

# Automatic Identification of Conceptual Metaphors with Limited Knowledge

Lisa Gandy<sup>1</sup> Nadji Allan<sup>2</sup> Mark Atallah<sup>2</sup> Ophir Frieder<sup>3</sup> Newton Howard<sup>4</sup>  
Sergey Kanareykin<sup>5</sup> Moshe Koppel<sup>6</sup> Mark Last<sup>7</sup> Yair Neuman<sup>7</sup> Shlomo Argamon<sup>1\*</sup>

<sup>1</sup> Linguistic Cognition Laboratory, Illinois Institute of Technology, Chicago, IL

<sup>2</sup> Center for Advanced Defense Studies, Washington, DC

<sup>3</sup> Department of Computer Science, Georgetown University, Washington, DC

<sup>4</sup> Massachusetts Institute of Technology, Cambridge, MA

<sup>5</sup> Brain Sciences Foundation, Providence, RI

<sup>6</sup> Department of Computer Science, Bar-Ilan University, Ramat Gan, Israel

<sup>7</sup> Departments of Education and Computer Science, Ben-Gurion University, Beersheva, Israel

## Abstract

Full natural language understanding requires identifying and analyzing the meanings of metaphors, which are ubiquitous in both text and speech. Over the last thirty years, linguistic metaphors have been shown to be based on more general *conceptual metaphors*, partial semantic mappings between disparate conceptual domains. Though some achievements have been made in identifying linguistic metaphors over the last decade or so, little work has been done to date on automatically identifying conceptual metaphors. This paper describes research on identifying conceptual metaphors based on corpus data. Our method uses as little background knowledge as possible, to ease transfer to new languages and to minimize any bias introduced by the knowledge base construction process. The method relies on general heuristics for identifying linguistic metaphors and statistical clustering (guided by Wordnet) to form conceptual metaphor candidates. Human experiments show the system effectively finds meaningful conceptual metaphors.

## Introduction

Metaphor is ubiquitous in human language. Thus, identifying and understanding metaphorical language is a key factor in understanding natural language text. Metaphorical expressions in language, called *linguistic metaphors*, have been shown to be derived from and justified by underlying *conceptual metaphors* (Lakoff and Johnson 1980), which are cognitive structures that map aspects of one conceptual domain (the *source concept*) to another, more abstract, domain (the *target concept*). A large and varied scientific literature has grown over the last few decades studying linguistic and conceptual metaphors, in fields including linguistics (e.g., (Johnson, Lakoff, and others 2002; Boers 2003)), psychology (e.g., (Gibbs Jr 1992; Allbritton, McKoon, and Gerrig 1995; Wickman et al. 1999)), cognitive science (e.g., (Rohrer 2007; Thibodeau and Boroditsky 2011)), literary studies (e.g., (Goguen and Harrell 2010; Freeman 2003)), and more.

\*Corresponding author. Email: [argamon@iit.edu](mailto:argamon@iit.edu)

Copyright © 2013, Association for the Advancement of Artificial Intelligence (www.aaai.org). All rights reserved.

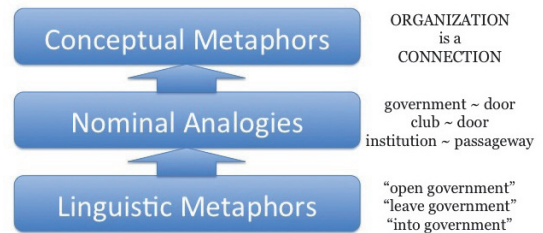


Figure 1: Three levels of metaphor processing.

In this paper, we describe a novel system for discovering conceptual metaphors using natural language corpus data. The system identifies linguistic metaphors automatically and uses them to find meaningful conceptual metaphors. A key design principle is to rely as little as possible on domain and language specific knowledge bases or ontologies, but rather more on analyzing statistical patterns of language use. We limit such use to generic lexical semantics resources such as Wordnet. Such a knowledge-limited, corpus-centric approach promises to be more cost effective in transferring to new domains and languages, and to potentially perform better when the underlying conceptual space is not easily known in advance, as, e.g., in the study of metaphor variation over time or in different cultures.

The system works at three interrelated levels of metaphor analysis (Figure 1): (i) linguistic metaphors (individual metaphoric expressions), (ii) nominal analogies (analogical mappings between specific nouns), and (iii) conceptual metaphors (mappings between concepts). For instance, the system, when given a text, might recognize the expression ‘open government’ as a linguistic metaphor. The system may then find the nominal analogy “government ~ door” based on that linguistic metaphor and others. Finally, by clustering this and other nominal analogies found in the data, the system may discover an underlying conceptual metaphor like “AN ORGANIZATION IS A CONNECTION”.

## Related Work

There has been a fair amount of computational work on identifying linguistic metaphors, but considerably less on finding conceptual metaphors. Most such work relies strongly on predefined semantic and domain knowledge. We give here a short summary of the most relevant previous work.

Birke and Sarkar (2006) approach the problem as a classical word sense disambiguation (WSD) task. They reduce nonliteral language recognition to WSD by considering literal and nonliteral usages to be different senses of a single word, and adapting an existing similarity-based word-sense disambiguation method to the task of separating verbs occurrences into literal and nonliteral usages.

Turney et al. (2011) identify metaphorical phrases by assuming that these phrases will consist of both a "concrete" term and an "abstract" term. In their work they derive an algorithm to define the abstractness of a term, and then use this algorithm to contrast the abstractness of adjective-noun phrases. The phrases where the abstractness of the adjective differs from the abstractness of the noun by a predetermined threshold are judged as metaphorical. This method achieved better results than Birke and Sarkar on the same data.

For interpreting metaphors researchers have used both rule and corpus based approaches. In one of the most successful recent corpus-based approaches, Shutova et al. (2012) find and paraphrase metaphors by co-clustering nouns and verbs. The CorMet system (Mason 2004) finds metaphorical mappings, given specific datasets in disparate conceptual domains. That is, if given a domain focused on lab work and one focused on finance it may find a conceptual metaphor saying that LIQUID is related to MONEY.

Nayak & Mukerjee (2012) have created a system which learns Container, Object and Subject Metaphors from observing a video demonstrating these metaphors visually and also observing a written commentary based on these metaphors. The system is a grounded cognitive model which uses a rule based approach to learn metaphor. The system works quite well, however it is apparently only currently capable of learning a small number conceptual metaphors.

## The Metaphor Analysis System

As discussed above, our metaphor analysis system works at three interrelated levels of metaphor analysis: finding linguistic metaphors, discovering nominal analogies and then clustering them to find meaningful conceptual metaphors. As shown in Figure 2, each level of metaphor analysis feeds into the next; all intermediate results are cached to avoid recomputing already-computed results. All levels of analysis rely on a large, preparsed corpus for statistical analysis (the experiments described below use the COCA n-grams database<sup>1</sup>). The system also currently uses a small amount of predefined semantic knowledge, including lexical semantics via Wordnet, dictionary senses via Wiktionary, and Turney's ratings of the abstractness of various terms (this last is derived automatically from corpus statistics). Input texts are also parsed to extract simple expressions as linguistic metaphor candidates to be classified by the system.

<sup>1</sup><http://corpus.byu.edu/coca/>

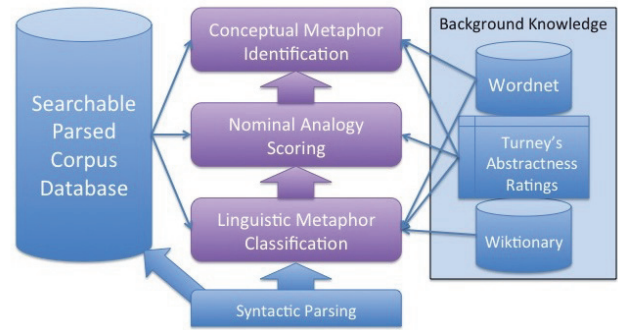


Figure 2: System Architecture.

Note that the only manually-created background knowledge sources used are Wordnet and the number of definitions of words in Wiktionary; these are used in very circumscribed ways in the system, as described below. Other than those resources, the only language-dependence in the system is the use of a parser to find expressions of certain forms as metaphor candidates, such as verb-object pairs.

## Linguistic Metaphors

The system uses a heuristic algorithm to identify linguistic metaphors. The method classifies expressions with specific syntactic forms as being either metaphorical or not based on consideration of the two focal words in the expression. The types of expressions we consider in this work are those delineated by (Krishnakumaran and Zhu 2007):

**Type 1** The subject and complement of an identifying clause e.g., "My **lawyer**<sub>T</sub> is a **shark**<sub>F</sub>."

**Type 2** A lexical verb together with its direct object, e.g., "He **entered**<sub>F</sub> **politics**<sub>T</sub>."

**Type 3** An adjective and the noun it describes, e.g., "**sweet**<sub>F</sub> **child**<sub>T</sub>," or "The **book**<sub>T</sub> is **dead**<sub>F</sub>."

We term the descriptive, possibly "metaphorical," term in each such expression the *facet* (marked *F* in the examples above), and the noun being described the *target* (marked *T*). Note that the facet expresses one aspect of a source domain used to understand the target metaphorically.

Our algorithm for identifying linguistic metaphors, an expansion of our earlier work (Assaf et al. 2013), is based on the key insight in Turney et al.'s (2011) Concrete-Abstract algorithm, i.e., the assumption that a metaphor usually involves a mapping from a relatively concrete domain to a relatively abstract domain. However, we also take into account the specific conceptual domains involved. Literal use of a concrete facet will tend to be more salient for certain categories of concrete objects and not others. For example, in its literal use, the adjective "dark" may be associated with certain semantic categories such as Substance (e.g. "wood") or Body Part (e.g. "skin").

To illustrate the idea, consider the case of Type 3 metaphors, consisting of an adjective-noun pair. The algorithm assumes that if the modified noun belongs to one of the

concrete categories associated with the literal use of the adjective then the phrase is probably non-metaphorical. Conversely, if the noun does not belong to one of the concrete categories associated with the literal use of the adjective then it is probably metaphorical. In a sense, this algorithm combines the notion of measuring concreteness and that of using selectional preferences, as has been well-explored in previous work on metaphor identification (Fass and Wilks 1983).

To illustrate the approach, consider analyzing an adjective-noun expression  $\langle A, N \rangle$  such as “open government”. First, if the adjective  $A$  has a single dictionary definition then the phrase is labeled as non-metaphorical, since metaphorical usage cannot exist for one sense only. Then, the algorithm verifies that the noun  $N$  belongs to at least one high-level semantic category (using the WordStat dictionary of semantic categories based on Wordnet<sup>2</sup>). If not, the algorithm cannot make a decision and stops. Otherwise, it identifies the  $n$  nouns most frequently collocated with  $A$ , and chooses the  $k$  most concrete nouns (using the abstractness scale of (Turney et al. 2011)). The high-level semantic categories represented in this list by at least  $i$  nouns each are selected; if  $N$  does not belong to the any of them, the phrase is labeled as metaphorical, otherwise it is labeled as non-metaphorical. The same method applies to other expression types, by varying how collocations are computed.

Based on analysis of a development dataset, disjoint from our test set, we set  $n = 1000$ ;  $k = 100$ ; and  $i = 16$ .

## Nominal Analogies

Once a set of metaphorical expressions  $S_M = \{\langle f, n_t \rangle\}$  (each a pair of a facet and a target noun) is identified, we seek nominal analogies which relate two specific nouns, a *source* and a *target*. For example, if the system finds many different linguistic metaphors which can be interpreted as viewing governments as doors, it may identify a nominal analogy “government  $\sim$  door.”

The basic idea behind the nominal analogy finding algorithm is that, since the metaphorical facets of a target noun are to be understood in reference to the source noun, we must find source/target pairs such that many of the metaphorical facets of the target are associated in literal senses with the source (see Figure 3). The stronger and more numerous these associations, the more likely the nominal analogy is to be real.

Note that the system will not only find frequent associations. Since the association strength of each facet-noun pairing is measured by PMI, not its raw frequency, infrequent pairings can (and do) pop up as significant.

**Candidate generation.** We first define the candidate set of source nouns  $S_C$  as the set of all nouns (in a given corpus) with strong non-metaphoric associations with the facets used in the linguistic metaphors in  $S_M$ . Start with the set of all facets used in  $S_M$  which have positive point-wise mutual information (PMI; cf. (Pantel and Lin 2002)) with at least one target term in the set.  $S_C$  is then defined as the set of all other (non-target) nouns (in a given large corpus) associ-

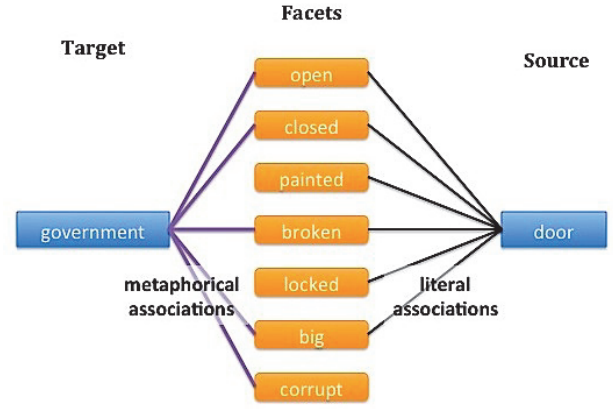


Figure 3: A schematic partial view of the nominal analogy “government  $\sim$  door.” Facets in the middle are associated with the target noun on the left metaphorically, and with the source noun on the right literally.

ated (in the appropriate syntactic relationships) with each of those facets with positive PMI.

For example, consider finding candidate source terms for the target term “government.” We first find all facets in linguistic metaphors identified by the system that are associated with “government” with a positive PMI. These include such terms as “better”, “big”, “small”, “divided”, “open”, “closed”, and “limited.” Each of these facets also are associated with other nouns in non-metaphorical senses. For instance, the terms “open” and “closed” are associated with the words “door”, “table”, “sky”, “arms”, and “house”.

**Identification.** To identify which pairs of candidate source terms in  $S_C$  and target terms are likely nominal analogies, the system computes a measurement of the similarity of the non-metaphorical facets of each candidate with the metaphorical facets of each target.

To begin, define for each pair of a facet  $f_i$  and noun (source or target)  $n_i$ , its *association score*  $a_{ij}$  to be its PMI, and the pair’s *metaphoricity score*,  $m_{ij}$  such that  $m_{ij} = 1$  if the pair is judged by the system to be metaphorical and  $m_{ij} = 0$  if it is judged to be non-metaphorical. Currently, all metaphoricity scores are 1 or 0, but we plan to generalize the scheme to allow different levels of confidence in linguistic metaphor classification.

We then define a *metaphoric facet distribution* (MFT) for facets given target terms by normalizing the product of association and metaphoricity scores:

$$P_M(f_i|n_j) = \frac{a_{ij}m_{ij}}{\sum_{ij} a_{ij}m_{ij}}$$

as well as a *literal facet distribution* (LFT), similarly, as:

$$P_L(f_i|n_j) = \frac{a_{ij}(1 - m_{ij})}{\sum_{ij} a_{ij}(1 - m_{ij})}$$

As noted above, we seek source term / target term pairs such that the metaphorical facets of the target term are likely

<sup>2</sup>See <http://provalisresearch.com/>.

to be literal facets of the source term and vice versa. A natural measure of this tendency is the Jensen-Shannon divergence between the LFT of  $n_s$  and the MFT of  $n_t$ ,

$$D_{JS}(P_L(\cdot|n_s) \parallel P_M(\cdot|n_t))$$

The larger the J-S divergence, the less likely the pair is to be a good nominal analogy.

We then compute a heuristic estimation of the probability that a noun pair is a nominal analogy:

$$P_{NA}(n_s, n_t) = e^{-\alpha(n_s)D_{JS}(P_L(\cdot|n_s) \parallel P_M(\cdot|n_t))^2}$$

where  $\alpha(n_s)$  is the abstraction score of  $n_s$ . Empirically, we have found it useful to multiply the J-S divergence by the abstractness score of the candidate source noun, since good source nouns tend to be more concrete (since metaphors typically map a concrete source domain to an abstract target domain).

The output of this process is a set of potential nominal analogies,  $S_A = \{\langle n_s, n_t \rangle\}$ , each a pair of a source noun and a target noun, with an associated confidence score given by  $P_{NA}$ . We note that at this stage recall is more important than precision, since many nominal analogies will be filtered out when nominal analogies are clustered to find conceptual metaphors.

## Conceptual Metaphors

After computing nominal analogies  $S_A$ , the system seeks a set of plausible conceptual metaphors that parsimoniously covers the most likely ones nominal analogies. A conceptual metaphor (CM) can be represented as a pair  $\langle C_s, C_t \rangle$  of a source concept  $C_s$  and a target concept  $C_t$ .

Which concept pairs most likely represent conceptual metaphors? We consider three key conditions indicating a likely conceptual metaphor: (i) consistency of the concept pair with linguistic metaphors (as identified by the system), (ii) the semantic relationship between the two concepts, and (iii) properties of each of the concepts on their own.

To address consistency with the corpus, consider a concept as a set of nouns, representing the set of possible instances of the given concept. A concept pair  $CP = \langle C_s, C_t \rangle$  thus may be viewed as the set of noun pairs  $NP(CP) = \{\langle n_s, n_t \rangle | n_s \in C_s, n_t \in C_t\}$ , each of which is a potential nominal analogy. The more such pairs are high-scoring NAs, and the higher their scores, the more likely  $CP$  is to be a conceptual metaphor. Furthermore, consider the facets that join these NAs; if the same facets are shared among many of the NAs for  $CP$ , it is more likely to be a conceptual metaphor.

For the second condition, the semantic relationship of the two concepts, we require that the target concept  $C_t$  be more abstract than the source concept  $C_s$ ; the larger the abstractness gap, the more likely  $CP$  is to be a conceptual metaphor.

Finally, the concepts should be semantically coherent, and the source concept should be comparatively concrete.

Currently, to ensure that concepts are semantically coherent we consider only sets of nouns which are all hyponyms

of a single Wordnet synset, by which synset we represent the concept<sup>3</sup>.

The rest of the conditions amount to the following hard and soft constraints:

- $\alpha(C_t) > \alpha(C_s)$ ;
- The larger  $\alpha(C_t) - \alpha(C_s)$ , the better;
- The more possible facets associated with covered nominal analogies, the better;
- The higher  $\sum_{\langle n_s, n_t \rangle \in NP(CP)} P_{NA}(n_s, n_t)$ , the better;
- The lower  $\alpha(C_s)$ , the better; and
- The more elements in  $NP(CP)$ , the better.

We thus define the overall score for each pair of concepts  $\langle C_s, C_t \rangle$  as:

$$\frac{(\alpha(C_t) - \alpha(C_s))\phi(CP) \log \sum_{\langle n_s, n_t \rangle \in NP(CP)} P_{NA}(n_s, n_t)}{\alpha(C_s)^2} \quad (1)$$

where

$$\phi(\langle C_s, C_t \rangle) = \frac{|NA(\langle C_s, C_t \rangle)|}{\max(|C_s|, |C_t|)^2}$$

such that  $NA(\langle C_s, C_t \rangle) = \{\langle n_s, n_t \rangle | P_{NA}(n_s, n_t) > \frac{1}{2}\}$  is the set of all likely nominal analogies for the concept pair. The function  $\phi$  thus gives the fraction of all possible such nominal analogies that are actually found.

The system then retains just the candidate concept pairs that satisfy the hard constraints above and have at least a given number of shared facets (in our experiments 5 worked well), ranking them by the scoring function in equation (1). It then removes any concept pairs subsumed by higher-scoring pairs, where  $\langle C_s, C_t \rangle$  subsumes  $\langle C'_s, C'_t \rangle$  iff:

$C'_s$  is a hyponym of  $C_s$  and  
 $(C'_t \text{ is a hyponym of } C_t \text{ or } C'_t = C_t)$

or

$C'_t$  is a hyponym of  $C_t$  and  
 $(C'_s \text{ is a hyponym of } C_s \text{ or } C'_s = C_s)$

The result is a ranked list of concept pairs, each with a list of facets that connect relevant source terms to target terms. Each pair of concepts constitutes a conceptual metaphor as found by the system, and its distribution of facets is a representation of the conceptual frame, i.e., the aspects of the source concept that are mapped to the target concept.

## Evaluation

### Linguistic Metaphors

We evaluated linguistic metaphor identification based on a selection of metaphors hand-annotated in the Reuters RCV1 dataset (Lewis et al. 2004).

For feasibility of annotation, we considered all expressions of Types 1, 2, and 3 for a small number of target terms, “God”, “governance”, “government”, “father”,

<sup>3</sup>Our approach here is similar to that of (Ahrens, Chung, and Huang 2004), though they rely much more strongly on Wordnet to determine conceptual mappings.

Table 1: Linguistic metaphor detection results for RCV1

| Metaphor Type | % metaphors | Precision | Recall |
|---------------|-------------|-----------|--------|
| Type 1        | 66.7%       | 83.9%     | 97.5%  |
| Type 2        | 51.4%       | 76.1%     | 82%    |
| Type 3        | 24.2%       | 54.4%     | 43.5%  |

and “mother”. From each of the 342,000 texts we identified the first sentence (if any) including one of the target words. Each such sentence was parsed by the Stanford Dependency Parser (De Marneffe, MacCartney, and Manning 2006), from which were automatically extracted expressions with the syntactic forms of Types 1, 2, and 3.

Four research assistants annotated each of these expressions to determine whether the key word in the expression is used in its most salient embodied/concrete sense or in a secondary, metaphorical sense. For instance, in the case of “bitter lemon” the first embodied definition of the adjective “bitter” is “Having an acrid taste.” Hence in this case, it is used literally. When asked to judge, on the other hand, whether the phrase “bitter relationship” is literal or metaphorical, that basic meaning of “bitter” is used to make a decision; as “relationship” cannot have an acrid taste, the expression is judged as metaphorical.

Annotators were given the entire sentence in which each expression appears, with the target word and related lexical units visually marked. Inter-annotator agreements, measured by Cronbach’s Alpha, were 0.78, 0.80 and 0.82 for expressions of Type 1, 2, and 3 respectively. The majority decision was used as the label for the expression in the testing corpus.

Results (precision and recall) are given in Table 1. Both precision and recall are significantly above baseline for all metaphor types. We do note that both precision and recall are notably lower for Type 3 metaphors, partly due to their much lower frequency in the data.

## Conceptual Metaphors

To evaluate the validity of nominal analogies and conceptual metaphors found by our system, we ran the system to generate a set of linguistic metaphors, nominal analogies, and conceptual metaphors starting with slightly larger but more focused set of target terms related to the concept of

Table 2: Top conceptual metaphors generated by the system for governance-related targets.

|        |       |            |
|--------|-------|------------|
| IDEA   | is a  | GARMENT    |
| POWER  | is a  | BODY PART  |
| POWER  | is a  | CONTAINER  |
| POWER  | is a  | CONTESTANT |
| POWER  | is an | ARMY UNIT  |
| POWER  | is a  | BUILDING   |
| POWER  | is a  | CHEMICAL   |
| REGIME | is a  | CONTAINER  |
| REGIME | is a  | ROOM       |

Table 3: Conceptual metaphor validation results. See text for explanation.

| Condition        | CMs validated by majority | Trials validated |
|------------------|---------------------------|------------------|
| All CMs          | 0.65                      | 0.60             |
| High-ranking CMs | 0.71                      | 0.64             |
| Foils            | 0.53                      | 0.60             |

Table 4: Framing results. Implied LMs are facet/target pairs from a single CM while foils are taken from different CMs.

| Condition   | Expressions validated by majority | Trials validated |
|-------------|-----------------------------------|------------------|
| Implied LMs | 0.77                              | 0.76             |
| Foils       | 0.66                              | 0.69             |

“governance”. These terms were: “governance,” “government,” “authority,” “power,” “administration,” “administrator,” “politics,” “leader,” and “regime.” The highest ranked conceptual metaphors found by the system are in Table 2.

For evaluation, we recruited 21 naive subjects, 11 female, 10 male, with native fluency in American English. Subject age ranged from 19 to 50, averaging 25.5 years old. Education ranged from subjects currently in college to possessing master’s degrees; just over half of all subjects (12) had bachelor’s degrees.

Each subject was first given an explanation of what conceptual metaphors are in lay terms. The subject was then given a sequence of conceptual metaphors and asked whether each conceptual metaphor reflects a valid and meaningful conceptual metaphor in the subject’s own language. Some of the conceptual metaphors presented were those generated by the system, and others were foils, generated by randomly pairing source and target concepts from the system-generated results.

Our evaluation metrics were the fraction of CMs validated by a majority of subjects and the fraction of the total number of trials that were validated by subjects. We also separately evaluated validation for the highest-scoring CMs, which were the top 19 out of the 32 total CMs found.

Results are given in Table 3. We note first of all that a clear majority of CMs produced by the system were judged to be meaningful by most subjects. Furthermore, CMs produced by the system are more likely to be validated by the majority of subjects than foils, which supports the hypothesis that system-produced CMs are meaningful; higher ranking CMs were more validated yet.

To evaluate the system’s framing of its conceptual metaphors, subjects were presented with a succession of word pairs and asked to judge the meaningfulness of each expression as a metaphor. Each primary stimulus represented a linguistic metaphor implied by a conceptual metaphor found by the system, comprising a pair of a facet word and a target word taken from the conceptual metaphor. (Note that such pairs are not necessarily linguistic metaphors

found originally in the data.) In addition, we included a number of foils, which were based on randomly matching facets from one conceptual metaphor with target terms from different conceptual metaphors. We constructed foils in this way so as to give a meaningful baseline, however, this method also tends to increase the rate at which subjects will accept the foils.

Results for framing are given in Table 4. A large majority of CM framing statements were judged valid by most subjects, supporting the validity of the systems framing of conceptual metaphors.

### Error analysis

To better understand what is needed to improve the system, we consider selected conceptual metaphors proposed by the system and their component framing statements.

Consider first the proposed conceptual metaphor AN IDEA IS A GARMENT. The framing statements given for it are as follows:

- An idea can be matching and a garment can be matching
- An idea can be light and a garment can be light
- An idea can be tight and a garment can be tight
- An idea can be blue and a garment can be blue
- An idea can be old and a garment can be old
- An idea can be net and a garment can be net
- An idea can be pink and a garment can be pink

While the notion of an idea being like a garment may be plausible, with someone being able to put the idea on (consider it true) or take it off (dismiss it), etc., many of the framing statements are not clearly relevant, and a number seem to be simply mistaken (e.g., “old idea” is likely literal).

We see a similar pattern for A REGIME IS A ROOM. The framing statements the system gives for this conceptual metaphor are as follows:

- A regime can be cold and a room can be cold
- A regime can be old and a room can be old
- A regime can be supported and a room can be supported
- A regime can be shaped and a room can be shaped
- A regime can be burning and a room can be burning
- A regime can be new and a room can be new
- A regime can be austere and a room can be austere

Again, the notion of thinking of a regime as a room is cogent, with the members of the regime in the room, and people can enter or leave, or be locked out, and so forth. However, not all of these framing statements make sense. One issue is that some of the facets are not specific to rooms per se, but apply to many physical objects (old, shape, burning, new).

The two main (interrelated) difficulties seem to be (a) finding the right level of abstraction for the concepts in the conceptual metaphor, as related to the framing statements, and (b) finding good framing statements.

One way to improve the system would be to first use the current system to find proposed CMs and associated framing statements, and then refine both the concepts in the CMs and the associated framing statements. Refining the framing

statements would start by finding additional facets related to the nominal analogies for each conceptual metaphor. For A REGIME IS A ROOM, for example, it would be nice to see facets such as “encompassing”, “inescapable”, or “affecting”. In addition we can cluster some of the facets such as “burning/cold” into larger categories such as “temperature”, which give clearer pictures of the different aspects of the source concept that are being mapped.

### Conclusions

We have described a system which, starting from corpus data, finds linguistic metaphors and clusters them to find conceptual metaphors. The system uses minimal background knowledge, consisting only of lexical semantics in the form of Wordnet and Wiktionary, and generic syntactic parsing and part-of-speech tagging. Even with this fairly knowledge-poor approach, experiments show the system to effectively find novel linguistic and conceptual metaphors.

Our future work will focus on two main goals—improving overall effectiveness of the system, and reducing our already low reliance on background knowledge.

One key area for improvement is in the underlying LM identification algorithms, by incorporating more linguistic cues and improving the underlying syntactic processing. Developing a good measure of system confidence in its classifications will also aid higher-level processing. More fundamentally, we will explore implementation of an iterative EM-like process to refine estimations at all three levels of metaphor representation. Hypothesized conceptual metaphors can help refine decisions about what expressions are linguistic metaphors, and vice versa. We will also apply clustering to the facets within conceptual metaphors to improve framing by finding the key aspects which are mapped between source and target domains.

To reduce dependence on background knowledge, we plan to use unsupervised clustering methods to find semantically coherent word groups based on corpus statistics, such as latent Dirichlet allocation (Blei, Ng, and Jordan 2003) or spectral clustering (Brew and Schulte im Walde 2002), which will reduce or eliminate our reliance on Wordnet. To eliminate our need for Wiktionary (or a similar resource) to determine polysemy, we will apply statistical techniques such as in (Sproat and van Santen 1998) to estimate the level of polysemy of different words from corpus data.

Finally, to explore application of these methods to languages where corpus data may be limited, we will examine methods for triangulating multiple partial data sources in a given language/culture (such as combining information from Dari, Classical Persian, and Quranic materials with Modern Iranian Persian) to improve performance where quantities of labeled or parsed training data may be limited.

### Acknowledgements

This work is supported by the Intelligence Advanced Research Projects Activity (IARPA) via Department of Defense US Army Research Laboratory contract number W911NF-12-C-0021. The U.S. Government is authorized to reproduce and distribute reprints for Governmental purposes notwithstanding any copyright annotation thereon.

**Disclaimer:** The views and conclusions contained herein are those of the authors and should not be interpreted as necessarily representing the official policies or endorsements, either expressed or implied, of IARPA, DoD/ARL, or the U.S. Government.

## References

- Ahrens, K.; Chung, S.; and Huang, C. 2004. From lexical semantics to conceptual metaphors: Mapping principle verification with Wordnet and SUMO. In *Recent Advancement in Chinese Lexical Semantics: Proceedings of 5th Chinese Lexical Semantics Workshop (CLSW-5)*, Singapore: COL-IPS, 99–106.
- Allbritton, D.; McKoon, G.; and Gerrig, R. 1995. Metaphor-based schemas and text representations: Making connections through conceptual metaphors. *Journal of Experimental Psychology: Learning, Memory, and Cognition* 21(3):612.
- Assaf, D.; Neuman, Y.; Cohen, Y.; Argamon, S.; Howard, N.; Last, M.; Frieder, O.; and Koppel, M. 2013. Why “dark thoughts” aren’t really dark: A novel algorithm for metaphor identification. In *Proc. IEEE Symposium Series on Computational Intelligence 2013*.
- Birke, J., and Sarkar, A. 2006. A clustering approach for the nearly unsupervised recognition of nonliteral language. In *Proceedings of EACL*, volume 6, 329–336.
- Blei, D.; Ng, A.; and Jordan, M. 2003. Latent dirichlet allocation. *the Journal of machine Learning research* 3:993–1022.
- Boers, F. 2003. Applied linguistics perspectives on cross-cultural variation in conceptual metaphor. *Metaphor and Symbol* 18(4):231–238.
- Brew, C., and Schulte im Walde, S. 2002. Spectral clustering for german verbs. In *Proceedings of the ACL-02 conference on Empirical methods in natural language processing-Volume 10*, 117–124. Association for Computational Linguistics.
- De Marneffe, M.; MacCartney, B.; and Manning, C. 2006. Generating typed dependency parses from phrase structure parses. In *Proceedings of LREC*, volume 6, 449–454.
- Fass, D., and Wilks, Y. 1983. Preference semantics, ill-formedness, and metaphor. *Computational Linguistics* 9(3-4):178–187.
- Freeman, M. 2003. Poetry and the scope of metaphor: Toward a cognitive theory of literature. *Metaphor and Metonymy at the Crossroads: A Cognitive Perspective* 30:253.
- Gibbs Jr, R. 1992. Categorization and metaphor understanding. *Psychological review* 99(3):572.
- Goguen, J., and Harrell, D. 2010. Style: A computational and conceptual blending-based approach. In Argamon, S.; Burns, K.; and Dubnov, S., eds., *The Structure of Style*. Springer-Verlag: Berlin, Heidelberg. 291–316.
- Johnson, M.; Lakoff, G.; et al. 2002. Why cognitive linguistics requires embodied realism. *Cognitive linguistics* 13(3):245–264.
- Krishnakumaran, S., and Zhu, X. 2007. Hunting elusive metaphors using lexical resources. In *Proceedings of the Workshop on Computational approaches to Figurative Language*, 13–20. Association for Computational Linguistics.
- Lakoff, G., and Johnson, M. 1980. *Metaphors we live by*. University of Chicago Press: London.
- Lewis, D.; Yang, Y.; Rose, T.; and Li, F. 2004. RCV1: A new benchmark collection for text categorization research. *The Journal of Machine Learning Research* 5:361–397.
- Mason, Z. 2004. CorMet: A computational, corpus-based conventional metaphor extraction system. *Computational Linguistics* 30(1):23–44.
- Nayak, S., and Mukerjee, A. 2012. A grounded cognitive model for metaphor acquisition. In *Twenty-Sixth AAAI Conference on Artificial Intelligence*.
- Pantel, P., and Lin, D. 2002. Discovering word senses from text. In *Proceedings of the eighth ACM SIGKDD international conference on Knowledge discovery and data mining*, 613–619. ACM.
- Rohrer, T. 2007. The cognitive science of metaphor from philosophy to neuropsychology. *Neuropsychological Trends* 2:31–40.
- Shutova, E.; Teufel, S.; and Korhonen, A. 2012. Statistical metaphor processing. *Computational Linguistics* 39(2):1–92.
- Sproat, R., and van Santen, J. 1998. Automatic ambiguity detection. In *Proceedings ICSLP*.
- Thibodeau, P., and Boroditsky, L. 2011. Metaphors we think with: The role of metaphor in reasoning. *PLoS One* 6(2):e16782.
- Turney, P.; Neuman, Y.; Assaf, D.; and Cohen, Y. 2011. Literal and metaphorical sense identification through concrete and abstract context. In *Proceedings of the 2011 Conference on the Empirical Methods in Natural Language Processing*, 680–690.
- Wickman, S.; Daniels, M.; White, L.; and Fesmire, S. 1999. A primer in conceptual metaphor for counselors. *Journal of Counseling & Development* 77(4):389–394.