

Modeling Factual Claims by Frames

Fatma Arslan

The University of Texas at Arlington
fatma.dogan@mavs.uta.edu

Damian Jimenez

The University of Texas at Arlington
damian.jimenez@mavs.uta.edu

Josue Caraballo*

State Farm
josue.caraballo@mavs.uta.edu

Gensheng Zhang[†]

Google, Inc.
gezhang@google.com

Chengkai Li

The University of Texas at Arlington
cli@uta.edu

ABSTRACT

In this paper we introduce an extension of FrameNet for structured and semantic modeling of factual claims and an adaptation of the frame detection algorithms in Open Sesame for identifying frames and extracting frame elements from text. This claim modeling capability can be leveraged in assisting a variety of steps for automating fact-checking, e.g., matching claims with fact-checks, translating claims to structured queries, and so on. Our preliminary results show that while many challenges remain, which we discuss, frames can potentially improve the aforementioned steps. Further studies will reveal the strength and weakness of this modeling approach in more detail, as well as how to incorporate it into the full pipeline of fact-checking automation.

1 INTRODUCTION

With the increasing demands for fact-checking, there is a growing interest in automating various fact-checking steps (e.g., [3]). Such automation calls for structured and semantic modeling of factual claims that can capture various aspects of a factual claim, including the domain and topic of the claim, the template of the fact being expressed, the entities involved and their relationships, quantities, points and intervals in time, comparisons, and aggregate structures.

With such modeling capability in place, we will be able to construct fact-checking automation tools that exploit the idiosyncrasies of different forms of factual claims. For instance, in determining if new claims are identical or opposite to fact-checks and past claims in a curated repository, the claim-matching algorithm can go beyond current methods for paraphrase detection, semantic similarity and textual entailment by direct, fine-grained comparison of claims' structured representations. For translating claims into verification queries over knowledge graphs and structured databases, query templates can be carefully crafted beforehand for different types of claims, and methods can be designed to replace the variables in the query templates by entities and elements from the structured representations.

This paper presents our latest progress along this direction. Our approach is to extend the Berkeley FrameNet project,¹ a lexical resource for English built on the theory of meaning called frame semantics [1]. This "theory asserts that people understand the meaning of words largely by virtue of the frames which they evoke" [8]. A frame is a schematic representation with which we can formalize a structure to describe a particular kind of event, situation, object,

or relation along with its participants. Our extension of FrameNet includes 13 new frames specifically tailored to the subject of fact-checking-factual claims. Each frame comes with a definition and descriptions of its elements as well as a set of example sentences which are annotated with the corresponding frame elements.

While these new frames enhance our capability in modeling factual claims, automatically identifying frames from text and further identifying the frames' elements is still non-trivial. To tackle this challenge, we used Open Sesame,² an open-source frame identification and frame element extraction tool based on recurrent neural networks. For training its machine learning models, we annotated 900 sentences which were gathered from fact-checks released by PolitiFact in the past. Though this process was laborious and required a lot of time, we now have machine learning models that incorporate our frames and can detect other sentences that fall within the scope of our new frames. Furthermore, these models also detect each constituent element of a frame.

We further discuss how the frames detected using these new models can assist in various fact-checking automation steps including, but not limited to, translating pieces of text into structured queries, and matching claims and fact-checks by their frames. We are conducting empirical study of this and the preliminary results of the study are presented in this paper.

2 MODELING FACTUAL CLAIMS

2.1 Background: FrameNet

In FrameNet, each frame is comprised of several components: frame definition, associated frame elements, lexical units, exemplified and annotated sentences, and frame-to-frame relations. The FrameNet database currently contains 1224 semantic frames, 13,640 LUs, and 202,000 annotated sentences.³ A frame element (FE) is a frame-specific semantic role which provides additional information to the semantic structure of a sentence. A lexical unit (LU) is a pairing of a lemma with its parts of speech. A frame has many LUs associated with it and an LU can be part of many frames because it can have different meanings in these frames. For example, the Taking Sides frame, illustrated in Table 1, describes a situation involving frame elements such as Cognizer, Issue, Action, and Side. The LUs include frame evoking words such as back, oppose, support, etc.

*Work done while at the University of Texas at Arlington.

[†]Work done while at the University of Texas at Arlington.

¹<https://framenet.icsi.berkeley.edu/fndrupal/>

²<https://github.com/swabhs/open-sesame>

³https://framenet.icsi.berkeley.edu/fndrupal/current_status

Table 1: An example frame from FrameNet

Frame: Taking Sides		
Definition	A Cognizer has a relatively fixed positive or negative point of view towards an Issue . A Side in a debate concerning an Issue or an Action of a Side may stand in for the Issue . The Cognizer 's Degree of alignment may also be specified.	
	Hilary Clinton OPPOSED an individual mandate ...	
FEs	Cognizer Rick Scott backed the federal shutdown. Action He opposed the president's decision to go Issue David Perdue supports Common Core . Side He has supported George Bush .	
LUs	against.prep, back.v, backing.n, believe (in).v, endorse.v, for.prep,in favor.prep, opponent.n, oppose.v, opposition [act].n, opposition [entity].n, part.n, pro.adv, side.n, side.v, support.v, supporter.n, supportive.a	

2.2 Methodology: Defining New Frames

We collected fact-checked claims from PolitiFact.⁴ We chose a number of sentences among these claims and thoroughly examined them one by one. We then grouped the claims by finding semantic and syntactic similarities between them. Throughout this process we tried to maintain some generality to the groups, as to avoid having many groups with a small number of sentences. Instead our goal during this step was to capture the essence of what certain types of sentences were trying to convey. At the end of this process we had 20 groups of claims. We then adopted FrameNet's methodology to further fine-tune our claim modeling approach. For the claim groups where we found a preexisting FrameNet frame, we simply used that frame. For the claim groups where we felt either FrameNet did not have a suitable frame, or its frames did not capture what we thought they should, then a new frame was created. As previously mentioned, at the end of this process we have 13 new frames that are novel in what they captured and are created with factual claims specifically in mind, and we also identified 7 existing frames from FrameNet that are prominent among factual claims.

2.3 Outcome: New Frames

We enumerate the 20 frames, including 7 existing ones and 13 new frames defined by us. For each frame, we make note regarding whether it is new. We briefly explain the frame elements in each frame and provide a couple of sample sentences annotated with the lexical units (in boldface) and frame elements (in square brackets).

1. Taking sides (existing). A "Cognizer" has a relatively fixed positive or negative point of view towards an "Issue". A "Side" in a debate concerning an "Issue" or an "Action" of a "Side" may stand in for the "Issue".

[Sen. Kamala Harris *COGNIZER*] is **supporting** [the animals of MS-13 *ISSUE*].

[By 2006 *TIME*], [the American people *COGNIZER*] were [overwhelmingly *DEGREE*] **against** [the Iraq War *ISSUE*].

2. Oppose and support consistency (new). This frame is about the consistency of an "Agent's" "Stance" towards an "Issue". The "Agent" either alters or maintains his/her "Stance". The "Stance" may not be explicitly stated.

[Israel Prime Minister Benjamin Netanyahu *AGENT*] didn't **change** [his position *STANCE*] [on a two-state solution *ISSUE*].

[Republicans Chuck Grassley, John Boehner and John Mica *AGENT*] **flip-flopped** [on providing end-of-life counseling for the elderly *ISSUE*].

3. Speech (new). A "Speaker" uses language in the written or spoken modality to communicate a "Message" to some "Addressee". A "Topic" may be stated instead of a "Message".

[Ronald Reagan *SPEAKER*] **talked** [about converting the United States to the metric system *TOPIC*].

[President Barack Obama *SPEAKER*] **said** [at the beginning of the negotiations *TIME*] [that the basic approach was to dismantle Iran's nuclear program in exchange for dismantling the sanctions *MESSAGE*].

4. Recurring action (new). The Recurring action frame describes a repetitive "Action" that is performed by an "Agent" at the interval of a "Time_span".

[Last year *TIME*], [Exxon *AGENT*] [pocketed nearly \$4.7 million *ACTION*] **every** [hour *TIME_SPAN*].

[Undocumented immigrants *AGENT*] [pay \$12 billion of taxes *ACTION*] **every** [single year *TIME_SPAN*].

5. Recurring action with frequency (new). This frame is about a repetitive "Action" that is performed by an "Agent" at a given "Frequency".

[Donald Trump *AGENT*] [was forced to file for bankruptcy *ACTION*] not once, not twice, [four *FREQUENCY*] **times**.

[Chemical weapons have been used *ACTION*] probably [20 *FREQUENCY*] **times** [since the Persian Gulf War *TIME*].

6. Vote (new). The Vote frame is about an "Agent's" "Position" on a voting decision for an "Action" or "Issue".

[Mitch McConnell *AGENT*] **voted** [three times *FREQUENCY*] [for *POSITION*] [corporate tax breaks that send Kentucky jobs overseas *ISSUE*].

[Tom Cotton *AGENT*] **voted** [against *POSITION*] [preparing America for pandemics like Ebola *ACTION*].

7. Causation (existing). A "Cause" causes an "Effect". Alternatively, an "Actor", a participant of a (implicit) "Cause", may stand in for the "Cause". The entity "Affected" by the causation may stand in for the overall "Effect" situation or event.

⁴<https://www.politifact.com>

[Obamacare *CAUSE*] has **caused** [millions of full-time jobs to become part-time *EFFECT*].

Due to [actions by President Barack Obama *CAUSE*], [the Burger King national headquarters announced this month that they will be pulling their franchises from our military bases *EFFECT*].

8. Capability (existing). An “Entity” meets the pre-conditions for participating in an “Event”. A “Degree” modifier may be included to indicate by how much the “Entity” exceeds or falls short of the minimum requirements.

[Western Europeans *ENTITY*] **can** [fly in the United States *EVENT*] [without even having a visa *CIRCUMSTANCES*].

[Former President George W. Bush and former Vice President Dick Cheney *ENTITY*] are **unable** [to visit Europe *EVENT*] [due to outstanding warrants *CIRCUMSTANCES*].

9. Conditional occurrence (new). A “Consequence” materializes if the “Conditional_event” occurs.

If [the Iran nuclear deal gets rejected *CONDITIONAL_EVENT*], [they still get \$150 billion *CONSEQUENCE*].

[We would create thousands of jobs in Colorado *CONSEQUENCE*], **if** [the Keystone Pipeline were to be built *CONDITIONAL_EVENT*].

10. Correlation (new). It shows the connection or relationship between the occurrences of “Event_1” and “Event_2”.

Every time [we’ve increased the minimum wage *EVENT_1*], [we’ve seen a growth in jobs *EVENT_2*].

Whenever [we raise the capital gains tax *EVENT_1*], [the economy has been damaged *EVENT_2*].

11. Cause change of position on a scale (existing). An “Agent” or a “Cause” affects the position of an “Item” on some scale (the “Attribute”) to change it from an initial value (“Value_1”) to an end value (“Value_2”). The magnitude of the change (“Difference”) can be encoded.

[Thom Tillis *AGENT*] has **cut** [\$500 million *DIFFERENCE*] [from public education *ITEM*].

[In the last two years *TIME*], [we *AGENT*] have **reduced** [the deficit *ATTRIBUTE*] [by \$2.5 trillion *DIFFERENCE*].

12. Change position on a scale (existing). The change of an “Item’s” position on a scale (the “Attribute”) from a starting point (“Initial_value”) to an end point (“Final_value”). The magnitude of the change (“Difference”) can be indicated.

[Since 2007 *TIME*], [Texas *ITEM*] has **gained** [440,000 people *DIFFERENCE*] while Maryland has lost 20,000.

[During Obama’s first five years as president *TIME*], [black unemployment *ITEM*] **increased** [42 percent *DIFFERENCE*].

13. Comparing two entities (new). This frame is about comparing two entities using a “Comparison_criterion” while qualifying with a “Degree”.

[Hillary Clinton *ENTITY_1*] [has been in office and in government longer *COMPARISON_CRITERION*] **than** [anybody else running here tonight *ENTITY_2*].

[This president *ENTITY_1*] [has offered fewer executive actions *COMPARISON_CRITERION*] **than** [almost any other president preceding his presidency in recent history *ENTITY_2*].

14. Comparing at two different points in time (new). This frame is about comparing an “Entity” with itself at two different points in time using a “Comparison_criterion” while qualifying with a “Degree”.

[More *DEGREE*] [private-sector jobs *ENTITY*] [were created *COMPARISON_CRITERION*] [in the second year of the Obama administration *FIRST_TIME_POINT*] **than** [in the eight years of the Bush administration *SECOND_TIME_POINT*].

[The average family *ENTITY*] is [now *FIRST_TIME_POINT*] [bringing home \$4,000 less *COMPARISON_CRITERION*] **than** they did [just five years ago *SECOND_TIME_POINT*].

15. Creating (existing). A “Cause” leads to the formation of a “Created_entity”.

[In the last 29 months *TIME*], [our economy *CREATOR*] has **produced** [about 4.5 million private-sector jobs *CREATED_ENTITY*].

[Ohio *CREATOR*] has **created** [45,000 new manufacturing jobs *CREATED_ENTITY*] [since 2010 *TIME*].

16. Occupy rank (existing). This frame is about “Items” in the state of occupying a certain “Rank” within a hierarchy.

[Under Gov. Tom Corbett *TIME*], [Pennsylvania *ITEM*] **ranked** [49th *RANK*] [in job creation *DIMENSION*].

[The U.S. *ITEM*] only **ranked** [25th *RANK*] [worldwide *COMPARISON_SET*] [on defense spending as a percentage of GDP *DIMENSION*].

17. Occupy rank via ordinal numbers (new). This frame is about “Items” in the state of occupying a certain “Rank” specified by an ordinal number within a hierarchy.

[New Mexico *ITEM*] moved “up to” [**sixth** *RANK*] [in the nation *COMPARISON_SET*] [in job growth *DIMENSION*].

[The United States *ITEM*] is [**65th** *RANK*] [out of 142 nations and other territories *COMPARISON_SET*] [on equal pay *DIMENSION*].

18. Occupy rank via superlatives (new). This frame is about “Items” in the state of occupying a certain “Rank” specified by a superlative within a hierarchy.

[Job growth in the United States *ITEM*] is [now *TIME*] at [the **fastest** *RANK*] [pace *DIMENSION*] [in this country’s history *COMPARISON_SET*].

[The state of New York *ITEM*] is [the **worst** *RANK*] [in the nation *COMPARISON_SET*] [in economic recovery *DIMENSION*].

19. Ratio (new). In this frame, a “Criterion” determines a “Ratio” that quantifies the size of the subset of a larger “Group”.

[Today *TIME*], [about 40 *RATIO*] **percent of** [guns *GROUP*] are [purchased without a background check *CRITERION*].

[More than 72 *RATIO*] **percent of** [children in the African-American community *GROUP*] are [born out of wedlock *CRITERION*].

20. Uniqueness of trait (new). This frame distinguishes a “Unique entity” from a “Generic entity” based on a specific “Trait” where a “Trait” is some property, quality, point-of-view, or an arbitrary construct which is generally understood to be an attribute of an entity.

[The United States *UNIQUE_ENTITY*] is the **only** [advanced country on Earth *GENERIC_ENTITY*] [that doesn’t guarantee paid maternity leave to our workers *TRAIT*].

[New Jersey *UNIQUE_ENTITY*] is the **only** [state in the union *GENERIC_ENTITY*] [that spent less on higher education than it did at the beginning of the decade *TRAIT*].

3 THE POTENTIAL USE OF FRAMES

3.1 Claim Detection

Claim detection, a necessary stage of the fact-checking process, identifies claims to be fact-checked. We can exploit frames that have been specifically created to address different factual claims to accomplish claim detection task. With factual frames in hand, we can remodel the claim detection task as identifying claims that have been found to be affiliated with at least one of the 20 frames. The claims identified can be used as input for the downstream operations in the automation pipeline of fact-checking. In Section 4 we demonstrate the use of frames in claim detection.

3.2 Claim Matching

Claim matching is the process of partially or fully matching a new factual claim with either supporting or refuting fact-checked claims stored in a repository. In the simplest case, a new factual claim may match completely with an existing identical factual claim. In such a case, the user may be presented with the professional fact-checker’s verdict to gauge the truthfulness of the claim. In the not-so-simple cases, where a new factual claim is completely or partially opposite of or partially similar to a stored fact-checked claim, we can still employ fact-checked claims and augment their verdicts. The frames proposed in this paper can help us address the aforementioned cases. We can compare the detailed elements of claims (entities, time point and interval, quantity, aggregate, grouping, comparison) with those of the fact-checked claims. Based on the similarity of their corresponding frame elements, we can conclude whether the new factual claims are fully or partially similar to or opposite of the ones that they are compared with.

3.3 Claim to Query Translation

Another potential use of frames is by leveraging them in translating a given input sentence into a structured query so that the outcome of the query can be compared with the information embedded in the claim itself, in order to reach an assessment of its truthfulness. We can envision a situation in which each frame is mapped to a rough template for a query. The process for mapping a sentence to a query template would involve first understanding the schema

of the database that was being worked with. A query template could be associated with a list of tables that were pertinent to it, and then each table could be checked for columns that were relevant to the sentence waiting to be translated. Then using some templates queries could be constructed to check these tables using the different parts of the sentence that the frame had captured as inputs to the template.

4 PRELIMINARY EXPERIMENTS

We conducted several preliminary experiments to assess the efficacy of our study. We used open-sesame [9], an open-source frame semantic parser for automatically identifying FrameNet frames and their frame-elements from sentences. Open-sesame, a syntax-free system, is built on softmax-margin segmental recurrent neural nets and executes an array of tasks: target identification, frame identification, and argument identification. Target identification is the identification of the words or expressions that evoke the frames. Frame identification is identifying the frame that each target evokes. Argument identification is the identification of the frame elements and their corresponding span of text for each of the frames that a sentence triggers.

We utilized open-sesame for a claim detection task, more specifically to identify the frames that were evoked by a factual claim under consideration. We added three new frames along with their lexical units and hand-engineered annotated sentences into the FrameNet 1.7 dataset. The inserted frames were the Vote, Uniqueness of trait, and Recurring action frames. We chose these frames as they have a large enough number of sentences associated with them as to allow a satisfactory training phase.

We retrained open-sesame on this extended FrameNet 1.7 dataset and followed this by an evaluation of the trained model with the help of the “Share the Facts” dataset⁵. Since this dataset included some claims that were irrelevant for our current task of political fact-checking, we removed any fact-checks from international organizations and those associated with Hollywood gossip magazine sections. We evaluated open-sesame’s frame identification performance for the three new frames (i.e., the Vote, Uniqueness of trait, and Recurring action frames) in addition to the 7 pre-existing FrameNet frames (i.e., the Taking sides, Occupy rank, Creating, Capability, Causation, Change position on a scale, and cause change of position on a scale frames). Table 2 depicts the performance results for each of these frames. From the results we see that the Vote and Uniqueness of trait frames performed in line with the other pre-established frames. The recurring action had a low precision and thus lower F1-score. We also see some low scores for some of the pre-established frames, but we expect that as we are able to create a more robust labeling process we will have more data to feed the neural network to improve its performance. We can also look at fine tuning frame elements and/or lexical units from what they are currently defined as. However, the latter should only be necessary if we do not see a noticeable improvement with the inclusion of more training data. It should also be noted that during training, we were not only training the model to detect these frames but the entirety

⁵The “Share the Facts” dataset is the result of a joint effort by several prominent fact-checking organizations that aims to create a standardized format of fact-checks. The dataset now contains around 20,000 fact-checks and counting. (<http://www.sharethefacts.org/>)

Table 2: Frame prediction performance, in terms of Precision (P), Recall (R) and F-measure (F_1). Where avg_w denotes the weighted average of corresponding measure across ten frames. The number in parentheses below the frame name is the number of sentences used for evaluating that frame.

	Cause change of position on a scale (156)	Capability (48)	Causation (256)	Creating (114)	Change position on a scale (167)	Occupy rank (35)	Recurring action (29)	Taking sides (129)	Uniqueness of trait (33)	Vote (104)	avg_w
P	0.73	0.32	0.41	0.92	0.45	0.82	0.23	0.52	0.54	0.61	0.56
R	0.60	0.79	0.48	0.31	0.74	0.69	0.66	0.46	0.88	0.94	0.60
F_1	0.66	0.46	0.44	0.46	0.56	0.75	0.34	0.49	0.67	0.74	0.54

of the FrameNet frames. Thus another possible option would be to train the model only on the 20 frames we will eventually focus on to see if that improves performance and produces sound results. Overall, for preliminary results, we are satisfied in seeing that two of our frames performed decently as that validates the direction we are heading in.

5 RELATED WORK

Recently efforts have gone into developing taxonomies of political claim. One such effort was from fact-checkers in HeroX fact-checking challenge [2]. During the course of this challenge, a taxonomy of political claims was presented. This taxonomy was comprised of four claim types: numerical claims, verification of quotes, position statements, and lastly objects, properties and events.

Fullfact researchers [4] proposed a claim annotation schema based on their proprietary fact-checked claims. Their annotation schema consists of seven different categories: Personal experience, Quantity in the past or present, Correlation or causation, Current laws or rules of operation, Prediction, Other types of claim, and Not a claim. They assured that these claim categories cover an entire gamut of sentences used in political TV shows that they have come across over several years. This work is based on PolitiTax,⁶ a work done by us earlier in IDIR lab. The annotation schema that we propose in this paper is an enhanced version of PolitiTax. The main difference between our work and Fullfact annotation is that while Fullfact assigns a claim to a singular category only, in our schema, a claim can belong to multiple categories concurrently.

This work is also loosely related to research in entity annotation and word embeddings. The latter is a tool that was developed in order to better represent words in a vector space. Typically a large text corpus is used in training word embeddings, this way the embeddings better represent words as they are encountered in day to day text. The most prominent embeddings currently are perhaps the Google News vectors [5], GloVe [6], and ELMo [7], with the latter claiming to achieve state of the art results when compared with the former two. Our work aims to give meaning to the words within a sentence as well, but by compartmentalizing them into different parts of a frame. Similarly, entity annotation aims to annotate entities that are found within a piece of text, and our work also annotates entities within the context of our frames. The frames in a way are a composite effort that aims to identify key parts of a sentence (e.g., like entities) and give them context and

meaning that can be used in algorithms (i.e., what word embeddings aim to do). Although we don't think this method can replace the aforementioned two, it can be a powerful tool to use in conjunction with various other methods to produce higher quality results in research or work that uses these.

6 CONCLUSIONS

In summary we present an extension to FrameNet that focuses on annotating factual claims with key components that can be used to algorithmically check these or perform various NLP related tasks. Our work introduces 13 new frames along with 7 preexisting frames that we found useful. Our frames were crafted with fact-checking specifically in mind, so we expect that these will provide a substantially different functionality when fully integrated into a model. With this foundation set we can focus on future endeavors and continue to expand this work in the near future.

REFERENCES

- [1] Collin F. Baker, Charles J. Fillmore, and John B. Lowe. 1998. The Berkeley FrameNet Project. In *Proceedings of the 36th Annual Meeting of the Association for Computational Linguistics and 17th International Conference on Computational Linguistics - Volume 1 (ACL '98)*. Association for Computational Linguistics, Stroudsburg, PA, USA, 86–90. <https://doi.org/10.3115/980845.980860>
- [2] Diane Francis and Full Fact. 2016. HeroX fact check challenge. <https://www.herox.com/factcheck/5-practise-claims>. Accessed: 2018-10-25.
- [3] Naeemul Hassan, Chengkai Li, and Mark Tremayne. 2015. Detecting Check-worthy Factual Claims in Presidential Debates. In *Proceedings of the 24th ACM International on Conference on Information and Knowledge Management (CIKM '15)*. ACM, New York, NY, USA, 1835–1838. <https://doi.org/10.1145/2806416.2806652>
- [4] Lev Konstantinovskiy, Oliver Price, Mevan Babakar, and Arkaitz Zubiaga. 2018. Towards Automated Factchecking: Developing an Annotation Schema and Benchmark for Consistent Automated Claim Detection. *arXiv:cs.CL/1809.08193*
- [5] Tomas Mikolov, Ilya Sutskever, Kai Chen, Greg S Corrado, and Jeff Dean. 2013. Distributed representations of words and phrases and their compositionality. In *Advances in neural information processing systems*. 3111–3119.
- [6] Jeffrey Pennington, Richard Socher, and Christopher Manning. 2014. Glove: Global vectors for word representation. In *Proceedings of the 2014 conference on empirical methods in natural language processing (EMNLP)*. 1532–1543.
- [7] Matthew Peters, Mark Neumann, Mohit Iyyer, Matt Gardner, Christopher Clark, Kenton Lee, and Luke Zettlemoyer. 2018. Deep Contextualized Word Representations. In *Proceedings of the 2018 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies, Volume 1 (Long Papers)*, Vol. 1. 2227–2237.
- [8] Josef Ruppenhofer, Michael Ellsworth, Miriam Petruck, Christopher R. Johnson, and Jan Scheffczyk. 2010. FrameNet II: Extended Theory and Practice. (01 2010).
- [9] Swabha Swayamdipta, Sam Thomson, Chris Dyer, and Noah A. Smith. 2017. Frame-Semantic Parsing with Softmax-Margin Segmental RNNs and a Syntactic Scaffold. *arXiv preprint arXiv:1706.09528* (2017).

⁶PolitiTax: A Taxonomy of Political Claims by IDIR Lab <https://perma.cc/4RQF-FCPV>