

A context-sensitive framework for lexical ontologies

TONY VEALE and YANFEN HAO

School of Computer Science and Informatics, University College Dublin, Ireland
e-mail: tony.veale@ucd.ie, yanfen.hao@ucd.ie

Abstract

Human categorization is neither a binary nor a context-free process. Rather, the criteria that govern the use and recognition of certain concepts may be satisfied to different degrees in different contexts. In light of this reality, the idealized, static structure of a lexical-ontology like WordNet appears both excessively rigid and unduly fragile when faced with real texts that draw upon different contexts to communicate different world-views. In this paper we describe a syntagmatic, corpus-based approach to redefining the concepts of a lexical-ontology like WordNet in a functional, gradable and context-sensitive fashion. We describe how the most diagnostic properties of concepts, on which these functional definitions are based, can be automatically acquired from the Web, and demonstrate how these properties are more predictive of how concepts are actually used and perceived than properties derived from other sources (such as WordNet itself).

1 Introduction

Different contexts encourage different ways of speaking. This variation comprises more than differences in terminology and vocabulary—perhaps the most obvious reason for designing context-sensitive ontologies—but more insidiously comprises subtle differences in how common terms and their underlying concepts are employed. Indeed, this variation in how ontological concepts are expressed linguistically can give rise to context-defining shibboleths; the plural term ‘ontologies’, for instance, is more likely to identify its user as a computer scientist than as a philosopher (Guarino, 1998).

An ontology is a formalized and highly structured system of concepts in which the meanings of semantic structures can be grounded. Guarino (1998) notes that such ‘an engineering artifact, constituted by a specific vocabulary’ is used to describe ‘a certain reality’, and so one expects this system to be fairly stable if it is to serve as a reliable bedrock of meaning. However, concepts are no more than perspectives on the world, and these perspectives can change from context to context. For instance, when speaking of man-made objects, one can distinguish between the perspectives of designed functionality and ad-hoc functionality (see Barsalou, 1983). Banks do not design credit-cards so that they may be surreptitiously used to open locks, but in the context of certain movies and genres of fiction, this is an apparently frequent usage. Likewise, lamp-stands are not designed to be used as blunt instruments, or dinner plates as projectiles, yet these can be contextually appropriate functions for such objects. Clearly then, the categorization of a concept depends not on the intrinsic classification of the concept as defined *a priori* in an ontology (though this will obviously be an important factor), but on how the concept is perceived in a particular context relative to a particular goal.

Once we accept that categorization is sensitive to context, all cognitive decisions that follow from categorization—such as the perception of similarity between concepts—become context-sensitive also. We see this effect in experiments performed by Morris (2006), which reveal that, in the context of an article about the effects of movies on suggestible teenagers, subjects reported a stronger semantic relationship between the terms ‘sex’, ‘drinking’ and ‘drag-racing’. The context in question served to highlight the danger inherent in each of these activities, prompting the subjects to lump these ideas together under an ad hoc and highly context-sensitive concept of ‘dangerous behaviors’ (see also Barsalou, 1983; Lakoff, 1987; Budanitsky & Hirst, 2006). This concept is context-sensitive insofar as the specific behaviors that comprise it are contextually determined. Smoking, for instance, is often considered a dangerous behavior in a medical context, but hardly seems to meet the diagnostic criteria for this concept when viewed from the contexts of bomb-disposal, undercover police-work or high-wire acrobatics.

A high-level division of labor between ontologies and contexts thus suggests itself: ontologies provide intensional definitions for those concepts that are meaningful and stable across many contexts, while contexts provide local evidence that these concepts can be specialized in different ways and to different degrees in specific frames of reference. As such, we see contexts and ontologies as comprising two complementary pieces of the larger knowledge-representation puzzle, a view consistent with that of Giunchiglia (1993), Ushold (2000) and Segev and Gal (2005).

This work is thus a computational exploration of the common intuition that language use reflects conceptual structure. As noted by De Leenheer and de Moor (2005), ontologies are, in the end, lexical representations of concepts, so we should expect that the effects of context on language use will closely reflect the effects of context on ontological structure. An understanding of the linguistic effects of context, as expressed through syntagmatic patterns of word usage, should lead therefore to the design of more flexible ontologies that naturally adapt to their contexts of use. Given this linguistic bias, we focus our attention in this paper to the class of ontology known as *lexical ontologies*. These are ontologies like WordNet (Fellbaum, 1998), HowNet (Dong and Dong, 2006) and the Generative Lexicon (Pustejovsky, 1995) that aim to serve as a formal ontological basis for a lexical semantics by combining knowledge of words with knowledge of the world. Since many words and word senses are inherently suited to some contexts of use more than others, the problem of context is one of particular importance to the proper working of such ontologies. Our focus on WordNet-like ontologies, lightweight as they are, is largely motivated by the fact that these ontologies have hitherto ignored the role of context in their design.

We begin in Section 2 by considering the interlocking roles of contexts and ontologies. We view each as a complementary kind of knowledge-representation, the primary distinction being one of stability: an ontology is a formal representation of concepts and their inter-relationships that is stable across different frames of reference, while a context is a changeable set of subsumption mappings from the core concepts of an ontology to the specific concepts employed in a given reference frame. The problem of ‘contextualizing an ontology’ (Bouquet *et al.*, 2003; Obrst & Nichols, 2005) is thus seen as one of local categorization, in which concepts are locally imbued with the properties needed to allow them to be subsumed by the context-independent definitions of a base ontology. In Section 3 we describe how this contextualization can be computationally realized for a lexical ontology, not by modeling contexts directly and explicitly, but by using representative text corpora as sources of indicative linguistic behavior. These corpora yield local knowledge in the form of syntagmatic patterns, whereby, for instance, the patterns ‘X-addicted’, ‘X-addled’ and ‘X-crazed’ suggest that the entity X is a kind of drug, the pattern ‘X-wielding’ suggests that X is a kind of weapon, and ‘barrage of X’ suggests that X is a kind of projectile. In Section 4 we describe how stable ontological definitions can be automatically constructed from syntagmatic associations that are distilled from patterns of textual data on the Web, in an approach that extends that of Almuhareb and Poesio (2005). We evaluate the reliability of these efforts in Section 5, before concluding with some final remarks in Section 6.

2 Context and ontologies

As with any modeling task, ontological description is as much a matter of representational choice as it is one of representational verisimilitude. An ontology (qua engineering artifact) does not capture ‘objective reality’, or even a small portion thereof, but merely, as Guarino (1998) is careful to point out, ‘a certain reality’. While a plurality of ‘realities’ may be as confounding as a plurality of ‘ontologies’, the term ‘reality’ is nonetheless appropriate insofar as an ontology is designed to encode a common world-view that is shared by multiple (if not all) parties (see Guarino, 1998; Patel-Schneider *et al.*, 2003). The representational choice inherent in ontological design reflects the wide range of perspectives, biases, levels of detail and subject-oriented divisions that are available (consciously or otherwise) to the knowledge engineer. Regardless of the label one uses to motivate these choices, the notion of ‘context’ seems to play a key role in defining the particular realities of different ontologies (e.g. see Bouquet *et al.*, 2003).

In distinguishing between ontologies and contexts, the former is often conceived as an inherently stable world-view, while the latter is conceived as an altogether more fluid and changeable frame of reference that is mapped to the former, or in which the former is applied. For instance, Obrst and Nichols (2005) conceive of contexts as user-dependent and task-dependent views on an underlying ontology, while Bouquet *et al.* (2003) similarly conceive of contexts as local and private perspectives onto a shared encoding of a domain. One role of context is to provide an additional layer of knowledge that better informs how an ontology can be used in a given set of circumstances. In particular, Obrst and Nichols suggest that context can serve to annotate or label the shared concepts and relations of an underlying ontology, for example, to express the security-level and provenance of those elements. Segev and Gal (2005) see ontologies and contexts as complementary views of a domain of discourse, where the ontology—a carefully engineered domain model—serves as a unified, global knowledge-representation onto which contexts—partial and task-specific or user-specific views—are mapped. These latter authors see the process of reconciling local contexts to global ontologies as a core part of applications as diverse as e-mail routing, opinion analysis and topic clustering.

Another role of context is to separate those parts of a knowledge-representation that are mutually inconsistent into different, but complementary, perspectives, each perhaps owned by a different agent. As used in the Cyc ontology (Lenat & Guha, 1990), these perspectives are called microtheories; their propositional content remains local and private unless explicitly inherited by other microtheories or made visible through a process of lifting (Guha, 1991). For instance, the concept Sherlock-Holmes can be ontologized as a kind of fictional character, and thus, a kind of mental product, or it might be ontologized as a kind of detective. WordNet opts for the former course rather than the latter, thus sacrificing the ability to reason about Holmes as if he were an actual detective, or even an actual person. In an ideal ontology, both ontological perspectives would be made available for reasoning purposes, perhaps by representing each in a separate microtheory, or by representing each in different ontologies and providing a detailed system of mappings between each (e.g. as in Bouquet *et al.*, 2003).

Each of these apparent roles sees context as a means of partitioning ontological content into alternate views of reality. Indeed, the microtheory labels used by Cyc, such as HealthMt, HistoryMt and so on, can be seen as annotations on propositions that allow Cyc’s inference processes to selectively include or exclude large swathes of the ontology’s content in a given reasoning process. In this vein, another related role of context is to provide a bridge between the stable definitions of an ontology and the contingent facts of a particular world-view. For instance, an ontology of chemical substances may be agnostic with respect to how those substances are used, so that the same substance might be categorized as a medicine in one context, an illegal drug in another, and a poison in yet another. It is this division of labor between ontologies and contexts that interests us most in this current work: how can we create ontologies as collections of stable definitions that apply in all contexts, yet which are realized differently, by different concepts and to different degrees, in specific contexts? Given the significant design and engineering efforts that are

employed in the construction of well-formed ontologies (e.g. see Gangemi *et al.*, 2001), this division of labor should be a clean one, so that the base ontology only posits relationships that are safe in all contexts, and each context only posits relationships that complement, rather than contradict, those of the base ontology.

This division of labor requires a solution to two related computational problems: how do we acquire and represent the stable concept definitions that comprise the ontology; and how do we acquire the local contextual distinctions that cause these definitions to be instantiated by different entities in different frames of reference? Almuhareb and Poesio (2005) describe a Web-based approach to acquiring the property structure of concepts via text analysis of Internet content, as indexed by a search engine like Google. Their approach indicates how both stable concept definitions and contingent realizations of those definitions can be inferred from simple processes of text analysis. Almuhareb and Poesio use highly diagnostic search queries such as ‘*the * of a|an|the C is|was*’ to identify property values for a given concept C in Web texts. By acquiring properties (such as the fact that beverages have an associated temperature and strength) as opposed to property values (such as ‘hot’ and ‘cold’ for coffee), these authors acquire a general frame structure for each concept that can be instantiated differently in different contexts.

We too employ a large-scale analysis of Web-text to acquire stable concept definitions that will transcend context boundaries. However, we do not currently focus on the acquisition of property structure, but on prototypical property values. While Almuhareb and Poesio (2005) demonstrate that generic properties such as Temperature and Color are more revealing about conceptual structure than specific values such as ‘hot’ and ‘red’ (since these values can change without affecting the nature of the concept), we do not collect arbitrary contingent attributions (such as the fact that coffee *can* be cold) but highly diagnostic and concept-defining attributions (e.g. that espresso *should* be strong, that surgeons *should* be delicate, that gurus *should* be wise). To identify which property values are truly central to the consensus definition of a concept, we use the highly specific comparison frame ‘as * as a an C’ to collect similes involving a given concept from the Web. Once acquired and validated, we articulate the prototypical properties for a given concept as a set of logical constraints that serves as a functional definition for that concept. The form of these functions is presented in Section 3, while the Web-based acquisition of each function’s content is described in Section 4.

2.1 *Texts and con-texts*

De Leenheer and de Moor (2005) see a context as a mapping from a set of lexical (and potentially ambiguous) labels to a set of language-neutral concept identifiers. In this view, the same words can denote specializations of different concepts in different contexts. For instance, ‘cocaine’ can denote a kind of local anesthetic in a medical context and a kind of illegal drug in a law-enforcement context. This is more than a matter of lexical ambiguity; the same sense is intended in each context (i.e. the same substance) but a different ontological categorization is implied in each. Combining a mapping-theoretic view of context (e.g. Bouquet *et al.*, 2003; Obrst & Nichols, 2005), with the lexical emphasis offered by De Leenheer and de Moor, it is possible to obtain much of the reasoning benefits of a context without an explicit logical representation of context. For example, the similarity between chocolate and a narcotic-like heroin will, in most contexts, simply reflect the ontological fact that both are kinds of substances; certainly, taxonomic measures of similarity as discussed in Budanitsky and Hirst (2006) would capture little more than this basic categorization. However, in a context in which the addictive properties of chocolate are highly salient (in an online dieting forum, for instance), chocolate is more likely to be categorized as a drug and thus be considered more similar to heroin. Look, for instance, at the similar ways in which these words can be used: one can be ‘chocolate-crazed’ or ‘chocolate-addicted’ and suffer ‘chocolate-induced’ symptoms (each of these uses is to be found in chocolate-related Wikipedia articles). In a context that gives rise to these syntagmatic patterns, it is unsurprising that chocolate should appear altogether more similar to a harmful narcotic.

A given corpus may employ syntagmatic patterns which reflect the fact that the corresponding context views chocolate as a kind of drug, or military robots as soldiers, or certain kinds of criminals as predators. By augmenting a base ontology with these categorizations, the ontology may become sufficiently contextualized to reason fluently in this context. The model of corpus-based ontology augmentation we describe in this paper is consistent with, and complementary to, the Theory of Norms and Exploitations (TNE) proposed by Hanks (2004), in which corpus analysis is used to identify both the syntagmatic *norms* of word usage (i.e. highly conventional and normative uses) and meaning-coercing *exploitations* of these norms.

This current work attempts a synthesis of ideas from the fields of cognitive science and cognitive linguistics (as exemplified by the work of Lakoff, 1987), corpus linguistics (as exemplified by the work of Hanks, 2004) and lexical semantics (as exemplified by the work of Pustejovsky, 1995). Throughout this paper, the term ‘concept’ is used in the set-theoretic sense employed by the ontology and description-logic literature (e.g. see Welty & Jenkins, 1999), which in turn corresponds to the use of the word ‘category’ in the cognitive science literature (e.g. see Barsalou, 1983). As such, we view a concept as a hierarchically positioned cognitive structure to which properties can be ascribed and in which membership can be asserted. As in a description-logic, membership in and subsumption by a concept can be stated explicitly, or computed as necessary (Welty & Jenkins, 1999) by an application of the logical criteria that define the concept to the specific properties of a putative member. These logical criteria that define a concept form part of the ontology proper, while additional properties of a concept may be provided by a specific context. For instance, an ontology may stipulate that the concept *Insurgent* is subsumed by the concept *Person*, so that all insurgents, in any context, must be human (e.g. see Gangemi *et al.*, 2001). Nonetheless, different contexts may ascribe different additional properties to *Insurgent*. In one context, insurgents may be seen as craven and evil, supporting a local categorization of *Insurgent* as a kind of *Terrorist* or *Criminal*; in another, insurgents may be seen as upright and noble, supporting a local categorization of *Insurgent* as a kind of *Champion* or *Defender*.

2.2 Deriving context-specific insights from text

A syntagmatic approach to deriving ontological insights from text is hardly novel. Hearst (1992) describes a syntagmatic technique for identifying hyponymy relations in free text by using frequently occurring genre-crossing patterns like NP_0 such as $NP_1, NP_1, \dots, NP_n$. Like the approach of Charniak and Berland (1999), Hearst’s patterns seek out explicit illustrations of inter-concept subsumption, as in the phrase ‘drugs like Prozac, Zoloft and Paxil’. Such techniques are useful because contexts frequently introduce new terms that are locally meaningful. Nonetheless, such techniques do not reveal the subtle and shifting nuances of concept usage that underpin a particular context. These differences are implicit precisely because the existence of a context presupposes the existence of a shared body of knowledge and a common world-view. Context-specific corpora only reveal this shared knowledge indirectly, insofar as it is presupposed in the way that language is used.

Closer to the current approach is that of Cimiano, Hotho and Staab (2005), who do not look for unambiguous ‘silver bullet’ patterns in a text, but who instead characterize each lexical term and its underlying concept according to the syntagmatic patterns in which it participates. These patterns include the use of the term as the subject, object or prepositional complement of a verb. The key intuition, expressed also in Weeds and Weir (2005), is that terms with similar distribution patterns will denote ideas that are themselves similar. Cimiano *et al.* exploit the phrasal dependencies of a term as features of that term that can be used, through a process of conceptual clustering called Formal Concept Analysis (FCA), as introduced by Ganter and Wille (1999), to determine subsumption relations between different terms. At no point are explicit expressions of these relations sought in a text. Rather, from a tabular mapping of terms to their syntagmatic properties (called a Formal Context), FCA is used to infer these relations by determining which terms possess property descriptions that are a superset or subset of other descriptions.

These attributive descriptions serve a dual purpose: they allow an extensional comparison of different concepts to determine which is more general and inclusive; but they also serve as an explicit intensional representation of the conceptual terms that are ontologized. For example, the term ‘bike’ is *rideable*, *bookable* and *rentable* because of its use as an object with the verbs ‘ride’, ‘book’ and ‘rent’, so the set *rideable*, *bookable*, *rentable* provides an intensional picture of Bike.

Segev and Gal (2005) likewise see contexts as domain perspectives that are discernable from representative texts. Since these authors see contexts as partial, user-specific domain views, they argue that contexts can be conveyed by texts as small as e-mail messages, and describe a means of extracting contexts from frequency-weighted bags of words. By discerning contexts in texts in this way, contexts can be mapped to the appropriate ontological concepts, so that the category, opinion or topic of the text can also be discerned. In mapping contexts to ontology concepts, the texts themselves are thus placed into appropriate ontological buckets.

In this current work, we do not aim to extract contexts (as first-order objects or formal representations of such) from texts, rather we see texts as a useful proxy for the knowledge provided by a context. As such, our goal is to employ the syntagmatic features of a text to infer the context-sensitive categorizations that are communicated by the text, and to augment the classification structures of the base ontology with these additional context-specific viewpoints. To achieve this requires an understanding of the diagnostic properties of concepts on which categorization in those concepts is based, as well as an understanding of how these properties are communicated in a text by particular word choices and syntagmatic patterns.

3 Conceptual norms and contextual exploitations

In this section we present a functional framework for defining concepts in terms of the linguistic cues used to signal their usage in text. These cues or expectations are articulated as syntagmatic norms (Hanks, 2004) that capture, for example, the most diagnostic adjectival modifiers that contribute to a lexical description of the concept, the kinds of verbs for which the concept typically acts as an agent or a patient, the kind of group terms (like ‘army’, ‘herd’, ‘flock’, etc.) that are typically used to describe aggregations of the concept, and so on. Each concept is assigned a different functional form that expresses the appropriate syntagmatic expectations. This functional form is not Boolean in construction, but yields a continuous output in the range $0 \dots 1$, where 1 indicates total satisfaction of the syntagmatic expectations (and thus, indirectly, the logical criteria) that comprise the concept definition.

A continuous or fuzzy categorization function allows concept definitions to be viewed as *radial category structures* in the sense of Lakoff (1987). For example, to the extent that a collocation like ‘army of X’ is found in a corpus, the associated context can be said to categorize X as a sub-type of Soldier. Likewise, to the extent that the syntagm ‘X-addicted’ has currency in a corpus, X should be seen as a kind of Drug. Interestingly, some of the most stable and unambiguous syntagmatic patterns are associated with metaphoric conceptualizations. Thus, the syntagmatic schema ‘barrage of X’ identifies X as a projectile, whether X is an arrow, a pointed question or an angry e-mail. The frequency of these patterns in a corpus yields a sliding scale of inter-concept subsumption in the associated context. Thus, something may be more representative of a particular concept in one context (e.g. Chocolate as a Narcotic in a dieting context) than in another.

We begin by supposing a function (*attr* arg_0 , arg_1) that returns a real number in the range $[0 \ 1]$ that reflects the frequency of arg_0 as an adjectival modifier for the noun arg_1 in a corpus. Suppose also a function (*%isa* arg_0 , arg_1) that returns a number in $[0 \ 1]$ reflecting the proportion of senses of arg_0 that are descendants of arg_1 in a base-ontology like WordNet (i.e. *%isa* is not a binary relation but a gradated function, and we mark it with ‘%’ to reflect this quantitative distinction). We can now define the concept Fundamentalism in a functional fashion.

That is, any extreme, violent or radical person or group that is either political or religious deserves to be categorized as a fundamentalist. The extent to which this person or group is a fundamentalist depends entirely on the contextual evidence for these criteria, as captured by the

use of *attr*. The precise workings of *attr* can be implemented in a number of ways, using any of a variety of corpus-based distributional similarity metrics, such as Dice’s coefficient or the Jaccard measure (see Lee, 1999; Weeds & Weir, 2005). Whatever measure is used, it must either return a value in the range [0 1] or be scaled to do so, so that each concept-defining function like that of *Fundamentalist* will similarly return a value in the [0 1] range. The value returned by each conceptual function thus corresponds to a context-sensitive degree of categorization in the corresponding radial category (Lakoff, 1987). For instance, in texts that are representative of a left-leaning liberal world-view, (*Fundamentalist evangelical*) should return a value closer to 1.0 than in texts with a right-leaning conservative bias. If (*Fundamentalist evangelical*) returns a value of 0.61 for a given corpus, one can consider evangelic believers to be highly representative, even exemplary, instances of Fundamentalists in the associated context.

The programmatic, LISP-like structure of these conceptual functions explicitly mirrors the logical structure of the concept’s intensional form, inasmuch as each function can be given an obvious logical interpretation. In the example of *Fundamentalist*, note how the mathematical functions *min* and *** (multiplication) are essentially used to encode a fuzzy-logic equivalent of the logical operator *and*, while the function *max* is used to encode a fuzzy-logic equivalent of the logical operator *or*. Conceptual functions can be built using all of the syntagmatic patterns employed by Cimiano *et al.* (2005), with some additions (see Table 1). For instance, the ‘GROUP of NOUN + plural’ pattern employs WordNet to identify group membership descriptions in a corpus, where GROUP is any group-denoting WordNet term (e.g. swarm, army) or group activity (e.g. barrage, invasion, influx). This syntagm is given functional form via the function (*group arg₀ arg₁*), which returns the extent (again in the range [0 1]) to which *arg₁* is described as a member of the group *arg₀* in a given corpus. For instance, using the text of the encyclopaedia Wikipedia as a corpus, and using Dice’s coefficient as a measure of association, we find (*group influx immigrant*) = 0.38 and (*group army mercenary*) = 0.31. The base functions of Table 1 thus serve as the interface between a context-independent intensional description, like that of *Fundamentalist* in Figure 1, and the specific linguistic evidence of a context-discerning corpus.

Some syntagmatic patterns are more directly associated with specific concepts than others. For instance, the pattern ‘mint-flavored’ clearly indicates that Mint is a flavor. Such hyphenated forms can be used to find figurative usage of concepts in context, as in:

(defineDrug(*arg₀*) (hyphenaddict*arg₀*)) %e.g., risk-addicted
(defineCausal-Agent(*arg₀*) (hypheninduce*arg₀*)) %e.g., drug-induced

A conceptual function may need to marshal different kinds of syntagmatic evidence to yield an overall categorization score, and the basic function *combine* allows us to combine this variety of contextual evidence into a score in the [0 1] range. If *e₀*, *e₁*, etc., are the scores associated with various pieces of evidence (as returned by the functions of Table 1), then *combine* adds these scores to yield another in the [0 1] range as follows:

$$\begin{aligned} (\text{combine } e_0 e_1) &= e_0 + e_1(1 - e_0) = e_0 + e_1 - e_0 e_1 \\ (\text{combine } e_0 e_1 \dots e_n) &= (\text{combine } e_0 (\text{combine } e_1 \dots e_n)) \end{aligned}$$

Table 1 Basic concept-defining functions and their syntagmatic correspondences

Function	Example	Syntagmatic Pattern(s)	Range
(<i>agent verb₀ noun₀</i>)	(agent kill robot)	“ <i>noun₀ verb₀</i> ”, “... <i>verb₀</i> + past by <i>noun₀</i> ”	[0 1]
(<i>patient verb₀ noun₀</i>)	(patient eat prey)	“ <i>noun₀ verb</i> + past by ...”, “... <i>verb₀ noun₀</i> ”	[0 1]
(<i>attr adj₀ noun₀</i>)	(attr knight brave)	“ <i>adj₀ noun₀</i> ”, “as <i>adj₀</i> as a an <i>noun₀</i> ”	[0 1]
(<i>group noun₀ noun₁</i>)	(group army grunt)	“ <i>noun₀ of noun₁ + plural</i> ”	[0 1]
(<i>of noun₀ noun₁</i>)	(of owner pet)	“ <i>noun₀ of noun₁</i> ”	[0 1]
(<i>hyphen verb₀ noun₀</i>)	(hyphen shape egg)	“ <i>noun₀- verb₀ + past</i> ” (e.g., egg-shaped, bite-sized)	[0 1]

```

(define Fundamentalist (arg0)
  (* (max (%isa arg0 Person)
    (%isa arg0 Group))
    (min
      (max (attr political arg0)
        (attr religious arg0))
      (max (attr extreme arg0)
        (attr violent arg0)
        (attr radical arg0))))
)

```

Figure 1 A functional description of the concept Fundamentalist

```

(define Invader (arg0)
  (combine (* 0.3 (max (%isa arg0 Person) (%isa arg0 Group)))
    (agent invade arg0)
    (attr invasive arg0)
    (group invasion arg0)
    (group influx arg0)
    ≥ 2
  )
)

```

Figure 2 A functional description of the concept Invader

This *combine* function is thus a naïve probabilistic *or* function, one that naively assumes independence among the evidence it combines to generate scores that asymptotically approach 1.0. If a piece of evidence is included multiple times (to reflect a greater diagnostic value), it is counted multiple times, but with a diminishing effect. Consider the use of *combine* in a concept definition for Invader in Figure 2. Note how four types of information are synthesized in this definition: general taxonomic knowledge (via the *%isa* function); adjectival modification (via the *attr* function); subject-verb knowledge (via *agent*); and group membership knowledge (via *group*).

We refer to the final clause of Figure 2, ≥ 2 , as a ‘quantitative cut’: it specifies the number of non-zero pieces of evidence that *combine* must have processed prior to this cut if it is to perform its normal function; if this threshold is not met, then *combine* aborts (i.e. cuts) early and simply returns a 0. Therefore, any concept in a given context that meets two or more of these intensional criteria (e.g. people or groups that invade, non-human invasive organisms that form an influx) is categorized as an Invader to a degree that reflects the linguistic evidence of the corpus. Note how the contribution of WordNet (or whatever ontology underpins the *%isa* function) is here scaled by a small multiplier of 0.3. This prevents the *%isa* clause—which merely serves as a soft taxonomic preference rather than a hard constraint—from making an undue contribution to the overall categorization score.

Consider a conceptual function for the concept Pet which, as formulated in Figure 3, combines several different types of evidence to diagnose ‘pet-hood’. The definition asks the following questions of each potential member: is it a kind of animal? Is it docile or domesticated? Is it cute? Is it something that one can own and care for? For those concepts that appear to meet two or more of these demands in a corpus, this definition can be used to introspectively explain why.

The definition of Figure 3 is also constructed to make animal-ness a soft-preference rather than a hard constraint for pet-hood, since one can conceive of human pets (favored children, slaves) and even artificial pets (toys, robots, etc.). Suppose, in a given context, the above function assigns

```

(define Pet (arg0)
  (combine (* 0.3 (%isa arg0 Animal))
    (max (of owner arg0) (of care arg0))
    (max(attr docile arg0) (attr domesticated arg0))
    (max(attr cute arg0) (attr cuddly arg0))
    ≥ 2 )
)

```

Figure 3 A conceptual function for the highly context-sensitive notion of Pet

a categorization score of 0.12 to the term Iguana. Introspecting over the symbolic structure of the definition, the system can explain why this score was assigned in this context, by pointing out that the associated corpus speaks of iguanas as cuddly or cute with this much frequency, and as docile with that much frequency, and so on. Now suppose a zero-categorization score is given to Piranha. The system can use a similar process to perform a *what-if* analysis, as in a spreadsheet. Looking at the taxonomic placement of Piranha in a base-ontology like WordNet, the system can determine which of the elements in the functional definition (such as animal-ness) are applicable to Piranha. Noting from the corpus that cuddliness, cuteness and docility are collocates of ‘animal’, it can then explain that Piranha is not a Pet because it is seen as neither cute, cuddly or docile in the associated context.

4 Web-based acquisition of conceptual functions

The functions of Figures 1, 2 and 3 make no reference to any kind of context. Rather, they encode diagnostic knowledge of a general character about individual concepts—what one might describe as the conventional wisdom about these concepts. Our definitions of Pet, Invader, Fundamentalist, Drug, etc., are intended to represent quantitative categorization functions for the correspondingly named concepts in a broad-scope lexical-ontology like WordNet. They should thus be seen not as comprising a local ontology of their own, but as additions to a base ontology (see Giunchiglia, 1993; Ushold, 2000). Nonetheless, these functions are inherently context-sensitive, in two crucial respects. First, they encode conventional wisdom about conceptual structure in a flexible manner, not as hard constraints but as soft preferences. In this way, they anticipate that certain contexts may observe certain diagnostic requirements and not others, for example, that invaders are not always human, or that pets may not always be animals.

Conventional wisdom has its own syntagmatic norms of expression. For instance, when one wishes to highlight a specific property in a given concept, it is commonplace to compare that concept to one for which that property is widely agreed to be diagnostic. Comparisons of the form ‘as ADJ as a/an NOUN’ work best when the exemplar that is used (e.g. ‘dry as *sand*’, ‘hot as the *sun*’) is familiar to the target audience and is truly exemplary of the given property in a context-independent manner. That is, such simile-based comparisons work best when they are generally self-evident and not dependent on a private or inaccessible context to give them meaning. By searching the Web for comparisons of this form, we achieve two important ends: we identify the exemplar concepts that are most frequently used as a basis of comparison, which one can expect to be most stable across varying contexts, and which are thus most deserving of representation in a base ontology; and, we identify the most salient properties of those concepts, thereby allowing us to assign to them a corresponding functional form.

As in Almuhareb and Poesio (2005), we use the Google API to find instances of our search patterns on the Web. We use two simile patterns, one in which the wildcard operator * substitutes for the adjectival property (where an exemplar noun is explicitly given), and one in which the wildcard operator substitutes for the noun (while the adjective is given). The first pattern collects salient

adjectival properties for a given noun, while the second collects the most common noun concepts that exemplify a given adjectival property. For purposes of radial category construction, we expect that adjectives which denote an end-point on a sliding scale, such as ‘brave’ (versus ‘cowardly’), ‘hot’ (versus ‘cold’) and ‘rich’ (versus ‘poor’) will be the most commonly used adjectives in comparative phrases, and will yield the most diagnostic properties for categorization. We initially limit our attention then to WordNet adjectives that are defined relative to an antonymous term. For every adjective ADJ on this list, the query ‘as ADJ as *’ is sent to Google and the first 200 snippets returned are scanned to extract different noun bindings (and their relative frequencies) for the wildcard *. The complete set of nouns extracted in this way is then used to drive a second phase of the search, in which the query template ‘as * as a NOUN’ is used to acquire similes that may have lain beyond the 200-snippet horizon of the original search, or that hinge on non-antonymous adjectives that were not included on the original list. Together, both phases collect a wide-ranging series of core samples (of 200 hits each) from across the Web, yielding a set of 74 704 simile instances (of 42 618 unique types) relating 3769 different adjectives to 9286 different nouns.

4.1 Property filtering

The simile frame ‘as ADJ as a NOUN’ is relatively unambiguous as such patterns go, but a non-trivial quantity of unwanted or noisy data is nonetheless retrieved. In some cases, the NOUN value forms part of a larger noun phrase that is not lexicalized in WordNet: it may be the modifier of a compound noun (e.g. ‘bread lover’), or the head of complex noun phrase (such as ‘gang of thieves’ or ‘wound that refuses to heal’). In other cases, the association between ADJ and NOUN is simply too ephemeral or under-specified to function well in the null context of a base ontology. As a general rule, if one must read the original document to make sense of the association, it is rejected. More surprisingly, perhaps, a substantial number of the retrieved similes are ironic, in which the literal meaning of the simile is contrary to the meaning dictated by common sense. For instance, ‘as hairy as a bowling ball’ (found once) is an ironic way of saying ‘as hairless as a bowling ball’ (also found just once). Many of the ironies we found exploit contingent world knowledge, such as ‘as sober as a Kennedy’ and ‘as tanned as an Irishman’.

Given the creativity involved in these constructions, one cannot imagine a reliable automatic filter to safely identify bona-fide similes. For this reason, the filtering task is performed by human judges, who annotated 30 991 of these simile instances (for 12 259 unique adjective/noun pairings) as non-ironic and meaningful in a null context; these similes relate a set of 2635 adjectives to a set of 4061 different nouns. In addition, the judges also annotated 4685 simile instances (of 2798 types) as ironic; these similes relate a set of 936 adjectives to a set of 1417 nouns. Surprisingly, ironic pairings account for over 13% of all annotated simile instances and over 20% of all annotated simile types.

4.2 Linking to WordNet

WordNet is used as a source for the adjectives that drive the simile retrieval process; it is also used to validate the nouns (unitary or multi-word) that are described by these similes. By sense-disambiguating these nouns relative to the noun senses found in WordNet, we can use their associated adjectival properties to assign functional forms to each of these WordNet senses. As such, we automatically construct context-sensitive categorization functions for the most commonly used concepts in the WordNet noun ontology.

Disambiguation is trivial for nouns with just a single sense in WordNet. For nouns with two or more fine-grained senses that are all taxonomically close, such as ‘gladiator’ (two senses: a boxer and a combatant), we consider each sense to be a suitable target. In some cases, the WordNet gloss for a particular sense will literally mention the adjective of the simile, and so this sense is chosen. In all other cases, we employ a strategy of mutual disambiguation to relate the noun vehicle in each simile to a specific sense. Two similes ‘as A_0 as N_1 ’ and ‘as A_0 as N_2 ’ are mutually disambiguating if N_1 and N_2 are synonyms in WordNet, or if some sense of N_1 is a hypernym or

hyponym of some sense of N_2 in WordNet. For instance, the adjective ‘scary’ is used to describe both the noun ‘rattler’ and the noun ‘rattlesnake’ in bona-fide (non-ironic) similes; since these nouns share a sense, we can assume that the intended sense of ‘rattler’ is that of a dangerous snake rather than a child’s toy. Similarly, the adjective ‘brittle’ is used to describe both saltines and crackers, suggesting that it is the bread sense of ‘cracker’ rather than the hacker, firework or hill-billy senses (all in WordNet) that is intended.

These heuristics allow us to automatically disambiguate 10 378 bona-fide simile types (85%), yielding a mapping of 2124 adjectival properties to 3778 different WordNet noun senses. Likewise, 77% (2164) of ironic simile types are disambiguated automatically. A remarkable stability is observed in the alignment of simile nouns to WordNet senses, which suggests that the disambiguation process is consistent and accurate: 100% of the ironic vehicles always denote the same sense, no matter the adjective involved, while 96% of bona-fide vehicles always denote the same sense.

4.3 From similes to membership functions

The above filtering and word-sense disambiguation processes associate the properties *stealthy*, *silent* and *agile* with the person sense of ‘ninja’ (denoted here as *Ninja.0*), leading to the following function shown in Figure 4.

As we cannot know which subset of these properties is sufficient for categorization, we use the quantitative cut ≥ 2 to ensure that more than one property is contextually present to support a categorization as a ninja. The more properties that are present, the higher the resulting categorization score (aggregated via the *combine* operator) will be. The hard constraint (*%isa arg₀ Person.0*) is chosen as the taxonomic constraint for all conceptual functions that represent a specialized kind of person in WordNet, ensuring that the context does not suggest categorizations that undermine the logical core of the ontology.

The nouns most commonly used in similes on the Web will typically provide an even richer set of properties on which to base categorization, as illustrated by the functional form of Snake in Figure 5.

Note that the taxonomic constraints (*%isa arg₀ Person.0*) and (*%isa arg₀ Animal.0*) serve to ensure that the resulting in-context categorizations are broadly literal w.r.t. WordNet. By weakening (i.e. generalizing) or removing these constraints, one could allow for contextually appropriate metaphoric categorizations to be made, for example, that agile animals, stealthy viruses or silent and clandestine organizations might be seen as ninjas, or that cunning and slippery people might be seen as snakes. The distinction between literal and metaphoric categorization in a given context is often blurred, and may, in principle, be impossible to delineate. Is chocolate really an addictive drug in some dieting contexts, or is such a categorization a creative over-use of the word ‘drug’? While this rather vexing question falls outside the scope of the current paper, we note that the framework of conceptual functions described here provides an ideal mechanism for exploring the contextual boundaries of literal and metaphoric categorization in future research.

```
(define Ninja.0 (arg0)
  (* (%isa arg0 Person.0)
    (combine (attr stealthy arg0)
              (attr silent arg0)
              (attr agile arg0)
              ≥2)
  )
)
```

Figure 4 A Web-derived functional form for the concept Ninja

```

(define Snake.0 (arg0)
  (* (%isa arg0 Animal.0)
    (combine (attr cunning arg0) (attr slippery arg0)
              (attr slim arg0) (attr flexible arg0)
              (attr sinuous arg0) (attr crooked arg0)
              (attr deadly arg0) (attr poised arg0)
              ≥2)
  )
)

```

Figure 5 A Web-derived functional form for the animal sense of ‘snake’ (snake.0)

5 Empirical evaluation

In this section we provide empirical support for the two main claims of this paper. The first is the relatively uncontroversial claim that syntagmatic patterns of usage at the textual level reflect distinctions in concept usage at the ontological level (e.g. Cimiano *et al.*, 2005; Hanks, 2006), so that the syntagmatic patterns in a given corpus can be taken to be indicative of categorization patterns in the corresponding context. The second is the more novel claim that Web similes are sufficiently revealing about the diagnostic properties of concepts to allow accurate categorization functions to be constructed for each.

We test the first claim using the HowNet ontology of Dong and Dong (2006). HowNet differs from WordNet in many respects (e.g. the former is bilingual, linking the same definitions to both English and Chinese labels) but the key difference is that HowNet defines the meaning of each word sense via a simple conceptual graph. For instance, HowNet specifies that a Knight is the agent of the activity Fight, while Assassin is the agent of the activity Kill. In addition, it states (in explicit logical terms) that the killing performed by an Assassin has the property *means* = *unlawful*. Each of these logical definitions is hand-crafted, allowing us to test whether the syntagmatic patterns exhibited by a corpus can suggest much the same semantic distinctions as made by an ontologist. For simplicity, we focus here on those concepts in HowNet that are defined as the agent of a given activity, like Knight and Assassin. Using the complete text of Wikipedia as our corpus (2 GB, from a June 2005 download), we find 1626 different nouns that have at least one sense that fills the agent role of a HowNet activity concept. In all, HowNet uses 262 unique verbs, such as *kill*, *buy* and *repair*, to describe those activities. Using Dice’s coefficient (Lee, 1999) to measure the association between each noun and each verb for which the noun is used as an active subject, we find that for 69% of nouns, the highest rating is given to the verb that is used to capture the noun’s meaning in HowNet. Although one could not argue from this result that an ontology could be automatically constructed from simple syntagmatic evidence alone, a result of 69% does strongly suggest that the semantic criteria that guide ontology construction are readily discernable from patterns of language used in a given context.

Our second claim concerns the simile-gathering process of the last section, which, aided by Google’s practice of ranking pages according to popularity, should reveal the most frequently used nouns in comparisons on the Web, and thus, the most useful concepts to functionally describe in a lexical-ontology like WordNet. But the descriptive sufficiency of these functional forms is not guaranteed unless the diagnostic properties employed by each can be shown to be collectively rich enough, and individually salient enough, to predict how each lexical concept is perceived and used by members of a language community. If similes are indeed a good basis for mining the most salient and diagnostic properties of concepts, we should expect the set of properties for each concept to accurately predict how the concept is perceived as a whole. One measurable clue as to how a concept is perceived is its affective rating.

For instance, humans—unlike computers—tend to associate certain positive or negative feelings, or affective values, with particular concepts. Unsavory activities, people and substances

generally possess a negative affect, while pleasant activities and people possess a positive affect. Whissell (1989) reduces the notion of affect to a single numeric dimension, to produce a *dictionary of affect* that associates a numeric value in the range 1.0 (most unpleasant) to 3.0 (most pleasant) with over 8000 words in a range of syntactic categories (including adjectives, verbs and nouns). Therefore, to the extent that the adjectival properties yielded by processing similes paint an accurate picture of each noun concept, we should be able to predict the affective rating of each concept by using a weighted average of the affective ratings of the adjectival properties ascribed to these concepts (i.e. where the affect rating of each adjective contributes to the estimated rating of a noun concept in proportion to its frequency of co-occurrence with that concept in our Web-derived simile data). More specifically, we should expect that ratings estimated via these simile-derived properties to exhibit a higher correlation with the independent ratings of Whissell's dictionary than properties derived from other sources (such as WordNet itself) or from other syntagmatic patterns.

To determine if this is indeed the case, we calculate and compare this correlation between predicted and reported affect-ratings using the following data sources:

- A. Adjectives derived from annotated bona-fide (non-ironic) similes only.
- B. Adjectives derived from all annotated similes (both ironic and non-ironic).
- C. Adjectives derived from ironic similes only.
- D. All adjectives used to modify a given noun in a large corpus (e.g. all possible uses of the function (*attr adj₀ noun₀*) for a corpus). We use 2 GB of text from the online encyclopaedia Wikipedia as our corpus.
- E. Set of 63 935 unique property-of-noun pairings extracted via the shallow-parsing of WordNet glosses; for example, *strong* and *black* are extracted from the gloss for Espresso ('strong black coffee brewed by forcing steam under pressure ...').

Predictions of affective rating were made from each of these data sources and then correlated with the ratings reported in Whissell's dictionary of affect using a two-tailed Pearson test ($p < 0.01$). As expected, property values derived from bona-fide similes only (A) yielded the best correlation (+0.514) while property values derived from ironic similes only (C) yielded the worst (-0.243); a middling correlation coefficient of 0.347 was found for all similes together (B), reflecting the fact that bona-fide similes outnumber ironic similes by a ratio of 4 is to 1. A weaker correlation of 0.15 was found using the corpus-derived adjectival modifiers for each noun (D); while this data provide quite large value sets for each noun, these property values merely reflect the potential rather than the intrinsic properties of each concept and so do not reveal what is most diagnostic about the concept. As also noted by Almuhareb and Poesio (2005), such values reveal very little about the actual structure of a concept. Those authors address this problem by instead seeking to mine property types (such as Temperature) rather than their values (such as *hot*), while we address the problem by only mining the most diagnostic property values.

More surprisingly, perhaps, property values derived from WordNet glosses (E) are also poorly predictive, yielding a correlation with Whissell's affect-ratings of just 0.278. Our goal in this paper has been to describe a framework for augmenting WordNet's concepts with functional forms that both reflect the diagnostic properties of these concepts and that allow them to categorize different concepts in different contexts. These results suggest that the properties needed to construct these categorization functions are not to be found within WordNet itself, but must be acquired by observing how people actually use concepts to construct and convey meanings in everyday language.

6 Concluding remarks

In this paper we have presented a context-sensitive functional framework for concept description in WordNet and other lexical ontologies. This framework serves as a flexible interface between, on one hand, the need for ontological clarity and a commitment to explicit logical definitions, and on the other, the context-sensitive utilization of these definitions in a representative body of text.

These conceptual functions establish their own categorization boundaries based on the context, and—under the ontologist’s control—can blur the traditional line between the literal and metaphoric usage of a concept when it is ontologically useful to do so (e.g. see Hanks, 2006).

This programmatic approach to concept definition complements the syntagmatic approach to ontology construction outlined in Cimiano *et al.* (2005), since the ontologist is here given access to the syntagmatic features of a context via a flexible but powerful representation language. We have also described how concept definitions can, like those of Cimiano *et al.* and Poesio and Almuhareb (2005) be created automatically, by identifying the most diagnostic properties of each concept as expressed in similes on the Web. By associating conceptual functions with the most commonly used WordNet noun senses, we achieve a pair of related goals: WordNet is augmented with a robust, non-classical view of conceptual structure; and, more importantly in the context of this special issue, WordNet is remade in a context-sensitive form. Ultimately, these two goals are flip-sides of the same coin, for insofar as context alters the perceived boundaries of familiar concepts, classically structured ontologies like WordNet cannot be made context-sensitive without first being augmented with a flexible sense of categorization.

Much work remains to be done on the current framework, not least on the development of a more formal treatment of how our approach serves to augment WordNet (or WordNet-like resources) with concept descriptions that can be used both to categorize in context and to reason about those categorizations. Such a formal treatment would likely parallel that offered to classification structures by Giunchiglia *et al.* (2005), who formalize the notion of a classification system by expressing topic labels in a propositional concept language. As a lightweight lexical ontology, WordNet is itself little more than a classification hierarchy, and the conceptual functions we now assign to its lexical entries serve much the same purposes (i.e. categorization and introspective reasoning) as the concept-language labels employed by Giunchiglia *et al.*

More speculatively, more exploratory effort is clearly merited by the tantalizing issue of where literal categorization ends and metaphoric categorization begins, and the role of context in blurring this boundary.

References

- Almuhareb, A. and Poesio, M. 2005 Concept learning and categorization from the web. In *Proceedings of CogSci 2005, the 27th Annual Conference of the Cognitive Science Society*. New Jersey: Lawrence Erlbaum.
- Barsalou, L. W. 1983 Ad hoc categories. *Memory and Cognition* **11**, 211–227.
- Bouquet, P., Giunchiglia, F., van Harmelen, F., Serafini, L. and Stuckenschmidt, H. 2003 C-OWL: contextualizing ontologies. In *Proceedings of 2nd International Semantic Web Conference*, LNCS vol. 2870:164–179. Springer Verlag.
- Budanitsky, A. and Hirst, G. 2006 Evaluating WordNet-based measures of lexical semantic relatedness. *Computational Linguistics* **32**(1), 13–47.
- Charniak, E. and Berland, M. 1999 Finding parts in very large corpora. In *Proceedings of the 37th Annual Meeting of the Assoc. for Computational Linguistics*. pp. 57–64.
- Cimiano, P., Hotho, A. and Staab, S. 2005 Learning concept hierarchies from text Corpora using Formal Concept Analysis. *Journal of AI Research* **24**, 305–339.
- De Leenheer, P. and de Moor, A. 2005 Context-driven disambiguation in ontology elicitation. In Shvaiko, P. and Euzenat, J. (eds.), *Context and Ontologies: Theory, Practice and Applications*, AAAI Technical Report WS-05-01. AAAI Press, pp. 17–24.
- Dong, Z. and Dong, Q. 2006 *HowNet and the Computation of Meaning*. Singapore: World Scientific.
- Fellbaum, C (ed.). 1998 *WordNet: An Electronic Lexical Database*. Cambridge, MA: The MIT Press.
- Gangemi, A., Guarino, N. and Oltramari, A. 2001 Conceptual analysis of lexical taxonomies: the case of WordNet’s top-level. In Welty, C. and Barry, S. (eds.), *Formal Ontology in Information Systems*. In *Proceedings of FOIS2001*. ACM Press, pp. 285–296.
- Ganter, B. and Wille, R. 1999 *Formal Concept Analysis: Mathematical Foundations*. Berlin: Springer Verlag.
- Giunchiglia, F. 1993 Contextual reasoning. *Special issue on I Linguaggi e le Macchine XVI*, 345–364.
- Giunchiglia, F., Marchese, M. and Zaihrayeu, I. 2005 Towards a theory of formal classification. In Shvaiko, P. and Euzenat, J. (eds.), *Context and Ontologies: Theory, Practice and Applications*, AAAI Technical Report WS-05-01. AAAI Press pp. 25–32.

- Guarino, N. (ed.). 1998 Formal ontology and information systems. *Proceedings of FOIS1998*, June 6–8, Trento, Italy. Amsterdam: IOS Press.
- Guha, R. V. 1991 Contexts: a formalization and some applications. *Technical Report STAN-CS-91-1399*. Stanford, CA: Stanford Computer Science Dept.
- Hanks, P. 2004 The syntagmatics of metaphor. *International Journal of Lexicography* **17**(3), pp. 245–274.
- Hanks, P. 2006 Metaphoricity is gradable. In Stefanowitsch, A. and Gries, S. (eds.), *Corpora in Cognitive Linguistics. Vol. 1: Metaphor and Metonymy*. Berlin and New York: Mouton de Gruyter pp. 17–35.
- Hearst, M. 1992 Automatic acquisition of hyponyms from large text corpora. In *Proceedings of the 14th International Conference on Computational Linguistics*. pp. 539–545.
- Lakoff, G. 1987 *Women, Fire, and Dangerous Things: What Categories Reveal About the Mind*. Chicago: University of Chicago Press.
- Lee, L. 1999 Measures of distributional similarity. In *Proceedings of the 37th Annual Meeting of the Association for Computational Linguistics*. pp. 25–32.
- Lenat, D. and Guha, R. V. 1990 *Building Large Knowledge-Based Systems: Representation and Inference in the CYC project*. Reading, MA: Addison-Wesley.
- Morris, J. 2006 *Readers' Perceptions of Lexical Cohesion and Lexical Semantic Relations in Text*. PhD thesis, University of Toronto.
- Obrst, L. and Nichols, D. 2005 Context and ontologies: contextual indexing of ontological expressions. In *Proceedings of the AAAI 2005 Workshop on Context and Ontologies*, Pittsburgh, PA.
- Patel-Schneider, P.F., Hayes, P. and Horrocks, I. 2003 *Web Ontology Language (OWL) Abstract Syntax and Semantics*. Technical report, W3C, www.w3.org/TR/owl-semantics.
- Pustejovsky, J. 1995 *Generative Lexicon*. Cambridge, MA: The MIT Press.
- Segev, A. and Gal, A. 2005 Putting things in context: a topological approach to mapping contexts and ontologies. In Shvaiko, P. and Euzenat, J. (eds.), *Context and Ontologies: Theory, Practice and Applications*, AAAI Technical Report WS-05-01. AAAI Press pp. 9–16.
- Ushold, M. 2000 Creating, integrating and maintaining local and global ontologies. In *Proceedings of the 1st Workshop on Ontology Learning (OL-2000)*, part of ECAI 2000.
- Weeds, J. and Weir, D. 2005 Co-occurrence retrieval: a flexible framework for lexical distributional similarity. *Computational Linguistics* **31**(4), 433–475.
- Welty, C. and Jenkins, J. 1999 An ontology for subject. *Journal of Data and Knowledge Engineering* **31**(2), 155–181.
- Whissell, C. 1989 The dictionary of affect in language. In Plutchnik, R. and Kellerman, H. (eds.), *Emotion: Theory and Research*, New York: Harcourt Brace. pp. 113–131.