



Technology Radar Feature Paper

Edition II/2006

June 2006

Web 3.0: Convergence of Web 2.0 and the Semantic Web

By Wolfgang Wahlster and Andreas Dengel,
German Research Center for Artificial Intelligence (DFKI)

Web 3.0: Convergence of Web 2.0 and the Semantic Web

By Wolfgang Wahlster, Andreas Dengel

German Research Center for Artificial Intelligence (DFKI)
Saarbrücken and Kaiserslautern

with contributions by Dietmar Dengler, Dominik Heckmann, Malte Kiesel, Alexander Pfalzgraf,
Thomas Roth-Berghofer, Leo Sauermann, Eric Schwarzkopf, and Michael Sintek

The World Wide Web (WWW) has drastically improved access to digitally stored information. However, content in the WWW has so far only been machine-readable but not machine-understandable. Since information in the WWW is mostly represented in natural language, the available documents are only fully understandable by human beings. The Semantic Web is based on the content-oriented description of digital documents with standardized vocabularies that provide machine understandable semantics. The result is the transformation from a Web of Links into a Web of Meaning/Semantic Web [1], (see arrow A in Fig. 1). On the other hand, the traditional Web 1.0 has recently undergone an orthogonal shift into a Web of People/Web 2.0 where the focus is set on folksonomies, collective intelligence, and the wisdom of groups (see arrow B in Fig. 1). Only the combined muscle of semantic web technologies and broad user participation will ultimately lead to a Web 3.0, with completely new business opportunities in all segments of the ITC market.

Without Web 2.0 technologies and without activating the power of community-based semantic tagging, the emerging semantic web cannot be scaled and broadened to the level that is needed for a complete transformation of the current syntactic web. On the other hand, current Web 2.0 technologies cannot be used for automatic service composition and open domain query answering without adding machine-understandable content descriptions based on semantic web technologies. The ultimate worldwide knowledge infrastructure cannot be fully produced automatically but needs massive user participation based on open semantic platforms and standards.

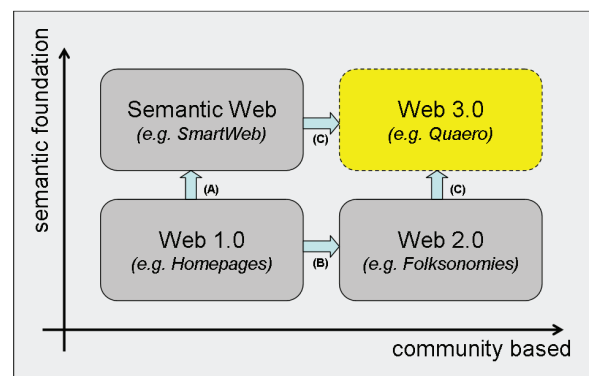


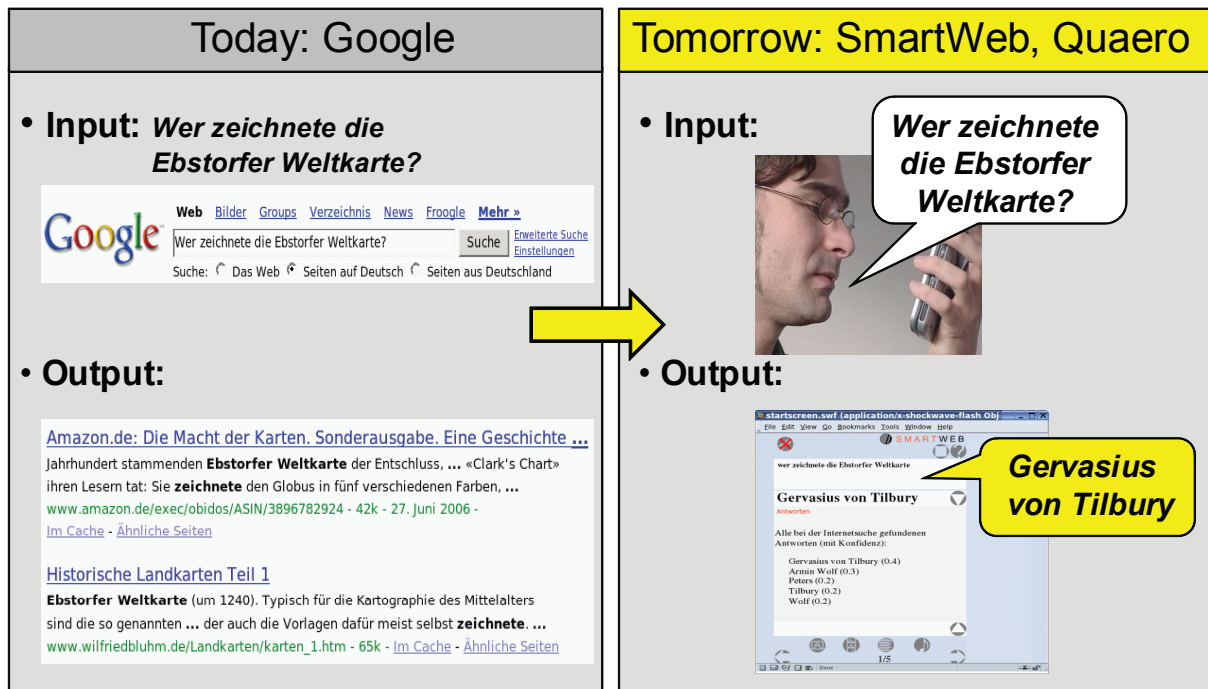
Figure 1: The upcoming trend of Web 3.0 in the World Wide Web

The interesting and urgent question that arises is: what happens when the emerging Semantic Web and Web 2.0 intersect with their full potential? We analyze this question throughout this feature paper and present the converging idea that we call Web 3.0. We use the following definition in this paper:

Web 3.0 = Semantic Web + Web 2.0.

A good example for developing Web 3.0 is the mobile personal information assistant (see Fig. 2). The user makes queries using natural language, and the assistant answers by extracting and combining information from the entire web, evaluating the information found while applying Semantic Web technologies. Today's second-generation search engines are based on keywords within the syntactic web, while open domain question answering engines are based on information extraction and the Semantic Web.

This paper is divided into three major parts. The first one presents a broad overview of the latest



[Figure 2: From search to question answering]

technologies, trends, and business models of Web 2.0. The second one deals with improvements and current research directions within Semantic Web. And in the third part we analyze this promising idea of bringing together the bottom-up approach of „Web 2.0“ and the top-down approach of „Semantic Web“.

Web 2.0

Web 2.0 describes a paradigm shift in the evolution of the Web: Web 1.0, the “Web of Companies”, has become Web 2.0, the “Web of People”. No new technologies have been introduced by this shift, but the role and “value” of the user has been changed significantly.

The term Web 2.0 was coined in spring 2004 by Tim O'Reilly and Dale Dougherty from O'Reilly Media in order to label new applications and emerging trends of the web. O'Reilly's definition of Web 2.0 [16] is comprised of the following seven principles:

- **The Web as a Platform:** The Web and all its connected devices are considered to be one global platform of reusable services and data where one can build on the work of others using open standards.
- **Collective Intelligence:** A Culture of Participation arose with the establishment of open systems that support and enable

cooperative creation of content following the maxim of “trust instead of control”. The observation is that the mass of users provides a mass of knowledge but that same mass of users prevents misuse of services and removes incorrect content. Wikipedia is a good example of harnessing the collective intelligence of people. It is an online encyclopedia where any user can publish some information and any other user is allowed to change that entry.

- **Power of Data:** The heart of Web 2.0 processes are databases where the data itself is much more important than the application or interface which uses it. In the market the race is on to own certain classes of core data, e.g. location, identity, product identifiers, and namespaces. Additionally, a critical mass of data can also be reached via user aggregation, and the aggregated data is then turned into a system service.
- **Service Operations and Open Source:** Software in the Web 2.0 era is delivered as a service and not as a product. The core competency of the involved companies is the daily operations of their services, rather than the algorithms used. The dynamic nature of the services requires a constant, cost-effective change by using an open source development style. This involves a “perpetual beta” status of many of the services over years, where new

features are added, sometimes on a daily basis. Successful Web 2.0 companies are experts in the real time monitoring of user behavior to see just which new features are used, how they are used, and how they should be adapted.

- **Lightweight Programming Models:** Since the syndication of data, not the bulkheading of its usage, has been recognized as a driving force in the market, the easy linking, extension, and mixing of data using simple open interfaces (APIs) allows for the formation of new, loosely coupled web services every day. So, the innovation of many Web 2.0 services is in assembling components in novel and effective ways.
- **Passing the Barriers of a PC:** A feature of Web 2.0 systems is that they are no longer limited to the PC platform. This is not news in respect to web applications, but Web 2.0 apps offer a fuller realization of the true potential of the web as a platform, not only encompassing the world of interconnected PCs, but also that of mobile phones and less powerful entertainment devices such as Apple's iPod.
- **Rich User Experiences:** It's common practice now that Web applications provide rich user interfaces and interactions with a server as only high-profile clients could implement before. Key components of the browser to reach such a behavior are standards-based display and interaction and an asynchronous data retrieval using JavaScript binding everything together.

From this definition, Web 2.0 is clearly not a new technology. It is a new principle of working with data and incorporating people into the process.

	Web 1.0	Web 2.0
Paradigms	▪ Web is unidirectional ▪ Home pages ▪ Web form interaction ▪ Services sold over Web ▪ API and IPR ownership	▪ Web as a platform ▪ Collective intelligence: „Wisdom of crowds“ ▪ Web services ▪ Real Web applications ▪ Power of Data ▪ Data management and enrichment ▪ Conversation
Key Success Factors	▪ Own high-quality content ▪ Strong brand ▪ Large user base	▪ Community leveraging ▪ Robust platforms ▪ Open systems

Web 2.0 is a human to human concept, social individualism and a retreat from the reliance on big brand software monopolies. Certainly Web 2.0 has plenty of hype, but it has enabled an evolution or maybe a revolution of new web applications. In the following sections, some aspects of Web 2.0 are analyzed in greater depth and links to key applications are given.

Mashups

Mashup of services is the concept of building composite online software made up of services from multiple Web sites. This way, the Web provides data sources connected by APIs, i.e., official external service interfaces of sites as well as community-managed scrape APIs where the official ones are missing.

Mashup of services is the concept of **building composite online software** made up of services from multiple Web sites. Mashup has its roots in music remixing by DJs. In Web 2.0, everything is remixed (news, images, videos, audios, blogs), and the appropriate tools are available online, according to the user-centered approach.

- Jumpcut is a recent service for video remixing which defines the future of online video as the following: upload media, grab shared media, create and remix movies, publish to friends, share with the world. The tool works in a standard Web browser augmented with a flash plugin.

Besides the remixing of multimedia as a form of new mass entertainment, the remix of information found in website mashups is particularly popular. The necessary external service interfaces of websites are provided as web services. REST and SOAP are the two competing technical approaches based on the HTTP protocol to specify a web service.

- REST (Representational State Transfer) is based on the concept of transferring state between two systems. The data that are transferred or manipulated are identified by a standard URI addressing scheme. REST is simplistic, lightweight and easy to use, and therefore the interface used most often in current web-service accesses.
- SOAP is a more formal and standardized solution to providing Web services. Like REST, SOAP is an open messaging framework based on XML but is backed by formal standards of

the World Wide Web Consortium (W3C). SOAP messages contain, in their body, the actual XML message being transmitted and carry in a header descriptive metadata information about security, transactions, or other useful context that does not change the message itself, but describes higher-level processes of which it is a part. SOAP is rather appropriate for larger, formal applications that require advanced capabilities.

The explosion of services around the open APIs of Google (for map access) or Amazon/A9's OpenSearch API (for export of search functionality of local sites) demonstrates the ease of combining services to create valuable new information. However, there are currently no more than two hundred official APIs of websites available for access [9]. As the web service ecosystem will evolve further, there is the need to use **unofficial APIs** – **scrAPIs (scrape APIs)**, i.e., open APIs for data sources that don't have them.

- A scrAPI consists of a scraper process on web sites (HTML parsing and some custom code for making sense of the data) and a REST interface, which provides controlled access to the scraper's functionality. The maintenance of a scrAPI is done by the community in order to assure service availability.

However, the scraped web sites can support the ecosystem by designing scrape friendly data. It is very helpful when the HTML-data are structured using **microformats**—a lightweight semantic annotation format which simply extends HTML by specific classes and attributes, e.g., the hCard format represents the vCard standard of contact information.

- Thus microformats are not specific to the service scraped but the **community specifies formats** for "events" which can then be used by all sites handling events.
- The widespread use of microformats is supported by structured blogging **plugins** for the well-established blogging tool WordPress. There are also approaches in web browser development (such as flock.com) which auto-detect microformats on a webpage, and provide a structured view on a page.

The "lowercase semantic web" (as it is called in the community) uses microformats to add simple semantics to today's web following a user-

centric design where people are helping to create metadata.

Blogs

A weblog is a website where posts of one or more bloggers are listed in newest-post-first order. A post may consist of text, hyperlinks, and arbitrary multimedia objects and is usually highly subjective, reflecting the opinion of its author. Blogging tools make blog creation very easy and allow community formation by supporting commenting and inter-blog linking. Blogs and the links between them form the **blogosphere**.

The term **weblog**, later shortened to **blog**, was coined in 1997 by Jorn Barger, who defined it as a „webpage where a weblogger (sometimes called a blogger, or a pre-surfer) ‚logs‘ all the other webpages she finds interesting“ [2]. A weblog page presents its entries in chronological order, with the most recent entry on top. Barger himself is still editor of Robot Wisdom (www.robotwisdom.com), one of the original weblogs.

While in the mid 1990s, there were only a handful of blogs, today Technorati, a blog-search engine, tracks about 40 million. The explosive growth of the **blogosphere**, the blogs and the links between them, began with the advent of user-friendly blogging systems in 1999. These systems simplified posting to a blog from manually editing HTML code to filling in a form and pushing a button. The most popular early blogging software was Blogger, and it changed weblog culture via its simple interface: A single text-input field into which users could type whatever they wanted. This led to **linkless posts** and the line between weblogs and web journals, web-based diaries, began to blur [11].

Then features appeared that supported the growth of blogger communities:

- In 2000, Blogger assigned to each individual post on a weblog a URL that did not change with modifications to the blog, thus introducing **permalinks** (permanent links). A blogger can use permalinks to reference blog posts of other bloggers. This greatly simplified crossblog discussions and is a fundamental building block of the blogosphere.
- With crossblog discussions becoming more prominent, some blogging systems allowed blog readers to add **comments** to blog posts.

- In 2001 MovableType, another weblog publishing system, introduced **TrackBack**. With Trackback, a blogger can inform a blogging system linked to one of its posts, allowing the system to create a **backlink** to the post of the blogger and hence supporting crossblog discussions.
- **Blogrolls**, lists of blogs of interest to the blogger and often presented on the blog's main page, and FOAF (Friend Of A Friend) metadata (see the section about social networks) are two additional means used by bloggers in order to create communities.

Bloggers were early adopters of **web syndication**: Blogs usually provide content not only in HTML format, but also as a web feed in an XML format such as **RSS** and **Atom**. Web feeds greatly simplify the use of blog posts in other websites and applications. **Blog aggregators**, such as Bloglines and Newsgator, allow users to subscribe to an arbitrary number of feeds and provide their contents via an interface similar to that of email readers.

What motivates people to create and maintain blogs? Nardi et al. [12] interviewed 23 bloggers living in California and New York and discovered five main motivations: documenting one's life; providing commentary and opinions; expressing deeply felt emotions; articulating ideas through writing; and forming and maintaining community forums.

Business blogs are used by companies to promote products, publish company news, react to bad press, or make the company more transparent to their clients. For example, Microsoft encourages their employees to create blogs and talk about new developments, technical issues and advise, and (almost) everything else that comes to their minds—according to the Microsoft bloggers directories [3,4], there are currently more than 3000 bloggers among their employees. A major motivation for Microsoft is to improve customer relationships and the general image of the company by giving the people behind the company a voice.

A hot topic in the search community is **blog search**, with the old-timers like Technorati, Feedster, and BlogDigger challenged by new entrants from, for example, Google and Sphere. The blogosphere poses some **new challenges to search engines**:

- **Blogs in general change much faster than other webpages**, and providing access to the newest posts is essential. Hence crawling a website every couple of days, typical for current web-search crawlers, is not good enough. Blog search engines tap into the infrastructure of the blogosphere to receive update notifications (blog software usually pings one or more ping servers when a new post is added), and rely solely on crawling for blogs that do not support notifications.
- Once a new post has been found, the search engine has to analyze it and **update its search index as fast as possible** while not reducing result quality, Technorati claims that it can index most new posts within 10 minutes of their publication.
- **Metadata about blogs and individual posts** presents another challenge. Web feeds provide, for example, the name of the author of a post, its publication date, a summary, and ever more often tags categorizing the post (see the section about social tagging). Some feeds even contain geographical information about where the blogger lives, used by BlogDigger to allow limited search results to posts from blogs in a specific geographical area. Trackback links, blogrolls, and FOAF data (see the section about social networks) can also be exploited to improve search results.

Judging from the quality of results from current blog search engines, there is still a lot of room for improvement.

Also of interest, in respect to search, is the **combination of blog search with web search** and search in other media. Currently, blog search engines are specialized to blogs, while web search engines regard blogs as web pages without exploiting the structure of the blogosphere or being as up-to-date on new posts. A hybrid search engine could show a list of web pages in order of relevance to the search terms, and for each page, a short list of blog posts referencing the page. Titles of and tags assigned to the posts could then be used by the searcher to better assess the relevance of the web page to current information needed.

In addition to traditional text blogs, there are also **audioblogs** (usually non-music MP3 files; called podcasts if provided as a web feed) and **videoblogs**. In general, **multimedia blogging** is gaining ever

more traction with the widespread availability of broadband Internet access.

Social Tagging

Users add tags to items like bookmarks or pictures and receive feedback about the tags used within the community. Feedback then leads to a coalescence of tag vocabularies and improved browsing of the tagged items. Social tagging systems provide immediate value to a user, for example, by allowing organization of personal data and by making things easier to find by others.

The general idea behind **social tagging** (also known as free tagging, ethnoclassification, or creating folksonomies) is to allow users to tag items with arbitrary words for their own use and to browse the tags and corresponding items of other users. Social tagging was popularized by **del.icio.us**, a social bookmarking application introduced in 2003 (now part of Yahoo!), which allows users to collect bookmarks and organize them via tags. A **tag** is usually used to categorize (tag as category) or to describe (tag as keyword) an item and represents user-created metadata. Possible tags are not predefined, so each user can create a personal vocabulary of tags and organize items in the way

that best fits the user's point of view. Thus, each user creates a unique personal information space.

Social tagging applications allow users to discover items tagged by other users with a specific tag, tags used by other users to label a specific item, or users using the same tags or collecting similar items—this is called **social browsing** or **pivot browsing**. It is an effective way to find related information about a subject, but it also provides user feedback about the communal use of tags: When adding a new item, the user sees which tags were used by other users for the item, and which items were associated with the tags chosen (assuming that the item and chosen tags are already in the system). Feedback about the metadata supplied by a user is important for a coalescence of the personal information spaces, which greatly increases the value of the application for each user, because with a common and well known vocabulary, browsing and finding is much easier.

Folksonomies are not information-management's silver bullet. They are conglomerations of different information spaces, which are at best partly integrated. From the point of view of a specific user, this conglomeration looks like an inaccurate and incomplete categorization of items, caused

Yahoo! Analyst Day 2006 [15]

Yahoo!'s mission is to 'To create the world's most relevant, vital and trusted internet services for consumers and businesses', and Web 2.0 concepts play a major role in their strategy. Their business model depends on advertiser money and user fees; both requiring a large user base. With content, community, personalization, search, and audience reach seen as competitive criteria, Yahoo! focuses on engaging users to create content, reaching niche audiences, social searching (leverage knowledge and metadata from user-generated content to improve search), and different social applications (Yahoo! 360, flickr, del.icio.us, etc., with a service similar to YouTube in development) in order to remain a big player in the new advertising environment. This environment is characterized by its global scale (billions of users, tens of millions of advertisers, billions of unique offers), advertiser impact (campaign optimization, relevant audiences, engaging audiences instead of just reaching them), shifting roles of agencies, advertisers, publishers, and consumers, and a change driven by data. Also important for Yahoo!'s goals of widening its reach and deepening the engagement of users is the mobile market, the target of the recently launched Yahoo! Go [6].

Google Analyst Day 2006 [17]

Google's core business lies in search and advertising, so it is not surprising that search and end-user traffic, quality of advertising as perceived by end-users, and building new products and services for publishers of information are among its strategic priorities in 2006. Tools like Google Mail and Google Earth are not seen as part of the core business, but as means to increase the user base and with it the reach of advertisements. Dimensions in which Google plans to improve its search service are user provided metadata, via Google Co-op and Google Base, and vertical searches (searches restricted to a specific domain, for example, health care).

by the differences in vocabularies, use of tags, intentions, and understanding between users—as well as simple categorization errors. But this issue is lessened by the social browsing aspect described above. Further, folksonomies do not provide any explicit relationships between tags (e.g., there is no ‘is-specialization-of’ relationship that would define a hierarchical order of tags), thus finding items about a very specific subject by relying solely on tags is difficult.

Searching and browsing the communal information space is but one facet of social tagging systems, the other being to provide to the users the means for adding and organizing items for themselves. The latter presents an **immediate value for each user** attainable without learning about the community, thus offering a low-barrier entry-point to the system, which can then **slowly introduce the user to the communal metadata**.

Somewhat related to social tagging are the services **Google Base** and **Google Co-op** provide. Both allow users to create metadata about web resources, but instead of personal information management and communities, Google’s focus is on improving its search service.

- Google Base lets users describe things that are online or offline via an attribute-value set. This metadata is then used by Google Search.
- Google Co-op is targeted at **vertical-search providers**, that is, companies, groups, or individuals who have special expertise in certain domains and want to offer a domain-specific search. A medical doctor might, for example, use Google Co-op to label a collection of good health-care sites for his patients, which they can use to look up health issues. The patients would then subscribe to the metadata and would receive an additional OnBox on top of the Google search-result list that shows results from their doctor’s collection if they submit a health-related query. Popular metadata will bubble up from the community of its users to all users of Google Search, so each user, whether subscriber or not, will see search results with the metadata taken into account. Here, the **motivation for domain experts** to tag pages is to share information and/or increase traffic (and thus commerce) on their website.

Social tagging can also be of use in an enterprise setting:

- Millen et al. [13] created a prototype of a social bookmarking application for a company called Dogear and did an early evaluation within IBM Research Cambridge, which showed promise for this technology. Two design decisions were to support different visibility levels of a bookmark list, so employees could make some of their bookmarks visible only to, for example, their colleagues in the same department and provide a REST interface to the data stored within Dogear, so programmatic integration into other enterprise applications was easy.
- Another example is **BrainFile** [8], a tool for knowledge workers that provides personal desktop search functionality, means for multicriterial classification and tagging of documents, and different views on document collections.

Social Networks

People, in addition to websites, have become entities on the web. Social networks connect these people together on the basis of individual profiles. These are personalized to express the user’s interests and tastes, thoughts of the day, and values. The friend network allows people to link to their friends and traverse the network via these profiles, and to give comments, votes, and recommendations on their content published.

There are several **data formats** available to specify facts about people. Besides the well-established vCard standard for electronic business card profiles (used in major messaging applications) there are formats that extend such a personal profile by specifying social relationships.

- FOAF (Friend-of-a-Friend) is a description of relationships using the Semantic Web format RDFS and allows for rich and powerful expression of personal information and relationships.
- XFN is a simpler description format for relationships that annotates already established links to resources. Links are extended by an indication of a personal relationship to the person responsible for the resource (e.g., that person’s home page) linked to.

As microformats can be easily integrated into HTML pages, the social network structure becomes part of any webpage and can therefore be used by

Properties Market Share [18]

Hitwise Intelligence analyst Bill Trancer published the following data about the current market share of Google, Yahoo and MSN properties on his blog in May 2006. The percentages represented in the right hand column are the percentages of visits in respect to visits to that category.

Portal Property Rankings and Market Share by Vertical Week Ending 5/13/06

Rank	Property	Market Share in Category
Computers & Internet – Search Engines (2322 sites, 7.3% of all Internet visits)		
1	Google	47.40%
2	Yahoo! Search	16.00%
3	MSN Search	11.50%
Computers & Internet – Email Services (1089 sites, 9.3% of all Internet visits)		
1	Yahoo! Mail	42.40%
2	MSN Hotmail	22.90%
3	MySpace Mail	19.50%
4	Gmail	02.54%
News & Media (6080 sites, 3.4% of all Internet visits)		
1	Yahoo! News	6.30%
5	Google News	1.90%
*MSN news results appear within search.msn.com domain		
Business & Finance – Business Information (1030 sites, 0.57% of all Internet visits)		
1	Yahoo! Finance	34.90%
2	MSN Money Central	13.40%
40	Google Finance	0.29%
Travel – Maps (164 sites, 0.47% of all Internet visits)		
1	Mapquest	56.30%
2	Yahoo! Maps	20.50%
3	Google Maps	07.50%
4	MSN Virtual Earth	04.30%
5	Google Earth	02.00%

appropriate tools or browser extensions. The social networks defined by the standards mentioned here form a kind of overlay network parallel to content-based Web 2.0 applications and strengthen the personalized online experience of their users.

Community-organizing platforms

- MySpace is an online community that allows you to meet the friends of your friends. By creating a private community on MySpace, one can share photos, journals, and interests with a growing network of mutual friends. At the

core are profiles that are connected by links to friends on the system. Profiles are personalized to express an individual's interests and tastes, thoughts of the day, and values. The friend network allows people to link to their friends and traverse the network via these profiles. People can comment on each other's profiles or photos, which are typically displayed publicly. "Who knows who", "how are you connected", and "check them out" are the main motivations for the mostly under-25 crowd who use the service; similar services are Friendster and Orkut (by Google), which is invitation-only, i.e., users must be invited to join the community by someone who is already part of it.

- Yahoo!, as another big player, has become the holder of some excellent Web 2.0 cards (acquisitions in 2005 include Flickr for photo sharing, De.licio.us for bookmarking, Blo.gs for pinging), started Yahoo!360 as an integrated environment of blogs, photos, music, voice chat, messaging, and reviews of favorite books and movies. People can connect with friends and invite them to their Yahoo! Groups.

However, the social software tools of the major players may not convince the entire community of web users. Important aspects for the enthusiastic user of community services are individuality, attention, lack of ads, exchange of less than polite messages, rather than mere recommendations of books and similar items. The standard tools do not match all of these criteria, but are considered to be sufficient for beginners. There are also some **niche segments** for social networks:

- LinkedIn is probably one of the best examples of serving a (quite large) niche market, namely **business people**. It allows the user to create a profile summarizing the users' professional accomplishments.
- Qype is a community platform that combines social networking with recommendations for sites, shops, etc., for specific locations. The idea behind Qype is that a **recommendation** is as worthwhile as the trust people have in the person making the recommendation; a valuable tip usually comes from a friend.
- CivicSpace follows a somewhat different approach to community management. It features an open-source **civic organizing platform** that provides a content management toolkit for organizing and mobilizing

communities through the web. It allows people to build communities online and offline that can communicate effectively, act collectively, and coordinate coherently with a network of other related organizations. CivicSpace enables bottom-up, people-powered campaigns that support grassroots democracy.

Currently, there are many social network services, but they remain disconnected from each other. So, a person has accounts on Orkut and LinkedIn, but the profiles cannot be connected, automatically extending the friend-of-a-friend space. Allowing the evolution to a semantic social network structure. This is one of the issues tackled by Identity 2.0 (see below).

Mobile Web 2.0

Access and integrated usage of Web 2.0 services from a restricted device (e.g., Smartphone) carried by the user and capable of media consumption as well as media production. The mobile context is the key to the enhancement of the services, e.g., the user's location, devices in the neighborhood, places recently visited as additional tags or filters to data.

Concerning mobile access to Web 2.0 services, the peculiarity is that the services are accessed from a restricted device (a Cellphone or Smartphone) carried by the user and capable of media consumption as well as media production.

Blogging services are one of the first mobile Web 2.0 applications. Moblogs (mobile blogs) are services that can be edited by their owner on the move, and everyone can browse these blogs from Internet or mobile devices. Updates to the blogs are usually sent by MMS or email.

- Using services such as moblogUK or FoneBlog, people can take photos, shoot videos, or

Roles	Key success factors	Players
Content provider and storage	▪ Integrate personal content ▪ Enable data enrichment ▪ Data quality management	▪ Google, Yahoo! ▪ YouTube
Trust provider	▪ Data of many users ▪ Become fast a quasi-standard ▪ Reputation as trustworthy on organizational level	▪ VeriSign ▪ Ebay (Recommendation) ▪ Microsoft (Passport)
Community provider	▪ Personalization ▪ Attractive community ▪ Support of information sharing	▪ MySpace ▪ Yahoo! ▪ Flickr, YouTube
Infrastructure provider	▪ Platform for platforms ▪ Authentication, Authorization, ▪ Accounting; Billing	▪ Main ISPs ▪ PayPal ▪ BitTorrent
Service composition provider	▪ Being fast, flexible	▪ Project "Smart Web" ▪ MSN, Yahoo!
Advertisement	▪ Standardization ▪ Personalization ▪ Critical size	▪ Google, Yahoo! ▪ MSN
Search Service	▪ Enabling of... - Tagging - Recommendations	▪ Flickr ▪ Google, Yahoo!
Advertisement	▪ Integrate mobile aspects ▪ Enable interaction interfaces ▪ Multimodal user interfaces	▪ Browser: Internet Explorer, Firefox ▪ Mobile devices: Nokia, ...

[Roles and players in the Web 2.0 domain]

capture audio with their camera phone and then email them to the service directly from the phone. The media file is then put online in an individual mobile multimedia blog, instantly sharing the moment and is immediately viewable using WAP or a standard web browser.

Social tagging in a mobile context extends social tagging as described above by adding **context information** to tags. For example, if the user tags a photo just taken with his mobile phone, the tag itself is metadata explicitly entered by the user, i.e., a traditional „web“ tag, but its context contains implicitly connected information about when the image was captured, the user's location, devices in the Bluetooth neighborhood, places recently visited, etc. Context-based tags will be enablers for enhanced social network services.

- Dodgeball (a Google company) is one of the first mobile social software tools. It sends registered participants text messages when other participants (or their friends) are nearby.
- Socialight is a more recent tool that allows users put virtual post-it notes (they call them StickyShadows) in places where other users can see them. A post-it note is made up of media, such as text and pictures, and information about who is allowed to see it, and when and where it's available.
- StreetHive combines the virtual post-it idea with social networking. It is a mobile social network that lets friends locate one another, send messages, and share location-tagged information right from their mobile devices.

Real world tagging will be further enhanced by **ubiquitous technologies** like RFID, color codes or barcodes that support a much finer granularity of resolution and interaction than standard GPS location coordinates.

Branded devices for social networking platforms also occurred recently.

- Helio is a branded device and mobile service that was launched around the well-established MySpace platform. The phones are geared towards the youth market and have integrated MySpace access, video messaging, news feeds, and some other multimedia features like game services. MySpace Helio members will also receive extra storage for their MySpace pictures, and be able to update their MySpace pics directly from their Helio handsets. Helio

is a high-priced device and service of specific carriers.

- A **mobile mashup** is a mashup service (see above) that runs on a mobile device and presents information that may be more relevant when people are not at home sitting in front of their PC. Further, it integrates services that are only accessible on the mobile. It corresponds to the successful, AJAX-based online software using a **web browser platform** to provide applications comparable to rich, responsive native software. Ajax has also entered the mobile phone-based.
- The mobile Opera platform is an AJAX framework that is fully designed for the mobile sector. In addition to the asynchronous XML-data exchange with a server, the native device APIs are wrapped by a set of Javascript APIs, and thus the developer gets access to the low level functions from the browser, i.e., the telephony API, phone book, text messages, etc. Therefore, services based on Google Maps can be used on the mobile, potentially extended by mobile device functionality.

However, the mobile browser is not the only platform for accessing mobile mashups; JAVA-based **J2ME clients** are also relevant players in this domain.

- MobileGlu supports access to Flickr and MoblogUK photos as well as events from upcoming.org, bookmarks from del.icio.us, blogger.com accounts, and RSS feeds from within a single J2ME application. The mobileGlu system optimizes all online data for the mobile screen and supports local photo snapping in order to update the favorite photo sharing sites.

Regarding these and other developments in mobile web development, it is obvious that the mobile web will converge at some not too distant point in the future with the rest of the web. Although the mobile web today will have to overcome some difficulties, which the web experienced 10 years ago, such as slow access and lack of interoperability, the situation is much better with respect to other dimensions: there are many potentially connected users and lots of potential content.

Identity 2.0

A user-centered approach to identity management. Instead of identity data being stored in enterprise-controlled data silos, the user is in possession and control of his identity data, showing as much or as little of it as required (like a driver's license). The jury is still out on which of the several new identity approaches will gain a critical mass of users.

With users at the center of Web 2.0, digital identity resurfaces as an important topic. **Digital identity** [10] is the digital representation of claims by one party about itself or another digital entity, where claims are assessments about attributes of an entity. For instance, attributes can be a user's name, age, credit-card information, and reputation – what others say about the user. Identity is not only important for e-commerce services, but wherever one party has to trust another, from providing access to web services to deciding whether to trust news posted on a blog.

Currently, **websites are at the center** of digital-identity management. Users have profiles on amazon.com, ebay, and Google, but they cannot transfer their identity from one site to another, nor is there any connection between their different profiles. All of which raises the question: who owns the identity data in the first place?

The phrase **Identity 2.0** was introduced by Dick Hardt [25] to denote, in reference to the user-centered Web 2.0, a **user-centered approach to identity management**.

Hardt's idea [22] is to provide a kind of web driver's license that can be shown by users to identify themselves to websites and web-services. Analogous to a physical license, the web license is kept by the user and the user knows exactly what data it contains. And just like a physical license, the web license is verified and certified by a trusted identity-certification authority, but the authority does not take an active role in the identification process between two parties.

This approach corresponds to what companies are doing today when they identify themselves to a client's browser before establishing a secure connection to receive, for example, credit card information.

Another approach to user-centered, transparent identity management is creating an open (user-controlled rather than enterprise-controlled) **infrastructure between existing identity silos**, as

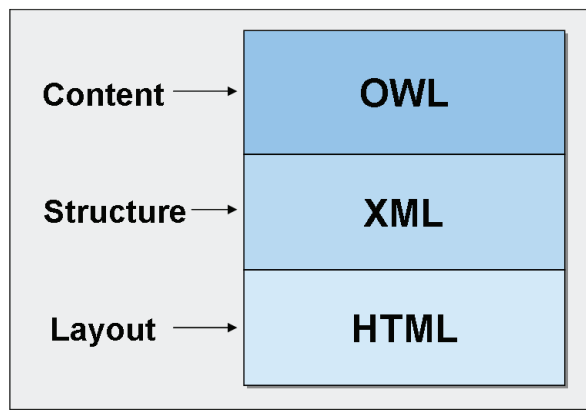
exemplified by the Higgins project [spwiki.editme.com/higgins].

Semantic Web

As Tim Berners-Lee, Jim Hendler, and Ora Lassila stated in their article in the Scientific American in 2001 [20] "The Semantic Web is not a separate Web, but an extension of the current one, in which information is given well-defined meaning, better enabling computers and people to work in cooperation." In order to do so, we need data about data, so called **metadata**, which may express meaning and which may be interchanged among computers. If metadata is stored in a standardized way, then machines may use it to both consume and produce data on the web. **Service agents** may use this information in order to search, filter, combine, structure, and rearrange it in new and prospective ways to assist us in solving tasks, preparing events, or just planning free time.

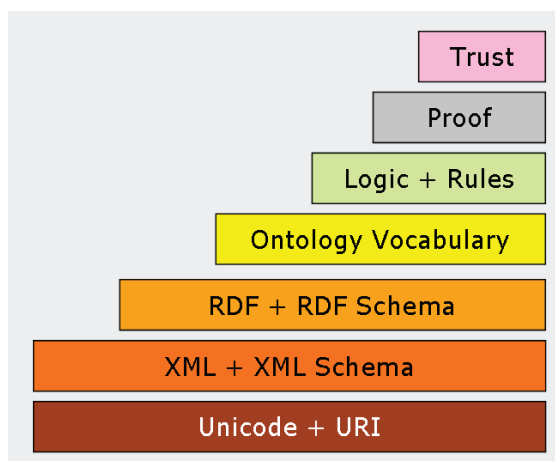
The core technology elements for the Semantic Web are markup languages with a formal syntax and semantics that provide a standardized concept for describing digital contents in the form of an ontology. Semantic markup languages like OWL (Web Ontology Language) allow for the world-wide distribution and shared usage of ontologies in the WWW. The semantic annotation with metadata forms the third layer of a document description, above the annotation layer of the structure (XML) and the annotation layer of the layout (HTML), see Figure 4.

	Web 1.0	Semantic Web
Paradigms	▪ syntactic markup ▪ web search ▪ inert data ▪ human-intelligible data ▪ information representation	▪ semantic markup ▪ question answering ▪ reusable data ▪ human- and machine-intelligible data ▪ knowledge representation
Key Success Factors	▪ simplicity ▪ open for everyone	▪ critical mass of semantically annotated data ▪ common ontologies ▪ return of investment in the creation of annotations by valuable semantic services



[Figure 4: Three Layers of Webpage Annotations]

These three layers of webpage annotations lead to Tim Berners-Lee's Semantic Web layer cake (shown in Figure 5). It arranges a mixture of emerging standards and concepts into the Semantic Web hierarchy. On the bottom layer Unicode and URI (Uniform Resource Identifier) form the base standard to denote resources in general. On the second layer, we find XML which is used to define the interchange format of underlying data models with a purely syntactic and structural nature. On the third layer, RDF (Resource Description Framework) and its schema form the basic and most widespread Semantic Web representation language that is based on triples. The next layer is made up of ontology vocabularies, which will be discussed shortly. The next layer brings functionality to the Semantic Web: with logic and rules, the systems are able to draw inferences and make decisions. However, in order to rely on these inferences, proof is necessary, which is the sixth layer of our cake. Finally, comes trust. The Semantic Web does not assert that all statements found on the web are true. However, all statements on the Web occur in some context and each application needs to evaluate the trustworthiness with the help of this context.



[Figure 5: The Semantic Web Layer Cake]

Furthermore, **security** and **privacy** are important for the final goal of a semantic „Web of Trust“.

Inferences (the layer of logic and rules) can improve search processes, while methods of machine learning, data mining and text mining, reduce the manual effort of creating and maintaining ontologies. A variety of software tools exist that support development, editing, evaluation, merging, and validation of ontologies. The immense problem of information overload can be treated with the Semantic Web. In the future, information will be provided on demand in a situation-aware and personalized fashion. Through the use of ontologies, the hyperlinks can be semantically classified, thus enabling semantic surfing and semantic-driven retrieval.

Ontologies

Ontologies define a vocabulary that can be used to create data models of a domain. Agents that commit to an ontology use the defined vocabulary in a way consistent with the ontology, and thus can exchange information and reason about objects in the modeled domain. Annotating data with ontological concepts therefore simplifies its automatic processing and provides a software agent with background knowledge, in the form of the ontology itself, to infer facts not explicitly part of the annotated data.

In computer science, ontologies are data models that represent a domain (for example, the domain of books in a bookstore) and are used to reason about objects in that domain. Ontologies represent **concepts** as well as **relationships between concepts** (for instance, there might be the concepts book and price in the bookstore ontology, with the relationship book ‚has price‘ price), and thus define a common vocabulary for talking about the domain. If an agent (either human or machine) **commits to an ontology**, it uses the concepts defined in the ontology to describe a domain instance or uses such a description in a way that is consistent with the ontology.

For example, assume there is a bookstore ontology that models a taxonomy of book types, such as Fictional Book with subconcepts Horror Book, Fantasy Book, and Sci-Fi Book. All book types are subconcepts of Book, which has relationships has ISBN and has price to the ontological concepts ISBN and Price, respectively — if this sounds similar to a database model, it is, but ontologies are more

expressive than the relational model. If a bookstore owner models its inventory of books according to the bookstore ontology, he assigns one or more book types to each of the books he sells, say Fantasy Book to Harry Potter, and then fills in the relationships has price and has ISBN required by the ontology.

A comparison-shopping agent committed to the bookstore ontology can then use this data to look up books in the store's inventory that are, for example, Fictional Books and cost less than a given price.

Because the bookstore owner and agent are committed to the same ontology, it is easy for them to exchange data in a meaningful way: both know that a Fantasy Book is a Fictional Book and that a Price is an amount of money the seller wants for a book. Without a common ontology, the agent would have to use brittle heuristics to extract the required data from a pure HTML representation of a book description.

Semantic Web Services

The development of Semantic Web Services results in a Web where programs act as autonomous agents to become the producers and consumers of information and enable automation of transactions. Semantic web service technology enriches the classical, purely syntactic web service markup languages, such as WSDL, with a semantic description of the contents of a service and its input and output parameters. Semantic markup of web services is the basis for the support of automatic discovery, selection, composition, monitoring, and invocation of services by intelligent software agents.

Recent progress in mobile broadband communication and Semantic Web technology is facilitating innovative Internet services that provide advanced **personalization** and **localization** features. The goal of the SmartWeb [www.smartweb-project.org] project is to lay the foundations for multimodal user interfaces to distributed and composable semantic Web services on mobile devices. The SmartWeb consortium brings together experts from various research communities: mobile services, intelligent user interfaces, language and speech technology, information extraction, and Semantic Web technologies. The appeal of being able to ask a question to a mobile Internet terminal and receive an answer immediately has been renewed by the broad availability of information on the Web. Ideally, a spoken dialogue system that uses the Web as its knowledge base would be able to answer a broad range of questions. Practically speaking, the size and dynamic nature of the Web, and the fact that the content of most web pages is encoded in natural language, makes this an extremely difficult task. However, SmartWeb exploits the machine-understandable content of semantic Web pages for intelligent question-answering as the next step beyond today's search engines.

SmartWeb provides a context-aware user interface, so that it can support the user in different roles, e.g., car driver, motorcycle rider, pedestrian or sports spectator. One realized example of SmartWeb is a personal guide to the 2006 FIFA World Cup in Germany, which provides mobile infotainment services to soccer fans, anywhere and anytime. Another demonstration of SmartWeb technology is based on P2P communication between a car and a motorcycle. When the car's sensors detect



The SmartWeb car version provides domain-specific information to the driver accessible by gesture and voice, e.g. to get answers for:

- Who has scored most goals at the FIFA World Cup?
- Where can I get the lowest price for diesel?
- Where are speed traps located today?

[Figure 6: The SmartWeb-Car demonstrator at CeBIT 2006 with Professor Wahlster]

hydroplaning, a passing motorcyclist is warned via SmartWeb „Hydroplaning danger in 200 meters!“. The motorcyclist can interact with SmartWeb through speech and haptic feedback; the car driver can use speech and gestures at the user interface. SmartWeb is a joint research project (consortium leader DFKI) of several academic and industrial partners (e.g., Deutsche Telekom) and is funded by the German Federal Ministry of Education and Research (BMBF) as a lead innovation project in the grant program “IT 2006”.

One core effort of systems like SmartWeb is the development of semantic Web services. The composition of more than one service is required if a certain goal cannot be achieved by invoking a single service but by chaining different services together. Composition and invocation has to be transparent for the user within an interactive scenario, so it must be done automatically. The representation of results achieved through web service execution also has to be semantic, in order to enable a dialog system for flexible and multimodal result presentation.

Semantic Search

Semantic Search attempts to augment and improve traditional search approaches by using data from the Semantic Web. This means that ontology-based metadata about documents acting as intrinsic descriptors, as well as the involvement of extrinsic descriptions, and by incorporating the inquirer's context of information, will allow for more complex queries to be asked, and more specific answers to be received.

Today, the pressure on organizations to **learn faster than their competitors** is more important for survival than ever before. The information to be processed rapidly increases and can hardly be mastered with respect to the collective benefit of an organization. The increasing complexity of data volume in company networks worldwide reveals the common dilemma:

- How to prioritize and avoid information overload.
- How to exchange the interesting bits for improved collaboration.
- How to select the right tools which allow knowledge sharing with obvious advantages for the participants.

Making information accessible at any time and from any place is no longer a problem. Without knowing if certain information is available and where it is stored, data inquiry often results in a cumbersome chain of queries whose success is highly dependent on chance.

Supporting information needs is a difficult task. A query system usually doesn't know more than a few words and some common constraining categories that can be used in the data inquiry task. The role or the individual interests of the inquirer as well as the context of the information needs are not visible for the query system (e.g., the actual task or a related process). The information needs are simultaneously spontaneous, specific, and subjective. At the same time, storing information or documents in the right place is a difficult task. The inherent “who”, “what”, “where”, and “when” facilitates a perspective query driven by **situational demands** with different and ambiguous terminologies. Our background, interests, roles, and tasks are the basic reasons that we consider documents as belonging to subjective categories, such as individually relevant events, people, topics, or projects relevant to each user.

Today, most search engines, such as Google or Yahoo!, allow an advanced, full text search with subject directory support and a ranking based on popularity. They allow searching not only on web pages but offer desktop or business search as well. However, when talking about documents, the intrinsic text features are not sufficient and additional metadata as extrinsic descriptors are needed in order to allow a perspective “who”, “what”, “where”, and “when” given by the subjective categories of a user.

One common technique is tagging, which allows users to describe the world in their own terms as taxonomies (see the section on “Social Tagging” above). Tagging may be seen as a promising way to organize documents, but tags, like taxonomies, are all about finding data. Tagging can deal well with attributes but cannot provide values or intrinsic relationships among terms. Moreover, it remains, in part, a **labour intensive annotation** process with problems in scaling up to the full free-text Web. Thus tagging cannot be seen as a competitive technique for search engines but complementary to them.

Modern advanced Web technologies and personalized information retrieval can make search peer-to-peer and semantic-based, with

synonyms automatically generated and updated, and documents are filled with metadata about the situational “w-questions”. These correspond at least to a shared vocabulary of semantic subjects being partially maintained by the entities of the native structures, such as file folders, e-mail, and Outlook contacts. Such a consideration is of interest especially if the document tags might be produced automatically either on intrinsic conceptual keywords, by the folder names, or by the existing attributes of the office applications. Semantic search also implies the combined search within informal and formal sources. This may be accomplished by allowing different levels of complexity. Besides the traditional keyword-search, more sophisticated techniques can be applied to the ranked results, such as heuristics in terms of business rules using ontology-based metadata from the “who”, the “what”, the “where”, and the “when” categories. For example:

- IF an event name or acronym was found, THEN include the names of the participants, or
- IF documents were found, THEN identify my personal categories they may be classified with, rank, and/or cluster documents according to those categories”.

However, the classical keyword search may also be enriched by specifying the respective resources, e.g., „who:Dengel where:DFKI“, or more complex combinations with relations, such as „who <workingFor> where:T-Labs AND who <emailedMe> when:2006“ may be used.

Semantic Grid

Semantic technology allows machines to cooperate more easily in grid environments. The concept of Semantic Grid is especially interesting for system administrators managing computer grids, since they want to be able to deploy applications without worrying about which individual servers or storage devices are involved. Semantic technologies can describe grid resources such that each device in the network can understand what's available, negotiate for resources, and execute application logic.

The Semantic Grid [24] is an extension of the current Grid in which information and services are given well-defined meaning, better enabling computers and people to work in cooperation, i.e., the application of the principles of the Semantic Web to the grid environment (read about Oracle's

position in [19]). There are really only two aspects of the Semantic Grid:

- The discovery of available resources for processing the data and the ability to integrate the data. The discovery side of the Semantic Web is designed to make it easier for grids to be discovered across the Internet. A detailed definition of the capabilities required to make use of the grid helps grid users to **reuse existing resources and technology** for their grid requirements. The complexity is in the way the grid services and capabilities are described. This is where the nature of the Semantic Grid, with the heavy classification and description of abilities and facilities, will be employed to more easily determine what facilities a specific grid can provide.
- The second main aspect of the Semantic Grid is the way in which users and applications can link and cooperate with the information stored and available within a grid. For example, resource grids (those sharing disk and storage space, instead of providing CPU power) use grid technology (Web services, security, etc.) to provide links and connectivity between information: a semantic grid component that stores photos, combined with a semantic grid that stores video material, allows the user to find videos related to photos of a subject searched for.

Within more complex scenarios, multiple grids can be used to solve complex calculations by feeding results generated by one grid into the other. Such processes are possible since the data and its structure are known and usable by the individual grids.

Queries, Inferences, and Rules

Query and inferencing languages are necessary to support machines that process data in an intelligent way. The available systems range from logic programming based SQL-like query languages to description logic based reasoning systems able to prove specific formal properties of rule systems. There are also markup languages for rules that support the interchange and reuse of rule systems over the Web.

The main goal of the Semantic Web is to allow machines to process data in an intelligent way. Several **processing languages** have been defined

so far, ranging from support for simple queries to complex inferencing. For example, SPARQL is a query language for RDF documents that is accompanied by a Web protocol allowing applications to send queries and receive answers across the web. The syntax of SPARQL resembles SQL database queries, but apart from SELECT queries returning variable bindings, SPARQL also has CONSTRUCT queries (which return RDF graphs), DESCRIBE queries (returning RDF graphs that describe the resources found), and ASK queries (returning a boolean indicating whether a query pattern matches or not). Since SPARQL does not support the definition of complex rules it can only be used to retrieve existing knowledge, not to derive new information through inferencing.

One of the first **inferencing languages** which has been developed for the Semantic Web was the logic programming based SiLRI (Simple Logic-based RDF Interpreter) [3], which has both an open-source and a commercial successor: TRIPLE [5], started as a collaboration between DFKI and Stanford University, and OntoBroker [www.ontoprise.com]. Rule systems of this kind are best used as a data and ontology transformation and mapping language, e.g., supporting mediator and matchmaking services for Semantic Web data and services.

For ontology languages based on description logics (DL) like OWL Lite and OWL DL, different kinds of reasoning are supported. Unlike the above mentioned rule systems DL-based reasoners provide much more basic inferencing capabilities, but with much cleaner semantics. These capabilities are usually consistency checking (determines if some class definition is inconsistent), finding implicit subclass relationships, finding synonyms (of classes and instances), and determining which class(es) a given instance belongs to. Just as in the case of rule languages, both open-source implementations like Pellet [26] and commercial systems like RACER [www.racer-systems.com] are available.

Apart from concrete rule and inference systems, markup languages that allow the exchange of rules are foreseen for the Semantic Web. One such existing and widely used rule markup language is RuleML, which was started at DFKI. RuleML defines an XML syntax for various kinds of rules, including forward (bottom-up) and backward (top-down) rules. The importance of rule exchange has also been realized by the W3C, which started a working group on this topic (Rule Interchange Format RIF), taking RuleML as input. The result of the RIF working group will be an industry standard allowing rules (which according to the charter will be semantically compatible to RDF and OWL) to be exchanged between agents/services on the web

Roles	Key success factors	Players
Ontology Provider	- Domain Knowledge - Reputation	- Cycorp - Dublin Core Metadata Initiative - IEEE
Inference Engine Provider	- Efficiency - Ability to process large datasets	- Cycorp - Ontoprise
Development Tool Provider	- Standard compliance - Support for knowledge-engineering processes	- Cycorp - Ontoprise
Semantic Service Provider	- Standard compliance - High-quality annotations of service and data	
Semantic Service Composition Provider	- Automatic service composition - Effective and efficient combination of services	Project "SmartWeb"
Semantic Search Service	- High-quality information extraction - Automatic semantic annotation - Integration of user-generated metadata	Project "Quaero"

[Roles and players in the Semantic Web domain]

using rule systems from different vendors, thus making knowledge encoded in rules reusable. This is of great importance since the acquisition of rules is often very expensive.

The RIF Use Cases document [27] enumerates several interesting scenarios showing how rules will be exchanged in industry, e.g., „Negotiating eBusiness Contracts Across Rule Platforms“, which „illustrates a fundamental use of the RIF: to supply a vendor-neutral representation of rules, so that rule-system developers and stakeholders can do their work and make product investments without concern about vendor lock-in, and in particular without concern that a business partner does not have the same vendor technology. It also illustrates the fact that the RIF can be used to foster collaborative work. Each developer and stakeholder can make a contribution to the joint effort without being forced to adopt the tools or platforms of the other contributors.“

Semantic Desktop

A Semantic Desktop is a virtual device in which an individual stores all personal digital information like documents, multimedia, and messages that are interpreted as Semantic Web resources and can be accessed as such. Ontologies allow the user to express personal mental models and form the semantic glue interconnecting information and systems. The Semantic Desktop is an enlarged supplement to the user's memory.

The Semantic Desktop [7] specifies a driving paradigm for desktop computing on the Semantic Web. Based on the needs and expectations of users today the software industry will evolve to a future way of computing, semantic desktop computing. The main task at hand is to transfer the Semantic Web to desktop computers, and this transfer will not only consist of the technology, but also of the philosophy and the people involved. Developers that today concentrate on services for the Semantic Web will need a complete RDF and ontology based environment to create applications on desktop computers. End users will benefit from these applications, as they integrate and also communicate better than (based on ontologies and Semantic Web standards) today's desktop applications.

A Semantic Desktop could be defined as a device in which an individual stores all personal digital information like documents, multimedia, and

messages. These are interpreted as Semantic Web resources, each identified by a Uniform Resource Identifier (URI) and accessible and queryable as RDF structure. Resources from the web can be stored and authored content can be shared with others. Ontologies allow the user to express personal mental models and form the semantic glue interconnecting information and systems. Applications respect this and store, read, and communicate via ontologies and Semantic Web protocols. The Semantic Desktop is an enlarged supplement to the user's memory. The European IST Project NEPOMUK (led by DFKI) bundles academic, industrial, and open source community efforts to create a new technical and methodological platform: the **Social Semantic Desktop**. It enables users to build, maintain, and employ inter-workspace relations in large scale distributed scenarios. Thereby, new knowledge can be articulated in semantic structures, be connected with existing information items on the local and remote desktops, and knowledge items and their metadata can be shared spontaneously without a central infrastructure.

Today, desktop search engines are a major market, and tools like Google Desktop, Apple Spotlight, Yahoo! Desktop Search or Microsoft Windows Desktop Search are products in a competitive market. The features provided in these free tools are satisfying to most users, but far behind the state of the art commercial tools available, like Autonomy or Convera, and current proposals in research papers.

Gnowsis [28] is an Open Source project that realizes a Semantic Desktop environment. It allows users to use desktop computers like a small personal Semantic Web. Linking documents across applications and browsing through the personal information space is now possible. Emails, documents, addresses, photos, appointments that have been spread in the local data jungle can be linked conveniently, weaving a personal Semantic Web. Data structures are not changed, and existing applications are extended and not replaced. Programmers can build Semantic Web and knowledge management applications on top of Gnowsis. Structured data from common desktop applications (like MS-Outlook) can be accessed in a simple web protocol, allowing developers and researchers to leverage this information. The core of the Gnowsis is a desktop web server which integrates data from different standard applications through standardized adapters. These reusable

Enriching flat data by using Web 2.0 services

Matt Biddulph presents, on his blog, an example of how to automatically add directed links to a flat set of strings with the intent of simplifying navigation among the represented concepts [21]. The strings in his example are names of politicians, and the links to be discovered are to represent relationships between those politicians. The first step in the process is to associate the strings with URLs pointing to web resources that describe the concepts represented by the strings. To that end, Wikipedia entries are chosen using Yahoo!'s search engine API to find the corresponding entry for each string. Now, in order to discover relationships, the contents of the Wikipedia entries are sent to Yahoo!'s term extraction service, which returns 'significant words or phrases' from within the supplied text. In the case of politicians, significant words in the Wikipedia entries contain the names of other, related politicians. Thus, by comparing the words extracted from the Wikipedia entry describing a politician with each of the other names in the initial set of strings, one can discover (most of) the relationships looked for.

This method of identifying relationships between concepts by using externally available web-services and web content does neither lead to semantic annotations nor does it scale to large datasets, but it points out a nice way of automatically enriching existing flat data.

adapters implement the application specific interfaces and publish them as semantic Web services based on RDF. Common applications like MS-Outlook are extended through plugins, which integrate the Gnowsif functionality into existing user interfaces. This allows the user to define the necessary links between resources such as a photo and contact information of a person. Links are stored and then used to navigate from one resource to another, across application borders.

Work on the realization of the Semantic Desktop is going forward with the goal to improve the way users search, find, and experience information, e.g., by integrating sophisticated information retrieval techniques like automatic clustering to the system.

Web 3.0: Semantic Web meets Web 2.0

As introduced earlier, the vision of Web 3.0 is to synergistically integrate both the concept of Web 2.0 and the new technologies of the Semantic Web. In order to reduce the tremendous costs for building a full Semantic Web, the lightweight and community-based methods of Web 2.0 must be used, so that every web user can contribute a piece of meaning to the emerging Semantic Web. The drive which evolves from the provision of a platform on which others can build, use, and combine applications signals the way in which Semantic Web applications should be allowed to evolve. The social software evolution and the widespread usage of Web 2.0 can be used

as a catalyst for the establishment of semantic technologies. Today, the complexity and lack of a network effect by a huge community prevent the Semantic Web from quickly flooding the public Internet. The necessary semantic tagging, extracting, classifying, and organizing work is labor intensive, and there are an insufficient number of simple approaches which immediately satisfy the self-interest of the users and show the benefit of using the Semantic Web technologies. At least semiautomatic methods are necessary means to overcome the Semantic Gap. Web 2.0 builds on simple standards and simple „loosely-coupled“ integrations. Applications like mashups and Web 2.0 sites about social bookmarking are drivers but not complete solutions to knowledge creation and sharing. But they represent what is fun and cool and people invest immense amounts of (private) time in adaptation and usage of simple tools. So, learning semantic structures from communities is a worthwhile approach to bridge the gap. This way, **semantic blogging** [14] is an attempt to combine blogging with the Semantic Web by adding more formal metadata (e.g., concepts from an ontology commonly agreed on in the blogger's community) to posts. This would allow use of inferencing to improve, for example, blog search. Given a suitable ontology stating that blogging is a subtopic of Web2.0, the search application could imply that a post about blogging is also about Web 2.0 and thus offer it to a searcher looking for Web 2.0 posts. Semantic blogging struggles with the same issues as the Semantic Web in general. In particular, there has to be a common ontological approach well understood by bloggers and people searching for

posts, and adding metadata to posts has to be very easy, yet lead to accurate annotations.

The Social Semantic Desktop approach, as followed by projects like Nepomuk, (see the section on Semantic Desktop above) combines both social aspects of Web 2.0 (supports the interconnection and exchange of data with other desktops and their users) and Semantic Web principles (information items obtain a well-defined meaning). Nepomuk will provide tools for social relationship building and knowledge exchange within social communities. These tools will provide semantically rich recommendations, which allow members of a community to exchange not only documents and other isolated information chunks, but all relevant background information about their context and the participating community as well. Furthermore, techniques for distributed search and storage of information will be developed based on scalable extensions of the distributed hash table technology and grid infrastructures. This will allow efficient searches over distributed information resources and provide a shared knowledge pool within a community. The Social Semantic Desktop also integrates datasources for popular tagging websites, such as del.icio.us or flickr, i.e., this allows the user to reuse personal "folksonomy" tag spaces with the added advantage of converting what had been a flat tag list into a first-class ontology. Semantic Wikis will also complement the Social Semantic Desktop. As Wikis are Web 2.0 collaborative writing tools par excellence, Semantic Wikis additionally provide an underlying model of the knowledge described in its pages. Now, if a user can enter both plain text and formalized ontology-based knowledge, he needs some interactive support. For example, autocompletion will be provided on the RDF-level, when typing „Paul knows“, with „Paul“ being a „foaf:person“, the system automatically proposes a list of „foaf:persons“ defined in the wiki to complete the RDF triple, as only „foaf:persons“ are allowed as range of „foaf:knows“. For non-formalized texts, new wiki texts will be created providing links to other wiki pages through autocompletion, suggesting relevant page names.

What does this mean for the future? There are lots of ways in which our computers can use our „individualized“ web content when they can understand it. Considering that forthcoming mobile phones may gradually take over the role of the personal digital assistant. They will contain a flash

memory card and may store all kinds of information, such as maps, bus timetables, favorite recipes, or information about food allergies. Thus, mobile phones can be used for a wide variety of purposes of which have not even been considered yet.

Search Engines like Google today are not able to understand what is meant behind the query „Ristorante Firenze“. There is no way to distinguish whether we are talking about a restaurant in Florence, Italy, or whether we want the address of our favorite place to get Italian food. Moreover, it is not able to interpret the prices behind the numbers shown on the corresponding website. In the future, mobile phones will be equipped with an appropriate interface to arrange a lunch date with a business partner, automatically negotiating an adequate time by accessing both calendars. Location-based semantic services will allow people to share with other people. Systems will let people register, and then when they are at a particular venue, i.e., at a restaurant, „hook up“ with other guests to get reviews about the venue or about the dishes they just had. However, there will not only be the obvious services like booking restaurant tables or airline tickets, accessing the departure time of the next train, or showing vacant hotel rooms in town. In the future, mobile phones will help maintain our health. They will hold all relevant, medical details which can easily be accessed in emergency situations via Medicare services. There will be automatic interfaces with our doctor's system during each visit. Based on the diagnosis, the doctor will use Web services to route our prescription to the pharmacy, and our phone will identify us when it is time to pick it up. Besides this, mobile phones will be used as a payment method to buy tickets for busses, trains, and airlines. They will even be used to purchase lunch at a restaurant. The next generation phone will feature a credit payment system so that credit card companies like MasterCard, Visa, or American Express will adopt their payment methods as integrated services. Future mobile phones will be aware of our current position. They will be able to interpret time and thus furnish us with all information we have requested or we need. Mobile phones might even replace identity cards allowing, in combination with thin electronic paper technology, access to our personal workplace from anywhere in the world or replace the existing keying systems for your car and your house. However, new challenges will emerge with sophisticated authentication systems.

On the present Web, the lack of semantics prevents computer systems from being able to interpret Web information automatically. While the Semantic Web is still largely discussed within research groups, and enterprise applications have begun to incorporate semantic technologies, Web 2.0 has brought a significant amount of informal knowledge onto the Web based around users as content providers, tagging to form loose “folksonomies” and open APIs to allow reuse of data in different settings. A big challenge will be to structure and manage the user-generated annotation data in a Semantic Web friendly way. This is certainly a task that the big guys, like Google and Yahoo!, have to deal with in order to reach an appropriate penetration of the Web. Certainly, an important aspect of the future Web is that new data can be easily and readily annotated by metadata; but what about the interpretation of already existing data? A primarily automatic annotation process is necessary to make the data accessible. A further challenge for the future will be the extension of coverage of the semantic technologies to the so-called Deep Web, i.e., the thousands of databases and archives of partly analog data, currently unreachable by the crawlers of the search engines. The Deep Web holds an enormous potential of high quality text and multimedia information, which is worth accessing.

As we have shown in this report, many facets of the emerging Web 3.0 are already available, others will arise and new technology will amaze us. The ultimate worldwide knowledge infrastructure of the Web 3.0 will boost network traffic, generate new revenue streams for providers of paid content and web services, and ensure full interoperability of web-based applications.

Bibliography

- [1] Wolfgang Wahlster, „Semantisches Web“, In: Bullinger, H.-J. (ed): Trendbarometer Technik: Visionäre Produkte, Neue Werkstoffe, Fabriken der Zukunft, p. 62-63, München, Wien: Hanser, 2004
- [2] <http://www.robotwisdom.com/weblogs/>
- [3] <http://blogs.technet.com/Bloggers.aspx>
- [4] <http://blogs.msdn.com/>
- [5] Michael Sintek, Stefan Decker: TRIPLE—A Query, Inference, and Transformation Language for the Semantic Web. International Semantic Web Conference (ISWC), Sardinia, June 2002.
- [6] <http://go.connect.yahoo.com/>
- [7] L. Sauermann, A. Bernardi and A. Dengel, Overview and Outlook on the Semantic Desktop, Proceedings International Semantic Web Conference, Galway, Ireland (Nov. 2005), pp. 1-19
- [8] <http://brainbot.com>
- [9] <http://programmableweb.com>
- [10] <http://www.identitygang.org/Lexicon>
- [11] Rebecca Blood, „How blogging software reshapes the online community“, Communication of the ACM, Vol. 47, Num. 12, pages 53-55, ACM Press, 2004
- [12] Bonnie A. Nardi, Diane J. Schiano, Michelle Gumbrecht, Luke Swartz, „Why we blog“, Communication of the ACM, Vol. 47, Num. 12, pages 41-46, ACM Press, 2004
- [13] David Millen, Jonathan Feinberg, Bernard Kerr, „Social Bookmarking in the Enterprise“, Queue, Vol. 3, Num. 9, pages 28-35, ACM Press, 2005
- [14] Steve Cayzer, „Semantic Blogging and Decentralized Knowledge Management“, Communications of the ACM, Vol. 47, Num. 12, pages 47-52, ACM Press, 2004
- [15] Yahoo! Analyst Day 2006 Presentation Slides, available at <http://yhoo.client.shareholder.com/downloads/2006AnalystDay.pdf>
- [16] Tim O'Reilly, „Definition of Web 2.0“, <http://www.oreillynet.com/pub/a/oreilly/tim/news/2005/09/30/what-is-web-20.html>
- [17] Google Analyst Day 2006 Presentation Slides, available at http://investor.google.com/pdf/20060302_analyst_day.pdf
- [18] Bill Trancer, „Google, Yahoo! and MSN: Property Size-up“, blog post available at http://weblogs.hitwise.com/bill-trancer/2006/05/google_yahoo_and_msn_property.html
- [19] David Baum, „Semantic Breakthrough“, Oracle Magazine, May/June 2006, pages 42-46
- [20] Tim Berners-Lee, James Hendler, and Ora Lassila. „The Semantic Web“. Scientific American, 279, 2001.
- [21] Matt Biddulph, „Using Wikipedia and the Yahoo! API to give structure to flat lists“, blog post available at <http://www.hackdiary.com/archives/000070.html>
- [22] OSCON 2005 Keynote; available at <http://www.identity20.com/media/OSCON2005>
- [23] <http://spwiki.editme.com/higgins>
- [24] <http://www.semanticgrid.org>
- [25] <http://www.identity20.com>
- [26] <http://www.mindswap.org/2003/pellet/>
- [27] <http://www.w3.org/TR/rif-ucr/>
- [28] <http://www.gnowsis.org>

About the Authors

Wolfgang WAHLSTER is the Director and CEO of the German Research Center for Artificial Intelligence and a Professor of Computer Science at Saarland University. In 2000, he was coopted as a Professor of Computational Linguistics at the same university. In addition, he is the Head of the Intelligent User Interfaces Lab at DFKI. He was the Scientific Director of the VERBMOBIL consortium on spontaneous speech translation (1993-2000) as well as the SmartKom consortium on multimodal dialog systems (1999-2003) and currently serves as the Scientific Director of the SmartWeb consortium on mobile multimodal access to semantic web services (2004-2008). He has published more than 170 technical papers and 7 books on language technology and intelligent user interfaces. His current research includes multimodal and perceptive user interfaces, user modeling, ambient intelligence, embodied conversational agents, smart navigation systems, semantic web services, and resource-adaptive cognitive technologies. Professor Wahlster is a member of the Supervisory Board of various IT and VC companies. He has been and is serving on a number of international advisory boards, including the Advisory Board of Deutsche Telekom Laboratories (Berlin, Germany). In 2001, the President of the Federal Republic of Germany, Dr. Johannes Rau, presented the German Future Prize to Professor Wahlster for his work on language technology and intelligent user interfaces. He was the first computer scientist to receive Germany's highest scientific prize that is awarded each year for outstanding innovations in technology, engineering, or the natural sciences. In 2003, he was the first German computer scientist elected Foreign Member of the Royal Swedish Academy of Sciences, Stockholm.

Andreas DENGEL is professor of Computer Science at the University of Kaiserslautern and a member of the DFKI Management Board. He is the German representative in the IST Prize Executive Jury of the European Council of Applied Science (Euro-CASE). He has been elected a Fellow of the International Association for Pattern Recognition (IAPR) and has been honoured for his work several times. Most prominent prizes are the "ICDAR Young Investigator Award", the Technical Communication Award of the Alcatel SEL foundation as well as a Document Analysis Systems Achievement Award at Princeton University, NJ. He has published 8 books and more than 120 internationally reviewed papers. As an

entrepreneur, he has been involved with several start-ups. He received his doctoral degree from the University of Stuttgart. His visiting and regular past positions include Johannes Gutenberg University at Mainz, Xerox Parc, Siemens Research Labs, and IBM.