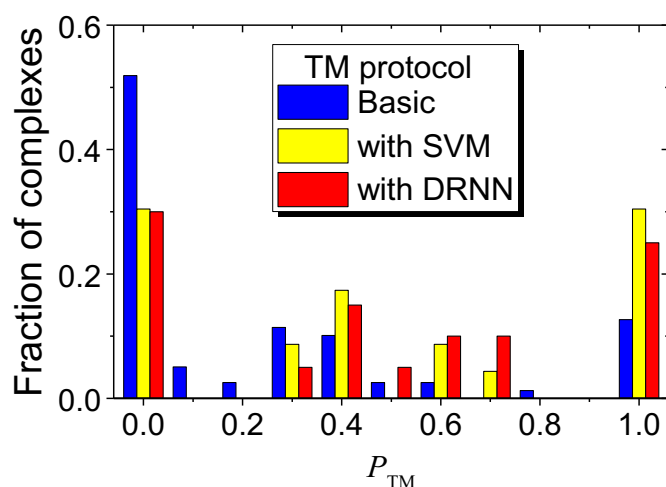


Figure S3. TM performance on the abstracts of PMC-OA test set of full-text articles with simplified residue filtering (basic TM) and with residue filtering by SVM model and DRNN. SVM model and DRNN were trained on the reduced full-text training set. The performance P_{TM} for individual complexes was calculated by Eq. 1 in the main text. The distribution is normalized to the total number of complexes for which residues were extracted (Table S3).

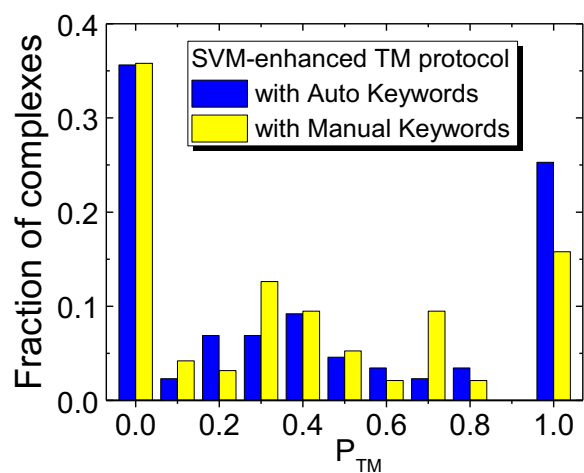


Figure S4. Comparison of the SVM model performance with automatic (Table S1) and manually selected (Badal, et al., 2018) keywords. The performance P_{TM} for individual complexes was calculated by Eq. 1 in the main text. The distribution is normalized by the total number of complexes for which residues were extracted (Table S2).

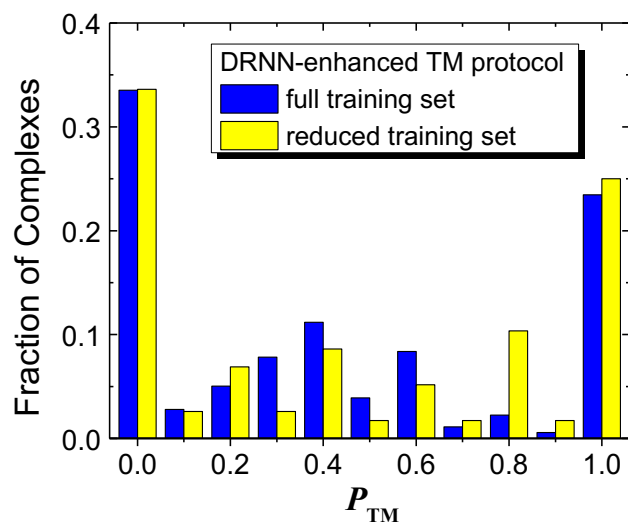


Figure S5. Influence of the training set on TM performance with residue filtering by DRNN. Full and reduced training sets consisted of 4,982 and 2,492 residue-containing sentences, respectively, from PMC-OA full-text articles. The performance P_{TM} for individual complexes was calculated by Eq. 1 in the main text. The distribution is normalized by the total number of complexes for which residues were identified (Table S4).

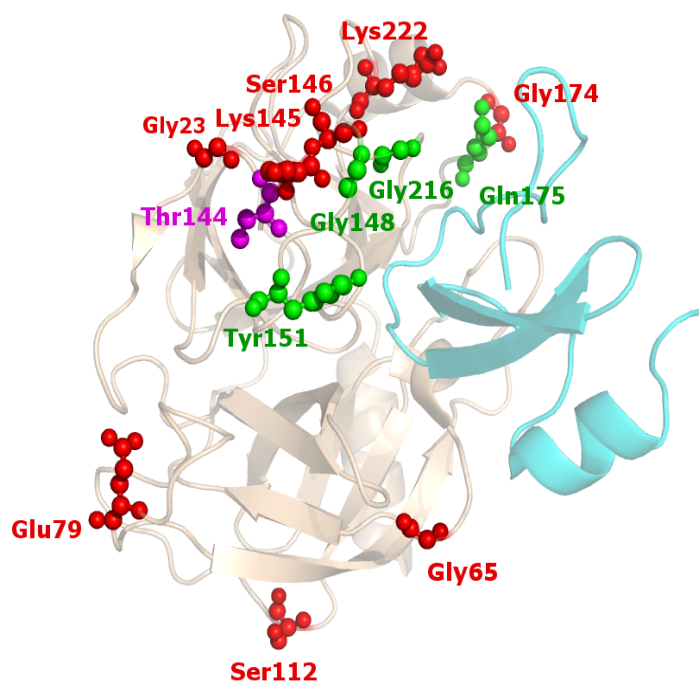


Figure S6. Example of residues mined from abstracts by DRNN-enhanced protocol in full sentence and with 7-words window. The structure is 2uuy chain A (gray) and B (cyan). Residues correctly identified at the interface are in green, and incorrectly identified in red. DRNN correctly identified two out of three residues in full sentences (the incorrectly identified residue is in magenta). With the window size of seven it correctly identified two residues (Gly216, Tyr151). Details are in Text S4.

Text S1: for Figure S2

Residue filtering by basic TM for the complex of reelin (chain A) and low-density lipoprotein receptor-related protein 8 (chain B) of 3a7q identified from PubMed abstracts and not from PMC–OA full text. Only AND-query identified 3 PubMed abstracts. Three residues passed the initial filters of the basic TM (see Methods in the main text), of which 2 are at the interface ($P_{\text{TM}} = 0.66$). All three residues belong to reelin (chain A of 3a7q). Direct citation of PDB entry identified Lys2467 and Lys2360 (both in chain A). These residues play an important role in interaction of reelin with apolipoprotein E receptor 2 (ApoER2) (Yasui, et al., 2010). Structure-guided alanine mutagenesis of fifth and sixth reelin repeats (R5-6) identified that residues Lys2467 and Lys2360 (both in chain A), are part of central binding site for low-density lipoprotein receptor (Yasui, et al., 2007). Another residue Ala2101 (chain A) was identified in mutant reelin where this mutant failed to assemble into multimers via disulfide bonds. However, it non-covalently associated high molecular weight oligomeric states in solution. The binding assay (surface plasmon resonance) showed that this mutant retained binding capability towards low density lipoprotein receptor (Yasui, et al., 2011).

Text S2: for Figure 6 (main text)

Residue filtering by basic TM for the complex of heat shock protein HSP82 (chain A) and AHA1 (chain B) of 1usu identified one residue from PubMed abstract. PMC (OA) abstracts did not identify any residues, while PMC (OA) full-text article identified six residues. The residue identified by PubMed is not at the interface ($P_{TM} = 0$). Four out of six residues identified by PMC (OA) full text are at the interface ($P_{TM} = 0.66$). NLP-hybrid method on PMC (OA) full-text further filtered the residues by identifying three residues, all of which are at the interface ($P_{TM} = 1$). All residues were identified by OR-queries.

TM on abstract from PubMed identified Glu90 in mutant study of Thr90 in Hsp90 α when investigating the phosphorylation impact of this residue, showing that Thr90 is involved in the regulation of the Hsp90 α chaperone machinery (Wang, et al., 2012) ($P_{TM} = 0$).

TM on PMC (OA) full-text articles identified six residues. The abstract of this full-text article, where the residues Leu66, Ile64 and Phe100 were identified as PPI-related using basic TM, is about a method to photo-crosslink interacting proteins using p-azido-L-phenylalanine (pAzpa). The non-canonical amino acid pAzpa was incorporated into a domain of Aha1 that was known to bind Hsp90 in vitro (Berg, et al., 2014). The abstract of the full-text article from which residue Asp145 was retrieved using basic TM is on the analysis of SGT1–HSP90 (Suppressor of G2 allele of skp1 and Heat-shock protein 90). The full text mentions this residue as a target for site-directed mutagenesis in HSP90 in wheat (Kadota, et al., 2008). Asp53 was identified in the full-text article on structural study of Aha1 binding with Hsp90 to modulate ATP hydrolysis cycle and client activity in vivo. When this residue is mutated in N-terminal domain in yeast Hch1p it impairs the ability to stimulate Hsp90 (Koulov, et al., 2010). The abstract where the residue Glu22 was identified by basic TM mentions the results of computational study of allosteric regulation in Hsp90 complexes with p23 and Aha1 as the co-chaperones. NLP-hybrid on full text correctly classified three interface residues (Leu66, Ile64 and Phe100). It rejected 3 residues (Asp145, Asp53 and Glu22) of which one (Asp53) was at the interface. Deep learning on full text yielded results identical to those of NLP-hybrid. DL with 7 words window accepted two residues which are at the interface (Leu66 and Ile64).

The interface residue Asp53 can be easily identified by a human reader from the following text (Koulov, et al., 2010) "*A previous study identified a mutation in the N-terminal domain of yeast Hch1p (D53K) that impaired its ability to stimulate yeast Hsp90 (Meyer et al., 2004). We made the equivalent mutation in human Aha1 (E67K) and tested for stimulation of Hsp90 ATPase activity. No stimulation was observed, similar to that observed previously for yeast D53K (Figure 5B).*" The failure of the automated protocols here is because this residue was mentioned not in relation to the proteins studied in the paper, but in a broader context of analogy between current and other studies.

Text S3: Examples of false negatives

Example 1. Residue Ser91 at the interface of 2yho chains A and B was found in the sentence “*Further, wild-type Otub1 and its catalytic mutant (Otub1(C91S)), but not Otub1(D88A), bind to the MDM2 cognate E2, UbcH5, and suppress its Ub-conjugating activity in vitro*” from a full-text publication (Sun, et al., 2012).

Example 2. Residue Glu67 at the interface of 2uyz chains A and B was found in the sentence “*Surprisingly, DPP9 binds to SUMO1 independent of the well-known SUMO interacting motif, but instead interacts with a loop involving Glu(67) of SUMO1*” from a full-text publication (Pilla, et al., 2012).

Both residues were rejected by DRNN and SVM models but can be easily identified by a human reader. We argue that this is due to the need to interpret a sequence of keywords (underlined in the above examples) rather than a single keyword as in our machine learning model, in order to access the correct context of the residue mentioning. Also, in these sentences, additional ambiguity is created by words “surprisingly”, “but” and the negation.

Text S4: for Figure S6

Residue filtering by basic TM for the complex of Cationic Trypsin (chain A) and Trypsin Inhibitor (chain B) of 2uuy identified 13 residues from 12 PubMed abstracts. Only 4 of 13 residues were correctly identified at the interface ($P_{TM} = 0.30$). DL using whole sentence identified 3 residues, of which only 2 were at the interface ($P_{TM} = 0.66$). DL using window of 7 identified 2 residues, both at the interface ($P_{TM} = 1$). All residues were identified by OR-queries.

Basic TM identified correctly Gly216 in a structural analysis of trypsin-BPTI interfaces (Kawamura, et al., 2011). Tyr151 was correctly identified by basic TM in trypsin forming hydrophobic interface in investigation of molecular specificity of Kunitz domain 1 (KD1) of tissue factor pathway inhibitor-2 (Schmidt, et al., 2005). Tyr151 was also identified in another abstract where crystal structure of trypsin were compared between different organism such as Atlantic salmon, chum salmon and bovine (Toyota, et al., 2002). Tyr151 again was correctly, and Ser146 incorrectly, identified in PubMed abstract where the residue is part of substrate activation binding site of bovine trypsin (Oliveira, et al., 1993).

Glu79 was incorrectly identified at interface by basic TM in the abstract describing E79K mutation in cationic trypsin causing increase in transactivation of anionic trypsinogen and used in-vitro analysis of recombinant wild and mutant enzymes (Teich, et al., 2004). Gly65 and Gly23 were incorrectly identified by basic TM from an abstract that describes the specificity of papaya proteinase IV (PPIV) for cleaving glycyl bonds (Buttle, et al., 1990). Of Thr144 and Gly148 identified by TM only the latter is at the interface. The abstract is on the structure of complex of bdellastasin and porcine beta-trypsin (Rester, et al., 2000). Of Gly174, Gln175 and Gly216 identified by TM only the last two were correctly determined to be at the interface. The abstract is about a comparative study of structures of cyclotheonamide A (CtA) complexes of alpha-thrombin and beta-trypsin (Ganesh, et al., 1996). Ser112 was incorrectly identified by basic TM at interface, from abstract about supercharged mutant (variant) of serine protease human enteropeptidase light chain (Simeonov, et al., 2012).

Lys145 and Ser146 were incorrectly identified by basic TM at the interface, in abstract about a crystal structure of an active autolysate form of porcine alpha-trypsin (APT) (Johnson, et al., 1997). Lys145 and Ser146 were also identified as an autolysis position in the abstract where the crystal structure of Porcine epsilon-trypsin was studied by molecular replacement (Huang, et al., 1994). Lys222 was incorrectly identified by basic TM from an abstract about covalent binding of proteinases by human alpha 2-macroglobulin (Sottrup-Jensen, et al., 1990).

DL using full sentence identified 3 residues (Tyr151, Thr144 and Gly148), of which 2 (Tyr151, Gly148) were correctly determined at the interface. Thr144 is incorrectly identified at the interface. DL using window size 7 identified 2 residues (Gly216, Tyr151), both at the interface.

References

- Badal, V.D., Kundrotas, P.J. and Vakser, I.A. Natural language processing in text mining for structural modeling of protein complexes. *BMC Bioinformatics* 2018;19:84.
- Berg, M., *et al.* An in vivo photo-cross-linking approach reveals a homodimerization domain of Aha1 in *S. cerevisiae*. *PLoS one* 2014;9:e89436.
- Buttle, D.J., *et al.* Selective cleavage of glycol bonds by papaya proteinase IV. *FEBS Lett.* 1990;260:195-197.
- Ganesh, V., *et al.* Comparison of the structures of the cyclotheonamide A complexes of human alpha-thrombin and bovine beta-trypsin. *Protein Sci.* 1996;5:825-835.
- Huang, Q., *et al.* Refined 1.8 Å resolution crystal structure of the porcine epsilon-trypsin. *Biochim. Biophys. Acta* 1994;1209:77-82.
- Irsoy, O. and Cardie, C. Deep recursive neural networks for compositionality in language. In, *Adv. Neural Information Processing Systems*. 2014. p. 2096-2104.
- Johnson, A., *et al.* The first structure at 1.8 Å resolution of an active autolysate form of porcine alpha-trypsin. *Acta Crystallogr. D Biol. Crystallogr.* 1997;53:311-315.
- Kadota, Y., *et al.* Structural and functional analysis of SGT1-HSP90 core complex required for innate immunity in plants. *EMBO Rep.* 2008;9:1209-1215.
- Kawamura, K., *et al.* X-ray and neutron protein crystallographic analysis of the trypsin-BPTI complex. *Acta Crystallogr. D Biol. Crystallogr.* 2011;67:140-148.
- Koulov, A.V., *et al.* Biological and structural basis for Aha1 regulation of Hsp90 ATPase activity in maintaining proteostasis in the human disease cystic fibrosis. *Mol. Biol. Cell* 2010;21:871-884.
- Oliveira, M.G., *et al.* Tyrosine 151 is part of the substrate activation binding site of bovine trypsin. Identification by covalent labeling with p-diazoniumbenzamidine and kinetic characterization of Tyr-151-(p-benzamidino)-azo-beta-trypsin. *J. Biol. Chem.* 1993;268:26893-26903.
- Pilla, E., *et al.* A novel SUMO1-specific interacting motif in dipeptidyl peptidase 9 (DPP9) that is important for enzymatic regulation. *J Biol Chem* 2012;287(53):44320-44329.
- Rester, U., *et al.* L-Isoaspartate 115 of porcine beta-trypsin promotes crystallization of its complex with bdellastasin. *Acta Crystallogr. D Biol. Crystallogr.* 2000;56:581-588.
- Schmidt, A.E., *et al.* Crystal structure of Kunitz domain 1 (KD1) of tissue factor pathway inhibitor-2 in complex with trypsin. Implications for KD1 specificity of inhibition. *J. Biol. Chem.* 2005;280:27832-27838.
- Simeonov, P., *et al.* Crystal structure of a supercharged variant of the human enteropeptidase light chain. *Proteins* 2012;80:1907-1910.
- Sottrup-Jensen, L., *et al.* Localization of epsilon-lysyl-gamma-glutamyl cross-links in five human alpha 2-macroglobulin-proteinase complexes. Nature of the high molecular weight cross-linked products. *J. Biol. Chem.* 1990;265:17727-17737.
- Sun, X.X., Challagundla, K.B. and Dai, M.S. Positive regulation of p53 stability and activity by the deubiquitinating enzyme Otubain 1. *EMBO J* 2012;31:576-592.
- Teich, N., *et al.* Interaction between trypsinogen isoforms in genetically determined pancreatitis: mutation E79K in cationic trypsin (PRSS1) causes increased transactivation of anionic trypsinogen (PRSS2). *Hum. Mutat.* 2004;23:22-31.

Toyota, E., et al. Crystal structure and nucleotide sequence of an anionic trypsin from chum salmon (*Oncorhynchus keta*) in comparison with Atlantic salmon (*Salmo salar*) and bovine trypsin. *J. Mol. Biol.* 2002;324:391-397.

Wang, X., et al. Thr90 phosphorylation of Hsp90alpha by protein kinase A regulates its chaperone machinery. *Biochem. J.* 2012;441:387-397.

Yasui, N., et al. Functional importance of covalent homodimer of reelin protein linked via its central region. *J. Biol. Chem.* 2011;286:35247-35256.

Yasui, N., et al. Structure of a receptor-binding fragment of reelin and mutational analysis reveal a recognition mechanism similar to endocytic receptors. *Proc. Natl. Acad. Sci. USA* 2007;104:9988-9993.

Yasui, N., Nogi, T. and Takagi, J. Structural basis for specific recognition of reelin by its receptors. *Structure (London, England : 1993)* 2010;18:320-331.