

Using Wearable Computers to Construct Semantic Representations of Physical Spaces

Fernando Diaz
Department of Computer Science
University of Massachusetts
Amherst, MA 01003
fdiaz@cs.umass.edu

Abstract

The representation of physical space has traditionally focused on keyphrases such as “Computer Science Building” or “Physics Department” that help us in describing and navigating physical spaces. However, such keyphrases do not capture many properties of physical space. As with the assignment of a keyword to describe a piece of text, these constructs sacrifice meaningful information for abstraction. We propose a system of spatial representation based on richer, emergent language models that encode information lost in keyphrase approaches. We use a mix of wearable and ubiquitous computing environments for the construction of these models. Wearable computers infer language models of their hosts. These language models then act as semantic paint over spaces in a ubiquitous computing environment. Spaces collect this information and construct representations based on interactions with augmented humans. A prototype navigation system based on this theory is presented and compared to traditional representations.

1. Introduction

Traditionally, the semantic labeling of spaces with building or room names involves the manual task of assigning some keyphrase to a space. Unfortunately, these assignments do not constitute a rich representation of space. A computer science building is more than just “computer science”; it also encompasses, to varying degrees, “algorithms,” “artificial intelligence,” “machine learning,” and many other topics *depending on the occupants of the building*. That is, a person’s conception of a building includes more than its structure (e.g., floor plans, lighting). Especially when familiar with the objects and people occupying a building, a person might think of that building as something more abstract and meaningful than a collection of

generic objects and people. For example, the task of labeling a building would be quite difficult if we were only given access to its structural properties. Knowledge of the occupants provides insight when constructing a meaningful description. When we are given the names and homepages of the occupants, we can better assign a useful label to a building.

We describe the development of a system of spatial representation grounded in the interaction of people in space. Related work in representation has been conducted in information retrieval and collaborative filtering. In these areas, good document or item representations are measured by an ability to effectively rank a set of items with respect to a query or active user. Likewise, the task of finding relevant spaces can motivate the adoption of similar representations; we want to rank space with respect to a description of what we are looking for.

We explore approaches to spatial representation that rely upon occupant-derived representation. Such a model requires both a representation of the individual occupants as well as an algorithm for constructing a representation of the space from this information. Wearable computers provide an excellent platform for the first task. Indeed, traditional user modeling techniques deployed in a variety of domains serve as a lower bound on the performance of wearable computers in constructing representations of individuals. Already, wearable computer systems have been developed which demonstrate the ability to construct fine-grained models of individuals [5, 11]. To address the second task, we adopt a computational partitioning of space similar to Dataspace [7]. In this architecture, physical spaces such as buildings or rooms maintain computational resources providing a location for accumulating knowledge. Individuals with wearable computers passing through these spaces provide the personal information used to build spatial models. The problem of constructing a representation of a space reduces to reasoning about the collection of user models

which pass through a space.

This paper develops the idea of interaction-based spatial representations by starting from traditional approaches to representation. A general theory of interaction-based representation is developed in Section 2. In order to contextualize our approach to previous methods of representation, we have organized several existing architectures into a set of categories. Using this theory of interaction-based representations, we develop a method for constructing interaction-based spatial representations in Section 3. In the course of this process, we describe representational techniques previously not explored. Having developed an algorithm for building these representations, we describe the implementation of a prototype navigation system in Section 4. We conclude by placing our research in context and discussing future directions in Sections 5 and 6.

2. Interaction-based Representation

Being concerned with the representation of space, we begin by discussing the various methodologies of constructing representations. Our focus will not be on a personal representation of space as is explored in artificial intelligence. Instead, we will prefer an approach which focuses on the social definition of a space. So, rather than having an agent ask, “what is this space about?”, we have the space ask, “what am I about?” In particular, we are interested in the representation of spaces as a product of interaction with people. However, we will first develop a more general method for building interaction-based representations using previous work.

2.1 Simple Interaction

The first type of representation to consider is a simple description of what is interacting in a system of objects. There are many dimensions upon which interaction may occur: two people speaking (linguistic interaction), several people being in eyesight of each other (visual interaction), two cars colliding (physical interaction). When considering a particular dimension, some objects are more relevant than others. The two people in a room are more relevant during a dialog than, for example, the chairs these individuals are sitting on. Relevance is mentioned as a means to reduce the system which we will have to describe. Even though the door may be semi-relevant to a dialog, such a system can be described as two people speaking. Therefore, linguistic interactions can be described by an interaction matrix of partners in conversations. That wearable computers provide this type of human-level monitoring partially motivates this work and explains why many of the examples involve people. Beyond this, however, ubiquitous computing results in a similar potential in physical objects. For example, if cars

are augmented with collision sensors, an interaction matrix can describe the physical interaction.

Much previous research implicitly adopts this interaction based framework. For example, collaborative filtering and Chalmers’ path-based information retrieval abstract information objects (e.g. documents, movies, songs) and manipulate their representations with respect to the people they interact with [4, 1]. In these systems, people are represented by the objects they have read, watched, or heard. Likewise, the information objects are represented by the people who read, watch, or hear them. Both representations ignore descriptions of the components (i.e. people and items).

In terms of ubiquity, Davis, *et al.* develop a representation of nodes in an ad hoc wireless network with respect to communicative interaction [6]. In such an environment, nodes may have very limited communication range and high mobility. Globally, the resulting network can be partitioned, dynamic, and altogether difficult to navigate. The goal is to find a relatively short and reliable route from one node to another given only a destination’s identifier. Interaction is defined by two nodes being within communication range. Here, a particular node is represented by the history of other nodes with which it has had possible communication. As with collaborative filtering, a description of the components of this representation (i.e. other nodes) is lacking.

2.2 Interaction Described

A second type of representation is possible using the descriptive history of interactions an object has participated in. We believe that there is power in describing the interaction itself. When people are speaking, one can describe the system not just by who is speaking to whom but also by what words are spoken between these individuals. When two cars collide, one can use a range of values to describe the collision. Consequently, a person can be represented by the words he or she has read, written, heard, or spoken. Likewise, a car can be described by the severity of collisions it has been involved in.

Traditional information retrieval may be cast in this representation scheme. The population of objects consists of the users and their document collection. Interaction is defined by reading a document and, hence, can be described by what is read. That is, a document is only represented by the words that flow between it and a reader (i.e. the text). More recent information retrieval systems incorporate additional knowledge into representation. For example, hypertext retrieval adds inter-document interaction to representation [15, 9].

With respect to collaborative filtering, content-based schemes incorporate linguistic knowledge about the interactions between objects and observers [13]. So, in addition to being represented by the people who have interacted with

it, a particular item is also represented by a description of the interaction which may be the text of a document or the synopsis of a film.

Both of the representational schemes describe important aspects of the system. Simple interaction tells one a lot. But, while it gives insight into the identities of the peers, this set alone provides nothing beyond a social context. Knowing the ISBN numbers of books I read and social security number of the people I speak with tells one little about my interests. It may, as in collaborative filtering, be able to represent interests in the abstract sense of individuals' overlapping book or dialog-peer selections. Nevertheless, if given the text of all of the books and the linguistic histories of all of my dialog-peers, then one may be able to better describe my interests. Similarly, knowing who is passing through space can tell one a lot about popularity and groups of people. Knowledge about the set of interests of that group can go further even if we do not know what is of particular interest in that space.

3. Spatial Semantic Model

The focus in our examples on people and language is not accidental. First, people are readily monitored and represented by wearable computers. One of the advantages of wearable computers is their persistent existence with an individual. The interaction of these wearable computers and other physically-bound computers allows the exploration of the representational methodologies we described above. Second, language grounds representation in flexible, understandable primitives. For our ends, a linguistic representation is useful since it allows us to build systems that are queriable using traditional information retrieval techniques [12, 19]. Words provide a powerful interface potential and carry a history of academic research. Given these aspects, we would like to build a spatial representation system which incorporates both interaction as well as linguistic representations. We will first describe a framework for building linguistic representations. Using this grounding and the ideas developed in Section 2, we will develop a method to bind meaningful linguistic representations to physical space.

3.1 Linguistic Representation

Wearable computers have access to a wealth of linguistic information in the form of both text (email, web browsing, and document composition) and speech (through speech recognition technology). The result is a history of words which have passed over the user's lips, ears, and eyes. Complete histories are informative but not compact and certainly not immediately comparable.

We propose the use of information retrieval techniques for abstracting from collections of words. The information

retrieval community represents documents in any number of ways: keywords, subject headings, abstracts, term vectors. Recent advances in information retrieval have found representational power in language models of documents [16]. The intuition with language models is that there is an underlying generative model for some collection of words. Word collections act as a sample from this model and can be used to estimate the "true", underlying language model. To a certain extent, for an individual, this representation can serve to describe a set of interests. A person who is interested in computer science is more likely to speak, write, hear, or read about "algorithms" and "artificial intelligence" than a person who is not interested in these things at all.

3.2 Traditional Interaction-based Spatial Representations

Since we are basing our spatial representation on interaction, an investigation of methods described in Section 2 is appropriate.

First, we consider the a representation based on simple interaction. The resultant representations would be similar to items in the collaborative filtering example or nodes in the ad hoc routing example. A single space would be represented by a vector of individuals who have passed through it. A collection of these spatial representations would allow query by example like collaborative filtering or the search for a particular individual like ad hoc routing. However, neither of these features result in easy map-based interaction.

We have already described the advantages of linguistic representations in Section 3.1. Practically, though, where does the linguistic information about a space come from? Although it is natural to think about linguistic interactions between people, using language to characterize interaction with space is not obvious. Let us consider the options. If it is claimed that there are word-based representations of spaces, where do the representations come from? Perhaps a linguistic interaction with a space means a linguistic interaction is occurring *in* that space. A linguistic history of a space would be constructed from the history of words spoken within that space. Basically, anything communicated between people in a space is monitored and incorporated into its representation. Figure 1 provides an interpretation of the source of linguistic information. The resulting linguistic representations are based solely on the linguistic interactions happening in a space. The identities of the participants are ignored. A linguistic representation of space built like this is reasonable but not practical. The history of words spoken in a space is potentially sparse or misrepresentative. Even though many a word may be spoken in an office, it can still have a representation based on the people occupying it.

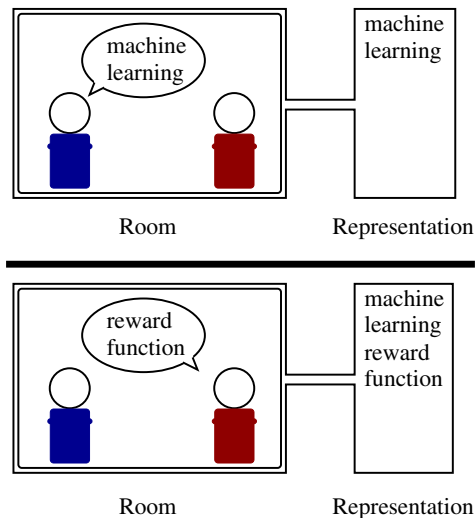


Figure 1. Linguistic information from linguistic interaction occurring in a space: The linguistic representation of a room is constructed solely from the words spoken in the room.

3.3 User-based Interaction-based Representations

The short-comings of traditional approaches to interaction-based representation lead us to consider novel methods of constructing such representations. One of the disadvantages described above was the potential sparsity in immediate linguistic information. In order to accumulate more data for our models, then, we advocate the construction of models based upon the linguistic representations of the people using that space. Wearable computers provide the ability to not only monitor speech but also a persistent monitoring of a user. It is this monitoring which allows the construction of rich models of user linguistic patterns. The information about the users occupying in the space is then exploited to construct a representation of the room itself. Figure 2 depicts the transmission of an entire linguistic representation to the space. This user representation includes the terms mentioned in the example in Figure 1 as well as an individual context for those terms.

A subtle distinction between this approach and previous approaches should be realized. Whereas the two interactive representations described place objects in either an immediate social or information context, our spatial representation attempts to combine the two. The linguistic representations are constructed not from immediate interaction such as text in a document but from the linguistic representations of the users. This would be akin to describing a document by the

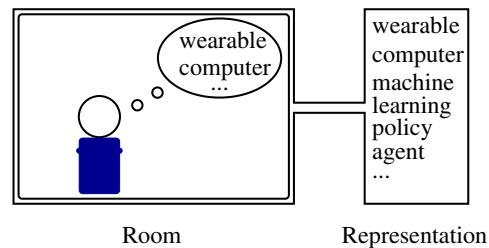


Figure 2. Linguistic information from linguistic representations of the individuals in a space: The linguistic representation of a room is constructed from an abstract representation of occupants in a room. This abstract representation describes a user's set of interests.

During each timestep,

1. wearable machines recompute linguistic representations of their hosts
2. if a user enters a space,
 - (a) the new user's wearable machine transmits its linguistic representation to computer associated with that space
 - (b) the spatially-bound machine recomputes a representation of itself based on the new collection of linguistic user representations

Figure 3. Representation construction algorithm

confluence of linguistic representations of its readers. This is an alternative representation we assume is a close approximation of immediate interaction. Even in cases where information about immediate interaction is provided, a reasoning about the linguistic representations of users is potentially exploitable. For example, if two statistics texts are equivalent with respect to content, knowing that one is far more often used by computer scientists perhaps tells us that this text is better suited for a computer science curriculum than the other.

A space, then, accumulates knowledge about its occupants in the form of these models from which composite representations can be constructed. Components of this distribution will be reinforced if people have common interests. So, even though several people in a particular space may be quite different, the composite representation should

For each space to be considered,

1. calculate the relevance of the space to the query
2. highlight the spaces according to this relevance measure.

Figure 4. Representation retrieval algorithm

encompass the similarities between those representations. In order to accomplish this with a collection of language models, we perform a uniform combination of between of people who have passed through a particular space. Figure 3 describes the behavior of the system construction algorithm during execution.

The querying of a collection of spatial representations can be thought of as analogous to reading a map for relevant areas. In this case, reading is substituted by natural language querying similar to modern information retrieval systems. Because our spatial representations are probabilistic models, we can compute the relevance of a space as the probability of that spatial representation generating the query. The relevance is related, then, to the likelihood of those words having been spoken by the occupants who passed through that space. Using our map metaphor, Figure 4 describes the retrieval algorithm.

3.4 Scenario: Alice's Day Out

Consider the following scenario. Alice, an undergraduate computer science student, owns, like everyone else, a wearable computer which infers a linguistic representation from her web browsing and document composing habits. Because Alice is interested in areas such as wearable computing and artificial intelligence, the language model allots a larger probability mass to words such as “wearable,” “mobile,” “learning,” and additional, related words. However, Alice is not one dimensional so her language model assigns relatively high probability to terms such as “guitar,” “tremelo,” and “flamenco.” Clearly Alice also has some interest in classical guitar.

Since Alice's environment is augmented with spatial representation machines, rooms in her department's building have representations associated with the language models of common foot traffic. So, as Alice enters her laboratory, her wearable communicates the inferred language model to the local spatial representation machine. The spatial representation machine then recalculates its representation based on this new information as well as the history of language models it has been transmitted by others. Because the majority of peers in Alice's laboratory also study machine learning, the combined language model reinforces terms such as

“learning,” “training,” and so on.

This afternoon, Alice visits a university campus she is considering for graduate school. Unfamiliar with the campus, Alice asks her wearable how to get to the machine learning laboratory. The wearable contacts the campus directory which maintains communication with all of the spatial representation machines on campus. This central directory then estimates the probability of Alice's query being satisfied by the different spaces. The campus directory constructs a campus map overlaid with color whose intensity is relative to this probability. The map is transmitted to Alice's wearable. Alice notices that there is a cluster of bright red spaces in the building next to her. A visit to the brightest spaces results in Alice finding the machine learning laboratory. Investigating other brightly colored spaces in the building, Alice discovers that the robotics laboratory also conducts interesting work in machine learning.

Satisfied with the machine learning research on campus, Alice asks about guitar playing on campus. Disappointed by the initial results (almost every building has a guitar player), Alice specifies *classical* guitar playing. A small cluster of rooms gets highlighted in a nearby building. Here, Alice finds a music department where, apparently, flamenco is embraced.

4. Prototype Interface

This spatial representation system is being deployed for the Computer Science Building at the University of Massachusetts. While a department-wide adoption of wearable computers is welcome, it is not feasible at the moment. Linguistic data for occupants of the building has been synthesized from publications and home pages. This data will serve to construct language models of the occupants of the building. However, before constructing the models, some preprocessing was conducted. First, text extracted from these documents was normalized by stemming according to the KStem algorithm and dropping a list of high-frequency, content-free stop words [10]. Terms occurring only once in the entire collection were omitted from calculations. For each individual in the department, his or her documents were used to build a vector of term-frequency pairs. Simple language models are built using the maximum likelihood estimate,

$$P_o(w) = \frac{tf_o(w)}{\sum_i tf_o(w_i)},$$

where $tf_o(w)$ is the count in the term vector for building occupant o . This gives us a naïve language model. Unfortunately, such an estimate assigns zero probability to unseen words. This problem is addressed by by mixing the maximum likelihood language model with a model of general English. In this case, a general English model is constructed

from the entire collection of documents for all users. Therefore,

$$P'_o(w_i) = \lambda P_o(w_i) + (1 - \lambda) P_{GE}(w_i),$$

where λ is a mixing parameter which we set to .80.

A model for the second and third floors of the Computer Science Building was then constructed to allow simulation of occupancy. The individuals were associated with their offices in the model. Many offices are shared, demanding a combination of the language models of the occupants. Composite language models were built by uniformly combining the individual language models,

$$P(w_i|M_s) = \frac{1}{|O_s|} \sum_{j \in O_s} P'_j(w_i)$$

where O_s is the set of occupants in the space s for which M_s is a model.

This semantic model of the building is constructed locally on a Xybernaut wearable computer [21]. Due to the relatively small number of potentially relevant spaces, no optimization of the indexing needed to be conducted. This would be necessary in very large buildings or sets of spaces. We designed the system for speech to allow the flexibility of traditional information retrieval querying without the overhead of learning to use traditional wearable keyboard alternatives. A user interacts with the system by issuing speech queries recognized by IBM ViaVoice runtime libraries [20]. The set of recognized words constitutes the query. The system then generates a relevance measure for all the spaces in the building based upon the probability of the space's semantic model, M_s , generating those words:

$$P(q|M_s) = \prod_i P(q_i|M_s),$$

where q is the sequence of query terms. These probabilities are then used to mark up a map of the building that is presented to the user on a touch panel display. Figure 5 shows this map for the query "robotics." The interface presents the user with the current state of the query, which provides context for the results. These results are displayed in two panels. The left-hand panel shows the relevance of spaces in the building. The right-hand panel displays the ranked list of relevant spaces using manually assigned labels. We found that this list helps in rapidly characterizing the space especially when used in conjunction with the highlighted map. In our example, the system highlights the robotics laboratory and offices of associated people. Interestingly, the system also detects the interest of machine learning and artificial intelligence laboratories in robotics.

5. User Experience

Having built a prototype system, we were interested in the application of this visualization to the task of navigation of space. Several computer science students with varying amounts of experience in the Computer Science building were given the system to use for exploring the space. Most users were enthusiastic about the system as a means of reducing the overhead when investigating a new building. Traditionally when trying to determine where relevant research is being conducted in a building, a coordination of web-browsing and physical maps is necessary. Our system combines this information into a single interface to allow more efficient navigation. Users were able to quickly find the offices and laboratories relevant to particular interests. Some also gained an awareness of previously unknown similarities between laboratories. Most participants were largely disappointed with speech recognition performance which resulted in longer search times. One user recommended the option for query reformulation so that the map state would change as terms were added to a query.

6. Related Work

With respect to representation, our approach is quite similar to *stigmergic* or pheromone-based algorithms [3]. These systems harness the distributed, socially-constructed representation of traffic on a network for problems such as finding shortest paths. Important to these algorithms is the notion of agent leaving markers at geographic locations and having representations emerge as a result of the marker accumulation. Our work in spatial representations demonstrates the application of this theory to domains outside of networking.

As an architecture, the Dataspace model comes closest to the system we describe [7]. While Dataspace describes at a high level how to partition and query spaces, the authors do not describe how the information in such a system comes to reside where it does. We consider our system to be an initial attempt at exploiting such an architecture in information retrieval.

Brown's work with stick-e notes is also related in the ascription of data to spaces [2]. Stick-e notes are text data stored in spaces. This text is broadcast to a computer user if certain contextual information is satisfied. Individuals may then leave similar notes for others traveling through such an augmented space. It is this latter part which we are automating so that instead of transmitting a text message, a user transmits a complex representation. Coincidentally, it is not impossible for a spatial machine in our system to broadcast its own representation to users passing through. This message encodes not only a spatial representation but also a potential user context. For example, many researchers have

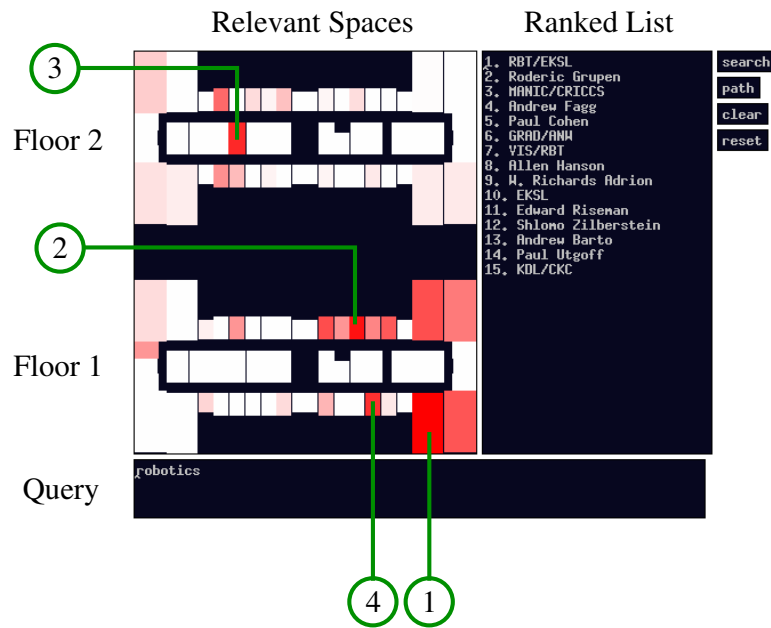


Figure 5. Spatial search results for the query “robotics”: The left-hand panel displays the relevant spaces graphically. The right-hand panel displays a ranking of relevant spaces using manually assigned labels.

described information retrieval systems which incorporate contextual information such as location or room occupants [8, 17]. This information retrieval system could seamlessly consider spatial context in the form of the models that we present.

Several recent artificial intelligence approaches to representation attempt to learn meaning based on the co-occurrence of spoken words and physical objects [14, 18]. These techniques reinforce specific word-sensor associations in an attempt to learn word meaning. The negotiated representation resides in the heads of the individual agents operating within the environment. In other words, the artificial intelligence community is interested in a vertical approach to intelligence by focusing on the construction of a highly sophisticated agent or group of agents acting in dynamic physical environments. In some ways, our work is an inversion of these artificial intelligence initiatives. The system attempts to learn object meaning by placing the representation into the object itself. Hence, we are interested in a horizontal approach to intelligence by focusing on the construction of sophisticated dynamic physical environments. Agents hold no privileged place.

7. Conclusion

We have presented a system to construct rich, emergent spatial representations. The representations result in meaningful spaces and aid in visualization and navigation. In designing the representational system, a novel method for approximating immediate linguistic representation was developed.

There are several extensions to the system we are currently considering. The temporal and dynamic aspects of these emergent representations remain unexplored. Realistic movement models would be necessary for these experiments. We are investigating the acquisition of empirical movement data for the faculty and students in our system. By incorporating movement into our model of the building, spatial representations can be constructed using different transformations on the interaction histories. For example, considering only the a short, recent history of people occupying a space may reduce the accuracy of the representation but will make the representation more robust to the dynamism of shared spaces.

While the prototype system holds promise, limitations exist. The type of queries possible is limited by the amount of representational power in text information related to a space. For example, it is unlikely that the system would work well on queries for subway stations, restaurants, or

other public places. The people occupying these places are too diverse. Inferring meaningful representations for these spaces from individuals' language data may not be possible but we believe useful linguistic histories exist somewhere in the environment.

Wearable computers provide the ability to model a vast amount of user interaction beyond words. Several initiatives to model context reveal the ability to model abstract states such as "walking" or "sitting" [5, 11]. One can imagine other abstract states such as "hammering". If such states were communicated to objects in the environment, then we could also imagine representing objects by the ways they have been used. For example, a hammer would most often be used for hammering though a shoe may also be used for the same task. An agent confronted with the need to hammer would not have to reason about the hammering properties of objects in the environment. Instead, it may merely seek those objects whose representations include hammering.

8. Acknowledgements

This work was supported in part by the Center for Intelligent Information Retrieval, in part by NSF grant #IIS-9907018, in part by DARPA/ITO #DABT63-99-1-0022 (SDR), and in part by DARPA/ITO #DABT63-99-1-0004 (MARS). Any opinions, findings and conclusions or recommendations expressed in this material are the author(s) and do not necessarily reflect those of the sponsor.

References

- [1] J. S. Breese, D. Heckerman, and C. Kadie. Empirical analysis of predictive algorithms for collaborative filtering. Technical Report MSRTR-98-12, Microsoft Research, 1998.
- [2] P. J. Brown. Triggering information by context. *Personal Technologies*, 2(1):1–9, September 1998.
- [3] G. D. Caro and M. Dorigo. Antnet: Distributed stigmergetic control for communications networks. *Journal of Artificial Intelligence Research*, 9:317–365, 1998.
- [4] M. Chalmers. Informatics, architecture and language. In *Social Navigation of Information Space*, London, 1999. Springer-Verlag.
- [5] B. Clarkson, K. Mase, and A. Pentland. Recognizing user context via wearable sensors. In *Proceedings of the Fourth International Symposium on Wearable Computers*, 2000.
- [6] J. A. Davis, A. H. Fagg, and B. N. Levine. Wearable computers as packet transport mechanisms in highly-partitioned ad-hoc networks. In *Proceedings of the International Symposium on Wearable Computing*, Zurich, Switzerland, 2001.
- [7] T. Imielinski and S. Goel. Dataspace - querying and monitoring deeply networked collections in physical space. In *IEEE Personal Communications Magazine, Special Issue on Networking the Physical World*, October 2000.
- [8] G. J. F. Jones and P. J. Brown. Information access for context-aware appliances. In *Proceedings of SIGIR*, 2000.
- [9] J. M. Kleinberg. Authoritative sources in a hyperlinked environment. *Journal of the ACM*, 46(5):604–632, 1999.
- [10] R. Krovetz. Viewing morphology as an inference process. In *Proceedings of SIGIR*, 1993.
- [11] K. V. Laerhoven and O. Cakmakci. What shall we teach our pants? In *Proceedings of the Fourth International Symposium on Wearable Computers*, 2000.
- [12] C. D. Manning and H. Schutze. *Foundations of Statistical Natural Language Processing*. MIT Press, Cambridge, MA, 1999.
- [13] P. Melville, R. Mooney, and R. Nagarajan. Content-boosted collaborative filtering. In *Proceedings of the 2001 SIGIR Workshop on Recommender Systems*, 2001.
- [14] T. Oates. *Grounding Knowledge in Sensors: Unsupervised Learning for Language and Planning*. PhD thesis, University of Massachusetts, Amherst, MA, 2001.
- [15] L. Page, S. Brin, R. Motwani, and T. Winograd. The pagerank citation ranking: Bringing order to the web. Technical report, Stanford Digital Library Technologies Project, 1998.
- [16] J. M. Ponte and W. B. Croft. A language modeling approach to information retrieval. In *Research and Development in Information Retrieval*, pages 275–281, 1998.
- [17] B. J. Rhodes. *Just-In-Time Information Retrieval*. PhD thesis, MIT Media Laboratory, Cambridge, MA, May 2000.
- [18] L. Steels and F. Kaplan. Situated grounded word semantics. In *IJCAI*, pages 862–867, 1999.
- [19] C. J. Van Rijsbergen. *Information Retrieval, 2nd edition*. Dept. of Computer Science, University of Glasgow, 1979.
- [20] http://www-3.ibm.com/software/speech/dev/sdk_linux.html.
- [21] <http://www.xybernaut.com>.