

Chapter 2

Literature Review

Abstract In this literature review, we have addressed the fundamentals of trust and reputation, in terms of an online environment. The literature review has progressed to defining the problem domain. We have realized that it is crucial to proactively carry out this research while online service technologies and platforms are still being developed. It will be very difficult to start integrating soft security solutions into online service technologies after the online service trust problems start emerging. Online systems are currently integrated into many information systems providing online services, where people can rate each other, or rate products and services offered and express their feedback as opinions about products and services offered either online or through the physical world. However, the field of online services trusts management in the recommender system is relatively new and different, addressing the demand for automated decision support in increasingly dynamic business to business (B2B) relationships. Some areas have not been reviewed in depth, such as how trust can be propagated through the web-based social network and how trust can be used to improve the automated recommendation quality. More investigation is required to develop trust management technology for online markets, including improvement in trust inference techniques and the methods for assessing the quality of recommendation related to trust and reliability.

Keywords Recommender systems • Trust • Inference • User tag • Web 2.0 • Social networks

In this literature review, we have addressed the fundamentals of trust and reputation, in terms of an online environment. The literature review has progressed to defining the problem domain. We have realized that it is crucial to proactively carry out this research while online service technologies and platforms are still being developed. It will be very difficult to start integrating soft security solutions into online service technologies after the online service trust problems start emerging. Online systems are currently integrated into many information systems providing online services, where people can rate each other, or rate products and

Table 2.1 Relationship between sections

Section	Contents
2.1	Reviews the state of the art in trust management
2.2	Summarizes recent studies on recommender systems including the most popular collaborative filtering techniques
2.3	Introduces the concept of personalized tagging in the Web 2.0 environment
2.4	Summarizes and highlights the implications from the literature affecting this study

services offered and express their feedback as opinions about products and services offered either online or through the physical world. However, the field of online services trusts management in the recommender system is relatively new and different, addressing the demand for automated decision support in increasingly dynamic business to business (B2B) relationships. Some areas have not been reviewed in depth, such as how trust can be propagated through the web-based social network and how trust can be used to improve the automated recommendation quality. More investigation is required to develop trust management technology for online markets, including improvement in trust inference techniques and the methods for assessing the quality of recommendation related to trust and reliability.

This book deals with the issue of trust management to improve the recommendation quality. An extensive literature review has conducted to investigate the current state of the art work on this issue in the existing literature. The organization of the literature review is presented in the Table 2.1.

In a recent work, Tavakolifard (2010) presents a categorization of the related research on trust management and similarity measurement which is closely related to our work. It has shown in the Table 2.2. We have discussed these works briefly including other state of the art research from the literature in the next 4 sections.

2.1 Fundamentals of Trust and Trust Network

2.1.1 Trust Overview

Trust has become an important topic of research in many fields including sociology, psychology, philosophy, economics, business, law and of course in IT. Far from being a new topic, trust has been the topic of hundreds of books and scholarly articles over a long period of time. Trust is a complex word with multiple dimensions. A vast body of literature on trust has grown in several area of research but it is relatively confusing and sometimes contradictory, because the term is being used with a variety of meanings (McKnight and Chervany 2002). Also a lack of coherence exists among researchers in the definition of trust. Though dozens of proposed definitions are available in the literature, we could not find a complete formal unambiguous definition of trust. In many instances, trust is used as a word

Table 2.2 Categorization of the related work

Paper	Web of trust	User-item rating	Rating-based similarity	Profile-based similarity
Papagelis et al. 2005		X	X	
Massa and Avesani 2004	X		X	
Massa and Bhattacharjee 2004	X		X	
Massa and Avesani 2007	X			
Avesani et al. 2004	X		X	
Avesani et al. 2005	X		X	
Weng et al. 2008	X			
Lathia et al. 2008	X			
Gal-Oz et al. 2008	X			
O'Donovan and Smyth 2005		X	X	X
Ziegler and Golbeck 2007		X		X
Ziegler and Lausen 2004		X		X
Golbeck 2006		X		X
Golbeck and Hendler 2006		X		X
Golbeck 2005	X			
Bedi and Kaur 2006	X			
Bedi et al. 2007	X			
Hwang and Chen 2007		X		
Kitisin and Neuman 2006	X			
Fu-guo and Sheng-hua 2007		X		X
Peng and Seng-cho 2009	X			X
Victor et al. 2008a	X			
Victor et al. 2008b	X			

or concept with no real definition. Sociologists such as Deutch and Gambetta have provided some working definitions and a theoretical background to understand the basic concepts of trust and how it operates in the real world (Folkerts 2005). Hussain and Chang (2007) present an overview of the definitions of the terms of trust from the existing literature. They have shown that none of these definitions is fully capable of satisfying all of the context dependence, time dependence and the dynamic nature of trust.

The most cited definition of trust is given by Dasgupta where he defines trust as “the expectation of one person about the actions of others that affects the first person’s choice, when an action must be taken before the actions of others are known” (Dasgupta 1990). This definition captures both the purpose of trust and its nature in a form that can be discussed. Deutsch (2004) states that “trusting behavior occurs when a person encounters a situation where she perceives an ambiguous path. The result of following the path can be good or bad and the occurrence of the good or bad result is contingent on the action of another person” (Hussain and Chang 2007). Another definition for trust by Gambetta is also often quoted in the literature: “trust (or, symmetrically, distrust) is a particular level of the subjective probability with which an agent assesses that another agent or group of agents will perform a particular action, both before he can monitor such action

(or independently of his capacity ever to be able to monitor it) and in a context in which it affects his own action” (Gambetta 2000). But trust can be more complex than these definitions. Trust is the root of almost any personal or economic interaction. Keser states “trust as the expectation of other persons goodwill and benign intent, implying that in certain situations those persons will place the interests of others before their own” (Keser 2003). Golbeck and Hendler (2006) defines trust as “trust in a person is a commitment to an action based on belief that the future actions of that person will lead to a good outcome”. This definition has a great limitation in that it considers trust as always leading to a positive outcome. But in reality, that may not be always true. Trust is such a concept that crosses disciplines and also domains. The focus of the definition differs on the basis of the goal and the scope of the projects. As mentioned earlier, in this work, we define trust as a subjective probability by which an agent can have a level of confidence in another agent in a given scope, to make a recommendation. Two general definitions of trust defined by Jøsang (Jøsang et al. 2007), which they called reliability trust (the term “evaluation trust” is more widely used by the other researchers, therefore this term is used) and decision trust respectively, are considered to define trust formally for this work. Evaluation trust can be interpreted as the reliability of something or somebody and decision trust captures a broader concept of trust.

Evaluation Trust: Trust is the subjective probability by which an individual, *A*, expects that another individual, *B*, performs a given action on which their welfare depends.

Decision Trust: Trust is the extent to which one party is willing to depend on something or somebody in a given situation with a feeling of relative security, even though negative consequences are possible.

Dimitrakos (2003) surveyed and analyzed the general properties of trust in e-services and listed the general properties of trust (and distrust) as follows:

- Trust is relevant to specific transactions only. *A* may trust *B* to drive her car but not to baby-sit.
- Trust is a measurable belief. *A* may trust *B* more than *A* trusts *C* for the same business.
- Trust is directed. *A* may trust *B* to be a profitable customer but *B* may distrust *A* to be a retailer worth buying from.
- Trust exists in time. The fact that *A* trusted *B* in the past does not in itself guarantee that *A* will trust *B* in the future. *B*’s performance and other relevant information may lead *A* to re-evaluate her trust in *B*.
- Trust evolves in time, even within the same transaction. During a business transaction, the more *A* realizes she can depend on *B* for a service *X*, the more *A* trusts *B*. On the other hand, *A*’s trust in *B* may decrease if *B* proves to be less dependable than *A* anticipated.
- Trust between collectives does not necessarily distribute to trust between their members. On the assumption that *A* trusts a group of contractors to deliver (as a group) in a collaborative project, one cannot conclude that *A* trusts each member of the team to deliver independently.

- Trust is reflexive, yet trust in oneself is measurable. *A* may trust her lawyer to win a case in court more than she trusts herself to do it. Self-assessment underlies the ability of an agent to delegate or offer a task to another agent in order to improve efficiency or reduce risk.
- Trust is a subjective belief. *A* may trust *B* more than *C* trusts *B* within the same trust scope.

Reputation systems are closely related to the concept of trust. The global trust in somebody or something can be referred as reputation. Wang and Vassileva (2007) identify that trust and reputation share some common characteristics such as being context specific, multi-faceted and dynamic. They argue that trust and reputation both depend on some context. Even in the same context, there is a need to develop differentiated trust in different aspects of a service. As they are dynamic in character, they say that trust and reputation can increase or decrease with further experiences of interactions or observations. Both of them also decay with time. Mui et al. (2002) differentiate the concepts of trust and reputation by defining reputation as the perception that an agent creates through past actions about its intentions and norms, and trust as a subjective expectation an agent has about another's future behavior based on the history of their encounters. The difference between trust and reputation can be illustrated by the following perfectly normal and plausible statements:

1. I trust you because of your good reputation.
2. I trust you despite your bad reputation.

Statement (1) reflects that the relying party is aware of the trustee's reputation, and bases his or her trust on that reputation. Statement (2) reflects that the relying party has some private knowledge about the trustee, e.g., through direct experience or intimate relationship, and that these factors overrule any reputation that the trustee might have. This observation reflects that trust ultimately is a personal and subjective phenomenon that is based on various factors or evidence, and that some of those carry more weight than others. Personal experience typically carries more weight than second hand recommendations or reputation, but in the absence of personal experience, trust often has to be based on reputation. Reputation can be considered as a collective measure of trustworthiness (in the sense of reliability) based on ratings from members in a community. Any individual's subjective trust in a given party can be derived from a combination of reputation and personal experience.

That an entity is trusted for a specific task does not necessarily mean that it can be trusted for everything. The scope defines the specific purpose and semantics of a given assessment of trust or reputation. A particular scope can be narrow or general. Although a particular reputation has a given scope, it can often be used as an estimate of the reputation of other scopes (Jøsang et al. 2007).



Fig. 2.1 The Amazon.com website

2.1.2 Previous Works on Trust

The issue of trust has been gaining an increasing amount of attention in a number of research communities including those focusing on online recommender systems. There are many different views of how to measure and use trust. Some researchers assign the same meaning to trust and reputation, while others do not. Though the meaning of trust is different to different people, we consider that a brief review of these models is a good starting point to research in the area of online trust management. While it is impossible to review all of the influential work related to trust management, the range of interest in the topic across nearly every academic field is impressive. These are just a few examples to give a sense of that scope.

As trust is a social phenomenon, the model of trust for an artificial world like the Internet should be based on how trust works between people in society (Abdul-Rahman and Hailes 2000). The rich and ever-growing selection of literature related to trust management systems for Internet transactions, as well as the implementations of trust and reputation systems in successful commercial applications such as eBay and Amazon (Fig. 2.1), give a strong indication that this is an important technology (Jøsang et al. 2007).

Commercial implementations seem to have settled around relatively simple principles, whereas a multitude of different systems with advanced features are

being proposed by the academic community. A general observation is that the proposals from the academic community so far lack coherence and are rarely evaluated in a commercial/industrial application environment. The systems being proposed are usually designed from scratch, and only in very few cases are authors building on proposals by other authors. The period in which we are now can therefore be seen as a period of pioneers.

Stephen Marsh (1994) is one of the pioneers to introduce a computational model for trust in the computing literature. For his PhD book, Marsh investigates the notions of trust in various contexts and develops a formal description of its use with distributed, intelligent agents. His model is based on social and psychological factors. He defines trust in three categories; namely basic trust, general trust and situational trust.

These trust values are used to help an agent to decide if it is worth it or not to cooperate with another agent. Besides trust, the decision mechanism takes into account the importance of the action to be performed, the risk associated with the situation and the perceived competence of the target agent. To calculate the risk and the perceived competence, different types of trust (basic, general and situational) are used. But the model is complex, mostly theoretical and difficult to implement. He did not consider reputation in his work.

Abdul-Rahman and Hailes (2000) proposed a model for supporting trust in virtual communities, based on direct experiences and reputation. They have proposed that the trust concept can be divided into direct and recommender trust. They represent direct trust as one of four agent-specified values about another agent which they called very trustworthy, trustworthy, untrustworthy, and very untrustworthy. Recommended trust can be derived from word-of-mouth recommendations, which they consider as reputation. However, there are certain aspects of their model that are ad-hoc which limits the applicability of the model in broader scope.

Schillo et al. (2000) proposed a trust model for scenarios where the interaction result is Boolean, either good or bad, between two agents' trust relationship. They did not consider the degrees of satisfaction. They use the formula to calculate the trust that an agent Q deserves to an agent A is

$$T(A, Q) = \frac{e}{n} \quad (2.1)$$

where n is the number of observed situations and e the number of times that the target agent was honest.

Two one-on-one trust acquisition mechanisms are proposed by Esfandiari and Chandrasekharan (2001) in their trust model. The first is based on observation. They proposed the use of Bayesian networks and to perform the trust acquisition by Bayesian learning. The second trust acquisition mechanism is based on interaction. A simple way to calculate the interaction-based trust during the exploratory stage is using the formula

$$T_{\text{inter}}(A, B) = \frac{\text{number_of_correct_replies}}{\text{total_number_of_replies}} \quad (2.2)$$

In the model proposed by Yu and Singh (2002), the information stored by an agent about direct interactions is a set of values that reflect the quality of these interactions. Only the most recent experiences with each concrete partner are considered for the calculations.

Yu and Singh then define their belief function to be m and can use that to gauge the trustworthiness of an agent in the community. They establish both lower and upper thresholds for trust. Each agent also maintains a quality of service (QoS) metric $0 \leq s_{jk} \leq 1$ that equates to an agent j 's rating of the last interaction with agent k .

This model does not combine direct information with witness information. When direct information is available, it is considered the only source to be used to determine the trust of the target agent. Only when the direct information is not available, the model appeals to witness information.

Papagelis et al. (2005) develop a model to establish trust between users by exploiting the transitive nature of trust. However, their model simply adopts similarity as trustworthiness which still possesses the limitations of similarity based collaborative filtering. Sabater and Sierra (2005) have proposed a modular trust and reputation system oriented to complex small/mid-size e-commerce environments which they called *ReGreT*, where social relations among individuals play an important role. The system takes into account three different sources of information: direct experiences, information from third party agents and social structures. This system maintains three knowledge bases. The outcomes data base (ODB) to store previous contacts and their result; the information data base (IDB), that is used as a container for the information received from other partners and finally the sociograms data base (SDB) to store the graphs (sociograms) that define the agent's social view of the world. These data bases feed the different modules of the system. The *direct trust* module deals with direct experiences and how these experiences can contribute to the trust associated with third party agents. Together with the reputation model they are the basis on which to calculate trust. The system incorporates a credibility module that allows the agent to measure the reliability of witnesses and their information. This module is extensively used in the calculation of *witness reputation*. All these modules work together to offer a complete trust model based on direct knowledge and reputation.

Mui et al. (2002) proposed a computational model based on sociological and biological understanding. The model can be used to calculate agent's trust and reputation scores. They also identified some weaknesses in the trust and reputation study, being the lack of differentiation of trust and reputation and that the mechanism for inference between them is not explicit. Trust and reputation are taken to be the same across multiple contexts or are treated as uniform across time and the existing computational models for trust and reputation are often not grounded on understood social characteristics of these quantities. They did not examine the effects of deception in this model.

In his research, Dimitrakos (2003) presented and analysed a service-oriented trust management framework based on the integration of role-based modelling and risk assessment in order to support trust management solutions. There is also a short survey on recent definitions of trust and the authors subsequently introduce a service-oriented definition of trust, and analyse some general properties of trust in e-services, emphasising properties underpinning the inference and transferability of trust. Dimitrakos suggested that risk analysis and role-based modelling can be combined to support the formation of trust intentions and the endorsement of dependable behaviour based on trust. In conclusions, they provided evidence of emerging methods, formalisms and conceptual frameworks which, if appropriately integrated, can bridge the gap between systems modelling, trust and risk management in e-commerce. Massa and Avesani (2004, 2006) first argued that recommender system can be more effective by incorporating trust than traditional collaborative filtering. However, they used trust metric with explicit trust rating which is not easily available in real world applications. In their subsequent work (Massa and Bhattacharjee 2004) they show that the incorporation of trust metric and similarity metric can increase the coverage. The limitations of their work include the binary relationship among users and trust inference is calculated only on distance between them. Avesani et al. (2004, 2005) apply their model in ski mountaineering domain. Their model requires direct feedback from the user. They also didn't evaluate their models effectiveness.

O'Donovan and Smyth (2005) distinguished between two types of profiles in the context of a given recommendation session or rating prediction. The consumer profile and the producer profile. They described "trust" as the reliability of a partner profile to deliver accurate recommendations in the past. They described two models of trust, called profile-level trust and item-level trust. According to

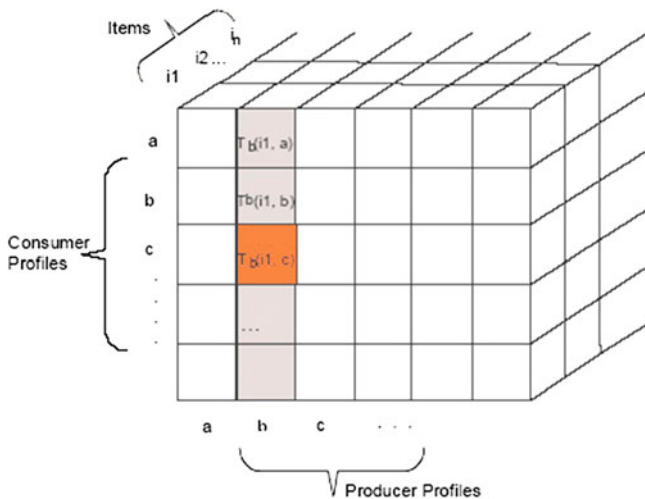


Fig. 2.2 Calculation of trust scores from rating data

them, a ratings prediction for an item, i , by a producer p for a consumer c , is correct if the predicted rating, $p(i)$, is within ϵ of c 's actual rating $c(i)$; see Eq. 2.1. Of course normally when a producer is involved in the recommendation process they are participating with a number of other recommendation partners and it may not be possible to judge whether the final recommendation is correct as a result of p 's contribution. Accordingly, when calculating the correctness of p 's recommendation we separately perform the recommendation process by using p as c 's sole recommendation partner.

For example, in Fig. 2.2, a trust score for item $i1$ is generated for producer b by using the information in profile b only to generate predictions for each consumer profile. Equation 2.8 shows how each box in Fig. 2.2 translates to a binary success/fail score depending on whether or not the generated rating is within a distance of ϵ from the actual rating a particular consumer has for that item. In a real-time recommender system, trust values for producers could be easily created on the fly, using a comparison between predicted rating (based only on one producer profile) and the actual rating entered by a user.

$$Correct(i, p, c) \Leftrightarrow |p(i) - c(i)| < \epsilon \quad (2.3)$$

$$T_p(i, c) = Correct(i, p, c) \quad (2.4)$$

From this they define two basic trust metrics based on the relative number of correct recommendations that a given producer has made. The full set of recommendations associated with a given producer, $RecSet(p)$, is given by Eq. 2.5. And the subset of these that are correct, $CorrectSet(p)$ is given by Eq. 2.6. The i values represent items and the c values are predicted ratings.

$$RecSet(p) = \{(c_1, i_1), \dots, (c_n, i_n)\} \quad (2.5)$$

$$CorrectSet(p) = \{(c_k, i_k) \in RecSet(p) : Correct(i_k, p, c_k)\} \quad (2.6)$$

The profile-level trust, $Trust^P$ for a producer is the percentage of correct recommendations contributed by this producer; see Eq. 2.7. For example, if a producer has been involved in 100 recommendations, that is they have served as a recommendation partner 100 times, and for 40 of these recommendations the producer was capable of predicting a correct rating, the profile level trust score for this user is 0.4.

$$Trust^P(p) = \frac{|CorrectSet(p)|}{|RecSet(p)|} \quad (2.7)$$

Selcuk et al. (2004) proposed a reputation-based trust management protocol for P2P networks where users rate the reliability of the parties they deal with and share this information with their peers.

Guha et al. (2004) proposed a method based on the *PageRank*TM algorithm for propagating both trust and distrust. They identified four different methods for propagating the net belief values, namely direct inference, co-citation, transpose

and coupling. Further to these, three different methods, which they called global rounding, local rounding and majority rounding were proposed to infer whether an element in the final belief matrix, corresponds to trust or distrust.

The Advogato *MaxFlow* (maximum flow) trust metric has been proposed by Levien (2004) in order to discover which users are trusted by members of an online community and which are not. Trust is computed through one centralised community server and considered relative to a seed of users enjoying supreme trust. Local group trust metrics compute sets of agents trusted by those being part of the trust seed. In *MaxFlow* model, its input is given by an integer number n , which is supposed to be equal to the number of members to trust, as well as the trust seed s , which is a subset of the entire set of users A . The output is a characteristic function that maps each member to a Boolean value indicating his trustworthiness:

$$Trust_M : 2^A \times N_0^+ \rightarrow (A \rightarrow \{true, false\}) \quad (2.8)$$

The *AppleSeed* trust metric was proposed by Ziegler (2005). *AppleSeed* is closely based on the *PageRank*TM algorithm. It allows rankings of agents with respect to trust accorded. *AppleSeed* works with partial trust graph information. Nodes are accessed only when needed, i.e., when reached by energy flow. Shmatikov (Shmatikov and Talcott 2005) proposed a reputation-based trust management model which allows mutually distrusting agents to develop a basis for interaction in the absence of a central authority. The model is proposed in the context of peer-to-peer applications, online games or military situations.

Teacy (2005) proposed a probabilistic framework for assessing trust based on direct observations of a trustee's behavior and indirect observations made by a third party. Their proposed mechanism can cope with the possibility of unreliable third party information in some contexts. They also identified some key issues to consider the assessment of the reliability of reputation, like (1) assessing reputation source accuracy, (2) combining different types of evidence, which is required when there is not enough of any one type of evidence for a truster to assess a trustee; (3) exploration of trustee behavior, which involves taking a calculated risk and interacting with certain agents, to better assess their true behavior etc.

Xiong (2005) also proposed a decentralized reputation-based trust-supporting framework called *PeerTrust* for a P2P environment. Xiong focused on models and techniques for resilient reputation management against feedback aggregation, feedback oscillation and loss of feedback privacy. Xue (Xue and Fan 2008) proposed a new trust model for the Semantic Web which allows agents to decide which among different sources of information to trust and thus act rationally on the semantic web. Tian et al. (2008) proposed a trust model for P2P networks in which the trust value of a given peer was computed using its local trust information and recommendations from other nodes. A generic method for quantifying and updating the credibility of a recommender was also put forward. With their experimental results they encountered security problems in open P2P networks. Lathia et al. (2008) outlined several techniques for performing collaborative filtering on the perspective of trust management. A model for computing trust

based reputation for communities of strangers is proposed by Gal-Oz et al. (2008). Bedi and Kaur (2006) proposed a model which incorporates the social recommendation process. In their latter work (Bedi et al. 2007) knowledge stored in the form of ontology is used. Hwang and Chen (2007) attempted to present their model of recommender system by incorporating trust. Victor et al. (2008a) attempted to solve the problem of cold start while making recommendation for a new user. They propose to connect a new user with underlying network for recommendation making. They also advocate the use of trust model in which trust scores are coupled (Victor et al. 2008b). Peng and Seng-cho (2009) applied their model in the domain of blog to help the readers of the blog to find desired articles easily. A recent work (Jamali et al. 2009) proposes a new algorithm called *TrustWalker* to combine trust-based and item-based recommendation. However, their proposed method is limited to centralized systems only. A trust based, effective and efficient model for recommender system is yet to develop and deserving.

2.2 Recommender Systems

Recommender systems intend to provide people with recommendations of items they might appreciate or be interested in. To deal with the ever-growing information overload on the Internet, recommender systems are widely used online to suggest to potential customers, items they may like or find useful. We use recommendations extensively in our daily lives. Examples of recommendations are not hard to find in our day-to-day activities. We may read movie reviews in a magazine or online hotel reviews on the Internet to decide what movies to watch or in which hotel to stay. Sometimes, we even accept recommendations from a librarian to decide which book to choose by discussing our interest and current mood. Usually, people like to seek recommendations from friends or associates when they do not have enough information to decide which books to read, movies to watch, hotels or restaurants to book etc. People love to share their preferences regarding books, movies, hotels or restaurants. Recommender systems attempt to create a technological proxy which produces the recommendation automatically

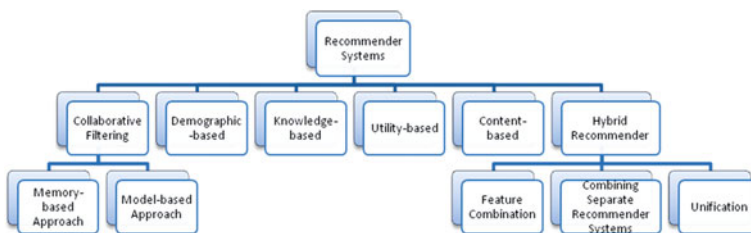


Fig. 2.3 Taxonomy of recommender systems

based on user's previous preferences. The assumption behind many recommender systems is that a good way to produce personalized recommendations for a user is to identify people with similar interests and recommend items that may interest these like-minded people. In this section, the popularly used recommendation approaches will be outlined.

Based on the information filtering techniques employed, recommender systems can be broadly categorized in six different ways (Fig. 2.3):

- Demographic-based,
- Knowledge-based,
- Utility-based,
- Content-based,
- Collaborative filtering, and
- Hybrid recommender systems.

2.2.1 Demographic-Based Filtering

Demographic filtering recommender systems aim to categorize the user based on personal attributes such as their education, age, occupation, and/or gender, to learn the relationship between a single item and the type of people who like it (Rich 1979; Krulwich 1997) and make recommendations based on demographic class. Grundy's system (Guttman et al. 1998) is an example of a demographic filtering recommender system which recommended books based on personal information gathered through an interactive dialogue. The users' responses were matched against a library of manually assembled user stereotypes. Hill et al. (1995) proposed a system which uses demographic groups from marketing research to suggest a range of products and services. A short survey was used to gather the data for user categorization. In Pazzani's (1999) proposed system, machine learning is used to arrive at a classifier based on demographic data. LifeStyle Finder (Krulwich 1997) is another good example of a purely demographic-filtering based recommender. LifeStyle Finder divided the population of the United States into 62 demographic clusters based on their lifestyle characteristics, purchasing history and survey responses. Hence, based on a given user's demographic information, LifeStyle Finder can deduce the user's lifestyle characteristics by finding which demographic cluster the user belongs to and make recommendations to the user. Demographic techniques form "people-to-people" correlations like collaborative filtering, but use different data. The main advantage of a demographic approach is that it may not require a history of user ratings of the type needed by collaborative and content-based techniques. However, there are some shortcomings of demographic filtering recommenders because they create user profiles by classifying users using stereotyped descriptors (Rich 1979) and recommend the same items to users with similar demographic profiles. As every user is different, these recommendations might be too general and poor in quality (Montaner et al. 2003). Purely demographic-filtering based recommenders do not provide any

individual adaptation regarding interest changes (Montaner et al. 2003). However, an individual user's interests tend to shift over time, so the user profile needs to adapt to change which is generally available in collaborative filtering and content-based recommenders as both of them take users' preference data as input for recommendation making.

2.2.2 Knowledge-Based Filtering

Knowledge-based recommender systems use knowledge about users and products to pursue a knowledge-based approach to generating a recommendation, reasoning about what products meet the user's requirements. The *PersonaLogic* recommender system offers a dialog that effectively walks the user down a discrimination tree of product features. Others have adapted quantitative decision support tools for this task (Bhargava et al. 1999). The restaurant recommender *Entree* (Burke et al. 1997) makes its recommendations by finding restaurants in a new city similar to restaurants the user knows and likes. The system allows users to navigate by stating their preferences with respect to a given restaurant, thereby refining their search criteria. A knowledge-based recommender system does not have a ramp-up problem since its recommendations do not depend on a base of user ratings. It does not have to gather information about a particular user because its judgments are independent of individual tastes. These characteristics make knowledge-based recommenders a valuable system on their own and also complementary to other types of recommender systems. Knowledge-based approaches are distinguished from other approaches in the way that they have functional knowledge such as how a particular item meets a particular user need, and can therefore reason about the relationship between a need and a possible recommendation. The user profile can be any knowledge structure that supports this inference. In the simplest case, as in Google, it may simply be the query formulated by the user. In other cases, it may be a more detailed representation of the user's needs (Towle and Quinn 2000). Some researchers have called knowledge-based recommendation the "editor's choice" method (Schafer et al. 2001; Konstan et al. 1997). The knowledge used by a knowledge-based recommender system can also take many forms. Google uses information about the links between web pages to infer popularity and authoritative value (Brin and Page 1998). Knowledge-based recommender systems actually help users explore and thereby understand an information space. Users are an integral part of the knowledge discovery process, elaborating their information needs in the course of interacting with the system. One need only have general knowledge about the set of items and only an informal knowledge of one's needs; the system knows about the tradeoffs, category boundaries, and useful search strategies in the domain. Utility-based approaches calculate a utility value for objects to be recommended, and in principle, such calculations could be based on functional

knowledge. However, existing systems do not use such inference, requiring users to do their own mapping between their needs and the features of products.

2.2.3 Utility-Based Filtering

Utility-based recommender systems make recommendations based on the computation of the utility of each item for the user. Utility-based recommendation techniques use features of items as background data, elicit utility functions over items from users to describe user preferences, and apply the function to determine the rank of items for a user (Burke 2002). The user profile therefore is the utility function that the system has derived for the user, and the system employs constraint satisfaction techniques to locate the best match. Utility-based recommenders do not attempt to build long-term generalisations about their users, but rather base their advice on an evaluation of the match between a user's need and the set of options available.

The e-commerce site *PersonaLogic* has different techniques for arriving at a user-specific utility function and applying it to the objects under consideration (Guttman et al. 1998). The advantage of utility-based recommendations is that they do not face problems involving new users, new items, and sparsity (Burke 2002). It can factor non-product attributes, such as vendor reliability and product availability, into the utility computation, making it possible, for example, to trade off price against delivery schedule for a user who has an immediate need. It attempts to suggest objects based on inferences about a user's needs and preferences. In some sense, all recommendation techniques could be described as doing some kind of inference. The main problem here is; how a utility function for each user should be created. The user must build a complete preference function and weigh each attribute's importance. This frequently leads to a significant burden of interaction. Therefore, determining how to make accurate recommendations with little user effort is a critical issue in designing utility-based recommender systems. Several utility-based recommender systems have been developed (Guan et al. 2002; Schmitt et al. 2002; Stolze and Stroebel 2003; Choi and Cho 2004; Lee 2004; Liu and Shih 2005; Schickel-Zuber and Faltings 2005; Manouselis and Costopoulou 2007). Most of them are based on Multi-Attribute Utility (MAU) theory. However, these methods need a considerable amount of user effort and they lack the evaluation of their preference-elicitation methods in different contexts.

2.2.4 Content-Based Filtering

Conventional techniques dealing with information overload use content-based filtering techniques. They analyse the similarity between items based on their

contents and recommend similar items based on users' previous preferences. (Jian et al. 2005; Pazzani and Billsus 2007; Malone et al. 1987). Typically, content-based filtering techniques match items to users through classifier-based approaches or nearest-neighbor methods. In classifier-based approaches each user is associated with a classifier as a profile. The classifier takes an item as its input and then concludes whether the item is preferred by associated users based on the item contents (Pazzani and Billsus 2007). On the other hand, content-based filtering techniques based on nearest-neighbor methods store all items a user has rated in their user profile. In order to determine the user's interests in an unseen item, one or more items in the user profile with contents that are closest to the unseen item are allocated, and based on the user's preferences to these discovered neighbor items the user's preference to the unseen item can be induced (Montaner et al. 2003; Pazzani and Billsus 2007). This technique is also less affected by the cold-start problem which is one of the major weaknesses of the collaborative filtering based recommenders.

Limitations of contents-based filtering: It has a number of limitations over the other techniques of recommender systems which are discussed below.

Limited applications: The item recommended in the content-based recommender systems should be expressed as a set of features. Information retrieval techniques can effectively extract keywords as the features from a text document. However, the problem is how to extract the features from multimedia data such as audio streaming, video streaming and graphic images. In many cases, it is not possible or practical to manually assign these attributes to the items due to limitation of resources (Shardanand and Maes 1995). Another problem is that if two different items are expressed by the same set of features, how does the content-based recommender system distinguish them. This also causes inaccuracies (Shardanand and Maes 1995).

New User Problem: When a new user logs on to the system, and has very few preferences known, how does the content-based recommendation work? The system cannot acquire sufficient user preference. The system would not recommend an accurate item to the user. With purely content-based filtering recommenders, a user's own ratings are the only factor influencing the recommenders' performance. Hence, recommendation quality will not be very precise for users with only a few ratings (Montaner et al. 2003).

Undue Specialisation: People's interests vary widely. For example, if a user likes three absolutely different areas such as sport, health and finance. The problem is that the systems find it hard to extract the commonalities from these absolutely different areas. Obviously, a content-based recommendation finds it hard to cope in this situation (Ma 2008). Therefore, this technique often suffers from the over-specialization problem. They have no inherent method for generating unanticipated suggestions and therefore, tend to recommend more of what a user has already seen (Resnick and Varian 1997; Schafer et al. 2001).

Lack of Subjective View: Content-based filtering techniques are based on objective information about the items such as the text description of an item or the price of a product (Montaner et al. 2003), whereas a user's selection is usually

based on the subjective information of the items such as the style, quality, or point-of-view of items (Goldberg et al. 1992). Moreover, many content-based filtering techniques represent item content information as word vectors and maintain no context and semantic relations among the words, therefore the resulting recommendations are usually very content centric and poor in quality (Adomavicius and Tuzhilin 2005; Burke 2002; Ferman et al. 2002; Schafer et al. 2001).

2.2.5 Collaborative Filtering

Collaborative filtering, one of the most popular technique for recommender systems, collects opinions from customers in the form of ratings on items, services or service providers. It is most known for its use on popular e-commerce sites such as [Amazon.com](#) or [NetFlix.com](#) (Linden et al. 2003). Essentially, a collaborative filtering based recommender automates the process of the ‘word-of-mouth’ paradigm: it makes recommendations to a target user by consulting the opinions or preferences of users with similar tastes to the target user (Breese et al. 1998; Schafer et al. 2001). Early recommender systems mostly utilized collaborative and content-based heuristics approaches. But over the past several years a broad range of statistical, machine learning, information retrieval and other techniques are used in recommender systems. Collaborative filtering offers technology to recommend items of potential interest to users based on information about similarities among different users’ tastes. Similarity measure plays a very important role in finding like-minded users, i.e., the target user’s neighborhood. The standard collaborative filtering based recommenders use a user’s item ratings as the data source to generate the user’s neighborhood. Because of the different taste of different people, they rate differently according to their subjective taste. If two people rate a set of items similarly, they share similar tastes. In the recommender system, this information is used to recommend items that one participant likes, to other persons in the same cluster.

There is a significant difference between collaborative filtering and reputation systems. In reputation systems, all members in the community should judge the performance of a transaction partner or the quality of a product or service consistently. Collaborative filtering takes ratings subject to taste as input, but reputation system take ratings assumed insensitive to taste as input. Collaborative filtering algorithms are generally classified into two classes: memory based and model based. The memory-based algorithm predicts the votes of the active user on a target item as a weighted average of the votes given to that item by other users. The model-based algorithm views the problem as calculating the expected value of a vote from a probabilistic perspective and uses the users’ preferences to learn a model (Xiong 2005).

Since conventional collaborative filtering based recommenders usually suffer from scalability and sparsity problems, some researchers (Badrul et al. 2001; Deshpande and Karypis 2004; Linden et al. 2003) suggested a modified

collaborative filtering paradigm to alleviate these problems, and this adapted approach is commonly referred to as ‘item-based collaborative filtering’. The conventional collaborative filtering technique (or user-based collaborative filtering) operates based on utilizing the performance correlations among users. Unlike the user-based collaborative filtering techniques, item-based collaborative filtering techniques look into the set of items the target user has rated and compute how similar they are to the target items that are to be recommended. While content-based filtering techniques compute item similarities based on the content information of items, item-based collaborative filtering techniques determine if two items are similar by checking if they are commonly rated together with similar ratings (Deshpande and Karypis 2004).

Item-based collaborative filtering usually offers better resistance to data sparsity problem than user-based collaborative filtering. It is because in practice there are more items being rated by common users than users’ rate common items (Badrul et al. 2001). Moreover because the relationship between items is relatively static, item-based collaborative filtering can pre-compute the item similarities offline whereas user-based collaborative filtering usually computes user similarities online to improve its computation efficiency. Therefore, item-based collaborative filtering is less sensitive to scalability problem (Badrul et al. 2001; Jun et al. 2006; Deshpande and Karypis 2004; Linden et al. 2003).

2.2.5.1 Benefits of Collaborative Filtering:

Subjective View: They usually incorporate subjective information about items (e.g., style, quality etc.) into their recommendations. Hence, in many cases, collaborative filtering based recommenders provide better recommendation quality than content-based recommenders, as they will be able to discriminate between a badly written and a well written article if both happen to use similar terms (Montaner et al. 2003; Goldberg et al. 1992).

Broader Scope: Collaborative filtering makes recommendations based on other users’ preferences, whereas content-based filtering solely uses the target user’s preference information. This, in turn, facilitates unanticipated recommendations because interesting items from other users can extend the target user’s scope of interest beyond his or her already seen items (Sarwar et al. 2000; Montaner et al. 2003).

Broader Applications: Collaborative filtering based recommenders are entirely independent of representations of the items being recommended and therefore, they can recommend items of almost any type including those items from which it is hard to extract semantic attributes automatically such as video and audio files (Shardanand and Maes 1995; Terveen et al. 1997). Hence, collaborative filtering based recommenders work well for complex items such as music and movies, where variations in taste are responsible for much of the variations in preference (Burke 2002).

2.2.5.2 Limitations of Collaborative Filtering

Cold-start: One challenge commonly encountered by collaborative filtering based recommenders is the cold-start problem. Based on different situations, the cold-start problem can be characterized into two types, namely ‘new-system-cold-start problem’ and ‘new-user-cold-start problem’. The new-system-cold-start problem refers to the circumstance where a new system has insufficient profiles of users. In this situation, collaborative filtering based recommenders have no basis upon which to recommend and hence perform poorly (Middleton et al. 2002). In the new-user-cold-start problem, recommenders are unable to make quality recommendations to new target users with no or limited rating information. This problem still happens for systems with a certain number of user profiles (Middleton et al. 2002). When a brand-new item appears in the system there is no way it can be recommended to a user until more information is obtained through another user rating it. This situation is commonly referred to as the ‘early-rater problem’ (Towle and Quinn 2000; Coster et al. 2002).

Sparsity: The coverage of user ratings can be sparse when the number of users is small relative to the number of items in the system. In other words, when there are too many items in the system, there might be many users with no or few common items shared with others. This problem is commonly referred to as the ‘sparsity problem’. The sparsity problem poses a real computational challenge as collaborative filtering based recommenders may find it harder to find neighbors and harder to recommend items since too few people have given ratings (Gui-Rong et al. 2005; Montaner et al. 2003).

Scalability: Scalability is another major challenge for collaborative filtering based recommenders. Collaborative filtering based recommenders require data from a large number of users before being effective, as well as requiring a large amount of data from each user while limiting their recommendations to the exact items specified by those users. The numbers of users and items in e-commerce sites might increase dynamically, consequently, the recommenders will inevitably encounter severe performance and scaling issues (Sarwar et al. 2000; Gui-Rong et al. 2005).

2.2.6 Hybrid Recommender System

From the recommendation techniques described in previous sections, it can be observed that different techniques have their own strengths and limitations and none of them is the single best solution for all users in all situations (Wei et al. 2005). A hybrid recommender system is composed of two or more diverse recommendation techniques, and the basic rationale is to gain better performance with fewer of the drawbacks of any individual technique, as well as to incorporate various datasets to produce recommendations with higher accuracy and quality (Schafer et al. 2001). The first hybrid recommender system *Fab* was developed in

mid 1990s (Balabanovi and Shoham 1997). The combination approaches can be classified into three categories according to their methods of combination: combining separate recommenders, combining features and unification (Adomavicius and Tuzhilin 2005). The active Web Museum, for instance, combines both collaborative filtering and content-based filtering to produce recommendations with appropriate aesthetic quality and content relevancy (Mira and Dong-Sub 2001). Burke (2002) has proposed taxonomy to classify hybrid recommendation approaches into seven categories, which are ‘weighted’, ‘mixed’, ‘switching’, ‘feature combination’, ‘cascade’, ‘feature argumentation’ and ‘meta-level’. The central idea of hybrid recommendation is that they usually comprise of strengths from various recommendation techniques. However, it also means they might potentially include the limitations from each of those techniques. Moreover, hybrid techniques usually are more resource intensive (in terms of computation efficiency and memory usage) than stand-alone techniques, as their resource requirements are accumulated from multiple recommendation techniques. For example, a ‘collaboration via content’ hybrid (Pazzani 1999) might need to process both item content information and user rating data to generate recommendations, therefore will require more CPU cycles and memory than any single content-based filtering or collaborative filtering techniques.

2.3 Web 2.0

2.3.1 *Social Tags*

For assigning a label or organizing items, users may provide one or more keywords as a *tag* in any web site associated with Web 2.0. It is a non-hierarchical keyword or term assigned to a piece of information such as an Internet bookmark or a file name. This kind of metadata helps describe an item and allows it to be found again by browsing or searching. Tags are chosen informally and personally by the item’s creator or by its viewer, depending on the system. In a tagging system, typically there is no information about the meaning or semantics of each tag. For example, the tag “tiger” might refer to the animal or the name of a football team, and this lack of semantic distinction can lead to inappropriate similarity calculation between items. People also often select different tags to describe the same item: for example, items related to the movie “Avatar” may be tagged “Movie”, “Film”, “Animated”, “Fiction” and a variety of other terms. This flexibility allows people to classify their collections of items in the way that they find useful, but the personalized variety of terms can make it difficult for people to find comprehensive information about a subject or item. Another common challenge in tagging systems is that people use both singular and plural words as tags. A user could tag an object with “dinosaur” or with “dinosaurs”, which can make finding similar objects more difficult for both that user and other users in the system. However,

user tags can be regarded as expressions of a user's personal opinion or considered as an implicit rating on the tagged item. Thus tagging information could be useful while making recommendations (Halpin et al. 2007).

Although there are a good number of works available in the field of recommender systems using collaborative filtering, very few researchers consider using tag information to make recommendation (Tso-Sutter et al. 2008; Liang et al. 2009a). Tso-Sutter et al. used tag information as a supplementary source to extend the rating data, not to replace the explicit rating information. Liang et al. (2009a) proposed to integrate social tags with item taxonomy to make personalised user recommendations. Other recent works include integrating tag information with content-based recommender systems (Gemmis et al. 2008), extending the user-item matrix to user-item-tag matrix to collaborative filtering item recommendations (Liang et al. 2009b), combining users' explicit ratings with the predicted users' preferences for items based on their inferred preferences for tags (Sen et al. 2009) etc. However, using tags for recommender systems is still in demand.

The Web 2.0 applications attract users to provide, create, and share more information online. The current popularly available online information that is contributed by users can be classified into textual content information, multimedia content information, and friends/network information. The popular textual user contributed content information includes tags, blogs, reviews, micro-blogs, comments, posts, documents, Wikipedia articles and others. Besides textual contents, user contributed video clips, audio clips, and photos are also popularly available on the web now. In addition, users' network relationship information such as "trust", "tweets", "friends", "followers", etc. becomes more and more popular. Compared with web logs, these kinds of new user information have the advantages of being lightweight, small sized and explicitly and proactively provided by users. They become important information sources to online users.

Tag information is kind of typical Web 2.0 information. Recently, social tags become an important research focus. Implying users' explicit topic interests, social tags can be used to improve searching (Bao et al. 2007; Begelman et al. 2006; Bindelli et al. 2008), clustering (Begelman et al. 2006; Shepitsen et al. 2008), and recommendation making (Niwa et al. 2006; Shepitsen et al. 2009; Jaschke et al. 2007; Zhang et al. 2010).

The research of tags is mainly focusing on how to build better collaborative tagging systems (Golder et al. 2006), item navigation and organisation (Bindelli et al. 2008), semantic cognitive modelling in social tagging systems (Fu et al. 2009), personalising searches using tag information (Bao et al. 2007) and recommending tags (Marinho and Schmidt-Thieme 2007) and items (Tso-Sutter et al. 2008) to users etc. In real life, collaborative tagging systems have not only been used in social sharing websites such as De.licio.us, and e-commerce websites such as Amazon.com, but also in other traditional websites or organizations such as libraries and enterprises (Millen et al. 2006, 2007).

In those early works, only the original tag set of a user was used to profile the topic interests of that user (Bogers and Bosch 2008). The binary weighting

approach (Bogers and Bosch 2008) and the *tf-idf* weighting approach (Salton and Buckley 1988) that borrowed from text mining are commonly used to assign different weights to tags. Tso-Sutter et al. (2008) proposed a tag expansion method to profile a user with tags and items. They proposed to convert the three-dimensional user-tag-item relationship into an expanded user-item implicit rating matrix. The tags of each user were regarded as special items to expand each user's item set. While the tags of each item were regarded as special users to expand each item's user set. This approach profiles users with their own tags and items.

Because of the noise of tags, some user profiling approaches proposed to use the related or expanded tags to profile each user or describe each item. In the work of Niwa et al. (2006) and Shepitsen et al. (2009), those tags of the same cluster were used to expand the tag-based user profiles. Yeung et al. (2008) proposed to find a cluster of tags to profile the topic interests of each user, called *personomy*. All the popular tags of the collected items of a user were used to profile the topic interests of that user. Moreover, association rule mining techniques were used to find the associated tags to profile users (Heymann et al. 2008). Wetzker et al. (2010) proposed a probability model to translate a user's tag vocabulary into another user's tag vocabulary. Thus, these approaches not only just profiled users with their own tags but also with a set of related or expanded tags.

Moreover, targeting the tag quality problem, some latent semantic models were proposed to process tags. The tag-based latent topics were used to profile users. Wetzker et al. (2009) proposed a hybrid probabilistic latent semantic model based on the binary user-item and tag-item matrixes. Siersdorfer and Sizov (2009) proposed a latent Dirichlet Analysis model to find the latent topics of tags. Thus, instead of using tags directly, these approaches used the latent topics of tags to profile users. Besides being used as a stand-alone information source, tags are also popularly used to combine with other information sources to profile users. In the next two subsections, firstly the user profiling approaches based on other popular Web 2.0 information sources will be briefly reviewed. Then, the hybrid user profiling approaches based on tags and other information sources will be reviewed.

Tags are also popularly used to combine with other information sources to profile users. Sen et al. (2009) proposed to combine tags, explicit ratings, implicit ratings, click streams and search logs to profile users and make personalized item recommendations. The work in (Heymann et al. 2008; Gemmis 2008) proposed to combine tags with the textural content information of tagged items to find users' topic interests. The work proposed approaches to combine tags with blogs (Hayes et al. 2007; Qu et al. 2008) to find users' opinions and topic interests. Recently, some approaches combined tags and images (Xu et al. 2010) and videos (Chen et al. 2010) to mine the semantic meaning of the multimedia items. The research of user profiling approaches that combine tags and micro-blogs such as tweets (e.g., hash tag) is one of the new research focuses (Huang et al. 2010; Efron 2010).

In conclusion, user profiling is an ongoing research area. The new user information provides new opportunities to profile users. At the same time, the noise and the new features of these kinds of user information bring challenges to the current user profiling approaches. How to reduce the noise, make use of the unique

features and the rich personal information of these kinds of user contributed information is vital to accurately profile users.

2.3.2 Web 2.0 and Social Networks

A Web 2.0 site allows users to interact and collaborate with each other in a social media dialogue as creators of user-generated content in a virtual community, in contrast to websites where users are limited to the passive viewing of content that was created for them. Examples of Web 2.0 include social networking sites, blogs, wikis, video sharing sites, hosted services and web applications. However, social networking has been around for some time. Facebook and MySpace have become iconic and other sites such as LinkedIn, hi5, Bebo, CyWorld and Orkut are becoming important as well. At the end of 2007, Microsoft paid \$240 million for a 1.6 % stake in Facebook, sparking a fierce debate about the theoretical valuation of Facebook. While few would go along with the \$15 billion price tag, nobody would deny the huge potential of Facebook. The relevance of social networking for advertisers is very high considering they want to invest their money where their potential customers are located on social networking sites.

The success of social networking should not come as a surprise. Social interaction is deeply rooted in human nature and is one of the most fundamental needs. Wireless and Internet technology act as enablers and facilitators for enhanced social interaction with a global reach. While social networking has been and still is dominated by teenagers and young adults, it is quickly spreading to all age groups and beyond the confines of consumer entertainment. Corporations are discovering the power of networking sites to enhance their brands, communities, and overall interaction with their customers by seamlessly linking corporate Web sites to public sites such as Facebook. And something even bigger is about to take place.

There has been dramatic growth in the number and size of Web-based social networks. The number of sites almost doubled over the two-year period from December 2004 to December 2006, growing from 125 to 223. Over the same period, the total number of members among all sites grew fourfold from 115 million to 490 million (Golbeck 2007). This growth has continued over the last five years to an even greater extent. A list of the current most popular social networks can be found at http://en.wikipedia.org/wiki/List_of_social_networking_websites and also at <http://trust.mindswap.org>.

The recent emergence of location-based mobile social networking services offered by providers such as Rumble, GyPSii, Whrrl and Loopt is revolutionising social networking, allowing users to share real-life experiences via geo-tagged user-generated multimedia content, see where their friends are and meet up with them. This new technology-enabled social geo-lifestyle will drive the uptake of location-based services and provide opportunities for location-based advertising in the future.

A location-based social network is a system in the form of a robust web service, used to build an open infrastructure to introduce and connect individuals based on

the intersection of physical location and other properties they might have in common. It is different from the wide range of existing social networking and instant messaging applications in terms of its basic activities. Location-based social networking is really all about the physical presence of a member, and feeds that information nicely into social networking applications such as Facebook, MySpace etc. Location-based social networking is a natural extension to the mobile devices of the Web-based versions of these major social networking sites. It is the mobile's ability to provide real-time location, and in turn actual presence ("I'm shopping in Brisbane") that will provide a real boost to the mobile versions of these sites. Instead of just being a cut-down version of the main site, the mobile version of a location-based social network adds real time value with presence from location services. The location-based social network can be a good application area of trust-based recommender system. Members of the network may receive recommendations from their trusted friends while they are on visiting a new place.

2.4 Chapter Summary

This chapter analyzed a wide range of literature relevant to trust, trust networks, recommender systems, and use of personalized tags in Web 2.0 environments. From the literature, it is understandable that trust is a way of doing social personalization. There are a good number of models proposed to calculate trust between users in a trust network while the trusts between users' data are present. It leads to research questions such as: if there is a social network with no trust values, how do we compute trust; how do we compute indirect trust, while finding the neighbors in trust networks; and how many hops should we traverse or how many of the raters should be considered, etc. If we use too few, the prediction is not based on a significant number of ratings. On the other hand, if we use too many, these raters may only be weakly trusted. In a large trust network, efficiency of exploring the trust network is another issue. It can be noted that all the relevant existing works have assumed that a trust network already exists or that the trust data is readily available, which is not very common in real world. We have considered such a situation in our proposed model which is described in detail in [Chap. 5](#).



<http://www.springer.com/978-1-4614-6894-3>

Trust for Intelligent Recommendation

Bhuiyan, T.

2013, XIV, 119 p. 34 illus., Softcover

ISBN: 978-1-4614-6894-3