

Trust and Privacy in Social Participatory Sensing

Author:

Amintoosi, Haleh

Publication Date:

2014

DOI:

<https://doi.org/10.26190/unsworks/16754>

License:

<https://creativecommons.org/licenses/by-nc-nd/3.0/au/>

Link to license to see what you are allowed to do with this resource.

Downloaded from <http://hdl.handle.net/1959.4/53434> in <https://unsworks.unsw.edu.au> on 2024-04-28

Trust and Privacy in Social Participatory Sensing

THE UNIVERSITY OF NEW SOUTH WALES



SYDNEY · AUSTRALIA

Dissertation submitted in fulfilment
of the requirements for the degree of

Doctor of Philosophy
in
School of Computer Science and Engineering

Haleh AMINTOOSI

Supervisor: Prof. Salil S. Kanhere

March 2014

COPYRIGHT STATEMENT

I hereby grant the University of New South Wales or its agents the right to archive and to make available my thesis or dissertation in whole or part in the University libraries in all forms of media, now or here after known, subject to the provisions of the Copyright Act 1968. I retain all proprietary rights, such as patent rights. I also retain the right to use in future works (such as articles or books) all or part of this thesis or dissertation. I also authorise University Microfilms to use the 350 word abstract of my thesis in Dissertation Abstract International (this is applicable to doctoral theses only). I have either used no substantial portions of copyright material in my thesis or I have obtained permission to use copyright material; where permission has not been granted I have applied/will apply for a partial restriction of the digital copy of my thesis or dissertation.

Signed *H. Amintoosi*

Haleh AMINTOOSI

March 2014

AUTHENTICITY STATEMENT

I certify that the Library deposit digital copy is a direct equivalent of the final officially approved version of my thesis. No emendation of content has occurred and if there are any minor variations in formatting, they are the result of the conversion to digital format.

Signed *H. Amintoosi*

Haleh AMINTOOSI

March 2014

ORIGINALITY STATEMENT

I hereby declare that this submission is my own work and to the best of my knowledge it contains no materials previously published or written by another person, or substantial proportions of material which have been accepted for the award of any other degree or diploma at UNSW or any other educational institution, except where due acknowledgement is made in the thesis. Any contribution made to the research by others, with whom I have worked at UNSW or elsewhere, is explicitly acknowledged in the thesis. I also declare that the intellectual content of this thesis is the product of my own work, except to the extent that assistance from others in the projects design and conception or in style, presentation and linguistic expression is acknowledged.

Signed *H. Amintoosi*

Haleh AMINTOOSI

March 2014

ACKNOWLEDGEMENTS

Working at the School of Computer Science and Engineering at the University of New South Wales (UNSW) has been a great pleasure and a wonderful privilege.

In the first place, I would like to express my sincere appreciation and deep gratitude to my supervisor, Prof. Salil S. Kanhere, for his exceptional support, encouragement and guidance during all stages of this research. Salil taught me how to do high quality research and helped me think creatively. His truly valuable academic excellence, scientific intuition and beautiful mind have made him a constant oasis of ideas and passions in science. This has inspired and enriched my growth as a student. I will be forever grateful for the time that he has spent for me and for the strong research vision that he has shared with me.

My sincere thanks goes to Iran's scientific counsellor, Dr. Hassan Khaleghi, for his help and support during my studies in Australia. I am also thankful to Professor Sanjay Jha and everyone in the Network Research Lab (NRL) group.

Most of all, I am deeply and forever indebted to my parents Kazem and Najmeh for their never-ending love, support and encouragement throughout my entire life and for providing me the high-quality education and the enthusiasm for higher studies. They are my source of strength and without their boundless support this thesis would never have been started. I would like to thank my sister Hedyeh, and my brothers Hamed and Hesam for their kindly and continuous support in my life.

Last but no means least, it gives me immense pleasure to thank my husband Masood and my daughters Zahra and Hosna, who have given me loving, caring and emotional support in all the difficult times. I thank them especially for their understanding and the acceptance that I have been so busy during my PhD studies that we have hardly been able to spend much time together.

Finally, I would like to thank the Iran Ministry of Science, Research and Technology and the School of Computer Science and Engineering at UNSW for their

financial support through scholarship and work opportunities.

Haleh AMINTOOSI

Sydney, Australia

March 2014

To my husband Masood and my daughters Zahra and Hosna,
and to my parents,
for their love, patience and understanding

ABSTRACT

Advances in the sensing capabilities of smartphones have resulted in the emergence of participatory sensing. In participatory sensing, ordinary citizens are recruited to collect sensor data from nearby environments which are then analysed to provide useful information. The information credibility is predominantly dependent on sufficient participation. There are however, various costs associated with contributing data including time, phone battery and bandwidth consumption and potential exposure to privacy threats. These issues may dissuade participants from contributing, thus decreasing the data quality. The integration of social networks with participatory sensing, referred to as social participatory sensing is a potential solution since it provides access to social network members as participants. This integration however, raises new challenges. First, is the potential sparseness of the requester's friendship graph which affects the ability to recruit sufficient contributors. Second, is the identification of well-suited participants who can fulfil the task's requirements. Third, is assessing the trustworthiness of provided contributions.

In this thesis, we propose an innovative framework comprising novel strategies that address the aforementioned issues. We first present a recruitment scheme that addresses the participation sufficiency issue by utilising friendship relations to provide access to adequate participants. The scheme also identifies credible communication paths to preserve the integrity and privacy of messages. Next, we design a participant selection scheme to select well-suited participants from a wider pool. Our scheme also prevents collusion among the selected group. Finally, we present a trust assessment scheme for comprehensive trust evaluation encompassing all personal and social influential parameters. The trust scores are then used to update the participants' reputations. The proposed ideas have been experimentally validated on real-world datasets. Results show that our framework is able to effectively address the participant sufficiency by recruiting the required participants with twice the suitability as that achieved by comparable methods. Moreover, it can accurately detect

83% of possible collusion instances. Our framework is also successful in increasing the overall trust to 90% which is 15% greater than that achieved by compared methods. To sum up, our proposed framework is successful in comprehensively addressing the challenges of social participatory sensing in an application-agnostic manner.

PUBLICATIONS

- Haleh Amintoosi, Salil S. Kanhere, “A trust framework for social participatory sensing systems”, in Proceedings of the 9th International Conference on Mobile and Ubiquitous Systems: Computing, Networking, and Services (MobiQuitous), pages 237-249, 2012.
- Haleh Amintoosi, Salil S. Kanhere, “A trust-based recruitment framework for multi-hop social participatory sensing”, in Proceedings of the 9th IEEE International Conference on Distributed Computing in Sensor Systems (DCOSS), pages 266-273, 2013.
- Haleh Amintoosi, Salil S. Kanhere, “A reputation framework for social participatory sensing systems”, in Journal of Mobile Networks and Applications (MONET), vol. 19, pages 88-100, 2014.
- Haleh Amintoosi, Salil S. Kanhere, “Privacy-aware trust-based recruitment in social participatory sensing”, in 10th International Conference on Mobile and Ubiquitous Systems: Computing, Networking, and Services (MobiQuitous), in press.
- Haleh Amintoosi, Salil S. Kanhere, “Providing trustworthy contributions via a reputation framework in social participatory sensing systems”, Technical Report UNSW-CSE-TR-201304, University of New South Wales, 2013.
- Haleh Amintoosi, Salil S. Kanhere, “Trust-based recruitment in multi-hop social participatory sensing”, Technical Report UNSW-CSE-TR-201312, University of New South Wales, 2013.
- Haleh Amintoosi, Salil S. Kanhere, “Trust and Privacy Considerations in Participant Selection for Social Participatory Sensing”, Technical Report UNSW-CSE-TR-201409, University of New South Wales, 2014.

- Haleh Amintoosi, Mohammad Allahbakhsh, Salil S. Kanhere, Masood Niazi Torshiz, “ Trust assessment in social participatory networks”, in Proceedings of the 3rd IEEE International eConference on Computer and Knowledge Engineering (ICCKE), pages 448-453, 2013.
- Haleh Amintoosi, Salil S. Kanhere, Mohammad Allahbakhsh, “ Trust-based privacy-aware participant selection in social participatory sensing”, submitted to the Elsevier Journal of Information Security and Applications (JISA).
- Haleh Amintoosi, Salil S. Kanhere, Masood Niazi Toshiz, “ Collusion Detection in Social Participatory Sensing”, submitted to the Elsevier Journal of Computer Communications (COMCOM).

Contents

1	Introduction	1
1.1	Social Participatory Sensing	4
1.1.1	Online Social Networks	5
1.1.2	Participatory Sensing	7
1.2	Key Research Issues	9
1.2.1	Sufficiency of Participants	9
1.2.2	Suitability of Participants	10
1.2.3	Assessing the Trust	11
1.3	Goal and Methodology	13
1.4	Overview of Contributions	14
1.4.1	A Recruitment Scheme to Ensure Sufficient Participation . . .	15
	Assessing the Credibility of the Communication Paths	15
	Selecting the Communication Paths	16
1.4.2	A Participant Selection Scheme to Ensure Participant Suitability	17
	Evaluating the Suitability of Participants	17
	Preventing Collusion	18
1.4.3	A Trust Scheme to Assess Trust and Reputation	19
	Assessing the Trust	19

Evaluating the Reputation of the Participants	20
1.5 Dissertation Organisation	21
2 Background and State of the Art	24
2.1 Introduction	24
2.2 Social Participatory Sensing	25
2.2.1 Background	25
2.2.2 Components	28
2.3 Trust Assessment in Social Participatory Sensing	34
2.3.1 Trust Assessment in Online Communities	34
2.3.2 Trust Assessment in Participatory Sensing	36
2.3.3 Reputation Management in Online Communities	39
2.3.4 Reputation Management in Participatory Sensing	40
2.4 Participant Selection in Social Participatory Sensing	43
2.4.1 Participant Selection in Online Communities	43
2.4.2 Participant Selection in Participatory Sensing	45
2.5 Recruitment in Social Participatory Sensing	50
2.5.1 Recruitment in Online Communities	50
2.5.2 Recruitment in Participatory Sensing	52
2.6 Summary	54
3 Ensuring Sufficient Participation	55
3.1 Introduction	55
3.2 Credibility Assessment Scheme	61
3.2.1 Trust Assessment	61
3.2.2 Privacy Assessment	62

3.2.3	Assessing the Path's Credibility Score	66
3.3	Path Selection Scheme	69
3.3.1	Credibility-based Selection	69
3.3.2	Efficient Selection	70
3.4	Experimental Evaluation	76
3.4.1	Simulation Set-up	76
3.4.2	Simulation Results	80
3.5	Conclusion	86
4	Selecting Well-Suited Participants	88
4.1	Introduction	88
4.2	Suitability Assessment Scheme	95
4.2.1	Assessing the Suitability Parameters	95
4.2.2	Computing the Suitability Score	98
4.2.3	Eligibility Assessment	99
4.3	Collusion Prevention Scheme	101
4.3.1	Identifying Potentially Colluding Groups	102
4.3.2	Collusion Indicators	103
4.3.3	Possibility of Collusion	107
4.4	Experimental Evaluation	108
4.4.1	Simulation Set-up	108
4.4.2	Performance Comparison	111
4.4.3	Sensitivity Analysis	115
4.4.4	Collusion Prevention Analysis	118
4.5	Conclusion	123

5	Assessing Trust and Reputation	126
5.1	Introduction	126
5.2	Trust Assessment Scheme	131
5.2.1	Quality of Contribution (QoC)	132
5.2.2	Trust of Participant (ToP)	133
5.2.3	Trust of Contribution (ToC)	138
5.3	Reputation Management Scheme	140
5.4	Experimental Evaluation	146
5.4.1	Simulation Set-up	147
5.4.2	Simulation Results	153
	QoC vs. ToP Set-up Results	154
	ToP Set-up Results	158
5.5	Conclusion	162
6	Conclusion and Future Work	164
6.1	Concluding Remarks	164
6.2	Future Directions	166
	Bibliography	169

List of Figures

1.1	Timeline of online social networks from past to present	5
1.2	Architecture of a typical participatory sensing system	7
1.3	The proposed framework for addressing the challenges of social participatory sensing	14
2.1	Social participatory sensing and its relation to other collective intelligence systems	25
2.2	Architecture of a typical online social network	29
2.3	Architecture of a typical participatory sensing system	31
3.1	Recruitment scheme at a glance	58
3.2	The sequence of steps in the recruitment of suitable participants . . .	59
3.3	A simple example for calculating the path's trust score	62
3.4	Membership function for the path's trust score (TS)	67
3.5	Membership function for the path's privacy score (PS)	67
3.6	Membership function for the path's credibility score (CS)	68
3.7	Evolution of average number of participants in different intervals . . .	80
3.8	Evolution of average mean path trust (MPT) in different intervals . .	81
3.9	Evolution of average overall trust in different intervals	82
3.10	Evolution of participation score in different intervals	83

3.11	Average privacy score of the selected paths for different privacy classes	84
3.12	Evolution of average overall trust in different intervals	84
3.13	Evolution of average path's credibility score passing from a malicious node	85
4.1	Participant selection scheme	92
4.2	The sequence of steps in the recruitment of suitable participants . . .	93
4.3	Membership functions of input and output linguistic variables	101
4.4	Performance of three methods in the case of requesters with few friends	112
4.5	Performance of three methods in the case of requesters with large number of friends	113
4.6	Performance of all methods regardless of requesters' number of friends	114
4.7	Evaluation of the proposed participant selection scheme with different scenarios	118
4.8	Evolution of number of groups and their maximum size according to the target size	119
4.9	Evolution of precision (%) in different rounds	122
4.10	Evolution of recall (%) in different rounds	122
4.11	Distribution of the values of indicators in collusion attacks	124
5.1	Trust scheme at a glance	128
5.2	Fuzzy inference system architecture	129
5.3	Gompertz function for friendship score	136
5.4	Inverse Gompertz function for time gap score	137
5.5	Membership function for quality of contribution (QoC) and trust of participant (ToP)	138

5.6	Membership function for trust of contribution (ToC)	138
5.7	A sample social graph of 4 members with mutual trust ratings	143
5.8	Sequence of steps in assessing the trust and reputation	145
5.9	Evolution of average overall trust for all methods, Scenario 1	154
5.10	Ranked lists provided by trust scheme for the requester, Scenario 1 .	155
5.11	Evolution of quality of contribution (QoC) & trust of contribution (ToC) for one participant, Scenario 2	155
5.12	Overall trust obtained in Fuzzy and Average methods in Scenario 2 .	157
5.13	Comparison of average overall trust for all methods for both scenarios	157
5.14	Evolution of average overall trust for all methods, Scenario 1	158
5.15	Evolution of overall trust, Fuzzy method	159
5.16	Reputation score for all 100 members after attending 5000 tasks, Fuzzy method, Scenario 1	159
5.17	Reputation score for all 100 members at 4000th campaign, Fuzzy method, Scenario 2	161
5.18	Evolution of reputation score for participant no.9 in all methods, Scenario 2	161

List of Tables

3.1	Fuzzy rule base for defining the path's credibility score according to its trust and privacy scores	68
4.1	Fuzzy rule base for defining eligibility score (ES) according to time suitability (TS) and selection score (SS)	100
4.2	Fuel prices of three different service stations uploaded by eight participants	104
5.1	Fuzzy rule base for defining trust of contribution (ToC) according to quality of contribution (QoC) and trust of participant (ToP)	139
5.2	Parameter settings for the calculation of trust of participant (ToP) .	153

Chapter 1

Introduction

The widespread prevalence of mobile computing devices such as sensor-rich smartphones has propelled the emergence of a novel sensing paradigm, known as *participatory sensing* [1]. In participatory sensing, ordinary citizens volunteer to use their mobile phones for collecting sensor data from their nearby environment. The aim of such sensor data gathering includes computing the aggregate statistics about a phenomenon, thus increasing the global awareness of issues of interest. A plethora of applications have been recently proposed based on this revolutionary paradigm. In PetrolWatch [2], a mobile phone is mounted on the dashboard and automatically captures photos from roadside fuel price boards when the car approaches a fuel station. The photos are then uploaded to a server which is responsible for extracting the fuel price via image processing techniques. Individuals query the server to obtain the cheapest fuel price in their vicinity. In LiveCompare [3], the participants are recruited to take pictures of a product's price tag and its barcode. The barcode is decoded into a textual representation on the mobile phone, and transferred to the server along with the picture displaying the current price. Other information such as the location/time of capture are also stored in the server. Users are then able to search for products in the application in order to compare prices. The server retrieves the corresponding price reports, selects the stores in proximity of the user's current location and displays the pictures of the corresponding price tags. In a series of other applications such as NoiseTube [4], Ear-Phone [5] and NoiseMap [6] mobile phone microphones are used to measure the surrounding noise level. The sound

samples are used to build representative noise pollution maps of urban spaces to enable specialists to understand the relationships between noise levels and behavioural problems.

The involvement of people in the sensing process, however, brings about new challenges. Accepting to contribute to a task will inherently require that the participant devotes some time and effort towards it. Moreover, collecting and uploading the sensor data consumes the mobile phone battery and communication bandwidth. Most importantly, engaging in such crowdsourcing activities may lead to potential privacy threat such as the disclosure of home/work address or private conversations [7, 8, 9]. Further, participatory sensing systems are based on voluntary participation and typically there is little incentive for contributors. With all these in mind, a participant may be hesitant to contribute to a sensing task. This may result in a lack of adequate number of participants, which in turn may compromise the fidelity of the obtained information and ultimately render the application to be not very useful.

Besides, some tasks may require that the participants have specific knowledge or expertise related to the task at hand [10]. For example, consider an application that is aimed at collecting the photos of rare plant species. In order to obtain high fidelity pictures, the requester may wish to recruit participants who have some knowledge of botany. In general, it is desirable to recruit suitable participants (who are those who can satisfy the task’s requirements at an acceptable level). To sum up, an important challenge in obtaining trustable results is the recruitment of participants who are (i) sufficient in number and (ii) well-suited to contribute to the task.

One potential solution to address this challenge is to leverage online social networks (constituting hundreds of millions of subscribers with various skills and expertise) as the underlying publish-subscribe infrastructure for participatory sensing applications [11, 12]. This new paradigm, referred to as *social participatory sensing*, offers the following advantages. First, it is possible to benefit from the dynamics in social networks and reconnoitre suitable participants according to different parameters specified in their profiles such as their specialisation, habits and interests, the geographical area of where they live, prior campaigns that they were involved

in, and their reputation based on past contributions. Second, social friendship relations often act as effective motivation to contribute to tasks created by friends since people generally like to be helpful to their friends [13]. They may help them to find out about cheapest prices by taking photos of grocery item price tags, informing about the traffic conditions, etc. Third, incentives in the form of e-coins [14] or reputation points can be awarded to well-behaved participants in recognition of their contributions, and can be made publicly available on their profiles. So, integrating online social networks with participatory sensing is a first step towards addressing the above mentioned challenge. In other words, social participatory sensing provides the foundation for addressing the issue of sufficient participation.

For social participatory sensing to be a success, one major challenge is identifying a sufficient number of participants via credible communication paths. Leveraging an online social network as the underlying substrate, while beneficial for participant recruitment and transferring the communication messages (such as tasks and contributions), raises new issues. First, the requester may have a sparse friendship network in the social network, thus making it difficult to recruit adequate participants. Second, the communication links that are used for transferring the communication messages may not be trustworthy enough to guarantee the integrity of messages. Messages may also contain sensitive information about the requester or participants (such as home/work address, interests, health information, etc.). The communication paths may not be secure enough to preserve the privacy of sensitive embedded information. Third, is the potential of collusion among participants with ulterior motives. A group of malicious members with a particular agenda may create a colluding group and strategically upload contributions to a task in a manner that will change the final outcome. For example, in a noise pollution mapping application (e.g., NoiseMap [6], NoiseTube [4], Ear-Phone [5]), a group of neighbours may collaborate to upload low decibel noise samples (gathered from other areas) in order to declare their location as a quiet area, which may potentially increase the property prices. The fourth challenge has to do with assessing the trustworthiness of provided contributions. Without confidence on the trustworthiness of sensor data, the obtained information will be of little use [15]. In the context of social participatory

sensing, new trust issues arise. People normally have more trust in contributions provided by their close friends than casual acquaintances, since interactions with close friends provides more emotional and informational support [13]. In particular, when data of the same quality is available from two social network contacts, one a close friend and the other a casual acquaintance, it is natural human tendency to put more credence in the data from the close friend. Hence, in social participatory sensing, it is crucial to consider the social trust of participants as a prominent factor in assessing the trustworthiness of contributions.

This dissertation aims to address the above mentioned issues in social participatory sensing. Particularly, we propose a framework including innovative schemes and techniques that address the participation sufficiency issue by leveraging social friendship relationships to provide access to a large number of participants. Such access is provided via credible communication paths to preserve the integrity and privacy of communication messages. We also design methodologies to select the most suitable participants who can fulfil the task's requirements, while simultaneously preventing the selection of colluding participants. Finally, we present solutions to assess trust and reputation in a comprehensive and precise manner in a way that encompasses all influential personal and social parameters.

The remainder of this chapter is organised as follows. In Section 1.1, we describe the basic concepts behind social participatory sensing. In Section 1.2, we outline the research issues studied in this dissertation. In Section 1.3, we describe the main objectives of our research. Section 1.4 summarises the contributions. The structure of this dissertation is described in Section 1.5.

1.1 Social Participatory Sensing

Social participatory sensing provides access to a pool of members with a vast range of social interconnections who collaborate in the tasks initiated by their friends. Specifically, basic participatory sensing procedures can be carried out via online social networks. These procedures include identifying and recruiting the suitable participants, distributing the tasks to them, and delivering their contributions to

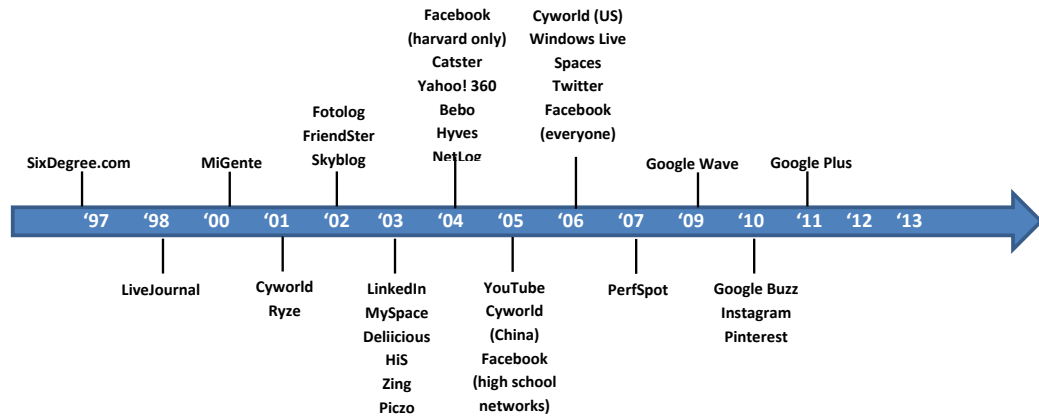


Figure 1.1: Timeline of online social networks from past to present

the requester [1, 16]. In this architecture, social network members can create a task and invite their friends and other members to act as participants and contribute to their tasks. Social ties between people are used for recruitment and the exchange of communication messages. Throughout the dissertation, the terms *task* and *campaign* may be used interchangeably. The same holds true for the terms *data* and *contribution*.

Social participatory sensing is a marriage of two important paradigms: online social networking and participatory sensing. In the following, we describe each paradigm in detail.

1.1.1 Online Social Networks

Since their introduction, online social networks (OSNs) have attracted millions of users, many of whom have integrated them into their daily practices [17]. Specifically, an OSN is made of a group of people, called users, who communicate with each other in an online basis in different ways. Most OSNs work in a publish-subscribe manner, that is, users publish their own generated or aggregated content, and let other users subscribe and get access to this published content. They are also able to add tags, reviews, comments and recommendations. It is also possible to create groups and communicate with group members according to common parameters such as interests or habits.

Figure 1.1 shows the evolution of OSNs as a timeline. As can be seen, the first

OSN, called SixDegrees.com, emerged in 1997 in which, users could create their profiles, add friends and browse through the friend lists. Post 2003, one can observe a noticeable spike in the establishment of OSN communities. This revolution has brought a dramatic shift in the business, the cultural and the research landscape of the World Wide Web [18].

OSNs like Facebook, ¹ MySpace, ² Flickr, ³ LinkedIn, ⁴ and YouTube ⁵ have achieved high growth in their user-base. For example, Facebook now has over 1 billion active users. Moreover, OSNs have brought about increased interactions among friends, and allowed people to reconnect with long-lost acquaintances and old classmates. Users join, establish social links to friends, and leverage their social links to share content, organise events, and search for specific users or shared resources.

In terms of their scope, OSNs can be categorised into two classes [18]: entertainment and business. Most online social networks are used as means for entertainment. They aim at providing online social communications to their users. Popular OSNs such as Facebook and Flickr are entertainment based. Business-oriented OSNs are those that intend to connect professionals from all around the world to share their experiences and knowledge. Users in these OSNs are able to present their professional expertise and achievements in their profiles. An indicative site in this class is LinkedIn.

From the point of community formation, OSNs can be classified into two categories: user oriented and content oriented. In user-oriented social networks, the users and their social relationships are the main emphasis, and the sharing of content is usually among the users in the same community or friendship graph. Examples are Facebook, MySpace and LinkedIn. In content-oriented OSNs, the users' networks are determined by their common interests rather than social relationships. Question-answering forums and video-sharing networks such as YouTube are examples of content oriented OSNs.

¹www.facebook.com

²www.myspace.com

³<https://flickr.com>

⁴www.linkedin.com

⁵www.youtube.com

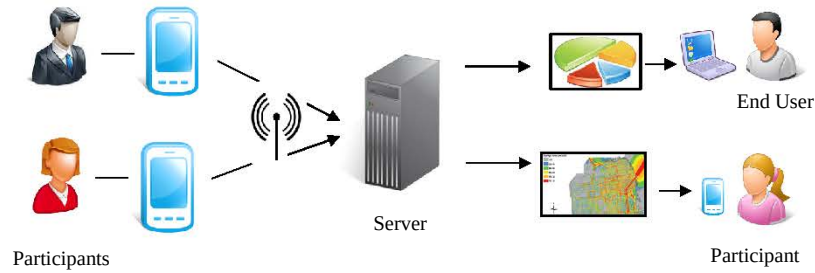


Figure 1.2: Architecture of a typical participatory sensing system

1.1.2 Participatory Sensing

Participatory sensing aims at utilising ordinary citizens to collect sensor data from their surroundings using their mobile phones [1, 19]. Though not built specifically for sensing, nowadays smartphones are able to work as sophisticated sensors. They have cameras for capturing images and video, and Global Positioning System (GPS) receivers that can provide location information. The microphone on the mobile phone can also act as an acoustic sensor. Other embedded sensors such as gyroscopes, accelerometers, and proximity sensors can collectively be used to estimate useful contextual information (e.g., if the user is walking or travelling on a bicycle).

Participatory sensing brings about several advantages that make it appropriate for urban life. First, since participatory sensing relies on existing sensing (i.e. mobile phone) and communication infrastructure (i.e. Wi-Fi or cellular), establishment and deployment costs are very low (almost zero). Second, the mobility of participants offers extensive spatiotemporal coverage which provides sampling diversity not possible with traditional static sensor networks. Third, by including people in the sensing loop, new applications can be designed in order to directly benefit the quality of life of individuals and communities.

The emergence of participatory sensing has resulted in several interesting applications, which can be broadly categorised as either people centric or environment centric. People-centric applications mainly aim at documenting activities (e.g., running, walking, etc.) and understanding the behaviour of individuals [19]. Examples of such applications are DietSense [20], BikeNet [21], PetrolWatch [2] and Live-

Compare [3]. Environment-centric applications, on the other hand, gather environmental parameters such as noise level or air quality. Examples are NoiseTube [4], NoiseMap [6] and Ear-Phone [5].

In participatory sensing, there are three different entities contributing to running and managing the system: application providers, participants and end users. Application providers are those who develop an application, and define its requirements, so that participants can perform it accurately. They are also responsible for setting up the application server to collect and process the data. Participants are those who contribute in gathering and sharing the sensor data. They download and install the application on their mobile phone and accomplish the sensing process according to the task specification. End users are those who can benefit from the information that has been extracted from the data gathered by participants. These end users can be participants themselves or their family members who want to have knowledge about, for example, the health condition of an old person. They might be scientists who want to obtain statistical information about the monitored phenomena. Organisations that have a role in making the application and are responsible for verifying actual contributions and results may also be the end users.

Figure 1.2 depicts the architecture of a typical participatory sensing system. The contributions produced by mobile phones of participants are reported (using wireless data communications) to a central server for processing. At the server-end, the contributions are analysed and made available in various forms, such as charts or maps demonstrating the results, which can be used by individuals as well as the community. Simultaneously, the results may be displayed locally on the users' mobile phones or accessed by the wider public through web-portals [22].

The success of participatory sensing mainly depends on recruiting a sufficient number of participants. Integration of online social networks with participatory sensing is expected to be helpful since it provides access to a pool of users who can be recruited as participants.

1.2 Key Research Issues

The main idea behind social participatory sensing is to recruit social network members to contribute to the tasks originating from their friends. Leveraging immediate friends for participatory sensing data collection is beneficial since friends usually tend to be useful to their friends. However, in the absence of adequate friendship relations, the challenge of participation *sufficiency* still exists. Without adequate participation, the resulting summary statistics may suffer from lack of validity. Another important challenge is the selection of participants with high *suitability*, since this leads to high quality data. The suitability of a participant should be evaluated in accordance with the task's requirements. Besides, participants who have previously collaborated in collusive attacks are not suitable to contribute, since their involvement may result in low quality or falsified contributions. The third important challenge in the success of social participatory sensing is in assessing *trust*. The existence of friendship relations in social participatory sensing brings up new dimensions, since more trust is normally placed on contributions obtained from close friends than those from acquaintances. The participant's social accountability is also a new aspect that should be considered in evaluating the participant's reputation. In this section, we highlight these key issues and explain their challenging aspects.

1.2.1 Sufficiency of Participants

The first challenge in the success of social participatory sensing is access to sufficient number of users, as potential participants, since lack of sufficient participation will inherently reduce the data reliability. Utilising direct friends as data contributors is beneficial, since people normally prefer to help their friends. However, with the lack of adequate friendship relations, the participation insufficiency is still a challenge.

In social participatory sensing, members who are not in a direct friendship relationship may be connected via two or more social links. The set of social links between these members forms a communication path between them. These communication paths can be used for the recruitment of participants and the exchange of communication messages (that are tasks and contributions). The communication

messages may contain the task specification, such as the location of the task or the required expertise. They may also contain sensitive information about the requester or participants which may reveal important information such as their residential address and interests. There are always people who are curious about others' private information and may try to access sensitive information such as a person's whereabouts, by eavesdropping on private communications. So, it is of great importance to recruit participants via the communication paths that better protect the integrity of messages. Furthermore, selecting the communication paths which better preserve the privacy of messages' embedded information is a serious concern.

1.2.2 Suitability of Participants

The second challenge is the selection of suitable participants. Leveraging well-suited participants is bound to increase the quality of obtained contributions, since suitable participants have better knowledge and expertise relevant to the task's requirements.

In the context of participatory sensing, tasks are normally location-based (i.e. contributions should be collected from a specific place) and are to be completed within a specific time period. As such, the suitability of the participant for a campaign is typically related to the participant's geographical and temporal availability as well as the participant's reputation [10]. The geographical and temporal availability is extracted from collecting and analysing the time-stamped location traces of the participants. The reputation of the participant is measured based on the quality of the contributions of the participant in the past.

In social participatory sensing, the existence of public profile information of participants (who are social network members as well), and the social links between them adds new dimensions to the evaluation of a participant's suitability. Through public profile information, access to the participant's interests, expertise and domain specific knowledge is possible. Moreover, the participant's social reputability can be derived from his social relations and interactions. These valuable pieces of information can be used to identify well-suited participants, and hence, overcome the challenge of suitability.

Another important issue that should be considered when evaluating the suitability of participants is the likelihood of their involvement in collusive groups. A group of malicious participants might form a colluding group such that they are recruited in preference to other potentially high-quality workers. The colluding group would then have the power to sway the outcome of the task in accordance with their agenda. So, it is important to identify potentially collusive members and prevent them from being selected as suitable participants.

In order to prevent collusion in participatory sensing, a series of works [23, 24] utilise a trusted platform module (TPM) [25]. TPM is a micro-controller provided with each sensor device to attest the integrity of sensor readings. This local integrity checking makes the system resistant to collusion. However, TPM chips are yet to be widely adopted in mobile devices. In other research that eschews TPM, such as [15, 26], the collusion detection is achieved by leveraging reputation management systems and outlier detection algorithms. The aim is to identify and revoke the colluders by investigating their behaviour and assigning a low reputation score to them.

In the context of social participatory sensing, the existence of social ties between members facilitates the formation of colluding groups. Colluders can easily communicate via the social network communication facilities. They are also able to establish social communities by creating groups in the social network and manage collusive attacks by collaboratively contributing to a series of tasks. They can also easily share their polluted contributions with other group members and hence propagate the bias. So, selecting the participants such that the probability of collusion among the selected members is very low is important for achieving high quality contributions.

1.2.3 Assessing the Trust

The third challenge in the success of social participatory sensing is assessing the trust of contributions and participants. As mentioned before, participatory sensing applications are potentially exposed to incorrect contributions due to their inherent

open nature. The incorrect or polluted contribution may be the result of inaccurate calibration of built-in sensors in mobile devices or careless/malicious behaviour of participants during the sensing action. Lack of trust in contributed data lowers the credibility of the resulting information [15]. So, addressing the issue of trust is considered to be an important challenge [15, 23, 26].

In the context of social participatory sensing, new issues arise. Nowadays, OSNs are not merely a medium to share users' opinions but have evolved to become a platform for disseminating information to a large user base. While beneficial from the point of providing vast amounts of information, the identification of the origin of the data and its credibility becomes more challenging. The reliability of the data is not solely related to its quality, but also dependent on the trustworthiness of its contributor. In other words, it is important to know who and with what level of reputability and trustworthiness produces the data. For instance, following the devastation of Hurricane Sandy in the US in October 2012, social media was flooded with misinformation and fake photos.⁶ While some of these were easy to identify as fake data (e.g., photoshopped images of sharks swimming in New York streets), several other fake pictures and reports were initially thought to be true (e.g., the photo of a storm brewing over lower Manhattan, which was not from Sandy but an April 2011 tornado). An investigation [27] of Twitter feeds collected during this event revealed that only a handful of users contributed to the majority of fake information dissemination. In another analytical study of tweets posted during the terrorist bomb blasts in India (Mumbai, July 2011) [28], it was observed that the majority of users who spread fake information had lower numbers of followers, which is an implicit measure of their credibility. This clearly highlights the importance of participant's social reputability as an influential factor in the reliability of the contributed data. In other words, there is a pressing need for a reputation management system which is responsible for performing necessary validations both from the perspective of data trustworthiness and also the reliability of data contributors.

In participatory sensing applications, the server typically has access to multiple contributions that characterise the same physical phenomenon but originate

⁶<http://news.yahoo.com/10-fake-photos-hurricane-sandy-075500934.html>

from different devices which are related spatially and temporally. In this case, the trustworthiness of contributions is measured by utilising outlier detection algorithms [29, 30]. An outlier detection scheme determines the likelihood of data coming from an untrusted source by measuring its distance to a common value (e.g., the average); the smaller the distance, the more trustworthy the data.

In social participatory sensing, the friendship relations between members add a new dimension to the issue of trust. Typically, more confidence is placed on the data provided by a close friend than the one collected by an acquaintance. So, it is desirable to consider the social trust relationships while evaluating the trustworthiness of contributions. Existing reputation management systems in participatory sensing such as [15, 22, 31, 32] aim at assigning a reputation score to each participant based on the quality of gathered sensor data. However, none of these works have considered participants' social accountability as an important factor in evaluating reputation. Long lasting friendship relations and continuous interactions between people normally translate to greater trust, which in turn, increases the reputation. As such, these solutions cannot be readily adopted in the social participatory sensing context.

1.3 Goal and Methodology

The goal of this dissertation is to address the three important but unresolved issues that are central to the success of social participatory sensing: the participation sufficiency, participant's suitability and trust. Towards this goal, we design novel and initiative schemes that can be utilised to improve social participatory sensing applications in terms of recruiting a sufficient number of well-suited participants and obtaining trustable results.

In order to achieve this goal, we have first conducted a literature review of these issues and identified the shortcomings of existing solutions. We have next designed novel schemes that address the identified challenges based on a comprehensive analysis of the design space. We conducted extensive simulations to evaluate the proposed schemes. We chose to evaluate the performance of our schemes based on

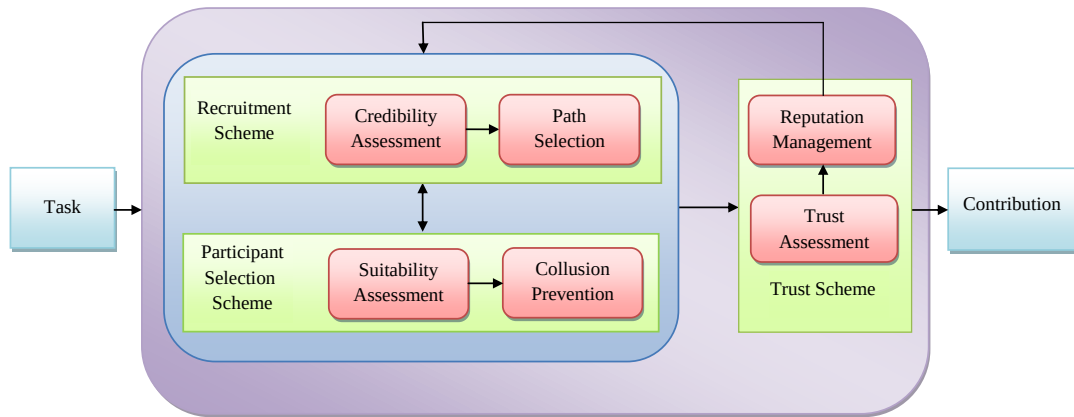


Figure 1.3: The proposed framework for addressing the challenges of social participatory sensing

simulations, as they allow a better exploration of the space of parameter values as compared to analytical modelling. Moreover, since there is a lack of available social participatory sensing systems, it was not possible to evaluate the performance of our developed schemes in real-world applications. However, we have established realistic conditions as a basis for our simulations in order to model a real-world social participatory sensing environment. For example, we have utilised the Advogato web of trust and Wikipedia adminship election datasets. Moreover, we created simulation scenarios that closely resemble situations that occur in real-world participatory sensing applications, e.g., trustable participants who provide low quality contributions for a short burst of time. Finally, we offer interesting insights into the performance of our proposed ideas.

1.4 Overview of Contributions

In order to address the identified challenges, we propose a framework (as depicted in Figure 1.3) to satisfy our research goals. Our contributions can be outlined as follows.

1.4.1 A Recruitment Scheme to Ensure Sufficient Participation

To address the issue of sufficient participation, we propose a trust-aware and privacy preserving recruitment scheme [33, 34, 35]. This scheme proposes the idea of crawling the social graph (starting from the requester) in order to provide access to a pool of participants among friends and friends-of-friends via credible communication paths. The contributions of our proposed recruitment scheme are as follows.

- **Assessing the Credibility of the Communication Paths**

The recruitment of participants requires the transmission of communication messages between the requester and selected participants. The communication messages may contain the task description and its requirements. Hence, it is imperative that the integrity of the messages is preserved. In addition, they may contain sensitive information about the requester or the participants (such as the home/work address, etc.). It is essential to preserve the privacy of the communication messages. So, the requester desires to recruit the participants for whom there exist trustable communication paths in order to preserve the messages' integrity. The requester also prefers that the messages are transmitted via privacy preserving communication paths in order to reduce the probability of sensitive information leakage via intermediate nodes.

Having these preferences in mind, we consider a credible communication path as a path that is (i) trustable and (ii) privacy preserving. We propose a credibility assessment methodology which evaluates the credibility of the communication paths between the requester and the participants [33, 34]. For each existing communication path, the credibility assessment is carried out by considering the trustworthiness of the communication path and its privacy. We propose the use of information entropy to quantify the privacy leak of sensitive information at each intermediate node. Entropy is a measure of the uncertainty in a random variable [36]. Maximising the entropy means the maximisation of the unpredictability of information for an adversary node.

Higher entropy means better privacy for the information contained inside a message. To evaluate the trustworthiness of the path, we assume that the friendship links are weighted by mutual trust rating, which is a dynamic value and is continuously updated according to the trustworthiness of the provided contributions and the requester's reputation score. The trustworthiness of the communication path is obtained by multiplying the mutual trust rates of all links along the communication path. The credibility of the path is then computed by the combining the trust and privacy scores.

• Selecting the Communication Paths

In order to recruit the participants via the best communication paths, we design a path selection scheme. For each selected participant, the path selection scheme identifies all the existing paths between the requester and the participant and selects the most credible communication path. The most credible path is then used for message transfer.

Identifying all the existing paths between two members and selecting the most credible path is beneficial since it offers the *best* path for participant recruitment. However, it leads to the potential increase in time and space complexity of the path selection process. An efficient alternative solution is to consider a customised random surfer for identifying the participants. The customised random surfer is based on the idea of the random walk [37]. Each random surfer begins its journey from the requester and selects the next intermediate node along the path randomly from the set of the requester's friends. The selection of the next node is done on-the-fly, which achieves significant savings in computational cycles and memory for finding the credible paths. In order to give better suited members a greater chance to be selected, the random surfer does not act in a purely random manner, but is biased such that it considers the suitability score of the participants and the pairwise trust scores along the path for the communication path selection.

1.4.2 A Participant Selection Scheme to Ensure Participant Suitability

In order to fully address the challenge of suitability of participants in social participatory sensing, we propose a novel participant selection scheme [33, 35]. Specifically, the contribution of our proposed scheme is as follows.

- **Evaluating the Suitability of Participants**

We carry out a thorough analysis of the design space and propose a suitability assessment scheme for evaluating the suitability of participants to a specific task. A suitable participant is one who is able to satisfy the task's requirements, mainly, the needed expertise and a minimum reputation level (as an indication of being a highly trustable participant). Moreover, the participant is considered as suitable if his recruitment does not impose any privacy threat to the requester's sensitive information. The requester may also prefer to give priority to some participants to be recruited (e.g., due to strong friendship relations). On the contrary, the requester may be reluctant to cooperate with some others due to poor behaviour in previous campaigns.

Keeping these in mind, the suitability assessment scheme identifies the participant's suitability according to the following parameters: (i) the participant's expertise (in order to satisfy the task requirements), (ii) his reputation score (as an indication of being a highly trustable participant), (iii) the pairwise privacy score between the requester and participant (to minimise the privacy breach of requester's sensitive information), (iv) the requester's list of preferred participants (to give priority to those who are preferred by the requester to be recruited), and (v) the requester's blocked list (those with whom, the requester is reluctant to contribute). These parameters are evaluated and combined to build a suitability score for the participant. In addition, the requester may desire to obtain timely contributions, especially in the cases of time critical tasks. In such instances, it is logical to select the participant that has shown timely behaviour in his past contributions, as this participant is most likely to

submit his contribution before the imminent deadline. So, we also take into account a set of time-aware parameters for each suitable participant. These parameters are (i) the selection score (that is the ratio of participants selected so far to the total number of required participants), (ii) the remaining time to the task deadline and (iii) the timeliness of the participant in previous tasks.

• Preventing Collusion

We propose a collusion prevention methodology, which is aimed at preventing the selection of colluding members as suitable participants. In other words, we intend to identify whether the addition of each new participant to the previously selected group will result in the formation of a group of colluders within the suitable participants.

Colluders are like-minded people who collaborate with each other on a specific agenda to obtain an objective by defrauding or gaining an unfair advantage. Their objective may be earning monetary or non-monetary profits. Colluders usually form a group which is large enough to make a considerable impact [38]. Moreover, group members usually target a considerable number of tasks and collaborate together in contributing to these tasks. Their contributions are typically similar to each other (in order to overwhelm the task with similar faulty contributions) and deviate from the other (genuine) participants (so as to change the task's outcome). Finally, the colluders may prefer to connect with each other in the form of social groups to facilitate their communications. Based on these collaborative behaviours, the collusion prevention scheme considers the following collusion indicators: (i) group size (i.e. number of colluders), (ii) group target size (i.e. number of tasks in which colluders have collaborated in the past), (iii) group deviation (i.e. an indicator to show the deviation of content produced by the colluders from those of other honest participants), (iv) group connectivity degree (an indicator to show to what extent the colluders are socially connected to each other), and (v) group content similarity (i.e. the degree of similarity of content produced by the group members). By considering all these indicators, the collusion prevention

scheme determines a collusion probability for each participant and prevents the selection of the colluding participants.

1.4.3 A Trust Scheme to Assess Trust and Reputation

In order to address the trust related challenges in social participatory sensing, we propose a trust scheme [39, 40] that offers a comprehensive view of trust and reputation. The major contributions are as follows.

- **Assessing the Trust**

To address the issue of trust evaluation in social participatory sensing, we propose a trust assessment methodology [39]. In fact, we aim at mimicking human perception of trust while assessing the trustworthiness of contributions. To decide whether to trust a contribution, the requester normally considers two factors. The requester considers the quality of contribution in terms of relevance to the task's specifications and requirements. The requester also takes into account the trustworthiness of the participant contributing the data. A participant who has demonstrated timely behaviour in the past, has relative expertise, or is acquainted with the task area should be considered more trustworthy. The same holds true for the participant who has close friendship relations with the requester, or has had frequent interactions with the requester in the past. So, for each received contribution, these influential factors are evaluated and quantified by the trust assessment scheme. The quality of contribution depends on the sensing modality. Current participatory sensing applications are related to capturing and transmitting a wide variety of sensing modalities such as location and time, pictures, sound samples, acceleration and environmental data. In order to evaluate the quality of contribution, the trust assessment scheme relies on the state-of-the-art methods such as image processing algorithms for image-based contributions and outlier detection algorithms for sound-based contributions. To evaluate the trustworthiness of a participant, in this scheme, we consider the set of personal and social factors

similar to the parameters typically considered by the requester (as mentioned above). In particular, we consider the participant's expertise, his locality (as a measure of his acquaintance with the task area) and his timely behaviour in previous campaigns as personal factors. We also consider the friendship duration between the requester and participant and the timegap between their successive interactions as social factors.

- **Evaluating the Reputation of the Participants**

In addition to the trustworthiness of contribution, it is also important to know the reputability level of the participant who has contributed the data. It is obvious that the data contributed by a highly reputable participant is more reliable than the data from a participant with a low reputation. So, in order to accurately manage the reputation evolution of participants in social participatory sensing, we propose a reputation management scheme [40, 41]. This scheme utilises the well-known Google PageRank algorithm [42] to update the reputation of participants based on their trustworthy behaviours. Specifically, we employ the following innovative ideas.

- We propose the concept of requester subjective evaluation which allows the requester to evaluate each contribution and ascertain how closely it satisfies his needs. Such kind of subjective evaluation is useful especially when it is difficult for the requester to express his real needs, desires or restrictions via task definition.
- In our proposed reputation scheme [40], the pairwise trust score is considered as an important factor in the quantification of a participant's reputation score. People normally have more trust in those who provide them with trustworthy contributions. So, the pairwise trust between the requester and the participant is updated based on the trustworthiness of the participant's contribution. The pairwise trust is increased with each trustable contribution, and decreased if the trustworthiness of a contribution is below a predefined threshold. The amount of increase/decrease

is influenced by the subjective trust rating that the requester assigns to the contribution and the reputation score of the requester.

- Finally, in order to evaluate and assign a reputation score, in our proposed reputation management scheme, we utilise the well-known Google PageRank algorithm. The participant's reputability is dependent on the amount of trust the requesters have of the participant. Besides, the amount of trust that a highly reputable requester has of the participant is more dependable than that of a low-reputable requester. So, it is rational to evaluate the participant's reputability based on the pairwise trusts and the reputability of requesters. The input to the algorithm is then, the pairwise trust scores between the requesters and a selected participant (which as mentioned above, is dependent on the trustworthiness of his contributions), and the requesters' reputation scores. Based on these, a reputation score is calculated and assigned to the participant. This score is further used as a criterion for assessing the suitability of a participant, once being selected for future tasks.

Once the task is defined, the recruitment scheme and the participant selection scheme are utilised to provide access to a pool of members as potential participants who are suitable to contribute to the task. In fact, these two schemes work in parallel to support the requester with a sufficient number of suitable participants who are invited to contribute. Once the selected participants report their contributions, assessing the trust of received contributions and updating the participants' reputation scores are carried out by employing the trust assessment and reputation management schemes.

1.5 Dissertation Organisation

The remainder of this dissertation is organised as follows. We start with a discussion of the current state-of-the-art of the above mentioned issues in Chapter 2. We first explain in more depth the fundamentals and basic concepts of social participatory sensing. We then study the related issues in trust and reputation, participant

selection and recruitment in online social participatory sensing.

In Chapter 3, we present the content of our proposed recruitment scheme. Our proposed scheme leverages multi-hop friendship relationships to identify participants via the most credible communication paths. We first explain the credibility assessment methodology which evaluates the credibility of each communication path between the requester and the participant. The credibility assessment methodology quantifies the trust and privacy scores of each communication path and combines them to assign a credibility score to the path. We then present the details of path selection scheme, which selects the most credible paths for message exchange. The path selection is performed with two configurations. In the first configuration, the network topology is known to the system (such as small-scale social networks within the organisations). This configuration results in obtaining the best communication paths. In the second configuration, we propose an efficient path selection approach and identify participants on-the-fly. In the latter case, the path selection scheme leverages a customised random surfer in order to crawl the requester’s social graph and identify the participants via trustworthy paths. The second configuration may not result in best paths, but it is successful in addressing the bootstrapping problem for the newcomers (who have recently joined the network) by giving them the chance of being selected in competition with more reputable participants. In the last section of the chapter, we present implementation details and results.

In Chapter 4, we present the specification of our participant selection scheme for social participatory sensing. We first present the details of a suitability assessment methodology which identifies the suitability of participants by evaluating the suitability parameters for each. It also considers the participant’s timeliness, the remaining time to the task deadline and the selection score to decide whether to select the member as an eligible participant. We then discuss the collusion prevention scheme which utilises the well-known Frequent Itemset Mining technique [43] and calculates a collusion probability for each eligible participant to prevent any possible collusion on the task. At the end of the chapter, we present the simulation results.

In Chapter 5, we present the details of our trust scheme. This scheme independently assesses the quality of the data and the trustworthiness of the participant and

combines these metrics using fuzzy logic to arrive at a comprehensive trust rating for each contribution. We then go through the details of our proposed reputation management scheme. We first explain the concept of a pairwise trust score and propose a method for updating the pairwise trust scores based on the trustworthiness of contributions. Then we explain the basics of the PageRank algorithm and its usage in the context of reputation management systems. We then discuss the modified version of PageRank that is used for calculating the reputation score of the participants. Finally, we present simulation details and evaluation results.

Finally, in Chapter 6, we make concluding remarks and discuss possible directions for future work.

Chapter 2

Background and State of the Art

2.1 Introduction

The ultimate goal of social participatory sensing is to address the serious challenge of insufficient participation by integrating participatory sensing and online social networks. In particular, this integration benefits from the considerably large numbers of existing friendship relations between members of online social networks, which can provide a pool of potential participants. Figure 2.1 illustrates the correlation of social participatory sensing with other collective intelligence systems. Social participatory sensing is a crowdsourcing system in which, workers are selected from social network friends who are then recruited to contribute to the tasks by utilising their mobile phones for data collection. As mentioned in Chapter 1, in order for the social participatory sensing to be successful, there are three main challenges to be addressed: the sufficiency of participation, selecting and recruiting well-suited participants, and assessing trust.

In this chapter, we present an overview of the state-of-the-art of the above mentioned issues in social participatory sensing. Particularly, in Section 2.2 we study the fundamental concepts of social participatory sensing and outline the general architecture and components. Then, in Section 2.3 we discuss the trust and reputation challenges and the proposed approaches which aim at addressing these issues. In Section 2.4, we outline the related work on participant selection methods in on-

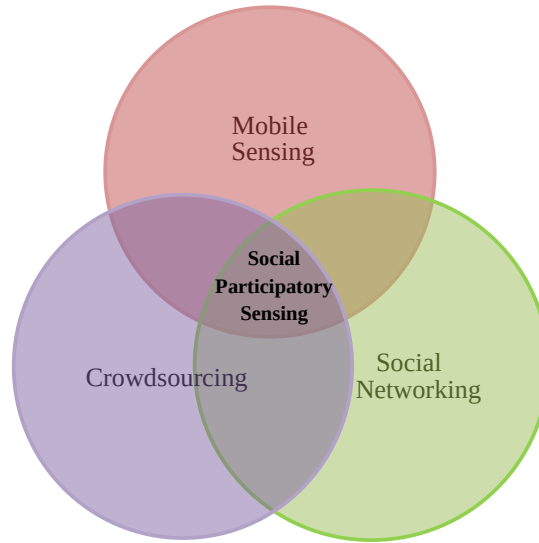


Figure 2.1: Social participatory sensing and its relation to other collective intelligence systems

line communities and participatory sensing. In Section 2.5, we study recruitment methods from the literature. We summarise this chapter in Section 2.6.

2.2 Social Participatory Sensing

2.2.1 Background

In recent years, we have witnessed tremendous improvement in mobile phone technologies in terms of processing power, storage capacities, sensing capabilities and network data rates. These improvements have transformed mobile phones into multifaceted devices that are capable of communicating, computing and sensing. It is thus no surprise that there are over 6.8 billion mobile users around the world which is equivalent to 97 percent of the world population.⁷ These advances in mobile phone technology in concert with their pervasiveness have given rise to an exciting new paradigm known as participatory sensing [1, 19, 44]. In participatory sensing, the key idea is to recruit ordinary people to contribute voluntarily in sensor data collection using their mobile phones. In fact, participatory sensing can readily compliment wireless sensor infrastructure deployments by the involvement of sensors already ex-

⁷<http://mobithinking.com/mobile-marketing-tools/latest-mobile-stats>

isting in people's hands. Moreover, participatory sensing allows people to focus on sensing their immediate surroundings, leading to several exciting applications that can directly benefit mankind.

For participatory sensing to be a success, a key challenge is the recruitment of sufficient volunteers. Typically, in participatory sensing campaigns there is no explicit incentive for participation and people contribute altruistically (we discuss incentive-based approaches in Section 2.4). Without adequate motivation, participants may not be willing to contribute, which in turn, reduces the data reliability. This is perhaps one of the main reasons why we are yet to see a killer participatory sensing application. While several novel applications have been proposed in recent years, most of them have only been tested with a small set of participants. Getting sufficient people to contribute to a participatory task is the *Achilles' heel* to its widespread acceptance.

The aim of social participatory sensing is to address these issues by integrating online social networks with participatory sensing systems [11, 12]. Social friendship relations can be used as a means to access and recruit friends and friends-of-friends as potential participants. By virtue of social networks, people are able to create an online identity, basically a public profile to share their photos, update their status and express their interests. This public profile can also be used to show the level of contribution of the participant to the task. In particular, this new paradigm offers the following advantages. First, social networks can enable organisers to identify and reach well-suited participants for data collections based on their geographic availability as well as their interests, expertise and habits. Indeed, participatory sensing applications are usually initiated for a specific geographic area, and desire to recruit people within a specific locality. Public information on member profiles could be used to select suitable members (or even whole groups) and invite them to participate. Second, the sense of community that exists in OSNs is known to be a motivator for contributing to online collaborations [45, 46]. Those who feel an attachment to a community are often willing to contribute altruistically for the benefit of that community. In the context of social participatory sensing, being part of social groups including family members, friends or colleagues may act as an

effective motivation for participation, since people normally like to be helpful to their friends [46]. Third, the desire for prestige is known to be another key motivation of individuals' contributions to the group [47]. High quality data may increase one's prestige in the community. In the domain of social participatory sensing, incentives in the form of e-coins [14] or reputation points could be devoted to well-behaved participants, and can be publicly available in participant profiles to express the degree of their contribution to the sensing campaign.

A pertinent example of such a system is Jelly⁸ which is built on top of existing social networks like Facebook and Twitter. When users encounter something unusual, they can take a picture of the object, formulate a query and submit it to their social network. Their friends, who also have the application, receive the query and respond to their friend's question with a link, by drawing on the original image, or simply using text. Another instantiation of the concept of social participatory sensing is found in [48] in which, Twitter was used as the underlying social network substrate to pave the way for ubiquitous crowdsensing and collaboration applications. The proposed system was tested in the context of two smartphone applications. The former, called weather radar, relied on textual tweets to collect data about the weather condition. Twitter members were asked to report the weather condition by tweeting 0 for sunny, 1 for cloudy, 2 for rainy, and 3 for snowy. The latter was devised to provide a noise level querying service over Twitter by aggregating the automatic noise-sensing updates from smartphones. The noise samples were mapped into three categories, low, medium and high and then, forwarded to Twitter. Through their experiments, the authors observed participation with around 15% reply rates and low reply latency (50% replies arrived in 30 minutes and 80% replies arrived in 2 hours) from Twitter members even without an incentive structure. These examples demonstrate the suitability of online social networks for tasking/utilising smartphones and pave the way for ubiquitous crowd-sourced sensing and social collaboration applications.

In a typical social participatory sensing system, a social network member may serve as a requester. The requester defines the task and specifies its requirements

⁸<http://blog.jelly.co/post/72563498393/introducing-jelly>

such as the location that the contributions should be collected, the required expertise to contribute to the task, etc. The requester then disseminates the task to his friends through friendship links (via email, private message or by writing as a post on their profiles (e.g., Facebook wall)) and asks them to provide contributions according to the task's needs and specifications. Some friends in turn, will accept to contribute, so they gather sensor data by their mobile phones and send it back to the requester. The sensor data is then aggregated and the required information will be available to the requester. Other entities can also make use of the collected data. Examples are participants willing to consult their own collected data, scientists attempting to gain insights about the monitored phenomena, health professionals checking patient data, or the general public [9].

2.2.2 Components

Social participatory sensing consists of two main components: online social network and participatory sensing.

Online Social Network. OSNs have become extremely popular in recent years. Already in 2014, Facebook has 1.23 billion monthly active users,⁹ 883 million users have registered with Twitter,¹⁰ and 277 million users have accounts on LinkedIn.¹¹ An online social network is defined as a web-based service that authorises people to (i) construct a public or semi-public profile within a bounded platform, (ii) create a list of other members with whom they share a connection, and (3) view and traverse their own list of connections and those established by others within the system [17]. The nature and nomenclature of these connections may vary in different OSNs.

A typical online social network has three main layers, data storage, control management and application layer [18] as shown in Figure 2.2. In the following, we present a short review of these layers.

- **Data Storage Layer.** This layer consists of two components. The storage

⁹<http://newsroom.fb.com/Key-Facts>

¹⁰<http://www.businessinsider.com.au/twitter-total-registered-users-v-monthly-active-users-2013-11>

¹¹<http://press.linkedin.com/about/>

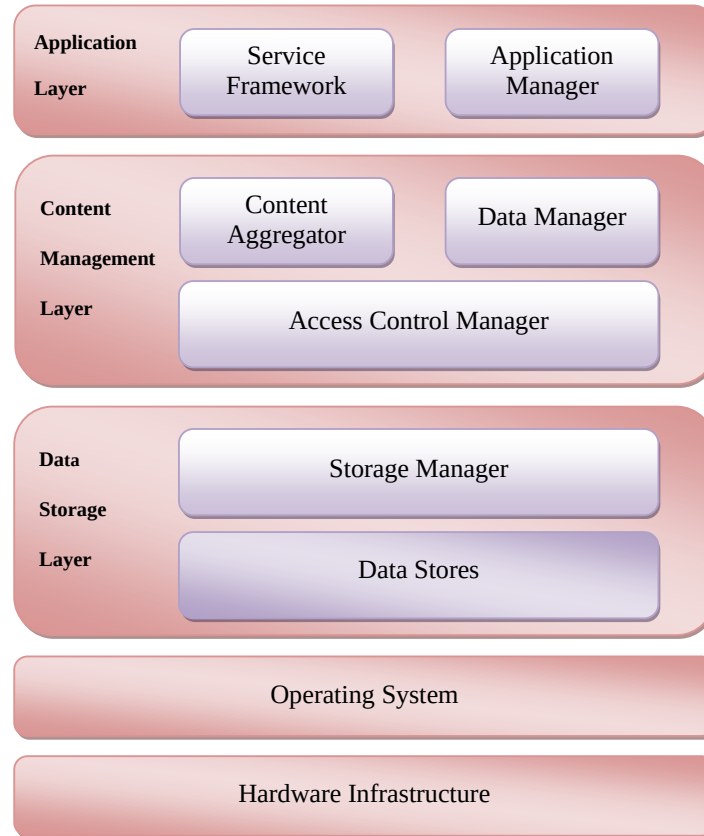


Figure 2.2: Architecture of a typical online social network

manager component is responsible for storing the information of social graphs on OSN databases. The data store component includes a set of storage elements such as user profile databases that accumulate information items of online social networks.

- **Content Management Layer.** This layer has three components. The content aggregator component is responsible for gathering and managing the content from remote OSNs. The data manager component facilitates the storage and retrieval of the social graph information. The access control component, as its name implies, controls access of social network members by establishing and managing an access control scheme.
- **Application Layer.** OSNs normally provide their members with various applications such as search, messaging, news feed, etc. The application layer

is responsible for managing the utilisation of different applications on social networks. This layer comprises two components. The application manager facilitates the interaction of members through a set of APIs. The service framework component enables users to develop their own applications and services on OSNs.

OSNs provide us with the ability to stay in contact with friends and family and allow us to maintain closer ties to our loved ones across long distances. Furthermore, they support us with the ability to increase our networking potential or work with others regardless of distance, presenting new ways for us to do business.

Participatory Sensing Systems. Mobile phones have now virtually become ubiquitous with almost every individual on the planet using one for basic connectivity. The new generation of smartphones contain a number of specialised sensors such as the: ambient light sensor, accelerometer, gyroscope, GPS, proximity sensor, and general purpose sensors including the microphone and camera.

These advances in mobile phone technology in concert with their ubiquity have brought about a new exciting platform, called participatory sensing. Participatory sensing emphasises the involvement of citizens and community groups in the process of sensing and documenting where they live, work, and play [9].

A participatory sensing system is generally organised in a client-server architecture and contains a set of components (as depicted in Figure 2.3) that interact to gather sensor data and provide the end users with the resulting information. Specifically, it comprises the following components:

- **Sensing Component.** This component resides on the participant's mobile phone and is responsible for performing the sensing process and providing necessary sensor data for the task undertaken. Most sensor readings are geo-tagged and are typically in the form of images (photos from the environment), sound (noise samples) or context data (heart beat rate, blood pressure). Some tasks may also require multi-modal sensing. The sensing process can be done automatically (in the background) without distracting the mobile user. This type of sensing is known as opportunistic sensing [19]. Opportunistic sensing

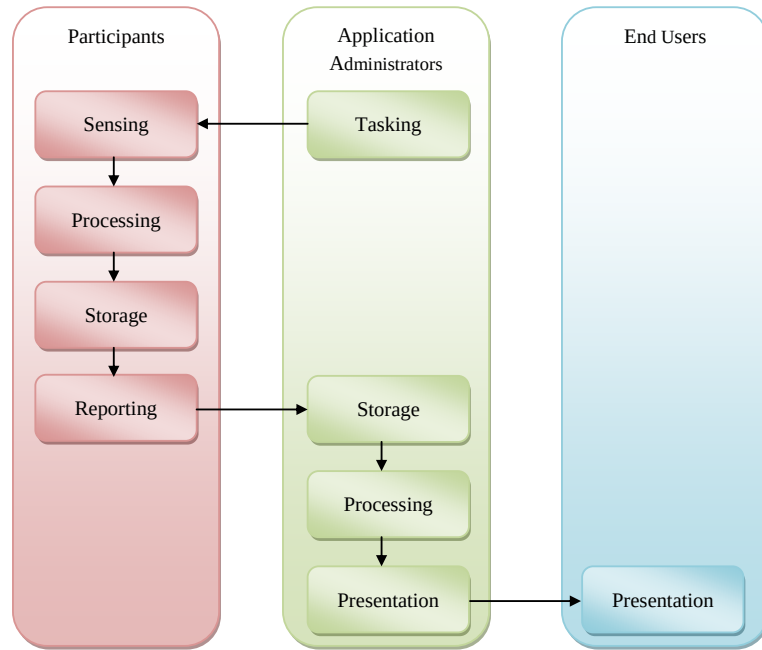


Figure 2.3: Architecture of a typical participatory sensing system

is appropriate for applications that require continuous sensor data collection. Examples are Ear-Phone [5] and PetrolWatch [2] in which, noise samples and fuel prices are automatically captured by mobile phones, respectively. The sensing process can also be done manually requiring some action by the participant (e.g. taking a photo). This sensing mode is normally used when the sensing action should be done upon the detection of the relevant event. For example in LiveCompare [3], the participants use their mobile phones to take pictures of the price tags, when they are out shopping. This is referred to as participatory sensing in the literature, since members actively participate in the sensing process [49].

- **Tasking Component.** This component is located at the application server and is responsible for disseminating the tasks to participants' mobile phones. The task is constructed according to the application requirements and should specify the time, location and other requirements of sensing process, such as sensor type and sampling frequency.
- **Reporting Component.** This component resides on the participant's mobile

phone and transfers the sensor readings collected by the sensing component to the application server. The existing communication infrastructures such as Wi-Fi and cellular networks can be used for data transmission.

- **Storage Component.** The sensor readings are stored temporarily on the mobile phone to be further processed and transmitted to the server. The server, on the other hand, manages the long-term storage of reported contributions. The storage component ensures the storage of collected sensor readings on the mobile phone. It also makes sure that the reported contribution has been stored in the server.
- **Processing Component.** The sensor readings may require processing so that useful information can be extracted from the data. The processing component is responsible for performing such processing on the mobile phone or the server. For example, in the Ear-Phone application [5], raw sound samples can be processed in the mobile phone in order to extract the noise level. At the server, the reported contributions from multiple users are combined to compute statistics and prepare results for end users and application providers.
- **Presentation Component.** This component is responsible for providing the results obtained from the processing component to the end users. The obtained results can be presented to end users in the form of graphs, maps, or even raw data which can be further analysed by themselves [50]. The results can be available for end users in different ways. They may be directly shown on the participants' mobile phones. Alternatively, they can be presented through web portals to a larger community.

Since the emergence of participatory sensing system, many applications have been developed that are aimed at collecting different kinds of sensor data and providing useful information for the users. As mentioned in Chapter 1, participatory sensing applications are typically classified into two categories: people centric or environment centric [9, 19, 49]. People-centric applications mainly focus on collecting sensor data about the participants and their behaviours. For example, DietSense [20]

allows participants to self-monitor their food and dietary options by taking photos of their everyday meals.

BikeNet [21] is an application for quantifying the cycling experience for the keen bicyclist. The data is collected by multiple body-area peripheral sensors including: microphone, magnetometer, pedal-speed sensor, inclinometer, lateral tilt, stress monitor, speedometer and a sensor for CO₂, all of which are connected wirelessly to a smartphone. The data collected can be physiological (heart rate, galvanic skin response), or may measure performance attributes (wheel speed, pedalling, and cadence). It is also possible to collect data about the environment around the route (such as pollution, noise levels, and irregularities of the roads). In PetrolWatch [2], a mobile camera phone is leveraged to take pictures of fuel price boards when the car approaches a fuel station. These images are then transported to a central server where computer vision algorithms are implemented for board detection and fuel price extraction. Participants can then make a query and ask the system to provide them the cheapest fuel station near a certain location.

On the other hand, environment-centric applications try to observe and collect sensing data from the environment. Ear-Phone [5] and NoiseTube [4] enable citizens to measure the noise level by their mobile phones from their everyday environment. Each user can participate in the creation of a collective map of noise pollution by sharing the geolocalised measurement data with the NoiseTube community. Ikarus [51] exploits sensor data collected during the flights of paraglider pilots to study thermal effects in the atmosphere. Cartel [50] is another example of this category of applications that utilises mobile phones carried in vehicles to collect information about traffic, quality of roadside Wi-Fi access points, and road conditions (e.g., presence of potholes).

The success of participatory sensing is largely dependent on adequate participation. The marriage of online social networks with participatory sensing (resulting in social participatory sensing) is helpful in lowering this barrier by supporting access to a pool of users who can be motivated to participate in participatory sensing campaigns. As mentioned in Chapter 1, this dissertation mainly focuses on addressing the issues and challenges in social participatory sensing, mainly, participation

sufficiency, suitability of participants, and trust related issues. In the following, we present an overview of the state-of-the-art methods used to address these challenges.

2.3 Trust Assessment in Social Participatory Sensing

Assessing trust is an important issue in virtually all crowdsourcing systems. According to the Crosby’s definition [52], the quality of the outcome of a task is “the extent to which the provided outcome fulfils the requirements of the requester”. Long-term trustworthy behaviour typically results in increased reputation. Reputation is an important criterion that can be used to identify suitable participants. So, accurate assessment of trust and comprehensive evaluation of a participant’s reputation score are challenging.

Since trust and reputation management have not been previously investigated in the context of social participatory sensing, we first provide a short review of the trust assessment methods in online communities and then present related work in participatory sensing. We adopt a similar approach with regards to related work on reputation management.

2.3.1 Trust Assessment in Online Communities

In online communities, there are multiple approaches to evaluate trust [53, 54]. In the following, we discuss the proposed methods in detail.

- **Expert Review**

In this approach, experts are hired to evaluate contributions, and the result of their evaluation is used by the requester to accept/reject the contribution. For example, in Wikipedia, administrators are assumed to be experts and have the authority to curate the articles [55]. Another example is Stack Overflow,¹²

¹²<http://stackoverflow.com>

which is a question and answer website that is widely used by computer programmers with 1,900,000 registered users ¹³ (as of August 2013). Users of Sack Overflow can earn reputation points (e.g., a person is awarded 10 reputation points for receiving an ‘up’ vote on an answer given to a question), and can receive badges for their valued contributions ¹⁴. In this context, only the users with specific levels of reputation can edit or delete posts. ¹⁵

Although this approach is widely used in practice, it can be time consuming for the tasks with a large number of contributions. Moreover, employing experts may incur additional cost.

- **Forced Agreement**

In this approach, the contribution evaluation is performed based on worker agreements [53]. The agreement may be obtained on the output of the task. In this case (which is called output agreement), two or more contributors work independently and simultaneously in different locations and receive the same input. The answer is only accepted if the pair can agree on the same description (i.e. the same *output*) for the input. A similar approach has been used by the ESP game for online image labelling [56]. Conversely, in the input agreement approach [57], two workers receive input that might or might not be the same. They must then describe the given input to each other. Based on the received descriptions, the workers decide whether they are dealing with the same input. If both workers agree on the *input* similarity, the description is inferred to be correct. This method has been used in Tag-a-Tune game, which collects descriptions of music clips. The players insert their description after listening to a sound file. If both players agree on whether or not the other’s descriptions describe the same sound clip, then the descriptions are considered to be relevant.

The advantage of these two approaches is providing the system with fast evaluations. However, they are only applicable to a limited range of simple tasks

¹³<http://stackoverflow.com/users>

¹⁴<http://stackoverflow.com/help/whats-reputation>

¹⁵<http://meta.stackoverflow.com/questions/5221/how-does-deleting-work-what-can-cause-a-post-to-be-deleted-and-what-does-that>

and cannot be used for more complex tasks.

- **Contributor Evaluation**

In the contributor evaluation approach, contributions are evaluated based on trustworthiness of their contributors. If contributing participants have the required reputation level, credentials, or experience, their contributions are considered as trustable. On the contrary, contributions received from contributors with low reputation may be revoked. In Amazon Mechanical Turk (Mturk), a worker who frequently submits low quality work may be assigned a low reputation and blocked from accessing future tasks [58].

- **Majority Consensus**

In the majority consensus approach, the requester submits several instances of the same task to the crowdsourcing platform. The requester then chooses the answer on which the majority of workers agree as the most appropriate one. This approach is commonly used in MTurk, where the task is submitted to several people and majority voting mechanism is used to choose the final answer. Majority consensus can also be used to define the result of image labelling tasks [59]. The annotation that receives the maximum number of votes is treated as the final aggregated label.

2.3.2 Trust Assessment in Participatory Sensing

In participatory sensing, trustworthiness can be viewed as the quality of the data that the participant produced by sensing via his mobile phone. The quality of sensor data is related to the extent it satisfies the task's requirements. Tasks typically specify the sensing modalities based on the application requirements, the sensors to be used, and the sampling frequency. They may also contain information about location and/or time frame of interest. In order to ascertain the trustworthiness of the data, it is highly desirable to ensure that sensor data has been captured from the prescribed location and time. In a series of works [60, 61], the authors propose a location/time attestation method in which, a tag is assigned to the sensor data by

content producers which is further used by a verification service to verify the location/time of data collection. In [60], a secure service has been proposed which allows participants to tag their content with a spatial time-stamp indicating its physical location, which is later used by a co-located infrastructure for location/time verification. A similar approach has been used in [61], in which, a proof of location is presented, which is a small piece of meta-data issued by a wireless infrastructure in coordination with a mobile device. Any device can request a time-stamped signed location proof from the infrastructure in its communication range. The above approaches rely on existing infrastructure which limits their scalability. Furthermore, such kind of verification limits the participatory sensing applications to where the required infrastructures have been readily installed.

The problem of verifying data received from user devices in participatory sensing has also been studied in [3, 23, 24, 62]. In [3], the data integrity issue is addressed by requesting manual validation from participants in cases of doubtful submissions. This approach can quickly become exhausting for users, if they are encountered with a considerable increase in the number of doubtful data. It may also potentially increase the network traffic as the number of applications grows. The works in [23, 24, 62] addressed the data integrity from a different point of view. They aim to ascertain that uploaded data exactly corresponds to the original data collected by the mobile phone sensors and has not been changed unintentionally or maliciously. Particularly, they assume that there exists a malicious user (or a malicious program) who is capable of tampering with software running on the phones and corrupting the sensor data. Their solutions rely on a Trusted Platform Module (TPM) [25], which is a micro-controller that resides on the mobile device and provides it with hardware-based cryptography as well as secure storage for sensitive credentials. In [23], each device has a trusted hardware element that implements cryptographic algorithms for content protection and prevention of software modification threats. In [24], two TPM-based design alternatives are presented: the first architecture relies on a piece of trusted code and the second design incorporates trusted computing primitives into sensors to enable them sign their readings. YouProve [62] is another TPM-based architecture that allows client applications to directly control the fidelity of

data they upload and services to verify that the meaning of source data is preserved. However, TPM-enabled mobile phones are not widely available as of yet. As such, their solutions are not readily deployable.

Common to all the above approaches is their focus on data integrity; that is, verifying and confirming that the contributed data is indeed from the participant device and was collected at the claimed location/time. However, unlike opportunistic sensing where the data comes from automatic readings from devices, in participatory sensing, contributed data is more subjective and includes users' participation. Therefore, there is a need for assessing the quality of contributions (in addition to data integrity) by factoring in the participants' behaviour.

In related research such as [15, 63], authors aim to address this issue. In particular, they consider scenarios in which some of the participants may deliberately affect the result of the readings from the sensor (e.g., by keeping the mobile phone in their pockets while collecting sound samples), and thus introduce poor quality information. They leverage majority consensus to assign a cooperative trust rating to each contributing device and use that rating to revoke untrusted devices. In the majority consensus, participants contribute to the same task and then the majority answer is taken as the correct result. Each device is then associated with a weight which is inversely proportional to the deviation between the device sample and the group consensus (e.g., a device which reports a value that is significantly different from the group consensus is assigned a low weight). However, this method is vulnerable to collusion. A group of participants may create a colluding group and collaboratively contribute the same (or nearly the same) malicious contribution, which leads to the selection of their polluted contribution as the final result.

To sum up, most of the aforementioned approaches do not focus on participants and their influence on the trustworthiness and quality of provided contributions. Besides, none of them has considered the participant's personal characteristics and his social accountability as the prominent factors in the quality of his contribution. Participant's personal specifications (such as his expertise or timely behaviour) as well as his social accountability (e.g., being a close friend, etc.) are influential factors on the trustworthiness of contributions. As such, these methods cannot be readily

used in the context of social participatory sensing.

A key issue in achieving trustable contributions is recruiting suitable participants. Reputation is a criterion that can be used to identify participant's suitability. So, accurate evaluation of participant's reputation score in a way that covers all effective issues is challenging. In the following, we first investigate this issue in online communities and then, discuss the related work in participatory sensing.

2.3.3 Reputation Management in Online Communities

The issue of trust and reputation has been omnipresent since the conception of the World Wide Web and is still a hotly debated problem. In online communities, people are being evaluated based on the quality of their contributions. These evaluations are aggregated to build a community-wide quality metric called *reputation score* for each person. Reliance on reputation scores is a popular method for people selection [64, 65].

Existing techniques for reputation management are classified into two categories: feedback based and content based. In feedback-based approaches, reputation score is computed using direct feedback explicitly received from other community members [66, 67]. The approach adopted by eBay [68] is an example of using direct feedback to build reputations. After completion of a transaction, the seller and the buyer can rate each other using the following scale: positive (+1), neutral (0), negative (-1) along with a descriptive comment. Based on these evaluations, the reputation score is calculated for every seller to be the sum of positive ratings (from unique users) minus the sum of negative ratings (from unique users). In order to provide information about a seller's more recent behaviour, the total of positive, negative and neutral ratings for the three different time windows (i.e. past 6 months, past month, past 7 days) along with his reputation score are publicly displayed [68]. Such information helps the buyers to buy from reputable sellers, in order to obtain high quality products.

In content-based approaches, feedback is implicit and extracted from the behaviour of evaluators [64, 69]. For example, in WikiTrust [69, 70], a reputation

management tool has been designed for assessing Wikipedia users. WikiTrust uses sophisticated data mining algorithms that assess the reputability of editors by tracking their contributions. It uses two criteria for awarding reputation: the survival of text, and the survival of edits. Text survival is the lifetime of a text fragment entered by an author (counted in terms of the number of subsequent edits that the text survives). In fact, it denotes how long the text entered by the author remains unchanged in a Wikipedia article. Edit survival captures how long the modifications performed by an author are kept in the article. If the contribution (or edition) survives for a long time, its quality is essentially proven since it shows that all subsequent editors to the same page have already implicitly voted on the contribution by leaving it in place. The amount of increase/decrease in the author's reputation score is thus based on the quality of the updates the author makes in the content. In particular, it is proportional to the amount of residual text and the similarity between the author's edit and the latest version of the article. If the update is preserved, then the user's reputation is increased. On the other hand, the user's reputation is decreased if the update is deemed to be incorrect.

The concept of reputation management has also been studied in peer-to-peer (P2P) networks [71]. The reputation of a peer is computed based on the opinion of its direct transacting partners as well as some third-party peers. In this approach, a peer A that wishes to know the reputation of another peer B, can ask some peers (e.g., its neighbours) to provide their opinion on B. A then combines the opinion from the peers to calculate B's reputation. In fact, in a fully distributed P2P system involving numerous peers, it is often impossible or too costly to obtain the opinions of all interacting peers with a given peer. Instead, the reputation score is based on a subset of opinions usually from the relying party's neighbourhood [65].

2.3.4 Reputation Management in Participatory Sensing

The issue of managing reputations of contributing volunteers in participatory sensing is an active area of research. The aim of the reputation management system is to associate a reputation score with each contributing device that reflects the level of

trust perceived by the application server about the data uploaded by that device over a period of time [15].

Authors in [26] propose a reputation-based framework which makes use of Bayesian reputation systems [65, 72, 73] to assign a reputation score to each sensor node in a wireless sensor network. Bayesian systems take binary ratings as input (i.e. positive or negative), and are based on computing reputation scores by statistical updating of Beta probability density functions. The proposed framework consists of a watchdog module which aims to detect the presence of invalid data resulting from compromised and faulty nodes. This module leverages outlier detection algorithms [29, 30] to assign a rating, inferred as the level of confidence, to each data reading. The framework also includes a reputation module which is responsible for maintaining the reputation of a node which reflects the level of trust given to its sensor data over a period of time. The authors adopted the Beta distribution framework [74] to compute node reputation. The resulting reputation scores are shown to be effective in isolating faulty sensors of various types.

A reputation framework for participatory sensing has been proposed in [15]. Similar to the previous work [26], it consists of a watchdog module which monitors the short-term behaviour of sensor devices to assign them a cooperative rating indicating the probability of the device being cooperative. These cooperative ratings are produced by executing a consensus-based outlier detection algorithm which works on group consistency and uses the deviations from a common consensus to identify outliers. The cooperative ratings are then fed to a reputation module which utilises a Gompertz function [75] to assign a reputation score to sensor devices. The reputation scores can then be used by the application server to compute the average statistic by weighting the sensor samples according to the device reputation scores, or they can be used as a feedback to filter contributions from untrustworthy devices.

Authors have extended their work in [31] and proposed a reputation anonymisation scheme which is aimed at preventing the privacy leakage due to the inherent relationship between reputation information. The reputation management system requires the system to know the history of user sensing actions to update the reputation score. This requirement is in conflict with the case wherein participants may

constantly change their pseudo-identities to preserve privacy, since transferring the reputation information between pseudonyms leads to the de-anonymisation of users. To address this issue, they propose an anonymisation scheme based on the concept of k -anonymity [76] to eliminate the uniqueness in the transitions of user reputation and prevent an adversary from de-anonymising users.

The work in [22] aims to address the same problem by proposing an anonymity-preserving reputation scheme which utilises blind signature [77] to provide a secure transfer of reputation scores between pseudonyms. It also cloaks exact reputation values into reputation groups. In theory, utilising the blind signatures ensures the authenticity of signed messages without revealing their content to the signing entity and prevents the signing entity from linking the message content with the identity of its creator. However, in practice, an adversary may still be able to track the reputation scores over several time intervals and link pseudonyms used in different periods. To prevent such attacks, authors propose that the participants cloak their reputation scores before their transfer. Their solution eliminates the assumption that the reputation and pseudonym managers must be trusted.

Authors in [32, 78] also address the problem of trust without identity by proposing a framework to compute the trustworthiness of sensing reports based on anonymous user reputation levels. Their proposed framework consists of a trust assessment module which evaluates the trustworthiness of contributions according to a set of contextual parameters such as location, time, sensor mode (i.e. whether it is a text report, voice clip, picture or video, etc.) and travelling mode. They propose a number of reputation levels to approximate the raw reputation score values. Based on the participant's obtained trust value and his current reputation level, the reputation feedback level is calculated for the participant. Their approach utilises blind signatures [77] to make the report submission and reputation update as two separate processes. Unlike the previous mentioned work that address the same challenge [31], their proposed solution does not require a trusted third party.

Common to all above solutions for reputation management in participatory sensing is their emphasis on data trustworthiness as the dominant factor in participant's reputation. However, in the context of social participatory sensing where social

relations play an important role in people's accountability, concentrating solely on data trustworthiness does not provide a comprehensive view of the reputation of individuals. Neglecting a participant's social accountability in terms of length and strength of friendship relations, interactions, helpfulness, etc. prevents achieving a comprehensive and accurate evaluation of the participant's reputation. As such, these solutions are not applicable in the context of social participatory sensing.

2.4 Participant Selection in Social Participatory Sensing

It is evident that the volume and the diversity of participants with different perspectives and backgrounds can lead to accurate high quality contributions. So, selecting suitable participants is of great importance in social participatory sensing.

To the best of our knowledge, the issue of participant selection in social participatory sensing hasn't been addressed in prior work. As such, we discuss related research focussing on selection issues in online communities and participatory sensing.

2.4.1 Participant Selection in Online Communities

Crowdsourcing systems commonly use one of three worker selection approaches: Open-call, qualification-based, and publish-subscribe [48, 54, 79, 80]. In the following, we discuss these approaches in detail.

- **Open-Call.** In this approach, there is no participant selection and each worker is able to contribute to the task. Wikipedia, Threadless,¹⁶ and the ESP game [81] use this approach, which is simple to implement and easy to use. For example, in Wikipedia, anyone can edit almost every page. In Threadless which is an online community of artists, designs are submitted online and are put to a public vote and each member is able to vote for his favourite

¹⁶<http://www.threadless.com>

design. Also in the ESP game, each user can take part in the image labelling game and receive incentives. While this openness, which can reach out to and attract members with different knowledge and interests, is an advantage of crowdsourcing, it may lead to selecting low quality workers and make quality assurance particularly challenging.

- **Qualification-Based.** This approach uses qualifications to select workers. The worker's qualification is affected by a series of parameters such as the quality of his previous work, number of his approved works, his skills, reputation, etc. Thus, smaller bespoke crowds can be assembled out of the workforce to complete highly specialized tasks [80]. Amazon Mechanical Turk is an example of the systems that utilise this approach [58, 82]. The requester can use qualifications to control which workers can perform his HITs (Human Intelligence Tasks). A HIT can have qualification requirements for a worker that must meet before the worker is allowed to accept the HIT. A requirement can also state that a worker must meet the requirement to see the HIT's question data when previewing the HIT.

Among the set of qualifications, reputation is probably the most important parameter. Reputation is in fact an overall estimation of a worker's quality and trustworthiness. An investigation of existing reputation management systems has been proposed in Section 2.3. Although reputation-based approaches have well-engineered foundations, they are prone to various types of attacks [54, 83]. The attacker may aim to boost or downgrade reputation scores of specific target objects. These objects either gain or lose advantage due to inaccurate reputation scores when competing with similar objects for users' attention or preference.

Expertise is another important parameter in the qualification of participant. Expertise-based participant selection consists of identifying users with relevant expertise or experience for a given topic. Expert finding has been extensively studied in social networks [84, 85, 86, 87]. Authors in [86] developed a Bayesian hierarchical model for expert finding that accounts for both social relationships and content. The model assumes that social links are determined by expertise

similarity between candidates. A propagation based approach for finding an expert in a social network has also been proposed in [87]. The approach consists of two steps. In the first step, an initial expert score for each person is estimated based on his local information, and then, the top ranked persons are selected as candidates. The selected persons are used to construct a sub-graph. In the second step, one's expert score is propagated to the persons with whom the person has relationships.

- **Publish-Subscribe.** In this approach, the task assignment is done on the basis of a publish/subscribe service. In particular, the participant (subscriber) shares his interests and preferences about a topic by subscribing to the server, and the requester (publisher) posts and forwards messages to the interested participants only [48, 79]. The main challenge in this approach is the lack of sufficiently qualified participants to attend to tasks that need specific knowledge or expertise [80, 88].

In addition to the shortcomings mentioned above, the main challenge in the above mentioned participant selection methods is their vulnerability to the colluding attacks. A group of malicious workers might form a colluding group with the aim of attacking the task, so that they are recruited as honest workers. The colluding group would then contribute polluted data with the aim of swaying the result of the task in accordance with their agenda.

2.4.2 Participant Selection in Participatory Sensing

In participatory sensing applications where participants are typically recruited to contribute to the tasks with defined spatiotemporal specifications, new selection criteria are brought up in addition to the general parameters (e.g., reputation). In [89, 10, 90], two indexes has been proposed for the identification of well-suited participants. Cross-campaign metrics provide a granular view of a participant's past performance across many campaigns. They include number of: previous campaigns volunteered, participated in, and abandoned. Campaign-specific metrics measure the quality and quantity of samples that can be expected for a specific data col-

lection. They include a set of parameters. The first parameter is timeliness, which represents the latency between when a phenomenon occurs and when the sample is available for a data processing unit. The second parameter is called capture, which describes the quality of a particular reading in terms of the ability in determining a particular feature. The third parameter is called relevancy, which specifies how well the sample describes the phenomenon that is sought for capture. The fourth parameter is called coverage, which represents the spatial and temporal availability associated with the coverage provided by the participant. The geographical and temporal availability is extracted from collecting and analysing the location traces in the form of latitude, longitude and time points for a period of time. The last parameter is called responsiveness which describes the probability of responding to a directed sensing request. These parameters are then combined to determine the overall reputation for the participant on a per campaign basis. However, in [10], they have limited reputation to considering participants' willingness (given the opportunity, is data collected) and diligence in collecting samples (timeliness, relevance and quality of data).

Also in [91], a distributed recruitment and data collection framework for opportunistic sensing has been proposed. The recruitment component exploits the suitability of user behaviours based on the mobility history information, and recruits only the nodes that are likely to be in the sensing area when the sensing activity is taking place. As a distributed recruitment framework, a set of recruiting nodes visit the sensor area before the campaign is launched and then, disseminate recruitment messages. In order to transfer collected sensor data to the requester, a collection of nodes called data sinks are used and participating nodes opportunistically exploit ad hoc encounters to reach data sinks that are temporarily deployed in the sensed area.

To sum up, the above mentioned methods aim at selecting the suitable participants according to the specifications and requirements of participatory sensing applications, mainly, spatiotemporal availability. However, they do not consider the social reputability as a main criterion for participant selection. The social accountability of a participant in terms of popularity, good history of helpfulness and

reciprocity, expertise and experience, etc. is of great importance in the suitability of participant. Thus, these methods cannot be readily utilised for participant selection in social participatory sensing.

Incentives as Motivations. As mentioned before, the success of participatory sensing strongly relies on user participation to provide a sufficient and continuous flow of contributions. On the other hand, typical participatory sensing applications are based on volunteer participation and participants normally do not earn *explicit* incentives in return for their contributions. While contributing to a task, a participant consumes his own private resource such as battery and computation power of his device. Also, the participant may expose himself to potential location privacy threats. Having these issues in mind, the participant may be unwilling to contribute to the sensing campaign without obtaining explicit benefits. A series of related works [92, 93, 94, 95, 96, 97] aim at addressing this issue by proposing incentive mechanisms as a driving force for user participation. In the following, we discuss these methods in detail.

A pilot study was conducted at UCLA [92] which investigated the effect of micro-payments, i.e. transactions in which small tasks are matched with small payments, as an incentive mechanism. In their study, 55 participants were recruited to capture photos and 5 different micro-payments were used: a lump sum payment (MACRO), medium micro-payment (MEDIUM), high micro-payment (HIGH), low micro-payment (LOW), and competition-based (COMPETE). The study found that monetary incentive is beneficial if combined with altruism and competitiveness. In particular, their results show that HIGH and MEDIUM were the most successful schemes, MACRO and LOW yielded poor results in terms of the number of photos submitted. Moreover, dynamic incentivising schemes such as competition might be better suited for short bursty data collections since several of the participants in the COMPETE group dropped out due to fatigue.

Lee and Hoh [93, 94] proposed a dynamic pricing mechanism that aims to keep the service requester's cost low while retaining a sufficient number of participants. They proposed an auction-based mechanism which allows participants to sell their sensing data by participating in a dynamic price reverse auction system. Specifi-

cally, in Reverse Auction-based Dynamic Pricing (RADP) incentive mechanism, the service provider (i.e. the requester) selects a predefined number of lower bid-price users, and the selected users receive their bid prices for their sensing data as a reward. Hence, the selling price in this mechanism dynamically changes based on bid prices of users. RADP incentive mechanism provides several inherent benefits over the auction mechanism. In RADP, users decide their own prices for selling their sensing data. This will simplify the pricing decision on incentive cost for the service provider. Moreover, users play more active roles in incentive negotiation, as compared to the auction mechanism. Additionally, RADP can adapt to dynamically changed data collection environments (e.g., geographic imbalance of collecting user sensing data) because when the number of participants decreases, the price increases to recruit more participants. The sensing data collecting mechanism in participatory sensing applications can be regarded as a reverse auction in which there are many bidders (i.e. participants) who want to sell their sensing data and one auctioneer (i.e. a requester) who wants to purchase some number of sensing data. The traded goods are the sensing data (e.g., environmental data such as traffic speed, temperature, etc.) in a certain geographic region for a specific time period. In order to retain and encourage the participants who drop out of the reverse auction, the service provider gives a virtual participation credit to the participants who lost in the previous reverse auction as a reward for their participation only.

Luo et. al. [95] proposed two incentive schemes to maximise the fairness and the social welfare of information service crowdsourcing scenarios, where the user gets the service as a payment. They defined fairness as the existence of balance between a user's received service quota and both his contribution level and his demand. Social welfare is defined as the aggregate user utility with respect to the service provided by the system. In fact, their solution is based on the assumption that users of such applications are both providers and consumers. Users who contribute sensor readings to the applications obtain access to further services provided by other users. This scheme, however, requires mutual relationships between providers and consumers that may not be applicable in all application scenarios.

Authors in [96] proposed two platform-centric and user-centric incentive mech-

anisms for participatory sensing. In a platform-centric incentive mechanism, the platform (i.e. the requester) has the absolute control over the total payment to participants, and participants can only tailor their actions to cater for the platform. Whereas in a user-centric incentive mechanism which is in fact an auction-based incentive mechanism, each user announces a suggested price, the lowest price at which it is willing to sell a service. The platform then selects a subset of users and pays them an amount that is no lower than the user's suggested price.

In [97] authors propose a reverse auction incentive mechanism, which is an enhancement to the method presented in [93, 94] (explained above). In fact, they aimed at incentivising the participants by considering a set of parameters such as the location of the participants, the coverage and budget constraints. The intuition behind considering these parameters is that it is preferable to buy the contributions that are better distributed throughout the area of interest, since users may all be physically located very close to each other, rendering pretty much the same information. Also, it is more practical to assume that the auctioneer has a limited budget to run the system. With these considerations, they leverage the Greedy Budgeted Maximum Coverage (GBMC) algorithm [98] to create a greedy auction-based incentive algorithm that not only includes provisions to retain the users but also to obtain the lowest cost samples that are best distributed to cover the area of interest within a given budget.

In the context of social participatory sensing, the sense of community and efficacy, as two powerful motivations, can be easily satisfied by leveraging social friends or community members. Moreover, a participant's sense of recognition (as another important motivation) can be satisfied through the publication of his reputation (as an indicator of valuable contribution) in the public profile. Incentives in the form of e-coins [14] can also be devoted to well-behaved participants. It should also be noted that explicit incentive mechanisms such as those discussed above may coexist with the contributions of this thesis and applicable to social participatory sensing.

2.5 Recruitment in Social Participatory Sensing

The recruitment of participants requires exchange of communication messages between the requester and participants. Messages may contain sensitive information about the requester or participant (e.g., residential address, interest, political view) which should be protected against privacy and integrity attacks. Thus, it is rational to use secure communication medium to ensure the safety and security of exchanged messages. In the following, we study the related research on recruitment in online communities and participatory sensing.

2.5.1 Recruitment in Online Communities

Preserving the privacy of participants is a crucial issue in participant recruitment. The message context is forwarded through various intermediate nodes that may not be trusted by the destination or the source. Moreover, trust relationships may be loose and therefore, people may want to keep a tight control over the access to their information by other nodes for privacy reasons.

In online communities such as peer-to-peer systems, delay tolerant networks and social networks, different methods have been proposed to transfer communication messages via a secure communication medium. A popular mechanism is *onion routing* [99], where packets are routed through a group of collaborating nodes, thus making it difficult to determine the source of a communication. In fact, onion routing helps to conceal relevant routing information from potential adversaries by using its interesting idea. Each relaying node decrypts (peels) one layer of the packet (onion) and sends it forward, while being unable to read the content of inner layers of the onion, to be decrypted with the secret keys of successor nodes. An instantiation of onion routing-based approaches is TOR ¹⁷ (previously an acronym for The Onion Router) [100]. TOR is aimed to route Internet traffic from source to destination by leveraging more than 5000 relay nodes ¹⁸ to conceal users' activities and locations from traffic analysers.

¹⁷<https://www.torproject.org>

¹⁸<http://torstatus.blutmagie.de>

Mix networks [101] are routing protocols that create anonymous communications. Anonymity in a communication context, also known as untraceability, prevents tracking back from a destination to the source. This is achieved by using a chain of proxy servers known as mixes which receive messages from multiple senders, permute them, and forward them in random order to the next destination (possibly another mix node). This prevents the linkage between the source of the request and the destination from being discoverable, making it more difficult for eavesdroppers to trace end-to-end communications. Moreover, mixes only know the node that it immediately received the message from, and the immediate destination to send the shuffled messages to, making the network resistant to malicious mix nodes. Moreover, each message is encrypted using public key cryptography in order to prevent attacks against message integrity. A complete survey of mix networks has been presented in [102].

In the domain of social networks, recent studies have focused on social relations and analysed the social network properties of these networks to assist the design of efficient routing algorithms. In [103], the social similarity (to detect nodes that are part of the same community) and ego-centric betweenness (to identify bridging nodes between different communities) have been used as two metrics to increase the performance of routing. When two nodes encounter each other, they calculate the joint utility function comprised of these two metrics for each of their destinations. The message is given to the node having higher utility for the message's destination. In [104], each node is assumed to have a global ranking which denotes the popularity of the node in the entire society, and a local ranking which denotes its popularity within its own community. Messages are forwarded to nodes having higher global ranking until a node in the destination's community is found. Then, messages are forwarded to nodes having higher local ranking within destination's community.

Random walk [37] is a popular concept that is extensively used in various branches such as physics, economics, psychology and brain research [105]. Specifically in computer science, random walk has numerous applications in graph theory [106]. Since a typical social network can be viewed as a graph, the concept of random walk is applicable for crawling the social graph and recruiting participants. Given a social

graph and a starting node as the requester, a random walk starts from the requester and randomly selects one of the requester's friends as the next step. The current state is then changed to the selected friend, from which, the next step should be taken. The (random) sequence of nodes selected in this way is a random walk on the graph. The concept of random walk has a wide range of applications in graph theory. Link prediction is an example of such applications. Given a large network, say Facebook, at time t , for each user, a link prediction algorithm is aimed at predicting what new edges (friendships) that user will create between t and some future time t_1 [107]. Similarly, link recommendation algorithms aim to suggest to each user a list of people that the user is likely to create new connections to [108]. Random walk has also been used for community detection in online communities. Authors in [109] leverage the idea that short length random walks on a graph tend to get trapped into densely connected parts corresponding to communities. In [110], authors leverage the idea of random walk for crowdsourcing and routing tasks that require people to collaborate and synchronise both in time and physical space.

To summarise, most of the above discussed methods rely on encryption-based solutions to provide a secure communication channel for reliable participant recruitment. However, encryption is not widely used by most online social networks. Only 3 of the top 5 online social networking services currently use HTTPS [111]. Moreover, they only make use of this security measure to protect login credentials. The rest of the communication happens unencrypted and is visible to everyone along the communication path [111]. Specifically, creating ad-hoc secure links between random nodes in a graph may involve complex key management, key distribution, etc., which may not be scalable in social networks. So, the above mentioned methods are not compatible with current settings and assumptions of existing social networks.

2.5.2 Recruitment in Participatory Sensing

As mentioned before, the recruitment process requires the exchange of messages (e.g., tasks and contributions) amongst the requester and participants. Typical participatory sensing applications operate in a centralised fashion [1, 19]. Specif-

ically, participants download the task from the application server (or a dedicated task server as proposed in [16]), and the sensor data collected by participants are reported (using wireless data communications) to a central server for processing.

In decentralised participatory sensing applications such as the one presented in [112], mobile phones autonomously disseminate tasks to devices in their proximity, depending on the availability of the required sensors. If the participant's mobile phone is not equipped with the sensors required by the tasks, task is transferred to other mobile phones embedding such sensors in the vicinity. The receiving participant can fulfil the transferred task as soon as entering the task's region of interest.

An alternative hybrid model is presented in [113, 114], which is a combination of centralised and decentralised task distribution schemes. Users create the task including the sensing area and task duration, and broadcast it to other users in a distributed manner. Other users may then attend in the specified task area, collect the sensor data and upload their contributions to the central server.

In the above presented approaches, task distribution and contribution uploading processes may impose a privacy threat to the participant. The task server may infer information about the location of the participant at specific times while the participant is downloading the tasks. The task's requirements may also reveal information about the participant. For example, the requester may design a task in a way that it requires a very specific device to be performed (e.g., an iPhone equipped with a heart-rate sensor); revealing the health status of the participant contributing data from such device.

In order to protect the privacy of users' sensitive information during the tasking and reporting processes, a set of privacy-preserving routing schemes have been proposed [16, 115]. These schemes hide the participant's location using well-known anonymisation techniques such as TOR [100]. For example, [16] utilises the TOR-based routing method and mix networks to anonymise the connections to the task and report servers and hide the Internet Protocol (IP) address, and thus location information of participants [116]. Also in [117], a decentralised routing mechanism to preserve location privacy during the collection of contributions has been proposed.

The main idea is to exchange the collected sensor readings between users in physical proximity. By exchanging their sensor readings, users combine their paths; the prior path of one participant becomes that of another participant and vice versa. This will result in the formation of paths composed of concatenated sub-paths from multiple participants. As a result, the reported contributions do not disclose the actual paths, but instead a path jumbled with other users.

In the context of social participatory sensing, disseminating tasks to participants in proximity is not efficient since it requires all participants to be in each other's vicinity to receive the task. Moreover, as explained in Section 2.5.1, encryption-based anonymisation schemes such as [16, 115] (discussed above) are not applicable to the context of social participatory sensing, since online social networks currently do not support these techniques for message transfer. So, proposing an efficient and reliable recruitment mechanism for social participatory sensing is challenging.

2.6 Summary

In this chapter, we have given an overview of existing concepts and techniques in the area of social participatory sensing. We studied the building blocks of a social participatory sensing system and investigated their representation and applications. Specifically, we discussed the architecture of online social networks. Then, we presented a thorough study on the components and applications of participatory sensing systems.

We discussed the state-of-the-art related to the three key challenges that are relevant to the scope of this thesis, namely, the trust assessment, participant selection and recruitment. For each of these challenges, we studied the existing related research in online communities and participatory sensing. We also highlighted some of the limitations and drawbacks of existing related work when dealing with each issue. In later chapters, we present our proposed approaches to address the above-mentioned issues.

Chapter 3

Ensuring Sufficient Participation

3.1 Introduction

As mentioned in Chapter 1, the goal of this dissertation is to address the three main challenges in social participatory sensing, namely, participation sufficiency, suitability of participants and assessing the trust. To achieve this goal, we propose an overall framework (depicted in Figure 1.3) comprising three key components: (i) the recruitment scheme to address the participation sufficiency issue, (ii) the participant selection scheme to address the suitability issue and (iii) the trust scheme to address the trust assessment issue. In this chapter, we present the details of our recruitment scheme. The participant selection and trust assessment schemes will be described in the subsequent chapters.

The first challenge in the success of social participatory sensing is participant sufficiency. User participation is one of the most important elements in participatory sensing applications for providing an adequate level of service quality, since application services are dependent on sensing data from multiple users.

Typically, there are several costs incurred by the participant when contributing to a task. User participation requires sensing and transmitting the measurements to an application server. During the sensing process (such as gathering sound samples), a participant consumes the mobile phone battery and computation power for data collection. Transferring the collected data to the server also consumes the user's

bandwidth (which may or may not be free depending on whether the data is uploaded over Wi-Fi or cellular networks) and energy. Moreover, the participant may expose himself to potential privacy threats. Possible threats include the recording of intimate discussions, taking photographs of private scenes, or tracing a user's path and monitoring the locations the user has visited. Besides, most participatory tasks are based on voluntary participation without any explicitly monetary rewards. In the face of such issues, the participants may be reluctant to participate. Such unwillingness would diminish the impact and relevance of sensing campaigns and consequently limit their benefits to the users [9].

In this chapter, we present a novel recruitment scheme which is intended to address the participation sufficiency challenge in a comprehensive manner. Our proposed scheme aims to solicit an adequate pool of participants for social participatory sensing applications. It should be noted that the recruitment scheme proposed herein and the participant selection scheme which is introduced in Chapter 4 work in parallel to support the requester with a sufficient number of suitable participants.

In addressing the issue of participation sufficiency, several limitations emerge. The first limitation is the potential sparseness of a requester's friendship graph. The requester may have few friends or may lack close friends who may be willing to contribute to tasks initiated by the requester. Empirical analysis has demonstrated that it is not uncommon for friendship graphs to be sparse in OSNs [118]. For example in Yahoo! Pulse¹⁹ which is an online social network involving hundreds of millions of users, almost half of the users only have a single friend [119]. While the above may represent a worst-case scenario, it exemplifies the challenge of identifying and recruiting sufficient number of participants.

The second limitation is the selection of secure and reliable communication paths. The aim is to prevent the privacy/trust threats that may incur in transferring the communication messages and (as mentioned earlier) may affect the user participation. In social participatory sensing, social links connect members to each other. We define a communication path to be a multi-hop path consisting of a number of so-

¹⁹pulse.yahoo.com

cial links. The communication paths can be used for the recruitment of participants and exchange of communication messages (i.e., tasks and contributions). The communication messages may contain the task specifications and requirements which should be preserved from manipulation. They may contain sensitive information about the requester or participants. The information may be a sensitive attribute such as participant's address or his telephone number, or it may be a combination of quasi-identifying attributes which would readily allow a malicious intermediary to infer the corresponding sensitive information. For example, according to a famous study [120] of the 1990 census data, 87% of the US population can be uniquely identified by gender, ZIP code and full date of birth. Access to such private information can potentially result in the leakage of user privacy. So, determining the most suitable paths to the selected participants is of great importance.

In order to fully address the above mentioned issues and limitations, we propose a novel privacy-aware trust-based recruitment scheme for social participatory sensing. The proposed scheme leverages multi-hop friendship relations to support the requester with a large number of potential participants. It also provides credible paths to participants with the aim of preserving the privacy and integrity of the exchanged communication messages.

The recruitment scheme utilises the social network as the substrate for supporting access to social network friends and friends of friends as potential participants. Specifically, the basic participatory sensing procedures (i.e., task distribution and uploading contributions) are carried out by utilising the social network communication primitives.

An online social network is best represented as an undirected graph with the set of nodes representing participants and the set of friendship relations between nodes. Each participant has a profile containing his attributes and related information. Some attributes represent the participant's personal information such as name and address. Others include the outcome of participant's social behaviour. Examples are the participant's reputation score, the history of his previous transactions, the pairwise trust scores, etc. A participatory task or simply a *task* is represented by θ_i , and Θ is the set of all the tasks to be solved ($\Theta = \{\theta_i\}$). The owner of the task

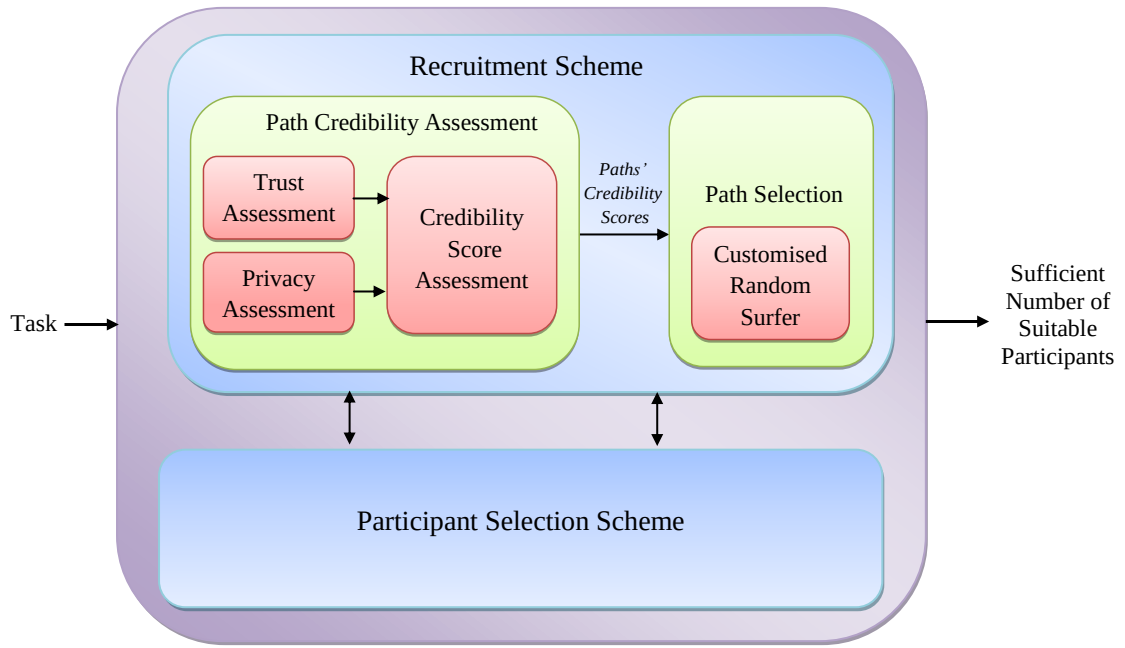


Figure 3.1: Recruitment scheme at a glance

is also called the *requester*. Ψ is the set of *participants* who contribute to the task ($\Psi = \{\psi_i\}$). They provide the requester with a set of contributions represented by κ .

Figure 3.1 illustrates the steps taken in the proposed recruitment scheme. A person wishing to start a participatory sensing campaign acts as a requester and defines the task, which includes the specification of task's main requirements such as needed expertise or location. Then, the recruitment scheme and the participant selection scheme (described in Chapter 4) are executed in parallel to select and recruit a sufficient number of suitable participants who will then contribute to the task. The recruitment scheme traverses the requester's social graph with the aim of supporting the requester with an adequate pool of potential participants. As mentioned above, constraining the graph crawling only to friends may lead to an insufficient number of participants, due to the potential sparseness of the friendship graph. Therefore, the proposed recruitment scheme extends crawling deeper in the social graph in order to maximise the possibility of finding potential participants. It also benefits from the suggested participants who have shown trustworthy behaviour in previous campaigns initiated by the requester (more details about the suggested

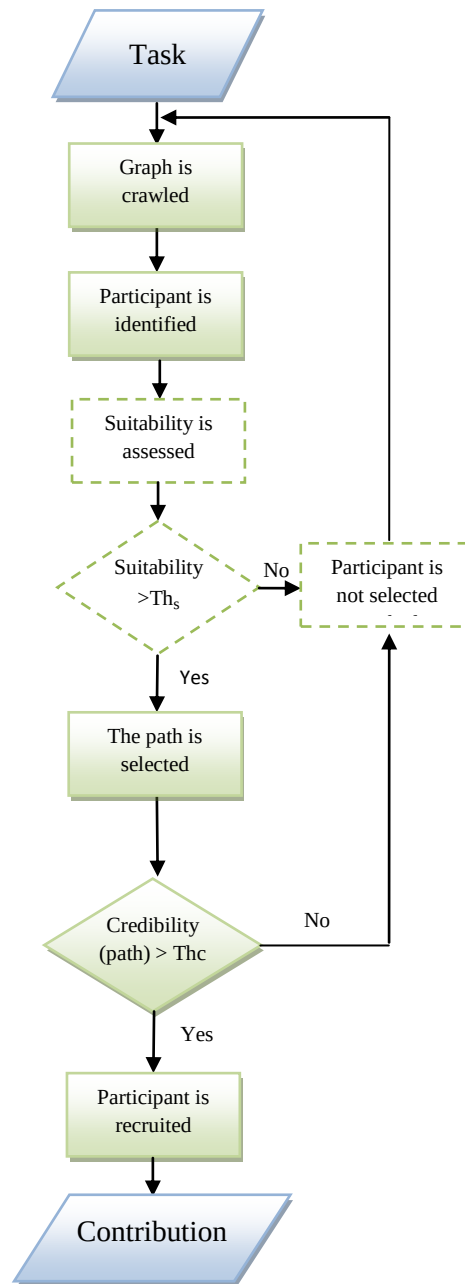


Figure 3.2: The sequence of steps in the recruitment of suitable participants

Steps shown in solid lines relate to the recruitment scheme; others suggest the participant selection scheme.

participants have been presented in Chapter 5). In particular, the recruitment scheme starts traversing the social graph from the requester and identifies potential participants (among friends and friends-of-friends). For each potential participant, the participant selection scheme is executed to assess his suitability (details described in Chapter 4). If the participant is identified as suitable, the path selection scheme is

then executed to select the most credible path for message exchange between the two parties. In particular, for each existing communication path between the requester and the participant, the credibility of the path is measured by making use of a novel credibility assessment scheme. This scheme constitutes the assessment of trust and privacy scores of the path, which are then combined to arrive at a credibility score for the communication path. The path with the maximum credibility score is then selected by the path selection scheme for further message exchange. The path selection scheme also enables the use of the customised random surfer method for fast and fair path selection. The participant is considered eligible for recruitment if there exists at least one credible communication path (with a minimum credibility score) between him and the requester. The participant is then recruited to contribute to the task. The sequence of steps and the exchanges and relations between the different schemes have been depicted in Figure 3.2. The steps displayed with dotted lines are carried out by the participant selection scheme. Other steps are performed by the recruitment scheme described herein.

In summary, the unique contributions of this chapter are as follows:

- We propose a privacy-aware trust-based recruitment scheme. The aim is to address the challenge of participation sufficiency in social participatory sensing by identifying the most credible communication paths to selected participants.
- We propose a credibility assessment scheme to obtain the level of credibility of the communication path to each participant. We consider a credible communication path as a path that is (i) trustable and (ii) privacy preserving. For each communication path, the proposed scheme measures the trust path's score based on the pairwise trust scores along the path. The path's privacy score is measured by utilising information theoretic formulations. A credibility score is then computed for each path by combining the trust and privacy scores.
- We propose a path selection scheme to select the best path to each suitable participant. This scheme considers the credibility score of all existing paths to the participant and selects the most credible path for further communications.

- We propose a novel customised random surfer to recruit the suitable participants in a fast and efficient manner. Our proposed random surfer crawls the requester’s social graph and selects the next to-be-visited-node based on the node’s suitability score and the trustworthiness of the path.

In Section 3.2, we present the details of our proposed credibility assessment scheme. We then discuss the path selection scheme and particularly our proposed random surfer in Section 3.3. Simulation set-up and results are discussed in Section 3.4. We conclude in Section 3.5.

3.2 Credibility Assessment Scheme

This section provides a detailed explanation of the credibility assessment scheme. In particular, we go through the procedures described in the previous section, i.e., assessing the trust and privacy of the communication path and calculating the path’s credibility score.

3.2.1 Trust Assessment

As mentioned in Section 3.1, the trustworthiness of the communication path is an effective factor in its credibility. In the following, we describe our proposed trust assessment method in detail.

In Section 3.1, we explained the structure of social participatory sensing. We assume that the edges of the social graph are labelled with weights equal to the trust score between the nodes. Consider the graph depicted in Figure 3.3, in which, *Req* is the requester and ψ_1 , ψ_2 , ψ_3 , and ψ_4 are participants. The weight of an edge from *Req* to ψ_1 is 0.7, showing the trust score of *Req* upon ψ_1 . If there are intermediate nodes in the path from requester to a participant, the trust score of the path is a combination of the trust scores of each of the pair nodes along the path. We leverage *multiplication* as the combination function since it has been shown in [121] to be an effectiveness strategy for trust propagation. For example, as shown in Figure 3.3, there are 5 paths to reach ψ_4 :

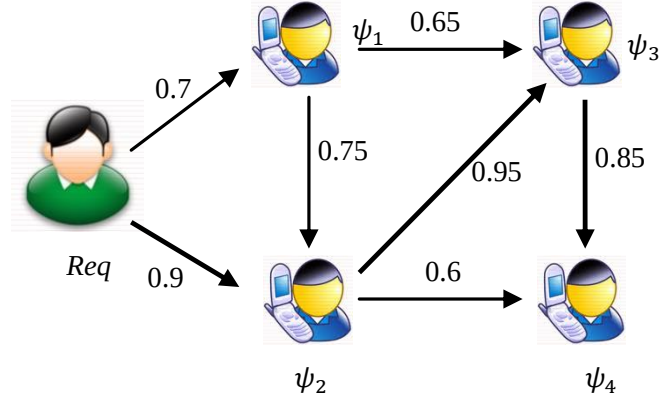


Figure 3.3: A simple example for calculating the path's trust score

$$r_1 = Req \rightarrow \psi_1 \rightarrow \psi_3 \rightarrow \psi_4 \quad \text{trust}(r_1) = 0.7 * 0.65 * 0.85 \simeq 0.39$$

$$r_2 = Req \rightarrow \psi_1 \rightarrow \psi_2 \rightarrow \psi_3 \rightarrow \psi_4 \quad \text{trust}(r_2) = 0.7 * 0.75 * 0.95 * 0.85 \simeq 0.42$$

$$r_3 = Req \rightarrow \psi_2 \rightarrow \psi_4 \quad \text{trust}(r_3) = 0.9 * 0.6 = 0.54$$

$$r_4 = Req \rightarrow \psi_2 \rightarrow \psi_3 \rightarrow \psi_4 \quad \text{trust}(r_4) = 0.9 * 0.95 * 0.85 \simeq 0.73$$

$$r_5 = Req \rightarrow \psi_1 \rightarrow \psi_2 \rightarrow \psi_4 \quad \text{trust}(r_5) = 0.7 * 0.75 * 0.6 \simeq 0.32$$

The most trustable path between Req and ψ_4 is r_4 .

In general, we assume the set R is the set of all possible paths between the requester and a specific participant. The path R_i ($R_i \in R$) has been defined with (N_{R_i}, E_{R_i}) in which, N_{R_i} is the set of nodes within this path and E_{R_i} is the set of edges of R_i . In that case, the trust score of R_i , denoted by $Trust(R_i)$ is calculated as:

$$Trust(R_i) = \prod_{k=1}^l w(e_k), e_k \in E_{R_i} \quad (3.1)$$

where l is the length of the R_i and $w(e_k)$ is the weight of the edge e_k .

3.2.2 Privacy Assessment

Now, we discuss the privacy assessment method in detail. A trivial solution to preserve privacy is encryption (e.g., HTTPS) which can be used to secure communication channels and protect against eavesdropping. However, as mentioned in

Chapter 2, this facility is not widely used by most online social networks. Only 3 of the top 5 online social networking services currently use HTTPS [111], and they only make use of this security measure to protect login credentials. The rest of the communication happens unencrypted and is visible to everyone along the communication path [111]. The primary reason for not using HTTPS for all communication is to minimise the hardware and connectivity costs. Moreover, public key cryptography needs additional computations and components for key management, which makes it computationally expensive for multi-hop social networks with extremely large numbers of nodes. In particular, setting up ad-hoc secure links between OSN members requires several complex tasks involved such as key creation, key distribution, etc. To sum up, encryption-based methods are most likely too complicated or expensive for general adoption.

In order to preserve the privacy of participant's information which is embedded in exchanged messages, it is desirable to consider potential privacy breaches in selecting the path between the requester and participant. In other words, when multiple paths exist, a reasonable approach is to select the most secure and trustable path in a way that the likelihood of privacy breaches in intermediate nodes is minimal.

Traditionally, two types of privacy breach have been studied: *identity disclosure* and *attribute disclosure*. Identity disclosure occurs when an adversary is able to map a profile in the social network to a specific real-world entity. Attribute disclosure, on the other hand, occurs when an adversary is able to determine the value of a sensitive user attribute, one that the user intended to stay private.

There are three sets of personal attributes [122]:

- Identifying attributes: attributes, such as social security number, which uniquely identify a person. To avoid identity disclosure, identifying attributes should be removed from profiles.
- Quasi-identifying attributes: a combination of attributes which can identify a person uniquely. As mentioned earlier, it has been observed that 87% of individuals in the U.S. can be uniquely identified based on their date of birth, gender and zip code [120].

- Sensitive attributes: those that users tend to keep hidden from the public, such as political view, location, and sexual orientation.

The messages exchanged between the requester and participant (including task or contribution) may contain private information such as sensitive attributes and quasi-identifiers, which may leak in intermediate nodes. To prevent such privacy leakage, it is reasonable to select the paths which contain intermediate nodes that are least likely to cause privacy breaches.

In order to quantify the privacy leak, we leverage the concept of entropy. Entropy is a measure of the uncertainty in a random variable [36]. We leverage this concept to quantify the privacy leakage of the communication messages. Our model aims to maximise the entropy which means the maximisation of the unpredictability of information for an adversary node. Higher entropy means better privacy for the information content inside a message. Since identifying attributes such as social security number are not normally kept in profiles, we assume that privacy breaches may happen if two types of information are leaked: sensitive attributes and quasi-identifiers. With this assumption, we aim at calculating the amount of uncertainty of a node about these two types of information inside a message.

For the intermediate node $m \in N_{R_i}$, we have the following definitions:

$M = \{M_1, M_2, M_3, \dots, M_n\}$ is the set of messages that have originated from a specific node and passed through the intermediate node m . Also $S = \{s_1, s_2, s_3, \dots, s_k\}$ is the set of sensitive attributes. Also assume $M(s_i)$ is the set of messages containing sensitive attribute s_i . Then, the probability of existence of sensitive attribute s_i in a message is $P(s_i) = n(M(s_i))/n(M)$, where $n(A)$ denotes the number of members in set A .

According to the definition of entropy [36], $H(P(S))$ is the amount of uncertainty about the existence of a sensitive attribute in a message:

$$H(P(S)) = - \sum_{i=1}^k P(s_i) \log_b(P(s_i)) \quad (3.2)$$

Similarly, assume that $Q = \{q_1, q_2, q_3, \dots, q_t\}$ is the set of quasi-identifiers, and

$M(q_i)$ is the set of messages containing quasi-identifier q_i . In that case, $P(q_i) = n(M(q_i))/n(M)$ is the probability of existence of quasi-identifier q_i in a message. Hence, the amount of uncertainty about the existence of quasi-identifiers in a message, denoted by $H(P(q)|Q)$ is:

$$H(P(q) | Q) = \sum_{i=1}^t P(q_i) H(P | Q = q_i) \quad (3.3)$$

So, the privacy breach (B) of a message in this intermediate node is:

$$B = |H(P(q) | Q) - H(P(S))| \quad (3.4)$$

We set the above definitions for a specific intermediate node m . In general, if n messages have been passed through intermediate node m , then the amount of privacy breach in node m , denoted by B_m , is calculated as:

$$B_m = \frac{1}{n} \sum_{j=1}^n B_{m,j} \quad (3.5)$$

in which, $B_{m,j}$ is the privacy breach of message M_j in intermediate node m . As this equation shows, B_m keeps a history of privacy breach for each message originator.

It is obvious that the privacy score of node m , ($Privacy(m)$), is inversely related to the privacy breach (B_m) in this node. However, in order to have the privacy score value of node m in the range of $[0,1]$, we divide the value of B_m by $\log(n)$. The reason is that $H(P(S))$ has a value less than $\log(n)$. So, the maximum value for $B_{m,n}$ and B_m will also be $\log(n)$, which results in $B_m/\log(n)$ in the range of $[0,1]$. To summarise, $Privacy(m)$ is calculated as:

$$Privacy(m) = 1 - \frac{B_m}{\log(n)} \quad (3.6)$$

in which, n is the number of messages that originate from a specific node and have passed through intermediate node m . So, for each path R_i consisting of a set

of nodes, the privacy of the path is:

$$Privacy(R_i) = \min(Privacy(m)) \quad \text{where } m \in N_{R_i} \quad (3.7)$$

3.2.3 Assessing the Path's Credibility Score

In order to assess the credibility of the communication path (R_i), we combine the privacy score of each path with its trust score via a combination function F to reach to a single value for the path's credibility. In other words,

$$Credibility(R_i) = \left(F(Trust(R_i), Privacy(R_i)) \right) \quad \text{where } R_i \in R \quad (3.8)$$

The selection of a proper combination function F is important. F should be efficient enough to handle possible conflicts between the trust score and the privacy score in a reasonable manner. The decision on how to combine these two independent factors affects the performance of the path selection scheme. In this scheme, two combination functions have been considered:

- **Geometric Mean**

The geometric mean is defined as the n_{th} root (where n is the count of numbers) of the product of the numbers. The geometric mean is often used for comparing different items and finding a single 'figure of merit' for these items, when each item has multiple properties. It is also appropriate for describing proportional growth. A geometric mean, unlike an arithmetic mean, tends to dampen the effect of very high or low values, which might bias the mean if a straight average (arithmetic mean) was calculated. This property makes geometric mean suitable for our situation since there may be conditions such that trust score and privacy score values are in opposite directions. The combination of a path's trust score and privacy score via the geometric mean is as follows:

$$Credibility(R_{best}) = \sqrt{Trust(R_i) * Privacy(R_i)} \quad (3.9)$$

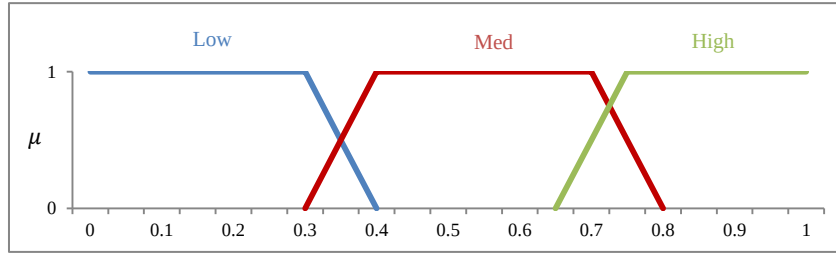


Figure 3.4: Membership function for the path's trust score (TS)

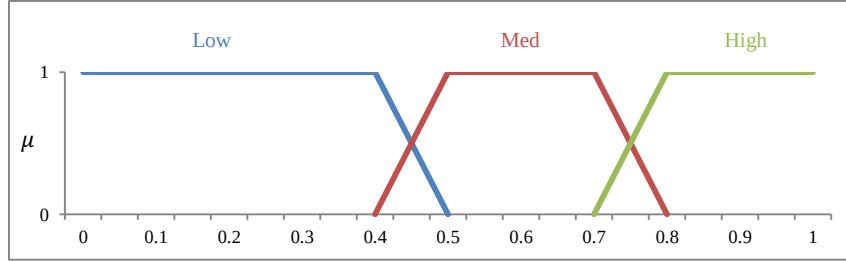


Figure 3.5: Membership function for the path's privacy score (PS)

- **Fuzzy Combination**

Another possible option is to employ fuzzy logic to calculate a comprehensive credibility score for the path. Consider a situation where the path's trust score is high but its privacy score is low. Leveraging fuzzy logic will help us make a meaningful balance between these two scores. The inputs to the fuzzy inference system are the crisp values of the trust score and privacy score of the communication path. In the following, we describe the fuzzy inference system components.

1. **Fuzzifier**

The fuzzifier converts the crisp values of input parameters into a linguistic variable according to membership functions. The fuzzy sets for the trust score (TS), privacy score (PS) and credibility score (CS) of the path are defined as:

$$T(TS) = T(PS) = \{\text{Low}, \text{Med}, \text{High}\},$$

$$T(CS) = \{L, M, H, VH\}.$$

For any set X , a membership function on X is any function from X to the real unit interval $[0,1]$. The membership function which represents a fuzzy set A is usually denoted by μ_A . The membership degree $\mu_A(x)$ quantifies the grade of membership of the element x to the fuzzy set A . Figure 3.4 and

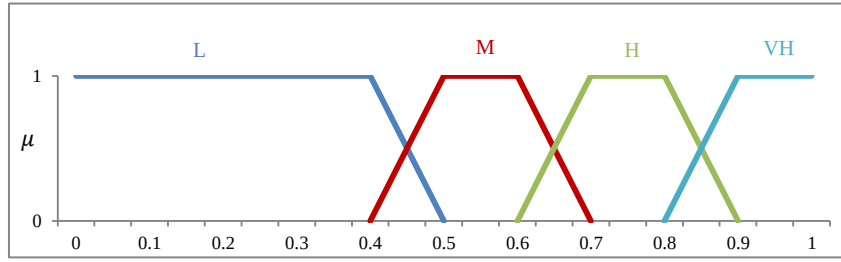


Figure 3.6: Membership function for the path's credibility score (CS)

Table 3.1: Fuzzy rule base for defining the path's credibility score according to its trust and privacy scores

Rule no.	trust score	privacy score	credibility score
1	Low	Low	L
2	Low	Med	L
3	Low	High	L
4	Med	Low	M
5	Med	Med	H
6	Med	High	H
7	High	Low	M
8	High	Med	VH
9	High	High	VH

Figure 3.5 represent the membership function of trust score and privacy score respectively. Figure 3.6 depicts the path's credibility membership function.

2. Inference Engine

The role of the inference engine is to convert fuzzy inputs to the fuzzy output (path's credibility score) by leveraging If-Then type fuzzy rules. The combination of the above mentioned fuzzy sets create $3 \times 3 = 9$ different states, addressed by 9 rules as shown in Table 3.1. The rule-base design is based on the experience of how the system should work [123] by leveraging the *max-min* composition method as follows.

$$\mu_{T(CS)}(CS) = \max \left[\min_{\substack{X \in T(TS), \\ Y \in T(PS)}} (\mu_X(TS), \mu_Y(PS)) \right] \quad (3.10)$$

The result is the credibility score which is a linguistic fuzzy value.

3. Defuzzifier

A defuzzifier converts the credibility fuzzy value to a crisp value in the range

of $[0, 1]$. We employed the Centre of Gravity (COG) defuzzification method, which is perhaps the most commonly used defuzzification technique.

In Section 3.4.2, the effect of each combination function on the performance of path selection scheme will also be investigated.

3.3 Path Selection Scheme

In the previous section, we discussed the operation of the credibility assessment scheme which measures the credibility of each communication path. When there exist multiple paths to a particular participant, only one communication path should be selected for message exchange. The path selection scheme is responsible for selecting the best communication path between the requester and the participant. In selecting the best path, the path selection scheme considers two different selection modes. In the following, we explain each selection mode in more detail.

3.3.1 Credibility-based Selection

In this selection mode, the aim is to find the *most credible* communication path to each potential participant. In order to do so, the path selection scheme considers *all* existing communication paths between the requester and the specified participant, measures the credibility of these paths one by one, and selects the most credible path as the best path for safe and reliable message exchange. Specifically, as stated in Equation 3.11, the path selection scheme identifies the best path to be the path with maximum credibility score.

$$Credibility(R_{best}) = \max \left(Credibility(R_i) \right) \text{ for all } R_i \in R \quad (3.11)$$

The most credible path is further used for any required communication with the selected participant.

3.3.2 Efficient Selection

Identifying all the existing paths between two members and selecting the most credible path is beneficial since it offers the optimal path for participant recruitment. However, for online social networks with large numbers of members and social links, this approach may lead to the potential increase in time and space complexity. Another challenge with this approach is the bootstrapping problem that occurs for newly-joined social network members. As explained in Chapter 5, when a participant provides a trustworthy contribution, the pairwise trust scores of all intermediate nodes along the path to this participant are increased, which results in the increase in the path's trust score. When a path is chosen for communication a relatively large number of times, its trust score is increased to a considerably high level. This, in turn, may increase its chance for being repeatedly selected as the best communication path in further recruitments (since as mentioned earlier, the path's trust score is effective in the selection). While this is valuable in terms of supporting highly trustable communications, it may result in a fixed set of communication paths (containing high trustable nodes) that are almost always selected for message exchange. So, low trustable members may have less chance to be leveraged in message exchange. This challenge is regarded as the bootstrapping problem. The low trustable friend, however, may be a new-comer. So, the new-comer friend should have the chance of being recruited or selected as the intermediate node. We claim that leveraging the random walk concept is able to address the bootstrapping issue, since it selects the intermediate nodes (almost) randomly.

In the following, we first have a short review on the concept of random walk and then, present the specifications and details of our proposed random surfer.

- **Random Walk**

Assume that we have a directed graph of nodes where some nodes have directed links to other nodes. One common approach to find the level of importance of each node in the set of all graph nodes is to use a random surfer [124]. The main idea of random surfing is as follows. One of the graph nodes is selected randomly as the starting node, from which, the surfer starts its journey. The

random surfer, then, picks one of the neighbouring nodes randomly and moves to that node. This process is repeated for a fixed period of time or till there is no further outgoing link. The level of importance of each node in the graph is proportional to the number of times it has been visited by the random surfer. In other words, the level of importance of node ψ_i is calculated using the following equation:

$$\text{Level of importance of node } \psi_i = \frac{\text{the number of times it has been visited}}{\text{total number of steps taken by the random surfer}}$$

The possibility of a particular node being visited by a surfer depends on the number of nodes that have outgoing links to this node. Recursively, for these neighbouring nodes, the possibility of being selected depends on the number of their incoming links. This implies that the importance of a node is greater if some other important nodes point to it. This is the main idea behind the PageRank algorithm for calculating ranks of web pages.

From the mathematical point of view, the random surfer concept is based on Markov chain theory. A Markov chain is a memoryless stochastic process in which, in each step, selecting the next state only depends on the current state of the process, and not on its history (i.e. those states visited earlier). A Markov chain is also called a ‘*random surfer*’ or a ‘*random walk*’.

The random surfer concept is widely used for graph processing such as node ranking and clustering (as discussed in Chapter 2). Assume that for a node ψ_i , the number of outgoing links is denoted by $|\psi_i|$. Then, the stochastic matrix Π representing the random surfer is defined as follows:

$$\Pi_{n \times n} = \begin{pmatrix} \pi_{11} & \pi_{12} & \cdots & \pi_{1n} \\ \pi_{21} & \pi_{22} & \cdots & \pi_{2n} \\ \vdots & \vdots & \ddots & \vdots \\ \pi_{n1} & \pi_{n2} & \cdots & \pi_{nn} \end{pmatrix}$$

in which, π_{ij} is the probability that the random surfer visits node ψ_j , assuming that it is visiting node ψ_i at the moment, and n is the number of nodes. In

the PageRank algorithm, these probabilities are the same in each row and calculated as follows:

$$\pi_{ij} = \frac{1}{|\psi_i|}, \text{ where } 1 \leq j \leq n \quad (3.12)$$

This means that in PageRank, the probability of moving from a node to each of its neighbours is the same. The matrix Π is the base for building the random surfer matrix in the PageRank algorithm. In this chapter, we propose a modified version of the random surfer to select the required participants for a task. The details of our proposed random surfer are discussed below.

- **Customised Random Surfer**

In order to provide sufficient number of participants in an efficient manner, the recruitment scheme leverages a customised random surfer. Our implementation of the random surfer is different from that employed in web page ranking and link recommendation systems such as [104, 107]. The typical random walk [124] is an iterative process wherein, in each iteration the next node to be visited is selected randomly and uniformly. The random walker may traverse each node multiple times based on the number of incoming links of a node. The importance of the node (i.e., the node's rank) is the number of times it has been visited by the random walker. It has been shown that in some instances convergence can take several iterations [124].

The proposed surfer is different since the probability of selecting a node as the next step is not the same for all available candidates. Some nodes have a higher chance than others to be visited by the random surfer. In other words, we propose a customised random surfer which does not act in a purely random manner, but is biased in a way that it considers the suitability score of the nodes for the selection (more details about the participant's suitability score are presented in Chapter 4). The intuition behind this strategy is to give better suited members a greater chance to be selected. Corresponding to a random surfer, we have a stochastic matrix Π which contains the probabilities of transitioning from one node to each of its direct neighbours. This matrix

is called the *transition probability matrix*. Each row π_i , called the *probability distribution row*, denotes the i^{th} row of the Π and contains the probability distribution row corresponding to the i^{th} member of the network.

Assume that in a social network, member ψ_i serves as the requester and intends to publish a task. Let $\varphi_i = \{\psi_j | \psi_j \text{ is friend with } \psi_i\}$ be the set of ψ_i 's friends. In order to find potential participants, we initiate K random surfers where $K = |\varphi_i|$. Each random surfer, denoted by ω_j , starts from the friend ψ_j and walks through the graph to find and nominate suitable participants. Assume that the current state of a random surfer ω_j is ψ_{cur} . The random surfer first checks the suitability of ψ_{cur} . If the suitability score is greater than a predefined threshold, ψ_{cur} will be invited to contribute. The surfer then continues its journey to find other participants from the list of ψ_{cur} 's friends and φ_i is updated accordingly. The next step will be selected from φ_{cur} based on the suitability scores. In other words, the probability of selecting ψ_{cur}^j (the j^{th} friend of ψ_{cur}) as the next step, denoted by $\pi_{cur,j}$, is:

$$\pi_{cur,j} = \frac{\sigma_j * \tau_{cur,j}}{\sum_{k:k \in \varphi_{cur}} \sigma_k * \tau_{cur,k}} \quad (3.13)$$

where σ_j is the ψ_{cur}^j 's suitability score and $\tau_{cur,j}$ is the pairwise trust of ψ_{cur} upon his j^{th} friend. It is evident that for each member ψ_i , the sum of probabilities of moving to his friends is equal to 1. In other words, $\sum_{j:1}^K \pi_{i,j} = 1$. Based on this, the stochastic matrix Π can be filled as $\Pi_{n \times n} = \{\pi_{i,j}\}$ in which, each element $\pi_{i,j}$ is the probability of selecting ψ_j as the next step for a random surfer that currently is in ψ_i . Π can be used by random surfers to determine the next steps.

The random surfer continues walking through the graph and identifies the participants. In order to control how far a random surfer can move from the requester, we define a parameter called the *propagation factor* and denote it by λ . The selection of an appropriate value for λ is challenging. A greater value of λ allows the random surfer to crawl deeper in the social graph, and thus increases the chance of finding more suitable workers. On the other

hand, it may increase the risk of privacy leakage as the exchanged messages may propagate several hops away from the requester's friendship network. In Chapter 4, we will go through extensive simulations for identifying the optimal value for λ .

In the following, we present the practical implementation of the process discussed above, in the form of an algorithm. Algorithm 1 presents our proposed customised random surfer. In this algorithm, W is a shared list which is accessible to all random surfers initiated for task θ_j , and includes the ID of all members who are identified by random surfers. Therefore, in each step, W contains the list of participants which have been identified so far. This list is used as a shared memory among random surfers to prevent them from selecting a member twice. The algorithm first initialises an empty list W . It also extracts the list of all ψ_i 's friends. Upon each friend ψ_i^j (j^{th} friend of ψ_i), a separate random surfer ω_j is initiated with the current state set to be ψ_i^j (lines 1 to 8 in the algorithm).

The lines 9 to 28 are the steps that each random surfer takes independently. Each random surfer, ω , checks to see if its current state is suitable to contribute to the task (details in Chapter 4). If the member is suitable to do the task, he will be identified. Then, the random surfer ω loads the row of Π corresponding to the current state of ω and then updates the transition probabilities. In order to do so, the pairwise trust between the current node and the participant is investigated. If less than a specific threshold, ψ_j 's suitability score will be set to zero. Otherwise, it will be updated with the value $\tau_{cur,j}$ stated in Equation 3.13. As can be seen, the probability of a particular member being visited by a surfer is in direct relationship with his suitability score.

The selection of each of the above explained methods depends on the settings of the application. In an organisational social network or a small social network (such as groups of acquaintances, co-workers, alumni, etc.) with less number of nodes and links, it is beneficial to use the credibility-based selection method. In other cases, utilising the proposed random walk will lead to an efficient and fast selection. In order to evaluate both methods, in the following experiments we make use of the

Algorithm 1 Customised Random Surfer

Input: Π as the transition probability matrix, ψ_i as the requester, NoP is the required number of participants, and θ_j as the advertised task

Output: W as the list of potential participants.

```

1:  $\varphi_i$  = list of  $\psi_i$ 's friends
2: Initialise  $W$  as an empty list.
3: for all  $f \in \varphi_i$  do
4:   Initiate a random surfer  $\omega_j$  from  $f$ 
5:   //current state of  $\omega_j$  is  $f$ 
6: end for
7:  $\Omega$  = set of all initiated random surfers
8: for all  $\omega \in \Omega$  do
9:    $L = \lambda$ 
10:  while true do
11:    //  $\psi_{cur}$  denotes the current state of the random surfer  $\omega$ 
12:    if  $\psi_{cur}$  is suitable for  $\theta_j$  then
13:      if  $|W| \leq 2 * NoP$  then
14:        Nominate  $\psi_{cur}$ 
15:        Add  $\psi_{cur}$  to  $W$ 
16:      end if
17:    else
18:      Stop random surfer  $\omega$ 
19:      Exit
20:    end if
21:    Load  $\pi = \Pi_{cur}$  //the row of  $\Pi$  corresponding to the current node
22:    Update  $\pi$  // see the algorithm description for details
23:    if  $\pi$  is empty then
24:      // there are no choices for next step
25:      Stop random surfer  $\omega$ 
26:      Exit
27:    end if
28:    Select a member of  $\pi$  as  $\psi_{cur}$  // see the algorithm description for details
29:     $L = L + 1$ 
30:    if  $L \geq \lambda$  then
31:      Stop random surfer  $\omega$ 
32:    end if
33:  end while
34: end for

```

credibility-based selection method. Since the random surfer is strictly dependent on the suitability of participants, we leave the evaluation of the efficient method (i.e., the customised random surfer) to Chapter 4.

3.4 Experimental Evaluation

This section presents simulation-based evaluation of the proposed recruitment system. The simulation set-up is outlined in Section 3.4.1 and the results are in Section 3.4.2.

3.4.1 Simulation Set-up

As mentioned in Chapter 1, we have chosen to evaluate the performance of our schemes based on simulations, as they allow a better exploration of the space of parameter values as compared to analytical modelling and measurement. Moreover, due to the lack of available social participatory sensing systems, it was not possible to evaluate the performance of our developed schemes in real-world applications. So, we developed a custom Java simulator for this purpose.

The dataset that we use for our experiment is the real web of trust of Advogato.org [125]. Advogato.org is a web-based community of open source software developers in which, site members rate each other in terms of their trustworthiness. Trust values are one of the three choices master, journeyer and apprentice, with master being the highest level in that order. The result of these ratings among members is a rich web of trust, which comprises of 14,019 users and 47,347 trust ratings. The distribution of trust values in the Advogato web of trust is as follows: master: 17,306, journeyer: 21,353, and apprentice: 8688. The instance of the Advogato web of trust referenced in this dissertation was retrieved on October 13, 2007. In order to conform the Advogato web of trust to our scheme, we map the textual ratings in the range of $[0, 1]$ as master = 0.8, journeyer = 0.6, and apprentice = 0.4. The Advogato web of trust may be viewed as a directed weighted graph, with users as the vertices and trust ratings as the directed weighted edges of the graph. So, it is in

perfect match with our assumptions related to participants and their trust relations in social participatory sensing.

Whenever a task is launched, one of the Advogato users is selected to be the requester. Then, the recruitment and the participant selection schemes (the latter is explained in Chapter 4) are executed to traverse the Advogato graph beginning from the requester until level L ($L = 3$) to find suitable participants via the most credible communication paths (details in Section 3.2). Tasks and contributions are then exchanged and the trust assessment scheme (explained in Chapter 5) is used to calculate the Trustworthiness of Contribution (ToC) for each receiving contribution. Pairwise trust ratings along the paths are then updated based on the ToC achieved. If above a threshold, all mutual trusts along the path from the requester to the participant are increased; otherwise, if less than a threshold, the pairwise trust between the participant and his immediate predecessor in the communication path is decreased. More details about the trust update process have been presented in Chapter 5; a short description has also been presented in Section 3.1.

The amount of private information contained in the exchanged messages (tasks and contributions) may vary. Some messages may contain more sensitive information than others. To simulate these differences, we assume that the total number of attributes contained in a message are 6 and a message may belong to one of three privacy classes: (i) *Class 0*: messages in this class contain 4 sensitive attributes; (ii) *Class 1*: messages in this class contain 2 sensitive attributes; (iii) *Class 2*: messages in this class contain 1 sensitive attribute. Greater number of sensitive attributes implies more private information. Whenever a message is created, a set of sensitive attributes, defined by numbers in the range of $[1,6]$, is randomly assigned to it. The credibility of the path is then computed via Equation 3.8. The path with highest credibility score is chosen for message exchange.

We run the simulation for 20 intervals, each consisting of launching 30 tasks. At the end of each interval, a list containing ‘suggested’ friends is provided with each requester, who are recruited in further tasks (details have been presented in Chapter 5). The list length is set to be 50.

In order to fully investigate the performance of our proposed recruitment scheme,

we consider two different set-ups for the scheme, as follows:

- **Trust-based Recruitment**

In this set-up, we assume that the path's credibility is calculated only based on the trustworthiness of the path. In this case, we compare the performance of our recruitment against one-hop recruitment, which broadcasts the tasks to all immediate friends. We also consider a multi-hop recruitment in which, no further friends are added. Specifically, we compare the following:

- (1) one-hop recruitment, which disseminates the task to all one-hop friends.
- (2) multi-hop recruitment without suggestion, in which, participants are selected through multi-hop friendship relations via most trustable paths. No further friendship establishment is done.
- (3) multi-hop with Friend Suggestion (FS): The same as (2) but with a suggestion scheme, which provides the requester with a set of suggested friends to be recruited via the most trustable paths in subsequent campaigns.

As we explained in Section 3.3, the path selection scheme searches for the most credible paths, which, in this scenario, is the path with greatest trust score. This path, however, may not be the path with shortest length. Accessing credible paths with shorter lengths is desirable since longer paths increase the task response time. Thus, in order to evaluate the performance of the recruitment scheme in terms of finding the best paths, we compute Mean Path Trust (MPT). MPT is defined as sum of trust ratings of all paths from the requester to the participants divided by sum of path lengths, where path length is denoted by the number of hops. Higher MPT demonstrates the better ability to find more reliable paths with shorter lengths. MPT is calculated after each campaign for all selected paths to all participants. MPT is a value in the range of $[0, 1]$.

As mentioned above, a ToC rating is calculated for each contribution by utilising the trust assessment scheme and those with ToC lower than a predefined threshold are revoked from further calculations. The ToCs for the non-revoked contributions are then combined to form an overall trust for that campaign.

In other words, $OverallTrust = \frac{\sum_{i=1}^n ToC}{n}$ in which, n is the number of non-revoked contributions. The revocation threshold is set to 0.5. We consider the overall trust as the evaluation metric. Greater overall trust demonstrates better ability to achieve highly trustable contributions and revoke untrusted ones. Overall trust has a value in the range of $[0, 1]$. The overall trust values obtained for all tasks will be averaged to arrive at a single value as the *average overall trust* for the entire simulation.

- **Privacy-aware Trust-based Recruitment**

In this set-up, we assume that the path's credibility is calculated based on both the trustworthiness of the path and its privacy score. So, in this set-up, the performance of our proposed recruitment scheme is compared against the one which only considers the trust (first set-up). To be more specific, we compare the following:

- (1) trust-based recruitment, in which, the path selection is based only on the trust score of the path.
- (2) privacy-trust recruitment, our proposed method in which, the best path is selected based on both privacy score and trust score.

In order to observe the performance of the scheme in the presence of noise, we artificially create situations in which, the privacy score of a specific node reduces for a period of time. This may happen in reality when a participant starts to reveal private information about another user. Our goal is to observe whether the system is able to rapidly detect such behavioural change and demonstrate a reasonable reaction accordingly. The duration of behavioural change has been set to be between the 5th and 10th intervals. We investigate the average credibility score of all paths passing through this malicious node. We expect to see that the proposed method is able to rapidly reduce the credibility score of these specific paths and thus, eliminate them from being selected for message exchange.

As explained in Section 3.1, the path selection scheme determines a set of potential participants (among all eligible participants) as those for whom, the

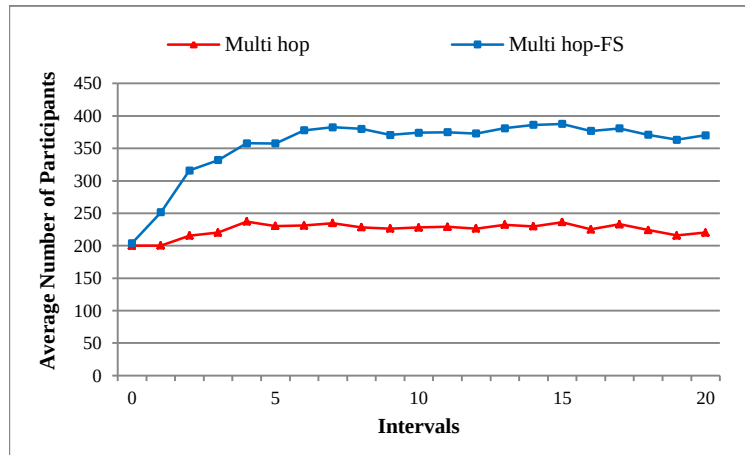


Figure 3.7: Evolution of average number of participants in different intervals

credibility score of the path is greater than a predefined credibility threshold (set to 0.6). Relevant to this selection process, we define the notion of *participation score* as the ratio of selected participants to eligible participants. A higher value of participation score implies better ability to recruit participants via optimum paths. Participation score has a value in the range of $[0, 1]$.

Since the proposed random surfer is strictly connected to the suitability of participants, and since the assessment of participants' suitability is thoroughly described in Chapter 4, we leave the experimental evaluation of the proposed random surfer to the next chapter.

3.4.2 Simulation Results

In this section, at first, we discuss the results obtained in the first set-up (trust-based recruitment), and then we will go through the outcome of the second set-up (privacy-aware trust-based recruitment).

- **Trust-based Recruitment Results**

Figure 3.7 demonstrates the average number of participants who have been identified in each interval for both multi-hop and multi-hop with FS recruitment methods. As the one-hop recruitment method has few participants in comparison with multi-hop, we omit its related data from this figure to better

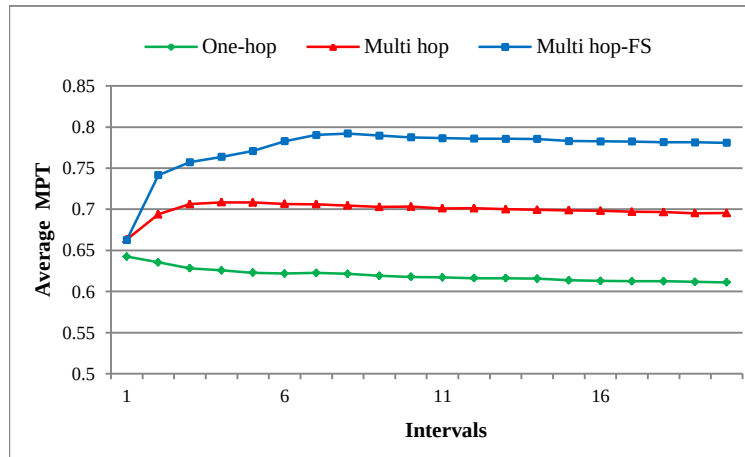


Figure 3.8: Evolution of average mean path trust (MPT) in different intervals

show the difference between two multi-hop recruitment methods. Note that interval 0 is in fact the first interval in which, no new friend has been identified yet and hence, both multi-hop methods recruit equal number of participants.

As this figure shows, multi-hop with FS trust based recruitment scheme is able to recruit participants with an average number twice as that of the multi-hop recruitment method, which demonstrates a better performance in terms of finding and recruiting more participants. Note that participant selection process in both multi-hop and multi-hop with FS recruitment methods is limited to L levels. But since our scheme adds a set of suitable participants as immediate friends, an easier access to friends of these newly added friends (who may be located beyond the L levels) becomes available, which results in an increase in potential participants. This demonstrates the effectiveness of our suggestion method in addressing the challenge of recruiting adequate participants.

Figure 3.8 depicts the evolution of average MPT for 20 intervals. As this figure shows, the average MPT of multi-hop with FS recruitment scheme is higher than the traditional multi-hop method. This means that our scheme demonstrates better performance in finding trustable paths which, at the same time, have shorter lengths. One may argue that average MPT for one-hop recruitment method should be higher given that the path length in such a method is less than path lengths in multi-hop recruitment methods. We would

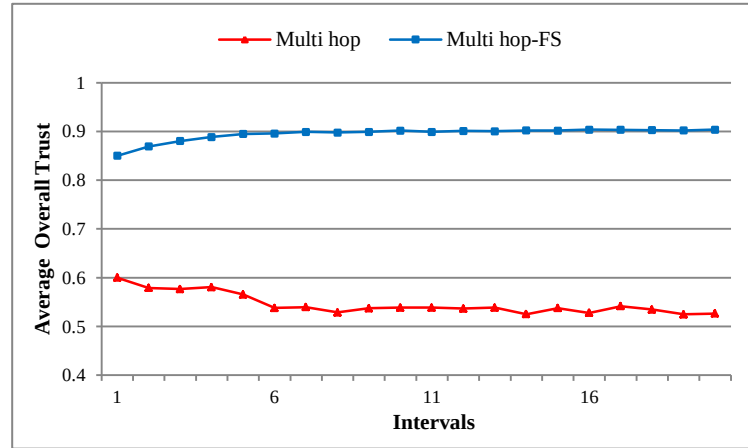


Figure 3.9: Evolution of average overall trust in different intervals

say that although path lengths in the one-hop method are shorter, the sum of trust ratings for selected paths in the multi-hop method is greater than that of the one-hop method. The reason is that, as described in Section 3.4.1, *all* the intermediate nodes along the path from the requester to participant are encountered with an increase in their pairwise trust ratings (in the case of participant's satisfiable behaviour). This means that a larger number of edges in multi-hop recruitment methods encounter such an increase, as compared to the one-hop method, which leads to higher MPT for multi-hop recruitment methods.

Finally, Figure 3.9 demonstrates the evolution of average overall trust obtained from contributions with multi-hop recruitment methods. As this figure depicts, the multi-hop with FS recruitment scheme obtains higher average overall trust in comparison with other method. This is because this method is able to identify more trustable paths to participants, which results in higher ToCs and hence, higher overall trusts. This clearly shows the success of the scheme in recruiting participants who produce trustworthy contributions.

To summarise, simulation results demonstrate a better performance for multi-hop with FS recruitment in terms of achieving higher overall trust (0.4 greater than multi-hop method) with higher MPT.

• Privacy-aware Trust-based Recruitment Results

In this section, we present the results for the second set-up in which, the

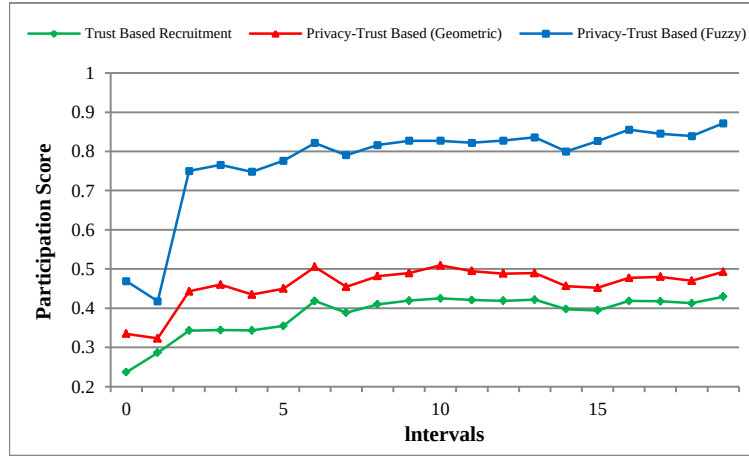


Figure 3.10: Evolution of participation score in different intervals

credibility score of the communication path is calculated considering both the trust and privacy scores of the path.

Figure 3.10 demonstrates the evolution of participation score for trust based and privacy-trust based recruitment methods with both geometric and fuzzy combination (details about combination functions can be found in Section 3.2.3). As this figure shows, our proposed recruitment scheme achieves a higher score in comparison with the trust based method, thus implying better performance in terms of recruiting more participants. Note that participant selection process in both methods is limited to L levels ($L = 3$). Since our proposed method takes participant's privacy into account, it results in the selection of a diverse set of paths which in turn, allows identifying participants from a broader group. In other words, our proposed scheme achieves greater diversity than the scheme that purely relies on trust. The figure also shows that fuzzy combination is more successful in achieving higher participation score in comparison with geometric mean. This is due to the adjustment of fuzzy rules (such as rule no.6 in Table 3.1) in a way that when the trust and privacy scores of the path are greater than the credibility threshold, the credibility score will be higher than the threshold, which results in leveraging more participants, thus increasing the participation score.

Figure 3.11 illustrates the average privacy score of the paths selected for message exchange between a specific requester and a set of selected participants,

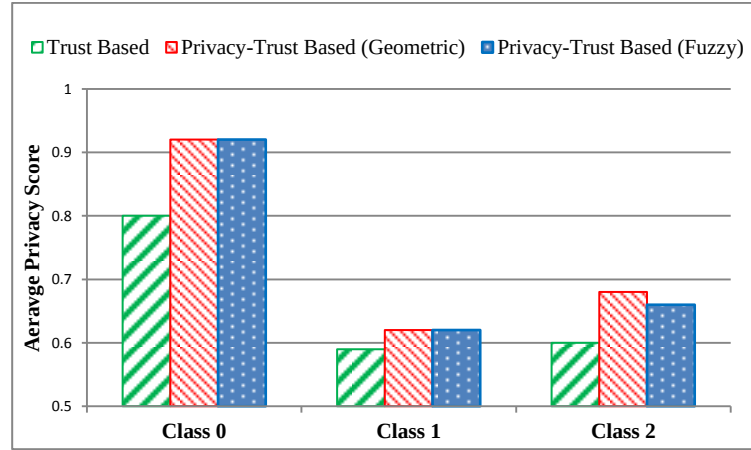


Figure 3.11: Average privacy score of the selected paths for different privacy classes

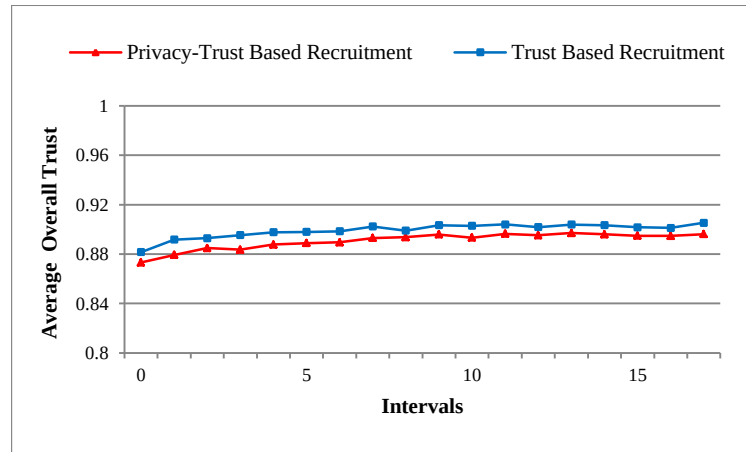


Figure 3.12: Evolution of average overall trust in different intervals

calculated separately for different privacy classes. As this chart demonstrates, our proposed method achieves a higher privacy score for all types of messages originating from a specific requester. For instance, the average privacy score obtained in our proposed method for class 0 messages is 12% higher than the one obtained in the trust-based method. This is because our scheme considers the privacy score of the path as an effective issue in evaluating the path's credibility. Note that this improved performance is consistent for all privacy classes (and more explicit for class 0), since the path selection method is identical for all types of messages. This implies that our proposed method is able to achieve higher privacy scores for all types of exchanged messages containing different levels of private information. The importance of this improved performance is better understood when obtained average overall trust is also

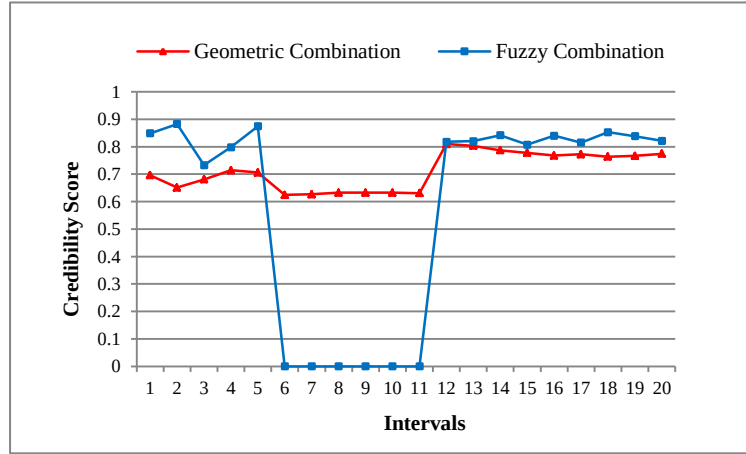


Figure 3.13: Evolution of average path's credibility score passing from a malicious node

considered in conjunction (details in Section 3.4.1). The average overall trust is 0.89 for trust-based recruitment method. Our scheme achieves a score of 0.89 with geometric mean and 0.84 with the fuzzy combination (as shown in Figure 3.12). This comparison shows that although our proposed method does not consider the trust score as the only determining factor for the path selection, the achieved overall trust is still similar to that of the trust based recruitment method, while better preserving the privacy of sensitive information.

As mentioned in Section 3.4.1, we also consider a scenario where the privacy score of an intermediate node decreases for a certain time interval (between 5th and 10th periods). The aim is to investigate the sensitivity of our recruitment method in promptly detecting and reacting to such a fluctuation. Figure 3.13 shows the evolution of average credibility score for all paths passing through this malicious node. Observe that the credibility score for both combination methods decreases in the transition period. However, the fuzzy method demonstrates better performance in early detection and severe punishment by a sudden decrease (to zero) in the credibility score. This is due to the adjustment of fuzzy rules such as rules no.1, 2 in Table 3.1. In other words, according to Equation 3.7, a low privacy score for a malicious node results in the low privacy score for all paths passing through it. We set the fuzzy rules in a way that when the privacy score of the path is low, the resulting credibility score will be L (Low). This will result in the exclusion of this path from

the set of candidate paths. This is not always true for the geometric mean method. There may be cases in the geometric mean combination (as observed in Figure 3.13) where the privacy score is low, but its credibility score is above the threshold, resulting in inclusion of this path for message exchange.

3.5 Conclusion

The three important challenges in the success of social participatory sensing are sufficient participation, suitability of participants, and the trust assessment. In this chapter, we addressed the first challenge by proposing a novel privacy-aware trust-based recruitment scheme for social participatory sensing. Our proposed recruitment scheme works in concert with the participant selection scheme (described in Chapter 4) to recruit participants who are selected (by participant selection scheme) as suitable to contribute. In particular, the proposed recruitment scheme crawls the social graph starting at the requester and leverages multi-hop friendship relations to identify the most credible paths to well-suited participants. The credible path is defined to be a trustable path to preserve the integrity of the exchanged messages. Moreover, the credible path is desired to be safe enough for the exchange of communication messages with embedded private information. The most credible paths are then used for recruiting the well-suited participants. The trust assessment scheme (proposed in Chapter 5) would then evaluate the trustworthiness of provided contributions.

Selecting the best communication path between the requester and the participant requires the evaluation of all existing paths and identifying the most credible one. This credibility-based path selection, while provides the best solution, requires extra time/space due to the need to check all potential paths between the pairs. It also suffers from the bootstrapping problem. So, in addition to the credibility-based path selection, we proposed an efficient path selection scheme by presenting a novel customised random surfer. The customised random surfer is aimed to support the requester with sufficient number of participants in an effective way and with less time/space requirements. Specifically, the proposed random surfer crawls the re-

requester's social graph beginning from the requester, selects the next to be selected participant among the requester's friends or friends of friends (almost) randomly. This would eliminate the need for examining all existing paths and at the same time, addresses the bootstrapping problem for the new-comers.

In order to evaluate the performance of our proposed scheme, we went through experimental simulation with various scenarios and settings. Results demonstrated that our scheme increases the overall trust as compared to other methods, and provides credible paths with shorter lengths for participant recruitment. Moreover, our scheme preserves better privacy for participants while achieving acceptable overall trust as compared to the trust-based method, and provides the system with more efficient recruitment of participants.

In the next chapter, we present a detailed description of our novel participant selection scheme.

Chapter 4

Selecting Well-Suited Participants

4.1 Introduction

We presented a framework in Chapter 1 to address the main challenges in social participatory sensing, which are participation sufficiency, suitability of participants and assessing the trust. The proposed framework includes three main components, each is aimed at addressing one challenge: (i) the recruitment scheme with the aim of addressing the participation sufficiency issue, (ii) the participant selection scheme to address the suitability issue and (iii) the trust scheme aiming at addressing the trust assessment issue. In Chapter 3, we presented our novel recruitment scheme. This scheme is aimed to address the sufficiency issue by traversing the multi-hop friendship relations and providing sufficient number of potential participants. These participants, however, may not all be well-suited to satisfy the requirements of the task. Selecting well-suited participants is important for acquiring high quality contributions since these participants have better knowledge and expertise relevant to the task's requirements. So, although an adequate number of participants may have been identified by executing the recruitment scheme, there still exists the challenge of selecting the well-suited participants among them. In this chapter, we address this challenge by presenting a novel participant selection scheme. Specifically, the participant selection scheme is aimed at addressing the suitability challenge by assessing the suitability of participants and selecting the well-suited ones. The collabora-

tion of these two schemes (which has been illustrated in Figure 4.2) results in the recruitment of a sufficient number of well-suited participants.

As described in Chapter 2, in the context of participatory sensing, the suitability of the participant for a campaign is typically related to the participant's geographical and temporal availability as well as the participant's reputation [10]. The geographical and temporal availability is extracted from collecting and analysing the timestamped location traces of the participants. The reputation of the participant is measured based on the quality of his contributions in the past.

In social participatory sensing, the existence of public profile information of participants (who are also members of online social networks), and the social links between them adds new dimensions to the evaluation of a participant's suitability. Through public profile data, additional information can be gleaned about participants such as interests, expertise and domain specific knowledge. This can provide interesting insights into the likelihood of the participant to contribute relevant and good quality data. For example, the participant who has academic knowledge on botany is more likely to provide high quality photos of rare plant species. The same is true for a participant who is local to the area of a task requesting for the traffic condition. This participant is more likely to contribute reliable data since this participant is more acquainted with that region and has a better understanding of its traffic conditions. The participant's social reputability can also be derived from his social relations and interactions. Reputability is in fact, an indication of trustworthy behaviour in the past. Selecting a reputable participant would invariably lead to high fidelity contributions and thus increase the likelihood of a successful campaign. To sum up, the valuable pieces of information that are normally available publicly in social networks can be effectively used as selection criteria for identifying well-suited participants, and thus overcome the challenge of suitability.

In addition to the above criteria, other time-varying parameters exist that are worth considering. Specifically, when participants are identified as suitable (according to the criteria mentioned above), they are all invited to participate (e.g., by receiving an email) and are able to accept/reject the invitation. For some participants, there may be a gap between the invitation time and the acceptance time

(e.g., due to a delay in checking email). During this time interval, the status of the task may change, e.g., other invited participants may have accepted the invitation sooner, and hence, fewer participants are now required. Moreover, as time goes by, the task’s deadline becomes closer, implying the need for more timely participation. So, it is rational to seek further information about these participants to see whether they are eligible to contribute. In general, with an eligibility check, we aim to improve our selection policy by selecting the rest of the required participants according to the current status of the task. Particularly, when the task’s deadline is imminent, it is rational to prioritise the recruitment of participants who are most likely to contribute in a timely manner, in order to ensure that with high probability, they contribute before the deadline. The prioritisation, however, depends on the number of participants selected so far. It is more stringent when we have selected most of the required participants (i.e. requires higher timeliness), otherwise it is more flexible (i.e. accepts participants with less timeliness score).

Moreover, some participants may create groups with the aim of colluding with other’s tasks and obtaining benefits. Specifically, they may collaborate to contribute malicious data to spoil the outcome of the tasks or sway them towards their goal. It is hence, essential to not select these participants.

In order to address the above mentioned challenges, we propose a novel time-aware collusion-free participant selection scheme. The proposed scheme aims to provide a comprehensive view of the participant’s suitability by identifying a set of influential parameters that can define the suitability of participants. It also aims at preventing the likelihood of collusion among selected participants. A participant is regarded as suitable if he is able to satisfy the task’s requirements. It is also preferable that the participant has shown trustworthy behaviour in previous campaigns. Moreover, the selection of the participant should not expose the requester to any privacy threat. The requester may also have a list of participants (created automatically by the application or manually by him) with whom the requester prefers to collaborate. Conversely, the requester may be reluctant to contribute to some other participants due to their poor behaviour in previous collaborations. It is beneficial to select the participant from the set of requester’s preferred list, and avoid selecting

those who reside in the requester’s blocked list. Based on these considerations, the proposed participant selection scheme considers the following parameters: (i) the participant’s expertise (in order to satisfy the requirements of the task), (ii) his reputation score (as a measure of being a trustworthy participant), (iii) the pair-wise privacy score between the requester and participant (to avoid the leakage of requester’s private information), (iv) the requester’s list of preferred participants (to select those who are preferred by the requester to be selected), and (v) the requester’s blocked list (those with whom, the requester is hesitant to contribute). Based on these factors, a suitability score is calculated and assigned to the participant.

In addition to the above factors, the participant selection scheme also considers a set of time-aware parameters. The main idea behind considering time-aware parameters is to improve the selection by including the current status of the campaign (in terms of number of selected participants, deadline, etc.) in the selection process. In particular, when the task deadline is close and a high percentage of the required participants have already been selected, it is reasonable to select the punctual participants since they are more likely to provide timely contributions (i.e. before the task deadline). With this idea in mind, the participant selection scheme takes into account the following three parameters for evaluating the eligibility of the participant: (i) the selection score (i.e. the ratio of participants selected so far to the total number of required participants), (ii) the remaining time to the task deadline, and (iii) the timeliness of the participant in previous tasks.

If a participant is eligible to be selected, a final check is carried out to ensure that the selection of this participant will not result in potential collusion. In particular, it should be identified whether the addition of this new participant to the previously selected group results in the formation of a group of colluders. This is performed by employing our novel collusion prevention scheme. The proposed scheme considers a set of indicators for detecting the potential formation of a colluding group such as size of the group and the similarity of the content contributed by group members. By considering these indicators, the collusion prevention scheme determines the possibility of collusion for each eligible participant.

Figure 4.1 depicts the sequence of steps in the selection of suitable participants.

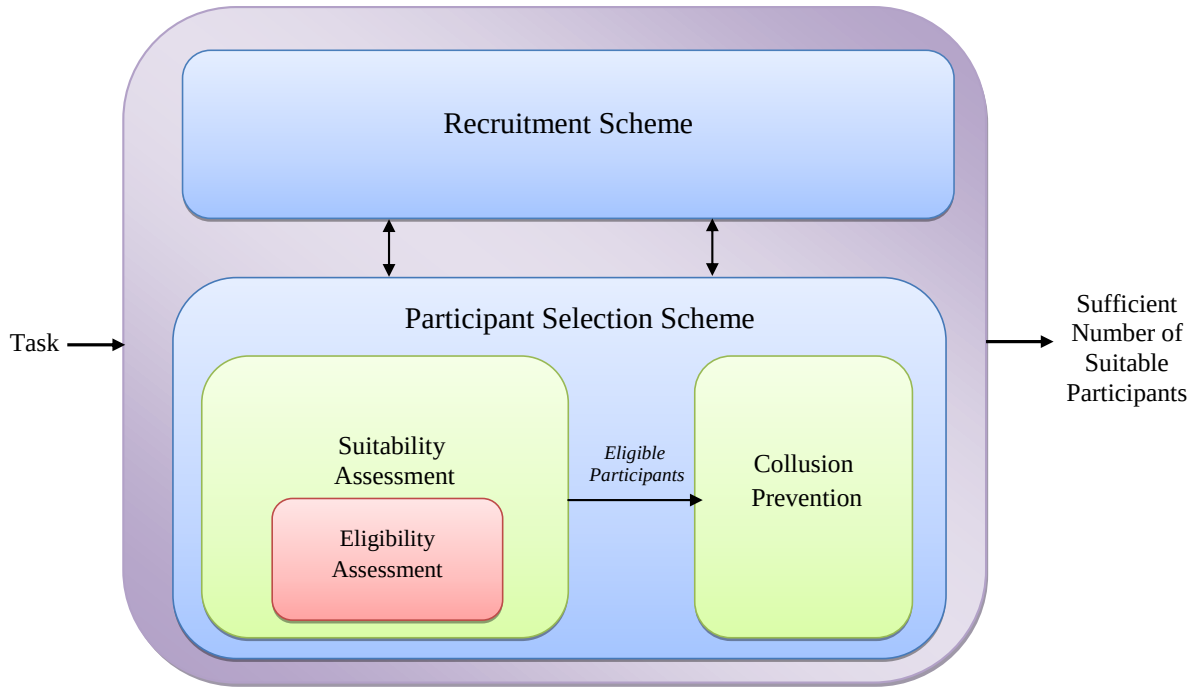


Figure 4.1: Participant selection scheme

The requester defines the task and specifies its requirements such as the needed expertise or reputation, the task’s deadline, number of required participants, etc. Then the participant selection scheme and the recruitment scheme (described in Chapter 3) co-operate to recruit a sufficient number of well-suited participants via the most credible paths. As mentioned in Chapter 3, the recruitment scheme traverses the social graph starting at the requester and identifies potential participants. At the same time, the participant selection scheme identifies whether this participant can be considered as suitable. If the participant is selected as well-suited, the recruitment scheme recruits the participant via the most credible path. In order to identify the participant suitability, the following steps are carried out by the participant selection scheme. At first, the suitability of the participant is assessed based on the parameters described earlier. If identified as suitable, the participant is invited to contribute. For each participant who accepts the invitation, the eligibility assessment scheme measures his eligibility based on the number of already selected participants and the task’s deadline, and assigns him an eligibility score. If the participant is considered as eligible, the collusion prevention scheme is executed

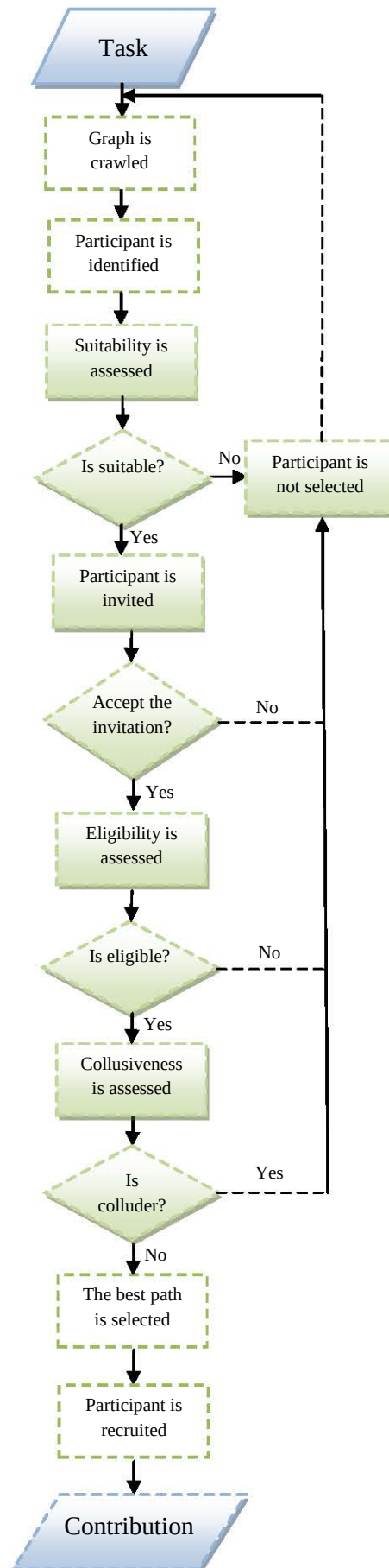


Figure 4.2: The sequence of steps in the recruitment of suitable participants
Steps shown in solid lines relate to the participant selection scheme; others suggest the recruitment scheme.

to identify the possibility of collusion for the participant. This is carried out by considering and measuring a set of collusion indicators. If the participant is not identified as collusive, he will be considered as a selected participant, who is then recruited to contribute. The above mentioned steps are illustrated in Figure 4.2. In this figure, the steps displayed with dotted lines are carried out by the recruitment scheme. Other steps are performed by the participant selection scheme described herein.

In summary, the main contributions of this chapter are as follows:

- We propose a suitability assessment technique to evaluate the suitability of a participant to contribute to a given task. In order to do so, the suitability assessment technique takes into account a set of parameters such as the participant's expertise, his reputation score and the pairwise privacy score between the requester and participant, and calculates a suitability score for the participant.
- We propose an eligibility assessment scheme to measure the eligibility of a suitable participant to satisfy the temporal requirements of the task. The eligibility score is measured for each participant who has accepted to contribute, based on the participant's timeliness, the remaining time to the task deadline and the number of participants selected so far.
- We propose a collusion prevention scheme to calculate a collusion possibility for each eligible participant and prevent any possible collusion.
- The accuracy and usability of the proposed techniques is tested using real world datasets from the Advogato social network and Wikipedia adminship election and simulated experiments. The evaluation results show superiority of our methods over the other common selection methods.

In Section 4.2, we present the details of our suitability assessment scheme. In Section 4.3, we discuss our novel collusion prevention scheme. Simulation set-up and results are discussed in Section 4.4. We conclude in Section 4.5.

4.2 Suitability Assessment Scheme

The suitability assessment scheme evaluates the suitability of participants as candidates for participation and identifies a set of eligible participants. In the following, we discuss these procedures in detail.

4.2.1 Assessing the Suitability Parameters

In order to evaluate the suitability of a participant, a set of parameters should be considered and evaluated. In the following, we first explain the definition and evaluation method of each parameter in detail and then discuss the calculation of suitability score.

- **Reputation**

The requester may specify a minimum level of reputation as a requirement for participation, in order to obtain high quality contributions. We assume that a reputation management system such as proposed in Chapter 5 is already in place, which calculates a reputation score ρ_i for each participant ψ_i . We also assume that the required reputation score of the task is denoted by ρ_{req} . In that case, the required level of reputation score for participant ψ_i to participate in task θ_j , denoted by $\Delta\rho_{ij}$ is as follows:

$$\Delta\rho_{ij} = \begin{cases} \rho_i & \text{if } \rho_i \geq \rho_{req} \\ 0 & \text{otherwise} \end{cases} \quad (4.1)$$

The higher the value of $\Delta\rho_{ij}$, the higher the eligibility of the participant to contribute to the task. We assume that $\Delta\rho_{ij}$ is a number in the range of $[0,1]$.

- **Expertise**

Expertise is defined as the measure of a participant's knowledge and is particularly important in tasks that require specific knowledge about a particular domain. In other words, participants may be asked to have specific exper-

tise such as programming skills, familiarity to a geographical area, proficiency with a particular language or so on. Greater credence is placed in contributions made by a participant who has expertise in the task. Expert finding systems such as [85, 84] may be employed for evaluating expertise. These systems employ social networks analysis and natural language processing (text mining, text classification, and semantic text similarity methods) to analyse explicit information such as public profile data and group memberships as well as implicit information such as textual posts to extract user interests and fields of expertise [84]. Expertise evaluation is done by incorporating text similarity analysis to find a match between the task keywords and participant's expertise. We denote the level of match between the i^{th} participant's expertise and the j^{th} task requirements by ΔE_{ij} . Assume that the E_j^t is the set of skills required by the task θ_j and E_i^ψ is the set of skills of the participant ψ_i , then:

$$\Delta E_{ij} = \frac{|E_j^t \cap E_i^\psi|}{|E_j^t|} \quad (4.2)$$

ΔE_{ij} is a number in the range of [0,1]. The higher the value of the ΔE_{ij} , the higher the match between the participant's profile and the task requirement.

• Privacy Requirements

As mentioned in Chapter 3, the communication messages exchanged between the requester and the participant may contain sensitive personal information. For example, if the requester asks for the gluten-free foods in a specific geographical area, it is probable that the requester has coeliac disease and lives in that place. So, it is desirable to maximise the privacy preservation of the requester in the selection process. Note that in Chapter 3, we took these privacy considerations into account with the purpose of selecting the communication path with the minimum likelihood of any privacy breach. Here, we focus on pairwise privacy issues with the aim of selecting the participants who do not expose the requester to these privacy threats.

We assume that the probability of a privacy breach in the one-hop neighbour-

hood of the requestor is zero. This is reasonable since friends are assumed to be trustworthy. For other non-friend nodes, the probability of privacy breach is greater in nodes who have been involved in a greater number of tasks initiated by that requester. The intuition behind this assumption is that the more ψ_i attends in tasks initiated by ψ_j , the more of ψ_j 's sensitive information will be revealed to him. So, the pairwise privacy score will decrease as the number of mutual tasks increases. With this intuition in mind, the pairwise privacy score of giving (the non-friend) ψ_i the permission to contribute to the task θ_j initiated by ψ_j , denoted by ΔPr_{ij} is calculated via the following function:

$$\Delta Pr_{ij} = \begin{cases} 1 - (\frac{t_{ij}}{T})^2 & \text{if } t_{ij} \leq T \\ 0 & \text{otherwise} \end{cases} \quad (4.3)$$

where, t_{ij} is the number of tasks ψ_i has done for ψ_j so far, and T is a system defined parameter which denotes the maximum number of the tasks initiated by ψ_j that ψ_i can participate in, following which, ψ_i 's privacy score reduces to zero.

- **Requester's Preferred and Blocked Lists**

The requester may have a list of preferred participants, whom the requester prefers to recruit in his future tasks. This list may be automatically generated by the application and would typically contain the requester's friends who have demonstrated trustworthy behaviour in tasks originated by the requester. The requester may also add some participants to the list manually, based on his trust upon them. It is clear that those who appear in this list should be assigned a higher suitability score. So, for the participant ψ_i who is being considered for selection for the task initiated by ψ_j , we define the parameter Pf_{ij} to be 1 if ψ_i belongs to the ψ_j 's preferred list, and zero otherwise. Note that belonging to a preferred list does not necessarily guarantee the selection of the participant for contributing in the task. Other suitability parameters

are also important and effective in identifying his suitability. In other words, there may be cases when a participant is a member of the requester's preferred list, but the participant is not considered as suitable to contribute to a specific task (e.g., due to not having the task's required expertise).

Similar to the requester's preferred list, a blocked list may also be available for the requester, which contains the list of those whom the requester desires to exclude from the list of contributors. This may be because of their poor behaviour in previous tasks, or due to privacy issues. It is obvious that the members belonging to this list should not be selected. So, for the participant ψ_i who is being considered for selection for the task initiated by ψ_j , we define the parameter B_{ij} to be 1 if ψ_i belongs to the ψ_j 's blocked list, and zero otherwise.

4.2.2 Computing the Suitability Score

Once the above parameters are evaluated, they should be combined to arrive at a single value for the participant's suitability score. To do so, the suitability score for a participant ψ_i to attend to task θ_j initiated by ψ_j , referred to as σ_i , is calculated as a weighted sum of parameters as:

$$\sigma_i = \begin{cases} 0 & \text{if } B_{ij} = 1 \\ w_1 * \Delta\rho_{ij} + w_2 * \Delta E_{ij} + w_3 * \Delta Pr_{ij} + w_4 * Pf_{ij} & \text{otherwise} \end{cases} \quad (4.4)$$

where w_i is the weight of each parameter, and $\sum_{i=1}^4 (w_i)$ equals to 1. The adjustment of the weights is application-dependant. For example, for privacy-aware applications, w_3 is set to be considerably high to give more importance to privacy parameter. Similarly, for tasks where expertise requirements are important, a higher weight may be associated with expertise (w_2). The suitability score is in the range of $[0,1]$.

4.2.3 Eligibility Assessment

Till now, we have a set of participants who are identified as suitable (by leveraging the suitability assessment scheme) and hence, invited to participate. The invited participants may accept the invitation in different times. During the time between the invitation of a specific participant and his acceptance, conditions may change. Some other participants may have accepted the invitation sooner, resulting in less number of required participants. The remaining time to the task deadline is also decreased which implies the need for recruiting those participants who provide timely contributions. So, in order to improve the participant selection policy to cover the variable situations, it is rational to carry out further investigation to evaluate the eligibility of the participant.

Specifically, for each participant who has accepted to contribute, the eligibility assessment scheme is executed to evaluate his eligibility according to a set of parameters. These parameters are: (i) the selection score (i.e. the ratio of participants selected so far to the total number of required participants), (ii) the remaining time to the task deadline and (iii) the timeliness of the participant in previous tasks. The first two parameters are combined via a geometric mean function to form a time suitability score. As explained in Chapter 3, geometric mean is useful for comparing different items and finding a single ‘figure of merit’ for these items, when each item has various multiple properties. So, for the participant ψ_i to be selected to attend in task θ_j , the time suitability will be as follows:

$$\text{Time suitability}(\psi_i) = \sqrt{\text{Timeliness}(\psi_i) \times \text{Remaining time}(\theta_j)} \quad (4.5)$$

The time suitability is then combined with the selection score via a fuzzy inference engine. The result is an eligibility score for the participant as follows.

$$\text{Eligibility Score}(\psi_i) = \text{Fuzzy}(\text{Time Suitability}(\psi_i), \text{Selection Score}(\theta_j)) \quad (4.6)$$

If greater than a predefined threshold, the participant will be considered to be eligible to participate.

Table 4.1: Fuzzy rule base for defining eligibility score (ES) according to time suitability (TS) and selection score (SS)

Rule no.	if TS	and SS	Then ES
1	Low	Low	M
2	Low	Med	L
3	Low	High	VL
4	Med	Low	H
5	Med	Med	M
6	Med	High	L
7	High	Low	VH
8	High	Med	H
9	High	High	M

Fuzzy inference system. Our proposed participant selection scheme employs fuzzy logic to calculate the eligibility score (ES) for each participant. The use of fuzzy logic allows us to achieve a meaningful balance between the time suitability and the selection score. We cover all possible combinations of time suitability (TS) and selection score (SS) and address them by leveraging fuzzy logic. Since the details of the leveraged fuzzy system is similar to the one employed in Chapter 3, we do not repeat the general definitions and concepts here and only explain the details specific to the eligibility assessment method.

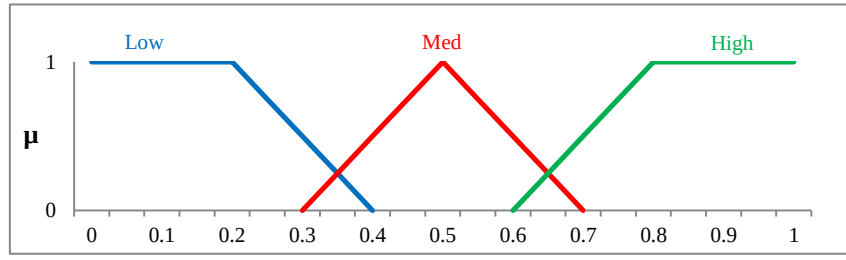
The inputs to the fuzzy inference system are the crisp values of TS and SS. The fuzzy sets for TS, SS and ES are defined as:

$$T(TS) = T(SS) = \{\text{Low, Med, High}\}$$

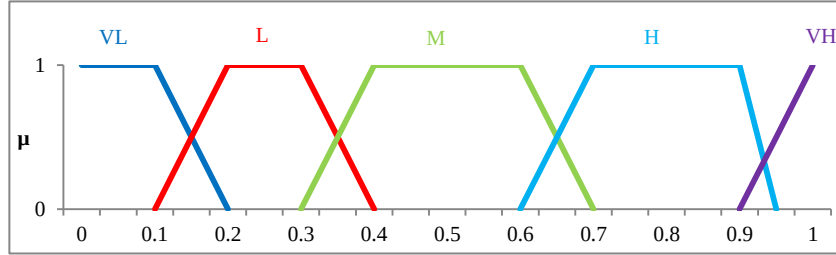
$$T(ES) = \{\text{VL, L, M, H, VH}\}.$$

Figure 4.3(a) represents the membership function of TS and SS and Figure 4.3(b) depicts the ES membership function.

The combination of the above mentioned fuzzy sets create $3 \times 3 = 9$ different states which have been addressed by 9 fuzzy rules as shown in Table 4.1. Fuzzy rules help in describing how we balance the various eligibility aspects. The rule base design has been done manually, based on the experience and beliefs on how the system should work [126]. To define the output zone, we used the *max-min* composition method. The result is ES which is a linguistic fuzzy value. We then



(a) Membership function for time suitability (TS) and selection score (SS)



(b) Membership function for eligibility score (ES)

Figure 4.3: Membership functions of input and output linguistic variables

employ the Centre of Gravity (COG) [127] method to defuzzify the linguistic value of ES and compute a crisp value for the eligibility score.

Once the crisp value for the eligibility score is computed, it is compared to a predefined threshold (has been set to 0.5 in the simulation). If greater than the threshold, the participant is considered as eligible.

4.3 Collusion Prevention Scheme

Once the participant ψ_i is considered to be eligible for being selected, a final check is done to ensure that the selection of ψ_i will not result in potential collusion. In particular, we aim to identify whether the addition of ψ_i to the set of previously selected participants will result in the formation of a group of colluders.

Collaborative attacks which are also called collusion attacks are those in which, a number of individuals form a clique and collaborate on changing the results of a task [38]. For example, colluders may collaborate as they wish to produce poor quality contributions that severely impact the goal of the task.

We define a *group* g consisting of a set of participants Ψ^g and a set of tasks Θ^g .

In other words, $g = \{\Psi^g, \Theta^g\}$. All the participants in Ψ^g have contributed to all the tasks in Θ^g .

Identifying the collusive groups requires two steps. In the first step, all existing collaborative groups (that fit within the above definition) are identified. In the second step, the potential collusive groups are detected among the identified groups. The detection of collusive groups is carried out based on a set of indicators. In the following, we discuss these steps in detail.

4.3.1 Identifying Potentially Colluding Groups

In order to identify all collaborative groups among the selected participants, the collusion prevention scheme employs the Frequent Itemset Mining (FIM) technique [128]. FIM is a method for market basket analysis. It aims at finding regularities in the shopping behaviour of customers of supermarkets, mail-order companies, on-line shops etc. More specifically, FIM intends to find sets of products that are frequently bought together. There are multiple applications for the identified frequent item sets such as improving arrangement of products in shelves, on a catalogue's pages etc., supporting cross-selling (suggestion of other products), product bundling and fraud detection [129, 130, 131]. Identified patterns are typically expressed as association rules, e.g., if a customer buys bread and butter, then this customer will probably buy cheese too. The performance and accuracy of the FIM technique is discussed in [43]. FIM is one of the major group detection algorithms which has been extensively used for collusion detection in online rating systems [38]. Hence, in our collusion prevention scheme, we make use of the FIM algorithm to find potential collusive groups.

The description of the FIM is as follows [43]: Let $I = \{i_1, i_2, \dots, i_n\}$ be a set of items and D be a multiset of transactions, where each transaction T is a set of items such that $T \subseteq I$. For any $X \subseteq I$, we say that a transaction T contains X if $X \subseteq T$. The set X is called an itemset. The count of an itemset X is the number of transactions in D that contain X . The support of an itemset X is the proportion of transactions in D that contain X . An itemset X is called *frequent* if

its support is greater than or equal to some given percentage s , where s is called the minimum support. In our context, the set of items (I) is the set of all selected participants for the current task. The set of transactions (D) is the set of all tasks that a participant has been involved in the past. By mining frequent itemsets, we find groups of participants who have contributed to multiple tasks together.

4.3.2 Collusion Indicators

Most existing collusion detection techniques rely on ‘behavioural’ indicators to identify colluding groups [38, 132, 133]. These indicators reflect suspicious behaviour from a group of participants which indicates the possibility of collusion. Colluders usually form a group which is typically large to gain the majority and make a considerable impact [38]. Moreover, group members usually target a considerable number of tasks and collaborate together in contributing to these tasks. We also claim that the colluders prefer to connect with each other in the form of groups (such as social groups in OSNs) to facilitate their communications.

Group connivance is also represented by some ‘content-related’ indicators. Colluders normally report contributions with typically similar (duplicate or near duplicate) contents in order to ensure that the task outcome is different from the true consensus. Moreover, their contributions deviate from the other (genuine) participants in order to change the task’s outcome. In order to have a better view of content-based collusion indicators, we provide an illustrative example. Recall the PetrolWatch application [2] in which, participants are recruited to take photos of fuel price billboards. The photos are then aggregated in the server and the fuel prices are extracted. The cheapest fuel price for each area is then identified (for example by leveraging majority consensus). People are then able to query the server to access the cheapest fuel price in their area of interest. Consider a situation in which, a service station operator is aware that there is a contest between the nearby stations to have more costumers. The operator is aware that PetrolWatch uses majority consensus and comes up with plan to game the system with the aim of attracting more customers to his business. The service station operator asks several of his

Table 4.2: Fuel prices of three different service stations uploaded by eight participants

The abbreviation “inc.” is used to denote incorrect prices (e.g., due to not being able to successfully recognise the price in the image).

Participants	Station 1	Station 2	Station 3
1	123.0	119.0	<i>inc.</i>
2	123.0	119.0	<i>inc.</i>
3	123.0	119.0	121.3
4	123.0	119.0	<i>inc.</i>
5	123.0	119.0	121.3
6 (m)	123.0	125.0	124.5
7 (m)	123.0	125.0	124.5
8 (m)	123.0	125.0	124.5
correct ¢	123.0	119.0	121.3
majority consensus ¢	123.0	119.0	124.5

social friends to collusively report false data for the competing service stations by uploading old pictures of higher fuel prices. If the false prices reported by his friends are more than the correct prices reported by other people, the collusion attack will be successful.

Table 4.2 (taken from [15]) is an example of this scenario that represents the fuel prices reported by 8 participants. We assume that the malicious operator owns service station 1 and that participants ψ_6 , ψ_7 , and ψ_8 (denoted by ‘m’ in the table) are his workers who have formed a collusive group and report higher prices for service stations 2 and 3. As shown in Table 4.2, all the colluding members report the same data (¢125.0 for station 2 and ¢124.5 for station 3) to set a higher price for these stations. Also, the price reported by the group members deviates from the prices forwarded by other genuine participants 1-5 in order to change the outcome of majority consensus. The result obtained from the majority consensus (the last row of Table 4.2) shows that the colluding group are successful in falsifying the genuine price of service station 3 since they constitute the majority and hence, their reported price is selected as the true price. This example illustrates the need to examine certain features that suggest the likelihood of the existence of a colluding group.

Similar to the concepts discussed above, in our collusion prevention scheme,

we consider a set of indicators. These indicators suggest that a colluding group is likely to exist among the selected participants. Note that these indicators reflect the likelihood of collusion only when they all occur together. In the following, we explain each indicator in detail and discuss how they identify possible collusive activities.

- **Group Size (GS).** The first indicator is the group size which is proportional to the *number of colluders* who have collaborated as a group in similar tasks. Group size (normalised) for a group g (GS^g) is calculated as follows.

$$GS^g = \frac{|\Psi^g|}{\max(|\Psi^g|)} \quad (4.7)$$

where $\max(|\Psi^g|)$ is the largest group size of all found groups. GS is a parameter in the range of $(0, 1]$, i.e. $0 < GS \leq 1$, showing how large the group is in comparison with other groups.

- **Group Target Size (GTS).** While the group size measures the number of group members, group target size measures the *number of tasks* in which the group members have targeted to collaborate in the past. Groups with a high value of target size are more likely to be colluding as the probability of a group of random people to have attended the same tasks together is rather small. For a group g , GTS^g is calculated as follows.

$$GTS^g = \frac{|\Theta^g|}{\max(|\Theta^g|)} \quad (4.8)$$

where $\max(|\Theta^g|)$ is the largest target size of all found groups. GTS is a number in range $(0, 1]$, i.e. $0 < GTS \leq 1$.

- **Group Deviation (GD).** The third indicator is group deviation which is an indicator to show the difference between the contents contributed by the colluders and those reported by other (honest) participants. In order to calculate the group deviation, we first calculate the deviation of the contents produced by group members from those of other participants for a single task $t \in \Theta^g$.

For each task $t \in \Theta^g$, the deviation of the group (GD_t^g) is calculated as follows:

$$GD_t^g = \left| \overline{\kappa_{i,t}}_{i \in \Psi^g} - \overline{\kappa_{j,t}}_{j \notin \Psi^g} \right|, \text{ for all } i, j \in \Psi^g \quad (4.9)$$

where $\overline{\kappa_{i,t}}$ and $\overline{\kappa_{j,t}}$ are the average of contents for task t given by members of group g and by other participants not in g , respectively.

Now, for a group g , the group deviation, denoted by GD^g , is the maximum of all group deviations for all tasks in Θ^g . In other words, GD^g is computed as:

$$GD^g = \max_{t \in \Theta^g} (GD_t^g) \quad (4.10)$$

GD is a number in range $(0, 1]$, i.e. $0 < GD \leq 1$.

- **Group Connectivity (GC).** The fourth indicator which is specifically suited for social communities is the group connectivity degree which is an indicator to show to what extent the colluders are connected to each other. For a group g , we first calculate the number of links between group members and denote it by link count (LC^g). LC^g is calculated as:

$$LC^g = \sum_{i \in \Psi^g} \sum_{j \in \Psi^g} T_{i,j} \text{ for all } i, j \in \Psi^g \quad (4.11)$$

where,

$$T_{i,j} = \begin{cases} 1 & \text{if } i \rightarrow j \text{ (there is a link from } i \text{ to } j) \\ 0 & \text{otherwise} \end{cases}$$

GC^g is then computed as follows.

$$GC^g = \frac{LC^g}{\max(LC^g)} \quad (4.12)$$

GC is a number in range $(0, 1]$, i.e. $0 < GC \leq 1$.

- **Group Content Similarity (GCS).** The fifth indicator is group content similarity which indicates the degree of similarity of contents produced by

group members. In order to evaluate this similarity, we first calculate the pairwise content similarity between every pair of members in the group. Pairwise content similarity between ψ_i and ψ_j , denoted by $GCS_{i,j}^g$, shows to what extent ψ_i and ψ_j have reported similar contents. In order to calculate the pairwise similarities between group members, we use the cosine similarity model, a well-known model for similarity detection [134]. Specifically, $GCS_{i,j}^g$ will be the cosine of the angle between two vectors containing the contribution contents of ψ_i and ψ_j and is a value in the range $(0, 1)$. The value 1 for $GCS_{i,j}^g$ means completely the same while 0 means completely different. $GCS_{i,j}^g$ is calculated as follows.

$$GCS_{i,j}^g = \frac{\sum_{t \in \Theta^g} \kappa_{i,t} \times \kappa_{j,t}}{\sqrt{\sum_{t \in \Theta^g} (\kappa_{i,t})^2} \times \sqrt{\sum_{t \in \Theta^g} (\kappa_{j,t})^2}} \quad (4.13)$$

We then calculate an overall degree of similarity for the group to show how all members are similar in terms of contents they have contributed. Group content similarity for every group g , denoted by GCS^g is the minimum amount of pairwise similarities between group members. In other words,

$$GCS^g = \min_{i,j \in \Psi^g} (GCS_{i,j}^g) \quad (4.14)$$

GCS is a number in range $(0, 1]$, i.e. $0 < GCS \leq 1$.

4.3.3 Possibility of Collusion

It is often difficult to determine with certainty whether a group is collusive [38]. Therefore, we define a metric called *Possibility of Collusion (PoC)* to show to what extent a group is potentially collusive. *PoC* is an aggregation of five collusion indicators. Since the importance of these indicators may be different in various applications, the collusion prevention scheme enables the applications to assign weight to each indicator based on its importance.

Suppose that W_{GS} , W_{GTS} , W_{GD} , W_{GC} and W_{GCS} are corresponding weights for indicators GS^g , GTS^g , GD^g , GC^g and GCS^g . The weights are initialised in a way

that: $W_{GS} + W_{GTS} + W_{GD} + W_{GC} + W_{GCS} = 1$. PoC is then calculated as:

$$PoC(g) = GS^g \times W_{GS} + GTS^g \times W_{GTS} + GD^g \times W_{GD} + GC^g \times W_{GC} + GCS^g \times W_{GCS} \quad (4.15)$$

PoC is a number in range $(0, 1]$. For each eligible participant ψ_i to be selected, we calculate the possibility of collusion ($PoC(g)$). If greater than a certain threshold, it implies that the selection of ψ_i may lead to potential collusion, and hence, the participant will not be selected.

4.4 Experimental Evaluation

In this section, we conduct a simulation-based evaluation to analyse the behaviour of our proposed schemes. First, we explain the experimentation set-up, the metrics we use for performance evaluation and the datasets we used in experiments in Section 4.4.1. Then, we compare our proposed schemes with other methods in Section 4.4.2. Then, we analyse the behaviour of our proposed schemes in Section 4.4.3 in order to find an optimum configuration. Finally, in Section 4.4.4, we investigate the efficiency of our proposed collusion prevention method.

4.4.1 Simulation Set-up

Our simulations have been conducted on a PC running Windows 7.0 Professional and having 4GB of RAM. We used Matlab R2012 for developing the simulator.

- **Datasets**

The dataset that we use for our experiment is the real web of trust of Advogato.org [125]. As explained in Chapter 3, Advogato members rate each other in terms of their trustworthiness. Advogato web of trust can be regarded as a social participatory sensing system with users as the potential participants and trust ratings as the friendship relations.

We pre-processed the dataset in order to remove the isolated nodes that have

no connection. 174 nodes were identified as isolated and were removed. We also have enriched the dataset in order to adapt it to our simulation scenario. To do so, we have computed a reputation score for each member by calculating the average of all pairwise trust scores each member receives from his friends. The reputation score is a number in the range of $[0, 1]$. We also assign each member a set of expertise attributes in order to use them for measuring the member's suitability score. We assume that there exist 10 different expertise attributes in the system. Each expertise attribute is an integer number in the range of $[1, 10]$. The total number of expertise attributes for each member is chosen randomly according to the reputation of the member. In other words, in the enriched dataset, those with higher reputation scores are likely to have greater number of expertise attributes. As mentioned before, the suitability assessment scheme calculates the privacy score based on the number of tasks a participant has attended for a particular requester. Therefore, we randomly and uniformly selected numbers in the range of $[1, 100]$ as the number of tasks completed by each member for each requester.

In order to evaluate the performance of our proposed collusion prevention method, we set two experiments. In the first experiment, we aimed at utilising a real dataset for which, there exists the possibility of collusion due to gaining benefits. Hence, we utilised the Wikipedia voting dataset. In Wikipedia, the voting process is used to elect administrators.²⁰ Every registered user can nominate himself or another user as an administrator in Wikipedia and initiate an election. The other users participate in the election and cast their votes on the eligibility of nominee. If the majority of users recognise a user as eligible, this user then will become a Wikipedia administrator. In order to incorporate this dataset in the context of our scheme, we employ the following mapping. The requester is the nominee, the worker is the voter, the task is evaluating the eligibility of the nominee as an administrator in Wikipedia and the contribution is the worker's vote. We use the log of Wikipedia Adminship Election²¹ which was collected by Leskovec et al. for behaviour prediction

²⁰http://en.wikipedia.org/wiki/Wikipedia:Requests_for_adminship

²¹<http://snap.stanford.edu/data/wiki-Elec.html>

in online social networks [135], referred to as WIKILog. WIKILog contains about 2,800 elections (tasks) with around 100,000 total votes and about 7,000 users participating in the elections either as a voter or a nominee. We use the WIKILog to demonstrate the efficacy of our proposed method to detect collusion.

In the second experiment, in order to better investigate the performance of our method, we artificially created collusive groups among Advogato members. We then investigated whether the proposed collusion prevention scheme is able to identify these groups.

- **Evaluation Method and Metrics**

In order to evaluate the performance of our proposed participant selection scheme, we run the experiment for a set of rounds. A simple experimentation round contains the following steps. In the first step, we choose a requester out of the members of the Advogato community. This selection is performed uniformly, meaning that all members have the same chance to be chosen as the requester. Then, a task is generated to be advertised to the community. Each task contains a set of attributes, mainly, a minimum accepted reputation score, a set of at most 5 required expertise attributes, and the maximum number of required participants. Once the requester is chosen and the task is generated, the random surfer scheme (proposed in Chapter 3) is executed in order to find the potential participants. Then the participant selection scheme chooses a subset of participants as selected participants to contribute to the task. We assumed that at least 50% of participants accept the invitation and apply to the task for contribution.

In order to evaluate the effectiveness of proposed schemes, we define four evaluation metrics. The first metric is the *number of suitable participants*. The ability to identify more suitable participants is a desirable property of the selection scheme. The second evaluation metric is the *overall suitability score of the suitable participants*, which is the average of all participants' suitability scores. A larger value for this metric suggests that the selection scheme is able to recruit well-suited participants. We have two similar metrics to evaluate

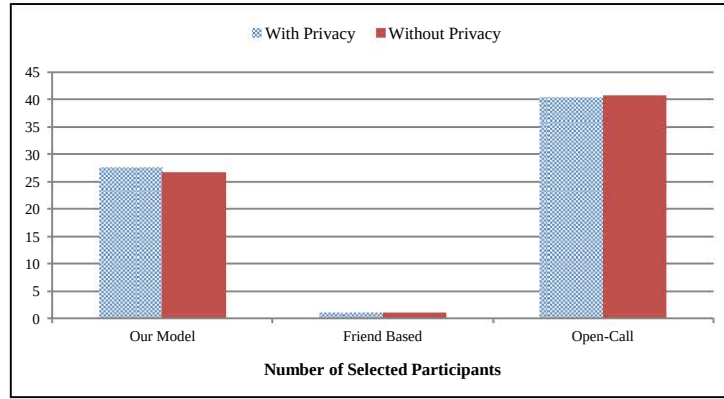
the performance of the eligibility assessment scheme: the *number of eligible participants* and the *overall eligibility score of eligible participants*. In the following, we will use these four metrics to evaluate the performance of proposed schemes. All results shown in charts are the average outcome of running the experiment for 1000 independent rounds.

4.4.2 Performance Comparison

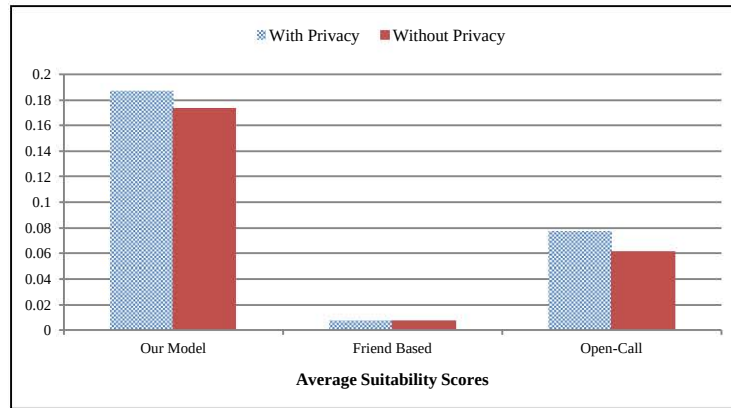
In this section, we compare the performance of our proposed participant selection scheme with two well-known selection methods: (i) *Open-call* which is used in most existing crowdsourcing platforms such as Amazon Mechanical Turk and CrowdFlower.²² In this scheme, the requester broadcasts the task to all members in the community and everyone is able to contribute to the task (details have been presented in Chapter 2), (ii) *Friend-based* which is widely used in social networks and related work such as [136, 137], wherein, the requester advertises the task to his friends.

It should be noted that neither of these methods consider the privacy preservation in their selection methods. So, in order to have a fair comparison, we consider each of the compared methods with two separate configurations: privacy-aware and non privacy-aware. In the privacy-aware configuration, the weights of reputation score, expertise, privacy score and requester’s preference in the computation of the suitability score are 0.5, 0.3, 0.1 and 0.1, respectively (refer to Section 4.2.1). In non privacy-aware recruitment, the weight of privacy score is zero and reputation, expertise and preference are taken into account with weights of 0.55, 0.35, and 0.1 respectively. Another important point is that the simulation results illustrated in the following figures have been scaled with the number of participants in order to have reasonable comparisons. For example, the number of selected participants for each of the aforementioned methods has been scaled with the corresponding maximum number of possible participants. In order to evaluate the performance of methods in real situations, we run the simulation in three different situations:

²²<http://crowdflower.com>



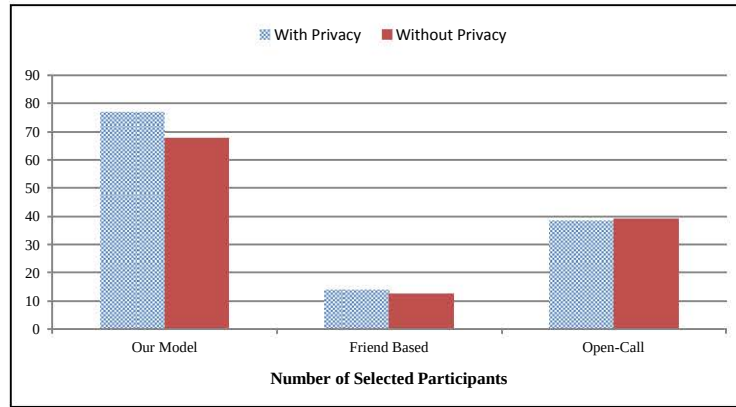
(a)



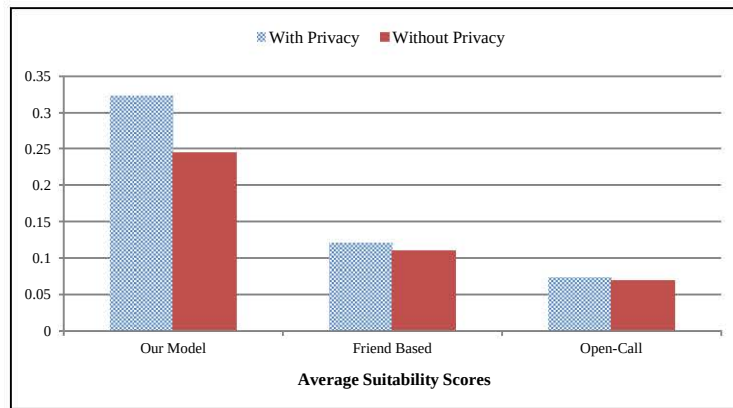
(b)

Figure 4.4: Performance of three methods in the case of requesters with few friends

(i) when the requester has few friends, (ii) when the requester has large number of friends, and (iii) when we select the requester randomly, regardless of his number of friends. Note that the concepts such as ‘few’ or ‘large’ are relative and depend on the characteristics of the underlying social network. In order to consider these situations, we first arrange all members (i.e. Advogato members) in ascending order according to their number of friends (outgoing links). For the first situation, the requester will be selected from the first one third of the members, and for the second situation, the requester will be selected from the last one third. The last situation will be the case when the requester is selected randomly from the unordered list. In Advogato, the range of number of friends for the first group is between 3 and 1000, and for the second group is between 3000 and 4000. To come up with dependable results, we run the simulation for 1000 rounds.



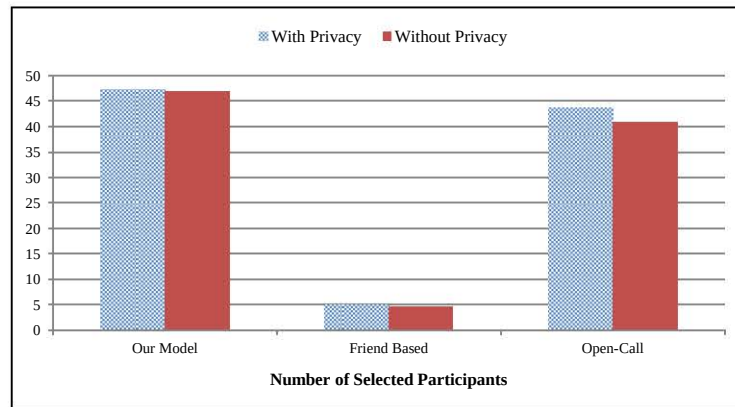
(a)



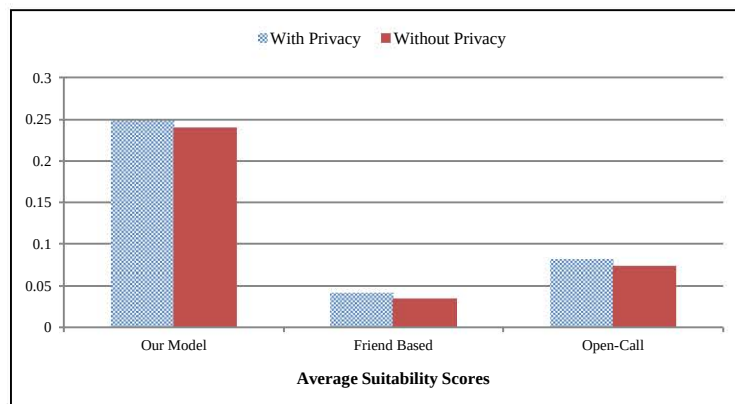
(b)

Figure 4.5: Performance of three methods in the case of requesters with large number of friends

Figures 4.4(a) and 4.4(b) depict the results of comparing three methods for the case in which, the requester has few friends. As it is evident from the charts, open-call outperforms the other two methods in terms of the number of selected participants. This is an expected result since in this method, it is possible to select the participants from everywhere inside the social graph, whereas there are restrictions in participant selection in the two other methods. In our proposed method, a potential limitation is due to λ which restricts the selection domain. Moreover, having few friends will result in less random surfers, which results in less selected participants. The friend-based method is also limited to recruiting one-hop friends. But when it comes to the overall suitability score of the selected participants, the best performance belongs to our scheme. This is because our method considers the suitability scores in the selection scheme and tries to assign higher



(a)



(b)

Figure 4.6: Performance of all methods regardless of requesters' number of friends selection probability to participants with higher suitability scores. This better performance is of great importance since it demonstrates a valuable achievement for the case of having a sparse friendship network, which, as mentioned in Chapter 3, is currently an issue in existing online social networks. The relative order of these three methods is consistent in both privacy-aware and non privacy-aware scenarios. Figures 4.5(a) and 4.5(b) illustrate the performance of the three methods for the case in which, the requester has a large number of friends. In this case, as expected, the performance of the friend-based method improves since the number of friends (potential participants) has increased for both (privacy-aware and non-privacy-aware) scenarios. The best performance still belongs to our scheme. Remember that in the previous case where the requester has few friends, the open-call method outperforms ours in terms of number of selected participants due to the limitation that occurred as a result of the propagation factor and the number of random surfers. Here, our

method outperforms the open-call method even in terms of the number of selected participants. This is due to the large number of friends in the requester's friendship graph, which in turn, increases the number of random walks and consequently, the number of selected participants. Finally, Figures 4.6(a) and 4.6(b) show the results of our experiments when we selected a requester from the community, regardless of his number of friends. In this case, as expected, due to the scarcity of the social network, the overall performance of the open-call is better than the friend-based method. Our proposed scheme is slightly better than the open-call method in terms of the number of selected participants. This small difference between our method and the open-call method is due to the improved performance of open-call in cases where the selected requester has few friends and better performance of our method in cases where the selected requester has a large number of friends.

As all above figures show, the number of selected and eligible participants decreases when privacy considerations are taken into account, since such considerations will result in tighter restrictions in the selection module. The important point is that the relative ordering of the methods in terms of performance remains unchanged in both the privacy-aware and non-privacy-aware scenarios.

4.4.3 Sensitivity Analysis

In this section, we run a series of experiments to reach an optimal setting for our proposed participant selection scheme. In particular, we first obtain the optimal value for λ (propagation factor), and then evaluate the performance of our scheme in the presence/absence of privacy considerations and eligibility assessment scheme.

- **Optimum Value of λ**

As mentioned in Chapter 3, one of the important parameters which impacts the performance of the random surfer is the propagation factor, denoted by λ . In fact, λ is a system-dependant parameter which denotes how deep the random surfer can explore the graph to find suitable participants. In order to assess the performance of our scheme, we need to find an optimum value for λ . Note that as a system-dependant parameter, the optimum value for λ totally

depends on the characteristics and the size of the social graph. We conducted an experiment to test the scheme on the Advogato graph with different values of λ in the range of 1 to 150. For each λ value, we generated 500 tasks and then investigated the outcomes. Based on various runs of this experiment, the highest value of overall suitability score for selected participants is obtained when λ is equal to 100. So, we select this value for λ as the optimum value for future experiments.

- **Performance Analysis of Scheme Components**

In addition to the value of λ , we investigate the impact of two other factors on the performance of our proposed selection scheme. The aim of these experiments is to obtain the best configuration for our proposed scheme.

At first, we try to investigate the effect of privacy score on the evaluation of suitability score. As mentioned before in Section 4.2.1, the probability of selecting a non-friend participant for further tasks of a particular requester has an inverse relation to the number of the tasks the participant has been involved in the past for that requester, due to the reduction in his privacy score. In other words, taking privacy into consideration, while valuable in terms of members' security, it will inherently decrease the number of potential participants. In the following experiments, we aim at investigating how the privacy score consideration will affect the scheme performance in terms of number of eligible and selected participants.

Next, we aim at observing the performance of our scheme with and without leveraging the eligibility assessment scheme. We expect that including the eligibility assessment scheme will increase the overall suitability score, but at the same time, will decrease the number of final participants, since it tightens the criteria of participant selection.

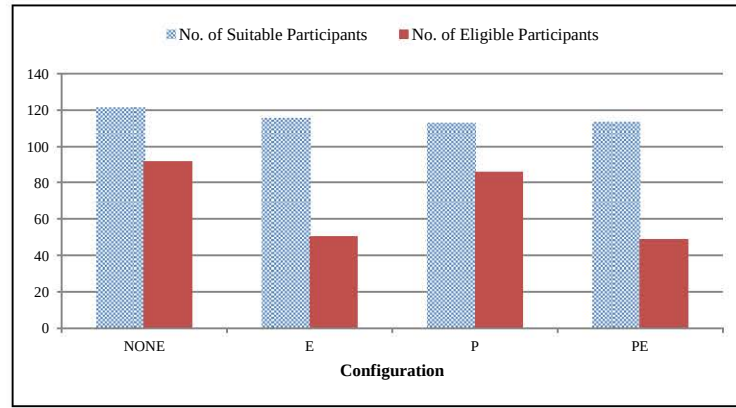
In order to evaluate the effect of these two components, we conducted an experiment in which, the performance of our scheme is evaluated with the following four scenarios:

1. In the first scenario, we neither take privacy nor the eligibility assessment

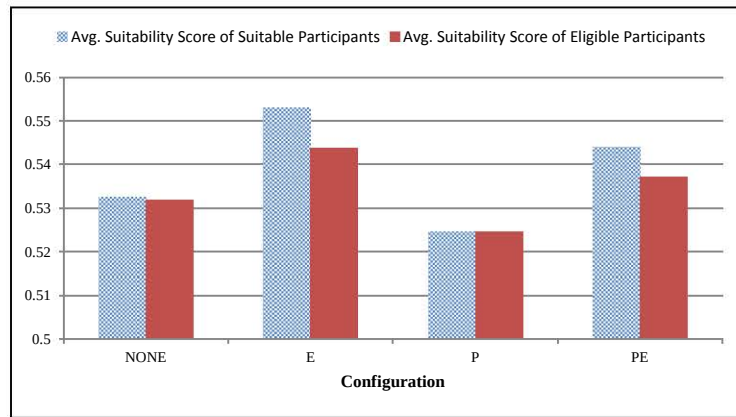
scheme into account. In other words, the suitability score of participants is only calculated based on their reputation, expertise and the requester's preferred and blocked list. Also, we do not consider the eligibility assessment scheme and its time-aware parameters. In our illustrations in Figures 4.7(a) and 4.7(b), we represent this scenario by '*NONE*'.

2. In the second scenario, the eligibility assessment scheme is executed. The privacy does not affect the suitability score. We denote this scenario by '*E*'.
3. In the third scenario, the privacy score is considered in the evaluation of suitability scores. The eligibility assessment scheme, however, is not included. This scenario is denoted by '*P*'.
4. The fourth scenario, is our proposed scheme where both privacy and eligibility assessment scheme are taken into account. We denote this scenario by '*PE*'.

The evaluation results are depicted in Figures 4.7(a) and 4.7(b). As shown in Figure 4.7(a), the overall number of suitable and eligible participants is the highest in the first scenario (when we have neither an eligibility assessment scheme nor privacy considerations). This is because both the privacy consideration and the involvement of an eligibility assessment scheme pose limitations on the number of suitable and eligible participants. However, as Figure 4.7(b) reveals, the overall suitability score in this scenario is too low, since there is no suitability check in the selection process. So, it cannot be deemed as a good configuration. The same argument can be applied to the third scenario as well. In this scenario, the number of suitable participants is greater than those methods which include the eligibility assessment scheme, since there is no limitation for participant selection. However, the average suitability score in this scenario is the least, compared to other scenarios, since it does not consider the suitability score as a dominant factor. So the optimum configuration is to be selected from the second and fourth scenarios (*E* and *PE* scenario). In both *E* and *PE* scenarios, the selection process is applied; but privacy is considered only in *PE*. The overall number of eligible and suitable



(a)



(b)

Figure 4.7: Evaluation of the proposed participant selection scheme with different scenarios

participants in both settings are approximately the same, but the overall suitability score in *E* scenario is slightly (about 0.009) higher than the suitability score in *PE*. Therefore, we conclude that *PE* configuration is the best for the privacy-aware systems and *E* configuration is appropriate for the rest.

4.4.4 Collusion Prevention Analysis

As mentioned in Section 4.4.1, to evaluate the performance of the collusion prevention scheme, we set two experiments. The experiments differ in their employed datasets. In the following, we explain the results of each experiment in detail.

Wikipedia Adminship Election Dataset

In the first experiment, we use the Wikipedia adminship election dataset to in-

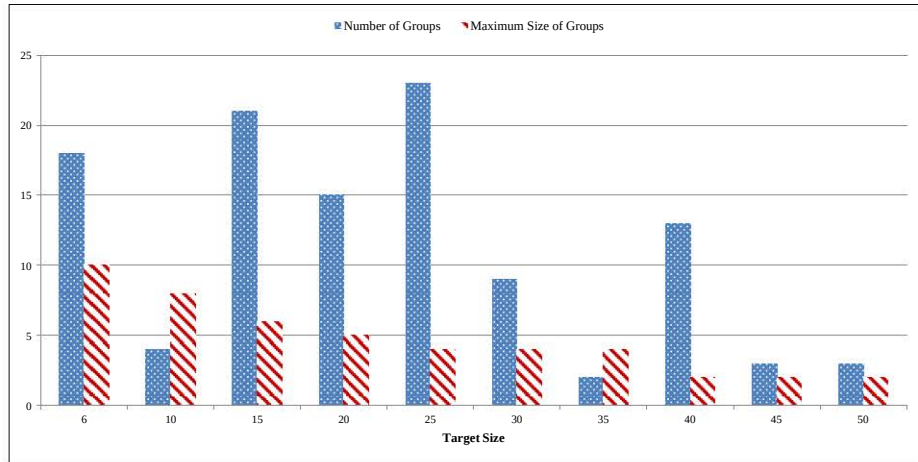


Figure 4.8: Evolution of number of groups and their maximum size according to the target size

investigate the performance of our proposed collusion prevention method. The dataset contains the information related to 2794 tasks. The average number of participants in these tasks is 40. In order to obtain reliable results, we consider the tasks with number of participants greater than the average as the sample data, and randomly select 100 tasks from these. We then test our proposed method to identify any potential colluding group among the participants. As mentioned in Section 4.3, we consider five indicators for detecting potential collusion. Among these indicators, the two indicators Group Size (GS) and Group Target Size (GTS) are the most important indicators as they are the basic conditions for the *formation* of a group. Basically, a group g is created when at least th_1 members of g have collaborated in at least th_2 tasks. So, we first run a short experiment to define the optimal values for th_1 and th_2 .

In order to find the optimum value for th_2 , we set an experiment in which, the target size (i.e. number of the tasks for which the group members have collaborated in the past) is changed. For each target size, we measure the number of groups identified, together with their size. As can be seen in Figure 4.8, the maximum size of identified groups decreases by increasing the target size. This is rational since the probability of finding groups whose members have collaborated in a greater number of tasks is smaller. We believe that the best setting is the one which results in the identification of the largest groups to make a considerable impact. As derived from

the figure, this situation is related to the case where the target size is 6. So, we set th_2 to be equal to 6. For the sake of simplicity, we assume that th_1 equals to th_2 .

In order to investigate the performance of our proposed collusion prevention method, we first utilise the FIM technique to find the candidate groups among the participants. The outcome is the discovery of 18 candidate groups with at least 10 members. We then employ our collusion prevention method and identify 9 of these 18 groups as collusive. To evaluate the efficiency and accuracy of our method, we examine a number of statistical metrics. At first, we measure the ratio of the tasks targeted with the colluding groups. The result shows that 14% of the tasks were affected by these 9 colluding groups. This means that our collusion prevention method is able to prevent these tasks from being targeted by the colluders. We then calculate the success ratio of the tasks targeted by the colluding groups as well as all 100 tasks. In the Wikipedia adminship election dataset, a task (an election) is successful if it results in the selection of the user as an administrator (note that the results are available in the dataset). By success ratio, we mean the ratio of the tasks that have resulted in a desired decision (i.e. resulted in the selection of a user as an admin), to the total number of tasks. We observe that overall success ratio of the tasks in our dataset is 71%. This ratio is 83% for the groups identified by our collusion detection method. This means that there is a high probability that the groups identified by our method are colluding groups, since their collaboration has resulted in a considerably high success ratio. This is a significant indication that the identified groups are much more likely to be collusive.

Advogato Dataset

In the second experiment, we use Advogato dataset. We first create a set of candidate groups among the Advogato members, and then, we define some of these candidate groups as collusive. In order to create candidate groups, we first select 90 Advogato members with at least 30 trust relations (i.e. 30 friends). Each of these members along with 20 out of his 30 friends forms a candidate group. When a task is released, a set of Advogato members are considered as eligible to contribute (by using the aforementioned suitability assessment and eligibility assessment techniques). Each candidate group with at least 10 eligible members is considered as

collusive. The collusive group members contribute polluted data while other eligible members contribute genuine data. Specifically, we assume that the genuine data (d) is a random number in $[0,1]$, while the polluted data is a random number in $(d - \mu, d + \mu)$. Greater values for μ result in polluted values with great deviation from the genuine values, which makes the collusion detection easier. In our experiments, we set μ to be 0.2. Note that for each task, all the collusive members report the polluted data, while others report the genuine data. We run the experiment for 10 rounds. In each round, 20 tasks are released. At the end of each round, we utilise the FIM technique to find the groups. The outcome is the set of all groups among the eligible participants (who have collaborated in at least 5 tasks). Then, for each group identified by FIM, the possibility of collusion (PoC) is computed by utilising Equation 4.15. Groups with $PoC > 0.5$ are identified as collusive (In Equation 4.15, we assume that all the indicator weights are equal to 0.2).

In order to evaluate the efficiency and the accuracy of our proposed method in identifying the colluding groups, we utilise two criteria. For the evaluation of accuracy, we use the well-known measures of precision and recall [138]. Precision measures the quality of the identification results, and is defined by the ratio of the correct identification of colluding groups, to the total number of groups identified by our method. Recall measures coverage of the identification results, and is defined by the ratio of the collusive groups identified correctly to the total number of all correct colluding groups that should be found. These two definitions are summarised in the following equations:

$$Precision = \frac{\text{number of collusive groups identified correctly}}{\text{total number of identified groups}}$$

$$Recall = \frac{\text{number of collusive groups identified correctly}}{\text{total number of existing collusive groups}}$$

These two measures are usually expressed as percentages. For an approach to be effective, it should achieve a high precision and high recall. However, in reality these two metrics tend to be inversely related [139]. This means that the improvements in precision come at a cost of reduction in recall, and vice versa.

Figure 4.9 shows the evolution of precision in different rounds. As displayed in

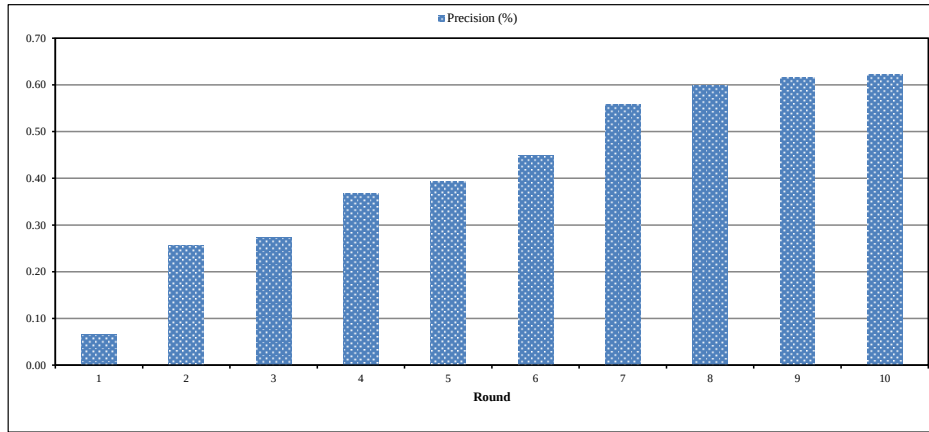


Figure 4.9: Evolution of precision (%) in different rounds

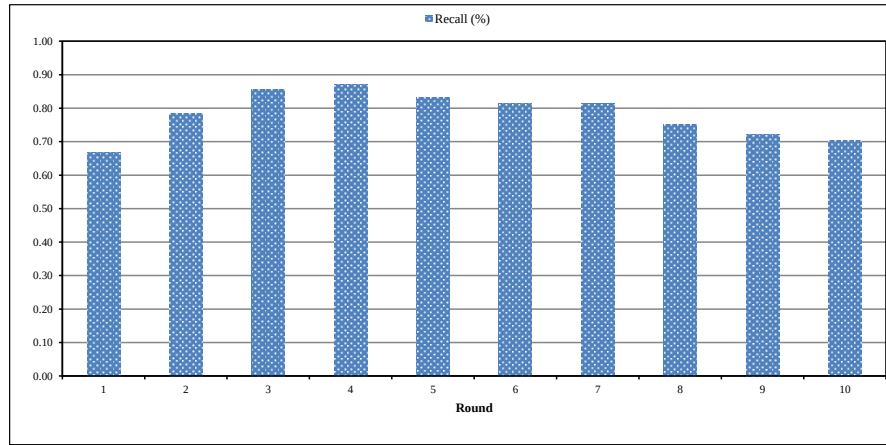


Figure 4.10: Evolution of recall (%) in different rounds

this figure, the collusion detection scheme achieves a precision of 63%. This means that our collusion prevention method is able to prevent 63% of the tasks from being targeted by the colluders. This is due to the suitability of the indicators, which correctly model the collusive behaviour of group members. Note that it may be possible to achieve greater precision but would result in a drop in recall. As can be observed in this figure, the precision values evolve in a constantly increasing manner. A lower value of precision in the first rounds is due to the lack of adequate history related to the colluders' behaviours. In other words, due to the small number of released tasks in the first rounds of the experiment, the collusion prevention scheme does not have the required information (e.g., content similarity, target size, etc.) at hand. As time goes by, collusive members collaborate in more tasks which results in the availability of more behavioural information such as number of the tasks they

have collaborated on, the contributions they have reported to these tasks, etc. This helps the collusion prevention method to better detect the collusive behavioural pattern.

Figure 4.10 depicts the evolution of recall in various rounds. As can be seen in this figure, our scheme also achieves a high percentage of recall (86%), which denotes that our collusion prevention scheme is successful in detecting 86% of the existing collusive groups. It can be observed that there is a slightly descending growth in recall after the fourth round which, as mentioned above, is natural in real systems, since precision and recall typically evolve inversely [139].

As mentioned above, a group is identified as collusive if the possibility of collusion (PoC) for this group is above 0.5. The possibility of collusion is obtained by averaging the indicator values. However, in order to ensure that the indicators are selected correctly, we calculate the distribution of values of each indicator in all collusive groups identified by our method. Figure 4.11 depicts the distribution of values calculated for collusion indicators. The values calculated for indicators are almost always higher than 0.5. This illustrates that the identified indicators are suitable and effective for detecting collusion in social participatory sensing.

To be brief, the results show that our proposed collusion prevention method is successful in preventing the formation of colluding groups among the selected participants with high accuracy.

4.5 Conclusion

In this chapter, we proposed a participant selection scheme for social participatory sensing systems. The aim was to address the challenge of identifying suitable participants in social participatory sensing applications. Our proposed participant selection scheme first assessed the suitability of a participant who has been identified via the utilisation of the recruitment scheme (presented in the previous chapter). The suitability assessment was carried out based on a set of factors such as participant's reputation, expertise level, etc. Then, a set of eligible participants were

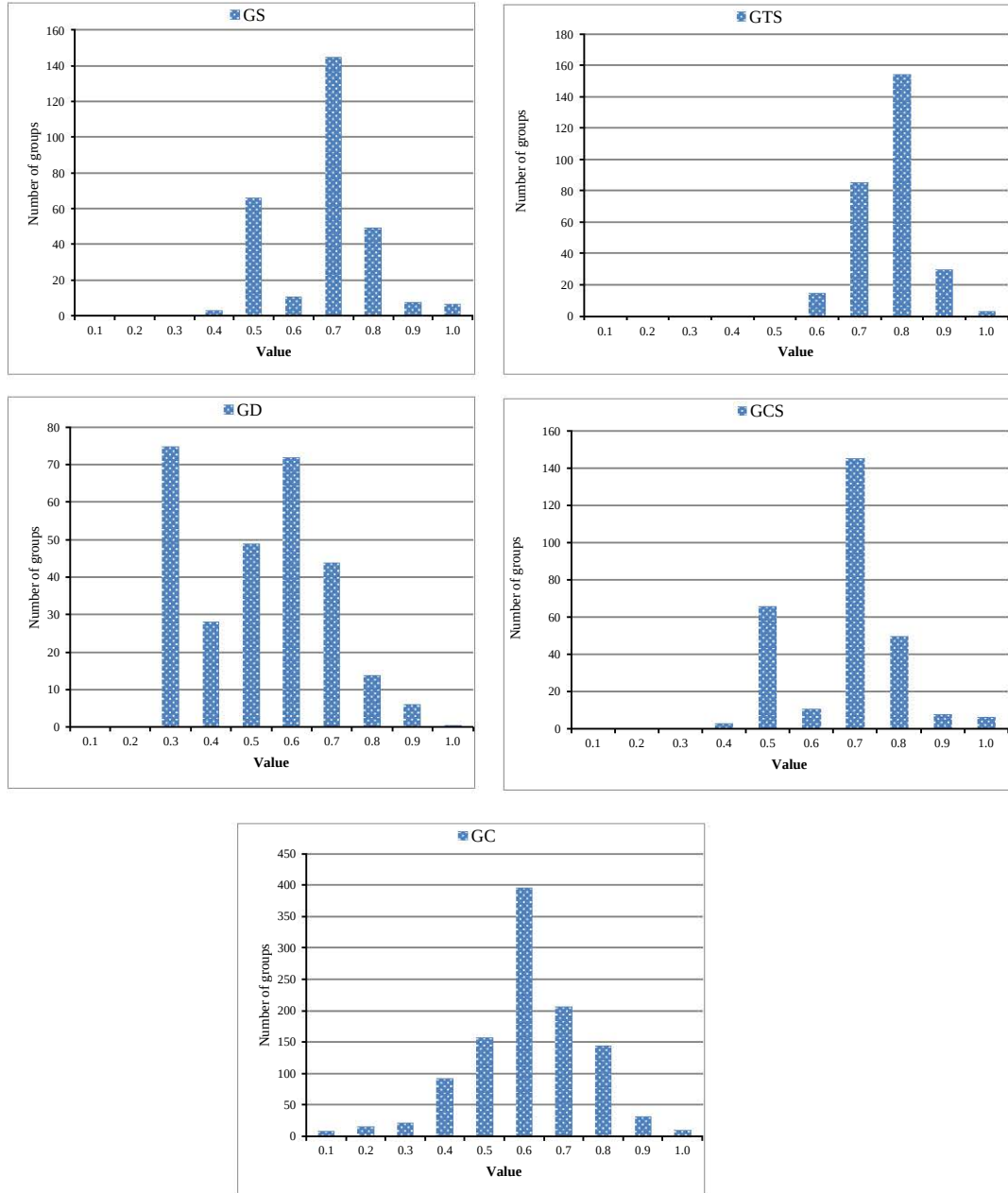


Figure 4.11: Distribution of the values of indicators in collusion attacks

selected who can satisfy the time limitations of the task. In particular, we proposed an eligibility assessment technique which considers a participant as eligible to contribute according to his timeliness as well as the remaining time of the task, to ensure that the participant would provide timely contributions. We went through extensive simulations to evaluate the effectiveness of our proposed participant selection scheme. The simulation results demonstrated that our scheme increases the number of participants who are reputable and well-suited to contribute.

We also proposed a collusion prevention scheme with the goal to prevent the formation of colluding groups within the selected suitable participants. The scheme investigates the possibility of collusion upon each eligible participant. This decision is made based on a set of indicators that are related to the common approaches utilised by colluders to arrange a collusive attack. Colluders normally form a large group and collectively collaborate on a large number of tasks. They normally contribute similar content which deviate from the genuine contributions provided by honest participants. They may also benefit from the social groups to better manage their communications. We then calculated the possibility of collusion based on these indicators. In order to measure the performance of the collusion prevention scheme, we set up two experiments in which, the datasets Wikipedia adminship election and Advogato were employed. The result of these experiments showed that our proposed scheme is able to detect the collusive groups with high precision. The results also demonstrated the correctness and effectiveness of proposed indicators.

The proposed participant selection scheme works in concert with the recruitment scheme described in the previous chapter to support the requester with sufficient number of well-suited participants. These participants are then recruited and provide multiple contributions. In the next chapter, we discuss the details of our novel trust scheme which is intended to assess the trustworthiness of reported contributions and update the reputation of participants accordingly.

Chapter 5

Assessing Trust and Reputation

5.1 Introduction

In Chapter 1, we proposed a framework to address the main challenges of social participatory sensing. As depicted in Fig 1.3, the framework consists of three main schemes. The recruitment scheme described in Chapter 3 aims at addressing the issue of sufficient participation. The participant selection scheme explained in Chapter 4, addresses the participant suitability issue. These two schemes collaborate hand in hand to provide a sufficient number of suitable participants who are then recruited for the task and provide contributions. In this chapter, we present the novel trust scheme to address the trust issue. In particular, the proposed trust scheme is intended to comprehensively assess the trustworthiness of contributions and update the reputation of participants accordingly.

The inherent openness of participatory sensing, while valuable for encouraging participation, also makes it easy for the propagation of erroneous and untrusted contributions. Typically, untrusted contributions originate from two types of behaviours: (i) inadvertent behaviours initiated by careless users and (ii) deliberate behaviours initiated by malicious users. In the former, we assume there exists careless users whose behaviours unintentionally cause the mobile phone sensors to produce corrupted data. For example in noise monitoring applications such as Ear-Phone [5], a careless user may keep the mobile phone in his pocket while collecting

sound samples. In the latter, we consider malicious participants who intentionally report false contributions to achieve personal gains. For example, in the Petrol-Watch [2] application in which, pricing billboards are captured by camera phones, a service station operator may intentionally report higher prices for other stations in order to attract more customers.

We have explained in Chapter 2 that the existing solutions for trust assessment in participatory sensing rely on the quality of contributions. However, we encounter new trust issues in the concept of social participatory sensing. People generally rely more on contributions provided by their close friends than those of others. Hence, in social participatory sensing, it is crucial to consider both, the participant's social trust and the data quality, as influential factors in evaluating the trustworthiness of contributions.

Another important issue in soliciting high quality contributions is the reputability of the contributors. Since trust and reputation are sometimes used interchangeably, we first present a formal definition for these concepts.

- We use the term ‘trust’ to represent the level of confidence about the reliability of a participant. In other words, trust is a *pair-wise* concept, which defines the belief of one node in the reliability, truth, ability, or strength of another node.
- The ‘reputation’ of a participant, denoted as ρ , is the overall quality or character as seen or judged by people in general. In other words, reputation is a *community-wide* opinion generally held about someone.

Our extensive review of the related research in reputation management schemes for participatory sensing in Chapter 2 has also shown that the reputability of a participant is typically measured based on the quality of his contributions in past campaigns [31, 22, 26]. During the Hurricane Sandy in and around USA in October 2012, social media was widely exploited by malicious entities to spread fake pictures (e.g., photoshopped images of sharks swimming in New York streets). Fake images and news, initially thought to be true, can become extremely viral on social networks and cause panic and chaos among the people affected by the incident. Recent

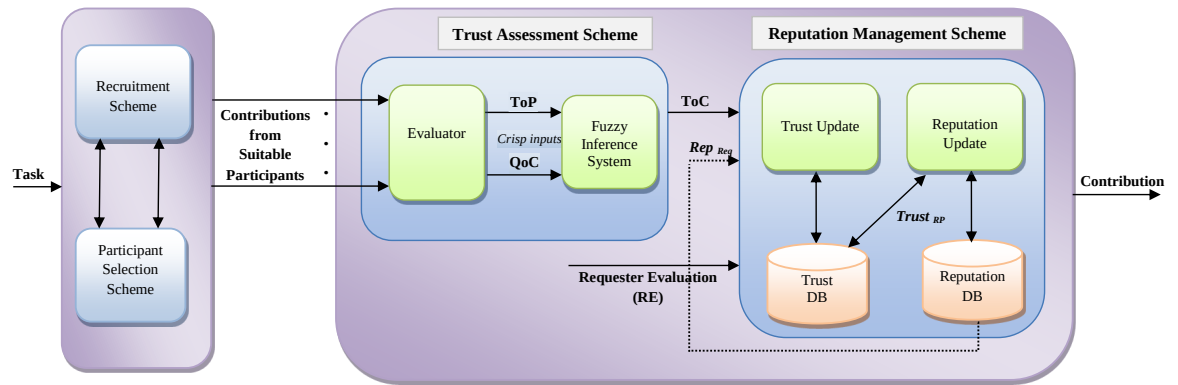


Figure 5.1: Trust scheme at a glance

analysis [27] of tweets collected during Hurricane Sandy revealed that a total of 10,350 unique tweets contained fake images. Retweets accounted for 86% of tweets spreading the fake images. In another study analysing tweets posted during the terrorist bomb blasts in India (Mumbai, July 2011) [28], it was noticed that the majority of Twitter members who disseminated fake information had a lower number of followers which is generally interpreted as being a low-reputable user [28]. This clearly demonstrates the relationship between the reputability of participants and the trustworthiness of their provided contributions.

In order to provide a thorough vision of trust and reputation, we propose a trust scheme which offers a comprehensive measure of trust by considering all possible influential factors. In our scheme, the trustworthiness of the participant, as an effective and significant aspect, is independently assessed and combined with the quality of the data using fuzzy logic to arrive at a comprehensive trust rating for each contribution. These trust ratings are then used by our proposed reputation management scheme to calculate and update the reputation score of participants. By adopting a fuzzy approach, our proposed scheme is able to concretely quantify uncertain and imprecise information, such as trust, which is normally expressed by linguistic terms rather than numerical values.

Figure 5.1 illustrates the data flow in the proposed trust scheme. The contributions received by the selected participants are given to the trust scheme which incorporates the proposed trust assessment technique and fuzzy inference engine

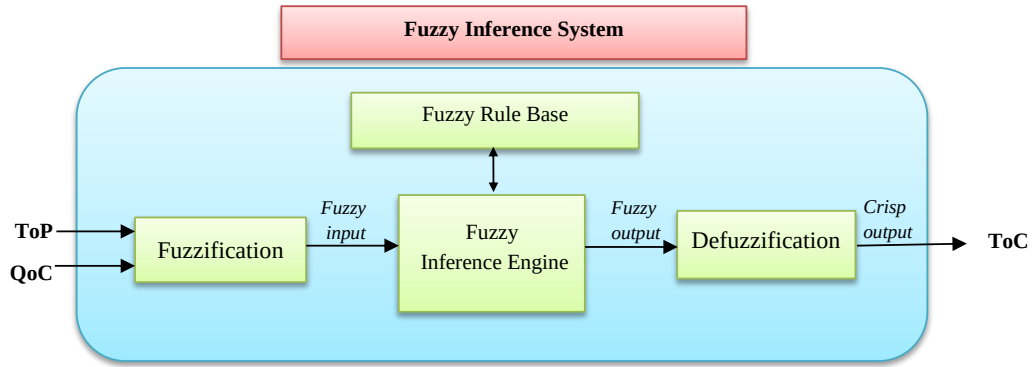


Figure 5.2: Fuzzy inference system architecture

(depicted in Figure 5.2) and arrives at a trust rating for each contribution. The trust assessment scheme maintains and evaluates a comprehensive trust rating for each contribution by considering two influencing factors: (1) Quality of Contribution (QoC) and (2) Trust of Participant (ToP). The application server maintains a trust database, which contains the required information about participants and the history of their past contributions. When a contribution is received to the trust assessment scheme, the effective parameters that contribute to the two aforementioned factors are evaluated by the evaluator component and then combined to arrive at a single quantitative value for each. The two measures serve as inputs for the fuzzy inference system, which computes the final trust rating. This cumulative trust rating is then used as a criterion to accept/reject the contribution by comparing it against a predefined threshold (Th_R).

At the end of each campaign, a cumulative objective trust rating, referred to as $Trust_{RP}$ is automatically updated for each participant, which denotes the level of the trustworthiness the requester can have on the participant. $Trust_{RP}$ is dependent on the trustworthiness of the contribution that the participant has provided for the requester.

For certain campaigns, depending on the nature of task, the requester may desire to add a subjective evaluation in order to indicate how much the contribution is compatible with his needs and expectations. In such a case, this subjective rating is combined with the system-computed rating to update $Trust_{RP}$.

At regular intervals, a reputation score is also calculated for each participant.

The reputation score of each node depends on (i) the trust ratings that other nodes have assigned to him , and (ii) the reputation of those nodes. The reputation score is further used as a weight for the participant's evaluations, ratings or reviews. Moreover, a suggested list is created for the requester, which contains a list of participants who have shown a satisfiable performance in multiple campaigns. The list is further used for recruitment or friendship establishment (more details in Section 5.3).

The novel contributions of this chapter are as follows:

- We propose a novel trust assessment scheme for social participatory sensing systems. The aim is to address the challenge of data trustworthiness by evaluating and assigning a trust score to each sensor data. This trust score can be further used by the requester to accept/revoke the corresponding contributions.
- We present two main influencing factors that affect the trustworthiness of a contribution: The Quality of Contribution (QoC) and the Trustworthiness of Participant (ToP). We then present effective parameters in each influencing factor and perform extensive investigations to set the parameters in a way that results in meaningful and valid values for each factor.
- We propose a fuzzy inference engine to combine these influencing factors together. Our methodology covers all possible combinations of trust influencing factors and combines them by leveraging fuzzy logic in an effort to closely align with the human decision-making process. The inputs to the fuzzy inference system are the crisp values of QoC and ToP. The output is a singular crisp value that denotes the trust score of the contribution.
- We introduce a new concept of subjective rating (called Requester Evaluation (RE)) which allows the requester to evaluate the contribution based on how closely it satisfies the task requirements.
- We design a novel scheme for calculating the reputation score for each participant that supports applications with accurate and inclusive reputation scores for participants. The reputation score can be used in different ways, depending

on the functionality of the scheme. In the design of the reputation scheme, we assume that each user acts as a participant in one campaign and as a requester in another campaign. While serving as a requester, the user's reputation score serves as a weight for the evaluation that this user assigns to each contribution. In other words, we assume that the requester's subjective evaluation has a weight, which is equal to the reputation score of the requester. As a participant, the reputation score of the user is used as a criterion while selecting the reputable contributors.

- We provide the requester with a connection suggestion list containing a set of participants who have shown trustworthy behaviour in previous campaigns and are proper candidates for further recruitment or friendship establishment.

The rest of this chapter is organised as follows. We present the details of our trust assessment scheme in Section 5.2. We then explain the reputation management scheme in Section 5.3. Simulation results are discussed in Section 5.4. Finally, Section 5.5 concludes the chapter.

5.2 Trust Assessment Scheme

Since the trust scheme attempts to mimic how human's perceive trust, we first present a simple illustrative example. Suppose John is a member of an online social network (e.g., Facebook). John has made a profile and has friended several people. John is a vegetarian. John is also on a budget and is keen to spend the least possible amount for his weekly groceries. John decides to leverage his social circle to find out the cheapest stores where he can buy vegetarian products. Specifically, John acts as the requester and asks his friends to capture geotagged photos of price labels of vegetarian food items when they are out shopping and to send these back to him. One of his friends, Alex decides to help out as a participant and provides him with several photos of price labels. In order to decide whether to rely on Alex's contributions, John would naturally take into account two influencing factors: (i) his personal trust perception of Alex, which would depend on various factors such as the

nature of friendship (close vs. distant), Alex’s awareness of vegetarian foods, Alex’s location, etc. and (ii) the quality of Alex’s data which would depend on the quality of the pictures, relevance of products, etc. In other words, John in his mind computes a trust rating for Alex’s contribution based on these two factors. Our proposed trust scheme provides a means to obtain such trust ratings by mimicking an approach similar to John’s perception of trustworthiness in a scalable and automated manner. This trust rating helps John to select trustable contributions and accordingly plan for his weekend shopping. Moreover, the reputation scheme provides a reputation score for each of the participating friends (such as Alex), according to the trustworthiness of their successive contributions. It also affords a list of trustable friends for the data consumer (e.g., John) for future recruitment.

The trust assessment scheme maintains and evaluates a comprehensive trust rating for each contribution. In particular, there are two influencing factors that need to be considered: (1) Quality of Contribution (QoC) and (2) Trust of Participant (ToP). In the following, we present a brief discussion about the underlying parameters and the evaluation methods.

5.2.1 Quality of Contribution (QoC)

In participatory sensing, contributions can be in any form, such as images or sounds. The quality of the data is affected not only by the fidelity of the embedded sensor but also the sensing action initiated by the participant. The in-built sensors in mobile devices can vary significantly in precision. Moreover, they may not be correctly calibrated or even worse not functioning correctly, thus providing erroneous data. Participants may also use the sensors improperly while collecting data (e.g., not focussing on the target when capturing images). Moreover, human-as-sensor applications such as weather radar in [48] are exposed to variability in the data quality due to subjectivity. For example, what is hot for one person may be comfortable for another. In order to quantify QoC, a group of parameters must be evaluated such as: relevance to the campaign (e.g., groceries in the above example), ability in determining a particular feature (e.g., price tag), fulfilment of task requirements

(e.g., specified diet restrictions), etc.

There already exists research that has proposed methods for evaluating the quality of data in participatory sensing, depending on the sensing modality. In image-based tasks such as PetrolWatch [2], the images taken from the fuel price boards are exposed to sophisticated image processing systems and computer vision algorithms in order to measure the quality of the contribution and extract the fuel prices. Also in DietSense [20], ImageScape which is a software tool for processing, clustering, and browsing large sets of images is used to investigate the photos taken from the meals in front of the users. It is also able to detect and remove the images which are too dark or blurry to be useful. In sound-based tasks such as those proposed in [63, 15], outlier detection algorithms [29] are used to evaluate the data quality and revoke unreliable contributions. In [15], a weight is assigned to each sensor device (by utilising the consensus based outlier detection algorithms), which is inversely proportional to the deviation between the device sensed data and the group consensus. The weight assigned to the sensor device reflects the quality of its sensed data.

Rather than reinventing the wheel, our system relies on the state-of-the-art methods for this evaluation. The result is a single value for QoC in the range of $[0, 1]$.

5.2.2 Trust of Participant (ToP)

ToP is a combination of personal and social factors. Personal factors consist of the following parameters:

- *Expertise (E)*

It is defined as the measure of a participant's knowledge and is particularly important in tasks that require domain expertise. Greater credence is placed on contributions made by a participant who has expertise in the campaign. We propose to use expert finding systems for evaluating expertise. These systems employ social networks analysis and natural language processing (text mining, text classification, and semantic text similarity methods) to analyse explicit information such as public profile data and group memberships as well as

implicit information such as textual posts to extract user interests and fields of expertise [84]. In particular, the Dmoz²³ open directory project is used for expertise classification. Expertise evaluation is done by incorporating text similarity analysis to find a match between the task keywords (e.g., vegetarian) and a participant's expertise.

We assume that the set TE contains the task's required expertise and PE is the set of participant's expertise attributes. In this case, the expertise score of each participant is defined as Equation 5.1:

$$E = \frac{n(TE \cap PE)}{n(TE)} \quad (5.1)$$

where $n(A)$ is the number of elements in set A .

- *Timeliness (T)*

Timeliness measures how promptly a participant performs prescribed tasks. It depends on the contribution response time (t) and the task deadline (d). To evaluate this parameter, we utilise the inverse Gompertz function, which is defined as $T(t) = 1 - e^{-be^{-ct}}$. The intuition behind utilising the Gompertz function is its compatibility with timeliness evolution [140]. The timeliness score is highest when the contribution is received immediately after the task release time. The score begins to decrease as the response time increases, reaching the minimum value when the response is received just before the deadline. In the original inverse Gompertz function, the lower asymptote is zero; it means that the curve approaches to zero in infinity. In our case, timeliness rate will only be zero if a contribution is received after the deadline; otherwise, a value between x and 1 is assigned to it. This means that the lowest timeliness rating will be x if contribution is received before the deadline, and is zero if received after the deadline. So, we modify the function as Equation 5.2 to calculate the timeliness (T):

²³<http://www.dmoz.org>

$$T(t) = \begin{cases} 1 - [(1 - x)e^{-be^{-ct}}] & \text{if } t < d \\ 0 & \text{otherwise} \end{cases} \quad (5.2)$$

- *Locality (L)*

Another significant parameter is locality, which is a measure of the participant's familiarity with the region where the task is to be performed. We argue that contributions received from people with high locality to the tasking region are more trustable than those received from participants who are not local, since the first group is more acquainted with and has better understanding of that region. In order to quantify the participant's locality, we rely on the results of the real experiment presented in [141]. In this experiment, authors implemented a mobile crowdsourcing platform that integrates location as a parameter for distributing tasks to workers. They asked the participants during the recruiting process to provide their home and office address, which they used to define different tasks. Particularly, they defined three types of location-based tasks: (i) tasks at/nearby the users' home locations, (ii) tasks at/nearby the users' office locations, and (iii) tasks in the city centre of their hometown. Participants were then recruited to attend in these tasks and provide contributions. According to their experimental results, people tend to perform tasks that are near to their home or work place (places that they are considered 'local' to them). This implies that if we log the location of participants' contributions, we can estimate their locality. A participant's locality would be highest at locations from where they make the maximum number of contributions. In order to evaluate locality, we assume that the sensing area has been divided to n regions, and a vector V with the length equal to n is defined for each participant, where, $V(i)$ is number of samples collected in region i . In this case, locality of a participant to region i is calculated by Equation 5.3:

$$L(i) = \frac{V(i)}{\sum_{i=0}^{n-1} V(i)} \quad (5.3)$$

Next, we explain the social factors that affect ToP:

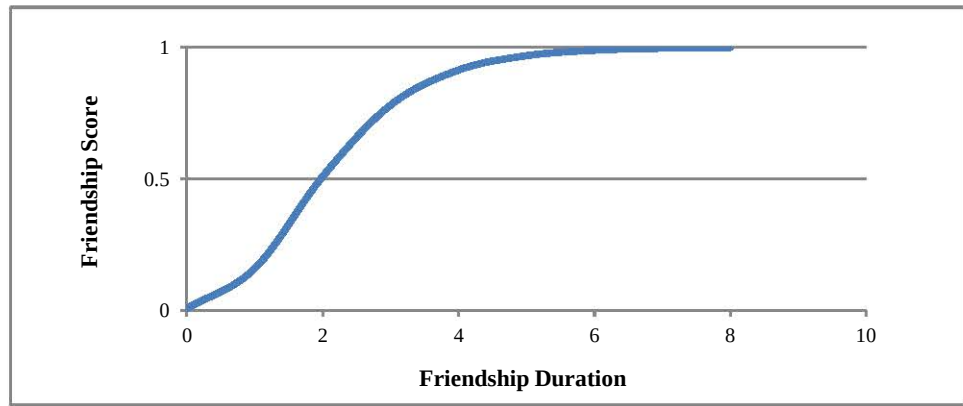


Figure 5.3: Gompertz function for friendship score

- *Friendship duration (F)*

In real as well as virtual communications, long-lasting friendship relations normally translate to greater trust between two friends [142]. So, friendship duration which is an estimation of friendship length is a prominent parameter in trust development [142]. We use the Gompertz function depicted in Figure 5.3 to quantify friendship duration due to its match with the friendship evolution. Slow growth at the start resembles the friendship gestation stage. This is followed by a period of accumulation where the relationship strengthens culminating in a steady stage.

The Gompertz function [75] is a well-known method for modelling a great variety of processes due to its flexibility. In particular, the Gompertz function provides the following important features [143]: (i) sigmoidal advancement; a monotonous increase in accuracy with increase in group size, (ii) the rate at which information is produced is smallest at the start and end of the process. (iii) asymmetry of the asymptotes, as for any value of t , the amount of information gathered in the first t time steps is greater than the amount gathered at the last t time steps. The Gompertz function is frequently used for modelling a great variety of processes such as population in a confined space [144] and growth of tumours [145]. As such, the friendship duration is evaluated according to Equation 5.4, in which, b and c are system-defined constants and

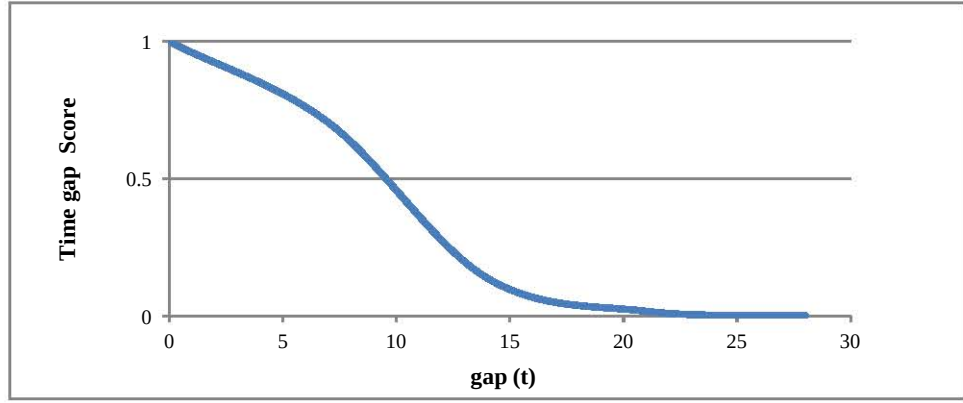


Figure 5.4: Inverse Gompertz function for time gap score

t is the time in years.

$$F(t) = e^{-be^{-ct}} \quad (5.4)$$

- *Interaction time gap (I)*

In every friendship relation, interactions happen in the form of sending requests and receiving responses. Interaction time gap measures the time between the consequent interactions and is a good indicator of the strength of friendship ties. If two individuals interact frequently, then it implies that they share a strong relationship, which translates to greater trust [140]. We propose to use the inverse Gompertz function depicted in Figure 5.4 to quantify the interaction time gap, since a smaller time gap indicates stronger relationship, which leads to high social trust and vice-versa. So, the interaction time gap is evaluated according to Equation 5.5, in which, b and c are system-defined constants and t is the gap (in days) between the current time and the latest interaction (LI) time.

$$I(t) = 1 - e^{-be^{-ct}} \quad (5.5)$$

The aforementioned parameters are combined by the Evaluator to arrive at a single value for ToP, as depicted in Equation 5.6,

$$ToP = w_1 \times E + w_2 \times T + w_3 \times L + w_4 \times F + w_5 \times I \quad (5.6)$$

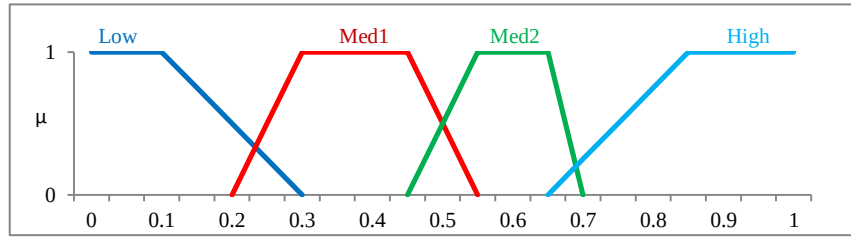


Figure 5.5: Membership function for quality of contribution (QoC) and trust of participant (ToP)

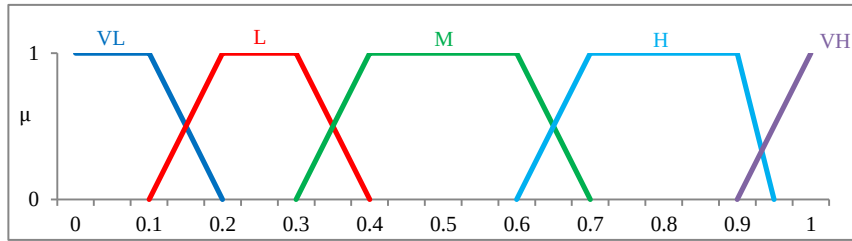


Figure 5.6: Membership function for trust of contribution (ToC)

where w_i is the weight of each parameter, and $\sum_{i=1}^5 (w_i)$ equals to 1. The adjustment of the weights depends on the nature of the task. For example, in location-based tasks, w_3 is set to be considerably high to give more weight to the locality parameter. Similarly, for tasks where real-time information is important, a higher weight may be associated with timeliness (w_2). ToP is in the range of $[0,1]$.

5.2.3 Trust of Contribution (ToC)

Our proposed scheme employs fuzzy logic to calculate a comprehensive trust rating for each contribution, referred to as the Trust of Contribution (ToC). The proposed scheme covers all possible combinations of influencing factors and combines them in an effort to closely align them with the human decision-making process. The inputs to the fuzzy inference system are the crisp values of QoC and ToP. Since the fuzzy inference system that is used in this chapter is similar to the one in Chapter 3, we omit the general explanations and definitions and emphasise solely the contents specific to this chapter.

As mentioned in Chapter 3, the fuzzifier converts the crisp values of input parameters into a linguistic variable according to their membership functions. The

Table 5.1: Fuzzy rule base for defining trust of contribution (ToC) according to quality of contribution (QoC) and trust of participant (ToP)

Rule no.	if QoC	and ToP	Then ToC
1	Low	Low	VL
2	Low	Med1	L
3	Low	Med2	L
4	Low	High	M
5	Med1	Low	L
6	Med1	Med1	L
7	Med1	Med2	M
8	Med1	High	M
9	Med2	Low	M
10	Med2	Med1	H
11	Med2	Med2	H
12	Med2	High	H
13	High	Low	H
14	High	Med1	H
15	High	Med2	VH
16	High	High	VH

fuzzy sets for QoC, ToP and ToC are defined as:

$$T(\text{QoC}) = T(\text{ToP}) = \{\text{Low}, \text{Med1}, \text{Med2}, \text{High}\}$$

$$T(\text{ToC}) = \{\text{VL}, \text{L}, \text{M}, \text{H}, \text{VH}\}.$$

Figure 5.5 also represents the membership function of QoC and ToP and Figure 5.6 depicts the ToC membership function.

The role of inference engine is to convert fuzzy inputs (QoC and ToP) to the fuzzy output (ToC) by leveraging If-Then type fuzzy rules. The combination of the above mentioned fuzzy sets create $4 \times 4 = 16$ different states which have been addressed by 16 fuzzy rules as shown in Table. 5.1. The result of the inference engine is ToC which is a linguistic fuzzy value. The defuzzifier is then used to convert the ToC fuzzy value to a crisp value in the range of $[0, 1]$ by employing the Centre of Gravity method (COG) [127, 146].

5.3 Reputation Management Scheme

Once the ToC is defined for a contribution, the corresponding requester-participant mutual trust is updated, which is then used to calculate/update the participant's reputation score. In the following, we describe these steps in detail.

As mentioned before, depending on the nature of the task, the requester may desire to add a subjective rating to a participant's contribution. Typically, in a participatory sensing system, there are two types of tasks. The first types of tasks are those where subjective rating is important. This is particularly relevant for campaigns where it is difficult for the requester to express his real needs, desires or restrictions via task definition. Subjective rating is also useful when the requester does not have enough knowledge about the task and needs an expert review to confirm the validity of the contributions. For example, assume a requester with a strict gluten-free diet who asks his friends to take photos of the price tags and ingredients of gluten-free products. The term gluten free is generally used to indicate a supposedly harmless level of gluten rather than a complete absence. For those with serious celiac disease, the maximum safe level of gluten in a finished product is even lower than the amount that exists in normal gluten-free products [147]. Hence, it is essential that the product ingredients are ratified by the requester himself and a nutritionist to ensure that it is safe to be consumed.

On the other hand, the second type of tasks do not benefit from subjective rating because the requester may not be in the best position to evaluate the quality due to not being aware of the ground truth. Ear-Phone [5] is an example of such a task in which, participants are recruited to gather noise samples. In such cases, the requester gives the authority to the trust assessment scheme and relies on the objective evaluation of the system, which automatically assigns a rating to the contributions by leveraging methods such as majority consensus [15].

In order to support both kinds of tasks, we denote the subjective rating as requester evaluation (RE) which implies the trustworthiness of the contribution from the requester's point of view. In our simulation in Section 5.4, we assume that RE has a value in the range of $(ToC - \mu, ToC + \mu)$, where $\mu = 1 - \rho_{Req}$ and ρ_{Req} is the

requester's reputation score. For a requester with a high reputation score, the value of μ is small, resulting in RE close to ToC. This means that a requester with a high reputation score is likely to assign a rating, which is close to the system-computed rating.

In the absence of subjective ratings, the requester simply relies on the objective ratings assigned by the trust assessment scheme. In this case, μ is simply set to zero, resulting in $RE = ToC$.

Based on the ToC assigned to each contribution, the trust of the requester upon the corresponding participant ($Trust_{RP}$) is updated. We adopt a reward/penalty policy for this update. A participant with ToC values greater than a predefined threshold (Th_1) is rewarded, and the amount of $|ToC - \rho_{Req} * RE|$ is added to $Trust_{RP}$. Similarly, a participant with ToC less than a predefined threshold (Th_2) is penalised, and the amount of $|ToC - \rho_{Req} * RE|$ is reduced from $Trust_{RP}$. This can be summarised in Equation 5.7. In our simulations in Section 5.4, we set $(Th_1) = 0.7$ and $(Th_2) = 0.3$.

$$Trust_{RP} = \begin{cases} Trust_{RP} + |ToC - \rho_{Req} * RE| & \text{if } ToC > Th_1 \\ Trust_{RP} - |ToC - \rho_{Req} * RE| & \text{if } ToC < Th_2 \end{cases} \quad (5.7)$$

Note that in this equation, we use the requester's reputation score (ρ_{Req}) as a weight for his evaluation (RE), since we believe an evaluation from a requester with a high reputation score is more trustworthy than an evaluation from a low reputation requester.

This process is repeated for all participants at the end of each sensing campaign, and $Trust_{RP}$ is updated for all of them.

After every n campaigns, $Trust_{RP}$ values upon each active participant act as inputs for the reputation scheme, which updates the participant's reputation score accordingly.

While there are already different crowdsourcing applications of online reputation systems [65] such as eBay, Epinions²⁴ and Amazon,²⁵ they are not applicable in this

²⁴<http://www.epinions.com>

²⁵<http://www.amazon.com>

context. For example, eBay is based on a direct feedback method, where the buyer assigns either a positive (+1), negative (-1) or neutral (0) rating to the seller based on his satisfaction with the transaction. The member's reputation score is then simply obtained by calculating the difference between the number of unique positive and negative ratings received in the past 12 months. While easy to understand and implement, the reputation lag of 12 months makes this approach unsuitable in our context, as a malicious user may be able to contribute bad data over an extended period of time before being detected.

In our proposed reputation scheme, we use web page ranking algorithms as the basis for computing reputation scores. We draw parallels between the rank of a page in a set of web pages and the reputation score of a member in a social network. Moreover, the weights of links from different pages to a specific page are considered to be equivalent to the trust ratings of one member as determined by the other members of the social network.

Having a set of objects, a ranking algorithm calculates a relative importance of all objects in the set and makes an ordered list according to the importance. Web page ranking algorithms such as PageRank [42] calculate and assign a rank to a web page by analysing the web graph. Broadly speaking, PageRank ranks a page according to how many other pages are pointing at it. This can be described as a reputation system, because the collection of hyperlinks to a given page can be seen as public information that can be combined to derive a reputation score. A single hyperlink to a given web page can be seen as a trust rating of that web page.

In PageRank, the rank of page P , denoted by $\rho(P)$ is defined as:

$$\rho(P) = \frac{\sum_{P_i \rightarrow P} (\rho(P_i))}{L(P_i)} \quad (5.8)$$

in which, P_i is the set of all pages which have an outgoing link to page P , and $L(P_i)$ is the number of outgoing links from page P_i .

In the original PageRank algorithm, it is assumed that all the outgoing links

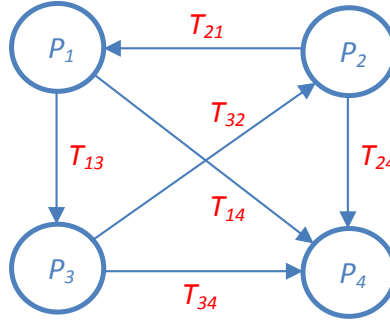


Figure 5.7: A sample social graph of 4 members with mutual trust ratings

have equal weights. This is not always true, since not all outgoing links from a web page are equally important. So, we adopted the extension offered in [148] which modifies the above equation as Equation 5.9,

$$\rho(P) = \sum_{P_i \rightarrow P} \frac{w_i}{\sum_{P_i \rightarrow P_j} w_j} \rho(P_i) \quad (5.9)$$

in which, w_i is the weight of the outgoing link, and the sum of weights of outgoing links is equal to 1.

We explain this further by presenting an illustrative example. Consider the graph in Figure 5.7 in which, P_1 , P_2 , P_3 and P_4 are the social network members. Links represent friendship relations with weights equal to the mutual trust between the pairs. In this case, according to Equation 5.9:

$$\rho(P_1) = T_{21} \times \rho(P_2)$$

$$\rho(P_2) = T_{32} \times \rho(P_3)$$

$$\rho(P_3) = T_{13} \times \rho(P_1)$$

$$\rho(P_4) = T_{14} \times \rho(P_1) + T_{24} \times \rho(P_2) + T_{34} \times \rho(P_3)$$

As can be seen in the above expressions, reputation calculation is an iterative process and continues until convergence is obtained. In our simulation in Section 5.4, we assume that the convergence occurs when $|\rho_k(P_i) - \rho_{k-1}(P_i)| \leq 10^{-10}$ for all P_i .

It should be noted that the reputation score calculation is independent of the campaign specification. As shown in Equation 5.9, the reputation score of each participant depends on the trust rating of the requester upon the participant ($Trust_{RP}$),

which itself, according to Equation 5.7, depends on two factors: (i) trustworthiness of contribution (ToC) and (ii) requester's evaluation (RE). All these parameters are related to the contribution itself and do not depend on the campaign.

The frequency with which reputations are updated is after every n campaigns. Updating the reputation scores in a regular manner results in more accurate perception of the suitability of participants and the accountability of requesters' evaluations. However, determining the optimum update interval of the reputation scores is challenging. One may argue that updating the reputation score after every campaign will allow the system to better reflect the behavioural changes of participants. However, as mentioned in Section 5.3, the reputation scheme utilises the PageRank algorithm to calculate and update reputation scores. The PageRank algorithm is a recursive algorithm wherein, the recursion continues until convergence happens. This can be quite time consuming, especially if the number of participants is large, which is typically the case in social networks. In other words, there is a trade-off between accurately reflecting the behavioural changes of the participants and the associated computational complexity of updating the reputation scores. The update period n , thus can be a system parameter that is configured by the system designer.

Participant Suggestion

When a participant ψ_i demonstrates reliable performance in multiple campaigns originated by the requester (Req), it would be beneficial if a direct relation is established between them, since: i) the time required for selecting ψ_i as a potential participant in further campaigns is reduced, (ii) less time and effort is consumed for task dissemination, since there is now only a one-hop distance to ψ_i , (iii) an easier access to ψ_i 's friendship network is now available.

To provide the requester with a suggestion list, the following process is carried out.

- For each participant ψ_i who is not an immediate friend of Req , a field called 'implicit trust' is kept. This field is initially set to zero and is updated whenever ψ_i contributes to a task originated from Req . The implicit trust update

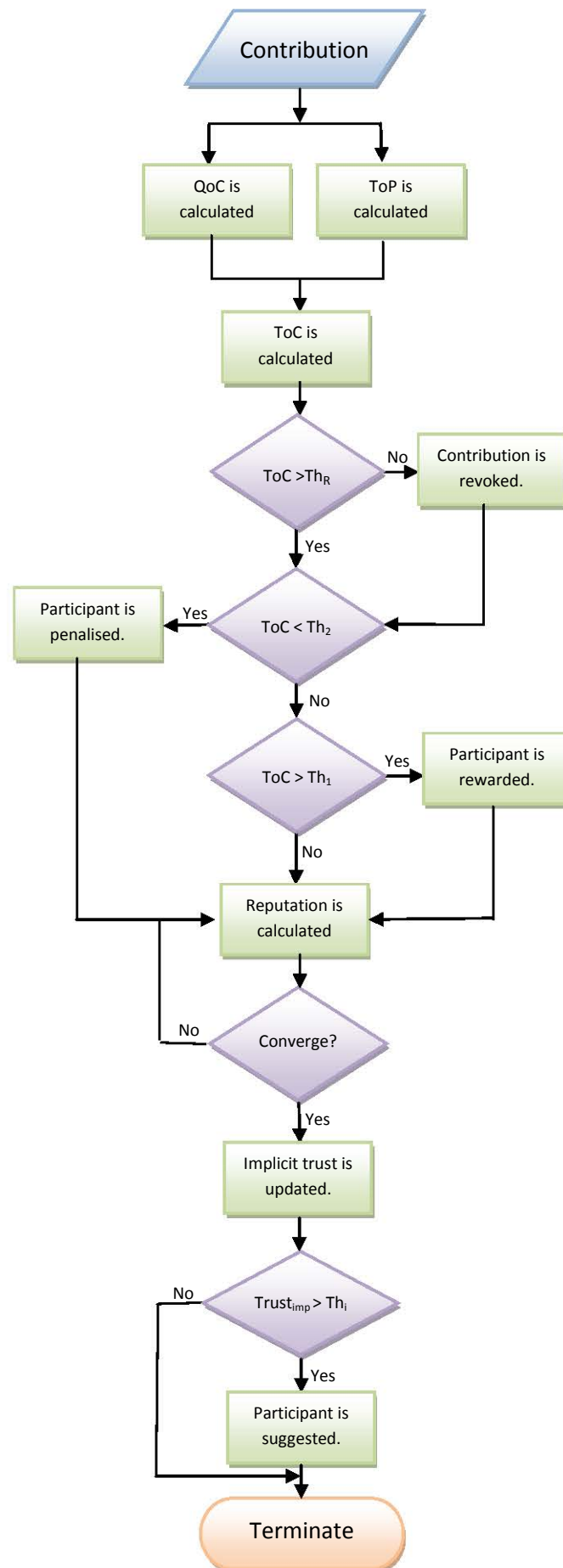


Figure 5.8: Sequence of steps in assessing the trust and reputation

process is the same as trust update performed in trust assessment scheme; i.e. it is increased by a constant amount $|ToC - \rho_{Req} * RE|$, if ψ_i provides a contribution with ToC greater than threshold1, and decreased by the same amount if ToC is less than threshold2.

- At certain intervals, implicit trust values are investigated to see whether ψ_i is eligible to be suggested for recruitment or friendship establishment. If above a threshold, ψ_i has such eligibility. In our simulation, we set this threshold to be 0.5.
- In the case of plenty of eligible participants, the best candidates are chosen from them. The best candidates are those participants who act as intermediate nodes in a larger number of paths. Adding such candidates as immediate friends will cause a considerable reduction in path lengths to other participants.
- The recruitment scheme is then provided with a suggestion list which consists of participants' IDs and their implicit trusts. The recruitment scheme is then able to utilise this list for further recruitment with initial trust value equal to implicit trust.

To summarise, once a campaign is launched, participants begin to send a series of contributions. For each contribution, the evaluator component computes a value for QoC and ToP. These values are fed to the fuzzy inference engine which calculates ToC for that contribution. The trust of the requester for each participant ($Trust_{RP}$) is updated according to his ToC. The server utilises $Trust_{RP}$ and ρ_{Req} to update the reputation score of each participant. At regular intervals, a list of suggested participants is prepared for the requester to be used in further campaigns. The sequence of steps is depicted in Figure 5.8.

5.4 Experimental Evaluation

This section presents simulation-based evaluation of the proposed trust scheme. We evaluate the performance of our proposed scheme via simulations since organising

real experiments in social participatory sensing is difficult. Moreover, simulations enable us to modify certain parameters and investigate the impact on the performance and outcome of the system. The simulation set-up is outlined in Section 5.4.1 and the results are in Section 5.4.2.

5.4.1 Simulation Set-up

In order to observe the performance of our proposed trust scheme, we considered two set-ups. In the first set-up (QoC vs. ToP set-up), we investigate the trade-offs between the ToP and QoC values. In the second set-up (ToP set-up), we study the trade-off between the personal and social trust factors of ToP. In the following, we explain the settings for each set-up in detail.

- **QoC vs. ToP Set-up**

We simulate an online social network where 50 members participate in 300 campaigns, producing one contribution for each. In the ideal case, for each contribution, we would have computed the value of each of the underlying parameters discussed in Section 5.2 based on some typical probabilistic distributions. However, this would digress from the primary objective of the evaluations, which is to evaluate if social trust is a useful contributor to the overall trust in social participatory sensing. For the sake of simplicity, we therefore, assign a random value of ToP to each participant and a random value of QoC for each contribution, both in the range of $[0, 1]$, based on criteria specific to the scenarios and leave extra investigation for the ToP set-up. In this set-up, we update ToP based on the quality of contributions. If below a specified threshold, the participant's trust will be decremented by α ; otherwise it will be incremented by β . Note that $\alpha > \beta$; since in typical social relations, trust in others is built up *gradually* after several trustworthy communications and torn down *rapidly* if dishonest behaviour is observed. We set α and β to 2 and 1, respectively.

Recall that, the goal of the trust scheme is to assign a trust rating to each contribution which is further used as a criterion to accept/reject the contribution.

As such, in the evaluations, we artificially create circumstances in which, some participants contribute poor quality data for a certain number of campaigns. We want to investigate if our trust scheme is able to identify this behaviour and revoke untrusted contributions in a robust manner. In order to create all possible combinations of QoC and ToP, we assume that participants belong to one of the following four categories, each of which resembles one type of friend in a typical social participatory sensing system:

Category 1: Participants with high ToP ($\text{ToP} \geq 0.5$) and high QoC ($\text{QoC} \geq 0.5$).

Category 2: Participants with low ToP ($\text{ToP} < 0.5$) but high QoC ($\text{QoC} \geq 0.5$).

Category 3: Participants with high ToP ($\text{ToP} \geq 0.5$) but low QoC ($\text{QoC} < 0.5$).

Category 4: Participants with low ToP ($\text{ToP} < 0.5$) and low QoC ($\text{QoC} < 0.5$).

The threshold 0.5 used above for a trustworthy participant/contribution has also been used previously in [15, 149]. Friends that belong to Category 1 would generally be more willing to volunteer and contribute data. As such, we assume that Category 1 contains more participants (20) than the other three categories, which contain 10 participants each. In the first scenario, we assume that participants do not alter their behaviour and thus QoCs follow the category settings throughout the entire simulation. In the second scenario, we assume that participants can transition from one category to another (details in Section 5.4.2).

- **ToP Set-up**

We simulate an online social network where 100 members participate in 5000 campaigns, producing one contribution for each. We assume that each member is connected to all others, similar to a social community. So, there are a total of 10000 friendship relations. All members can serve both as requesters who launch sensing campaigns and as participants who contribute data to these sensing campaigns.

In the previous set-up, we categorised participants according to the trade-offs between ToP and QoC. Our goal was to observe how accurately the system assigns trust ratings to contributions in the case of different ToP and QoC levels. Moreover, we artificially created scenarios where participants begin producing contributions with low QoC, which results in a decrease in ToC. We wanted to see if the system is able to quickly detect this transition and revoke low trustable contributions in an accurate and robust manner.

Here, instead of categorising the participants according to ToP and QoC, we designed the categories according to the trade-offs between personal factors and social factors within ToP, and simply assumed that QoC has a value in the range of $(ToP - \mu, ToP + \mu)$. This will allow us to observe how the system reacts to behavioural changes of participants and observe if it is successful in updating the reputation scores in case of such fluctuations. As mentioned in Section 5.2, ToP parameters can be divided into two groups: social factors which include friendship duration and interaction time gap, and personal factors which include timeliness, expertise and locality. In the real-world, there are often situations where a friend with a high rating of social factors (i.e. a very close friend with whom one has repeated interactions) has a low rating for personal factors for a period of time (i.e. does not have related expertise or does not produce timely contributions). In other words, we may have participants who have high social trust, but low personal trust, and vice versa. We thus define 4 different states based on the combination of different levels of personal and social trusts.

Specifically, we assumed that 60 members (out of 100) belong to Category A whereas the remaining 40 belong to Category B, adding the assumption that category A members have high personal trust, while category B members have low personal trust. We also assume that for each member P_A in category A, all other members score P_A with high social trust, and for each member P_B in category B, all other members score P_B with low social trust.

When P_A serves as requester, other members form two subcategories:

A-1: which includes 59 members from category A, excluding P_A . They have

high personal trust and score P_A with high social trust.

A-2: which includes 40 members from category B. They have low personal trust and score P_A with high social trust.

Similarly, when P_B serves as requester, other members form two subcategories:

B-1: which includes 60 members from category A. They have high personal trust and score P_B with low social trust.

B-2: which includes 39 workers from category B, excluding P_B . They have low personal trust and score P_B with low social trust.

It is natural that not all friends in a social network would contribute data to sensing campaigns. As such, we assume that 10% of the members in category A and 50% of the members in category B do not upload any data. The rationale for assuming unequal percentages is that the first group constitutes close friends and hence a higher percentage would be willing to contribute. The second group includes those who have low social connectivity and so, have less willingness to contribute.

Whenever a task is launched, one of the participants is selected to be the requester. Without loss of generality we assume that tasks are launched in sequential order by the social network members, i.e. member 1 launches the first campaign, member 2 launches the second campaign and so on.

ToP Parameter Settings

In the following, we discuss the initialisation of the various parameters introduced in Section 5.2. For a set of parameters, the assumptions and parameter settings are based on the statistics and results that have been proposed in other experiments. For the rest, we set the parameters in a way that allows us to configure user groups with different behavioural traits

In order to set the expertise value for a participant, we assume that there are a total of six expertise areas defined and that each task needs at most three expertise areas ($n(TE) = 3$). To calculate the expertise score for each participant, we assign a value to $n(PE)$ based on his category, as shown in Table 5.2. The expertise score E is then calculated using Equation 5.1.

For timeliness, we first set the response time (rt) for each participant. In order to initialise the response time (rt), we used the statistics presented in [141]. In this paper, the authors performed a real participatory sensing experiment and found that with the deadline of 1 day, 40% of the tasks were solved within the first 3 hours, 70% within 15 hours, and 90% within 20 hours. We have used the general trend from their observations to set the timeliness parameter in our simulation settings, where the deadline is 1 week. As seen in Table 5.2, for a participant P_A belonging to category A, with probability of 0.4, rt is at most one day, with the probability = 0.65, rt is at most half of a week, and with the probability of 0.9, rt is at most one week (Note that the greatest probability is 0.9, since with the probability of 0.1 (10%), P_A does not attend in sensing campaign). Response time then acts as the input value for Equation 5.2 which results in timeliness score T for participant. Other input parameters for Equation 5.2 have been set as $x = 0.3$, $b = 6$, $c = 0.6$, and $d = 7$ days.

For locality, we assume that there are a total of 25 regions and that each participant is local to three regions (i.e. locality score L for these three regions is 1). We also assume that when a participant has the maximum locality score in a region, this participant has a relatively high locality in its surrounding regions. We assume there are three surrounding regions N_1 , N_2 and N_3 , each representing a level of neighbourhood. Based on the participant's category, locality score L is assigned to each surrounding region, as shown in Table 5.2.

For friendship duration, as mentioned in Section 5.2, the input parameter (t) is the time (in years) elapsed since the beginning of the friendship establishment. The initial value of t is set according to the participant's category, as shown in Table 5.2 and a constant value of 0.02 is added to t after each participation. The friendship time (t) thus serves as the input value for Equation 5.4 which computes the friendship duration score F for the participant. Other input parameters for Equation 5.4 have been set as $b = 5$ and $c = 1$.

Finally, for the interaction time gap, as mentioned in Section 5.2, the input parameter t is the gap (in days) between the current time and the latest interaction (LI) time. We set LI based on the category of each participant, as

shown in Table 5.2, and calculate t accordingly. Gap time, t is then fed to Equation 5.5 which calculates the interaction time gap score I for the participant. Other input parameters for Equation 5.5 have been set as $b = 10$ and $c = 0.2$.

Once all of the aforementioned parameters are computed, ToP is calculated by simply averaging them. In other words, we simply assume that $w_i = 1/5$ in Equation 5.6. QoC is then assigned a value in a range of $(ToP - \mu, ToP + \mu)$ with $\mu = 0.1$.

ToC is then calculated and $Trust_{RP}$ is updated according to Equation 5.7. At intervals, reputation score is also updated for participants. We set the reputation interval to be after every 100 campaigns ($n = 100$).

In the first scenario, we assume that ToPs follow the category settings throughout the entire simulation. In the second scenario, we assume that ToP parameters change for a group of participants which results in a transition from one category to another (details in Section 5.4.2).

• Compared Methods and Evaluation Criteria

For both set-ups, we compare the performance of our scheme against a baseline system presented in [15], which does not consider the social trust as an effective factor in the trustworthiness of contribution. Then, in order to study the effect of other trust aspects, we incrementally add them to the baseline to see how considering each aspect influences trust. Specifically, we compare the following:

- Baseline-Rep: which follows the approach in [15] by calculating a reputation score for each participant according to the QoC of his successive contributions. This reputation score is used as a weight for QoC. In other words, $ToC = \sqrt{Rep * QoC}$
- Average: which includes ToP but computes the ToC simply as an average of ToP and QoC
- Fuzzy: our proposed scheme.

Table 5.2: Parameter settings for the calculation of trust of participant (ToP)

category A		category B	
param	value	param	value
n(PE)	4	n(PE)	2
rt	$\begin{cases} (0, 1] & prob = 0.4 \\ (1, 7/2] & prob = 0.65 \\ (7/2, 7] & prob = 0.9 \end{cases}$	rt	$\begin{cases} (0, 1] & prob = 0.1 \\ (1, 7/2] & prob = 0.3 \\ (7/2, 7] & prob = 0.5 \end{cases}$
N_1	random(0,1)	N_1	random(0,0.5)
N_2	random(0,0.9)	N_2	0
N_3	random(0,0.8)	N_3	0
t	rand[4,5]	t	rand[0,1]
LI	$\begin{cases} (0, d] & prob = 0.8 \\ 0 & prob = 0.2 \end{cases}$	LI	$\begin{cases} (0, d] & prob = 0.2 \\ 0 & prob = 0.8 \end{cases}$

As mentioned in Section 5.1, a ToC rating is calculated for each contribution and those contributions with ToC lower than a predefined threshold (Th_R) are revoked from further calculations. The ToCs for the non-revoked contributions are then combined to form an overall trust for that campaign. In other words,

$$OverallTrust = \frac{\sum_{i=1}^n ToC}{n} \quad (5.10)$$

in which, n is the number of non-revoked contributions. The revocation threshold is set to 0.5 ($Th_R = 0.5$). We consider the overall trust as the evaluation metric. The greater the overall trust the better the ability of the system to revoke untrusted contributions. Overall trust has a value in the range of $[0, 1]$. We also calculate the reputation scores for all participants to see whether they reflect the behaviour of participants in normal and transition settings. Reputation score value is a number in the range of $[0, 1]$ with initial value of 0.5 for each participant.

5.4.2 Simulation Results

In this section, we first present the simulation results for the first set-up. Following this, we elaborate on the findings for the second set-up.

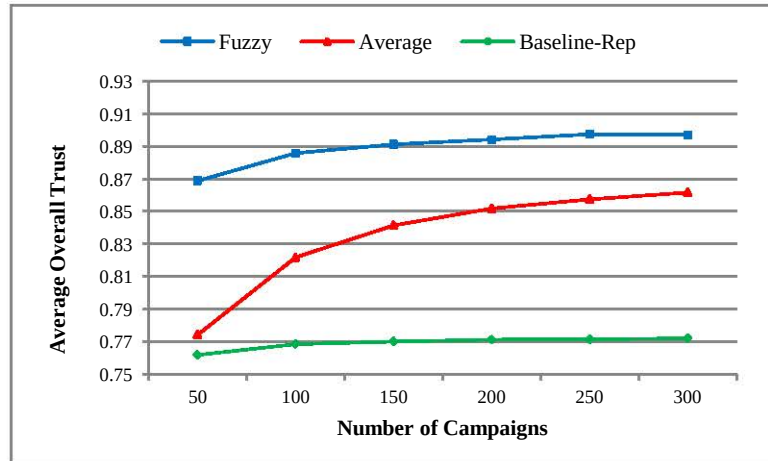


Figure 5.9: Evolution of average overall trust for all methods, Scenario 1

QoC vs. ToP Set-up Results

At first, we present results for the first scenario. Recall that in this scenario, participants do not alter their behaviour and thus QoCs follow the category settings. Figure 5.9 depicts the evolution of the average overall trust as a function of the number of campaigns. As shown in this figure, the Baseline-Rep method remains flat throughout the simulation. The Baseline-Rep method is totally based on QoC values and QoC values are purely random. This flat behaviour can be explained by the law of large numbers [150]. The law of large numbers describes the result of performing the same experiment a large number of times. According to this law, the average of the results obtained from a large number of trials will tend to become closer as more trials are performed. In fact, the law of large numbers guarantees stable long-term results for the averages of random values, which applies to averaging the random values of QoC in calculating the overall trust. Also, as this figure shows, we encounter a growth in the average method. The reason is that in this method, we increase ToP gradually by the amount of β , which inherently results in higher values of ToC. Moreover, we revoke the contributions with low ToC. Since the number of participants with high values of QoC and ToP (i.e. category 1 participants) is greater than other categories, we normally have a greater number of participants with monotonically increasing ToC values, which leads to the gradual increase in overall trust score. However, since this method computes the average QoC and ToP, the maximum value of ToC is less than the one obtained from the

Participant ID	0	8	13	9	1	14		41	28	44	45
ToP	1	1	1	0.98	0.95	0.93		0.2	0.16	0.13	0.05
Contribution ID	450	451	458	457	466	470		495	494	490	498
ToC	96.6	96.6	96.6	96.4	95.7	87.2		0.25	0.24	0.23	0.086

Figure 5.10: Ranked lists provided by trust scheme for the requester, Scenario 1

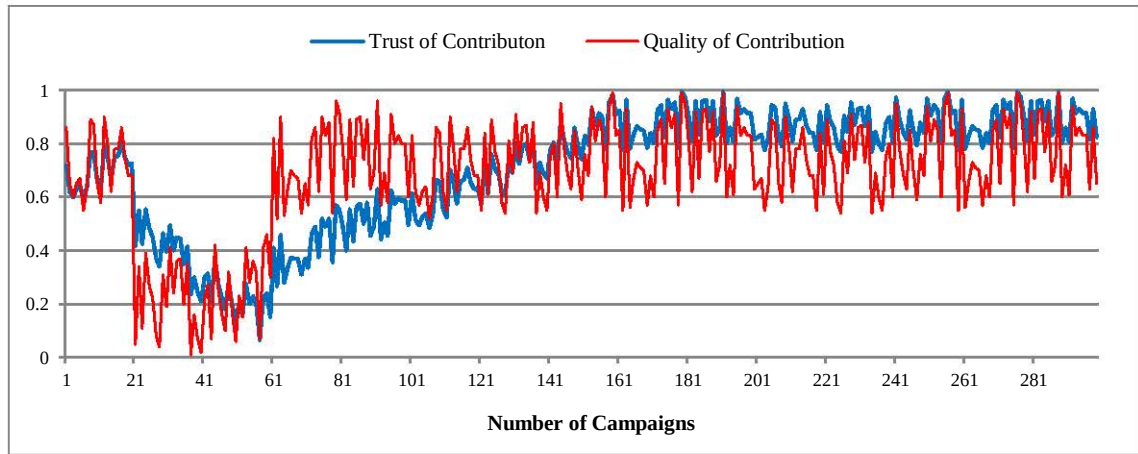


Figure 5.11: Evolution of quality of contribution (QoC) & trust of contribution (ToC) for one participant, Scenario 2

fuzzy method (regarding the rules no. 15 and 16 in Table. 5.1). So, although the average method approaches the fuzzy method, they would not converge. To sum up, our fuzzy trust method outperforms all the other methods.

Figure 5.10 depicts two ordered lists provided by the trust assessment scheme. The first list sorts the participants in a descending order of their ToPs. This can be used as a suggestion list for the data consumer for future recruitment of participants. The second list provides an ordered list of contributions according to the descending order of ToCs, which can help the data consumer to select the most trustable contributions based on a certain configurable threshold.

Now, we present results for the second scenario, wherein, the behaviour of the participants can change with time, which may result in a transition from one category to another. This scenario allows us to observe the performance of our method

in comparison with methods in the presence of noise. For example, consider a participant who is initially highly trusted and provides high quality data and thus belongs to category 1. After some time, this participant contributes low quality data for some campaigns. This may be because of incorrectly operating his mobile device for the purpose of the sensing task (e.g., capturing unfocussed pictures). In this scenario, we assume that 15 participants transition from category 1 to category 3. In other words, the total population of the 4 categories changes from (20, 10, 10, 10) to (5, 10, 25, 10). The transitionary period lasts from the 20th to 60th campaign. Following this, the 15 participants transition back to category 1 and we return to the initial population distribution.

Figure 5.11 demonstrates the evolution of QoC and ToC for one participant. As this figure shows, there is a sharp decrease for ToC and ToP after the 20th campaign and a gradual increase after the 60th campaign contribution, which resembles the real trust destruction and construction process. Note that in the period between the execution of 20th and 60th campaigns, ToC is even lower than QoC. This is related to the first rule in the fuzzy rule base, which defines a very low value for ToC in case of low values for QoC and ToP.

Figure 5.12 shows the evolution of overall trust as a function of the number of campaigns in the Average and Fuzzy methods (the baseline method is excluded, since we want to compare ToP related methods). There is a decrease in overall trust for both methods in the transition period, due to an increase in the number of category 3 participants, who produce low quality contributions. However, the fuzzy method is more robust at limiting the effect of these bad contributions and still achieves an acceptable level of trust. This is due to the correct adjustment of the fuzzy rules such as rule no. 6 in Table. 5.1 which assigns a low trust rating to low quality contributions, and which leads to their revocation. As can be seen in this figure, there is a small decrease in overall trust after the transitionary period. The reason is that when participants transition to category 3, they begin providing low quality contributions, which in turn, results in low ToP for them (Recall that ToP is updated according to QoC.). By transitioning back to category 1, they resume providing high quality contributions. But since ToP is still low, the obtained ToC is

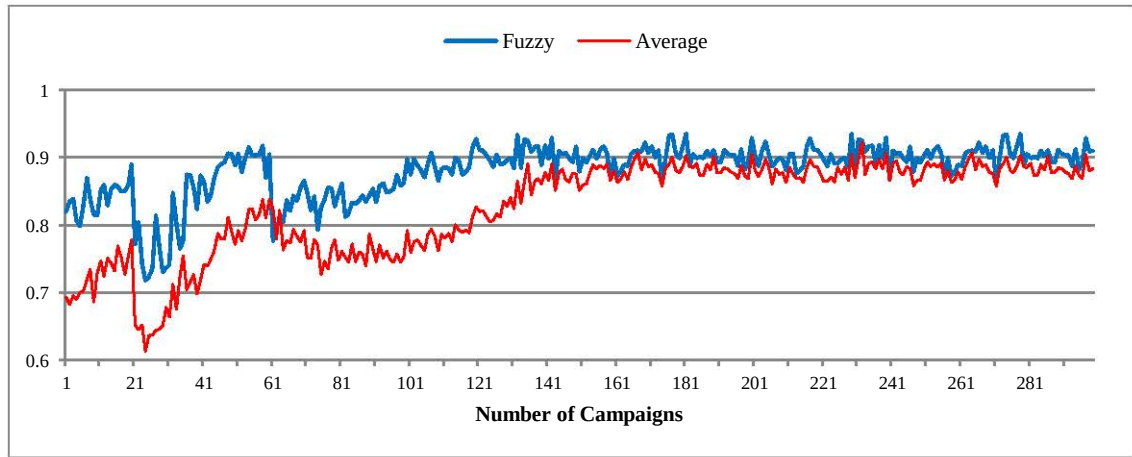


Figure 5.12: Overall trust obtained in Fuzzy and Average methods in Scenario 2

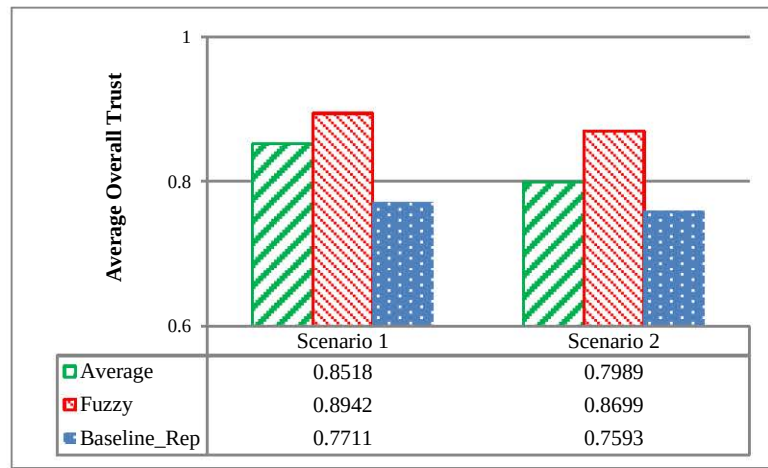


Figure 5.13: Comparison of average overall trust for all methods for both scenarios

a value that is lower than before, but greater than revocation threshold. So, these contributions are not revoked and considered in the overall trust calculation, which makes the aforementioned decrease.

Figure 5.13 presents a summary of the results for both scenarios, averaged over 300 campaigns. Observe that the proposed fuzzy trust scheme outperforms all other schemes in both scenarios. In particular, our scheme demonstrates high robustness to noisy contributions (scenario 2), as compared to the other schemes under consideration.

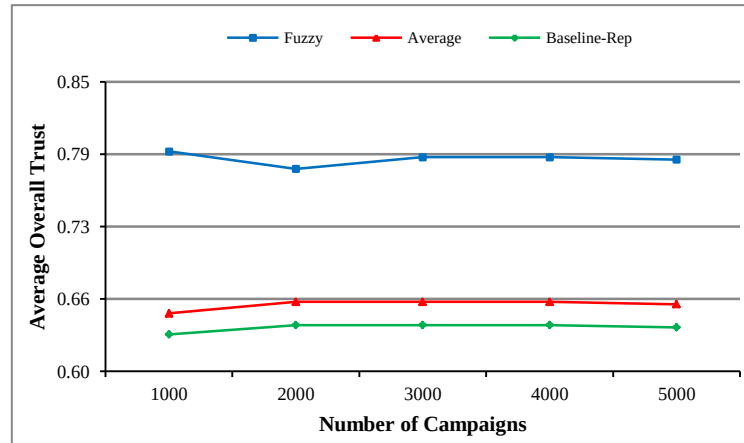


Figure 5.14: Evolution of average overall trust for all methods, Scenario 1

ToP Set-up Results

Now, the results for the second set-up are discussed in which, we study the trade-off between the personal and social factors of ToP and simply assume that QoC has a value in the range of ToP. We first present the simulation results for the first scenario. Recall that in the first scenario, we assume that ToPs follow the category settings throughout the entire simulation. Figure 5.14 depicts the evolution of the average overall trust as a function of the number of campaigns. As shown in the figure, our fuzzy trust method outperforms the other methods. This confirms its success in mimicking the human trust establishment process by correctly setting fuzzy rules. In particular, we have set the rules in a way that results in early detection and severe punishment of untrusted contributions and also put greater emphasis on highly trusted contributions. The former has been done by assigning a very low (VL) value to ToC in case of low ToP and QoC (i.e. Rule no. 1 in Table 5.1), whereas the latter has been obtained through assigning a very high (VH) value to ToC in the case of high QoC and above average ToP (i.e. Rule no. 15 and 16 in Table 5.1).

Figure 5.15 depicts the evolution of overall trust for 1000 contributions with Fuzzy method. As can be seen in this figure where, at each interval containing 100 contributions, two different levels of overall trust are achieved. Recall that the order of requesters is equal to the members' order. As observed in this figure, a higher level of overall trust is obtained when the requester is from category A. So,

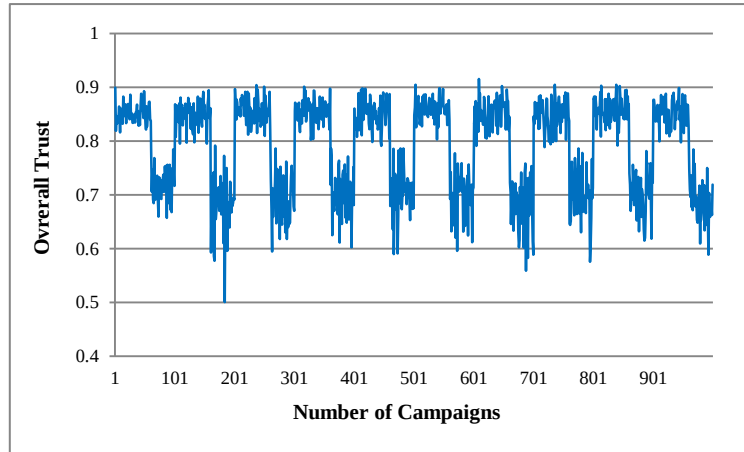


Figure 5.15: Evolution of overall trust, Fuzzy method

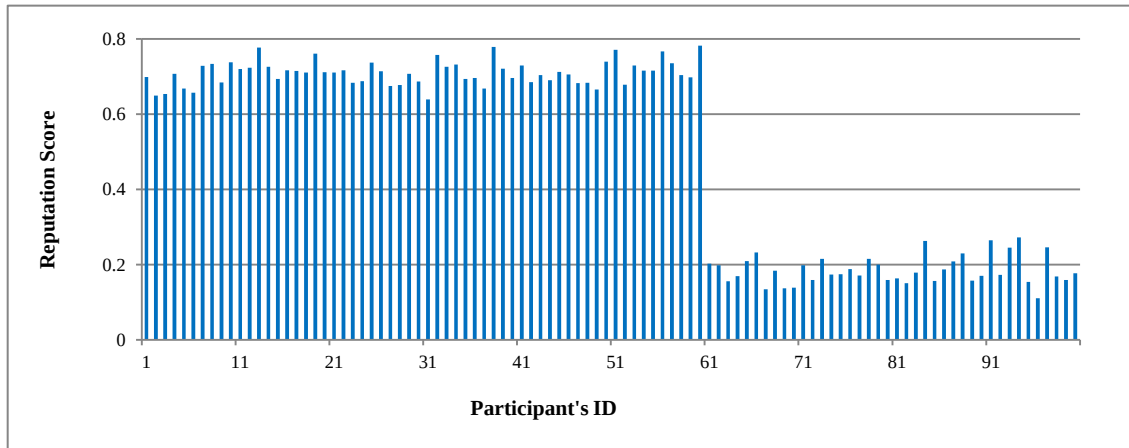


Figure 5.16: Reputation score for all 100 members after attending 5000 tasks, Fuzzy method, Scenario 1

participants are located either in subcategory A-1 or A-2. This will result in either high ToC values (when participants are from category A-1) or medium ToC values (when participants are from category A-2), which in turn, results in high overall trust. Similarly, lower level of overall trust is obtained when the requester is from category B. So workers are located either in category B-1 or B-2. This will lead to either medium ToC values (when participants are from category B-1) or low ToC values (when participants are from category B-2), which results in low overall trust. This variation is repeated regularly at each interval of 100 contributions.

Figure 5.16 presents the reputation of 100 participants after attending 5000 sensing campaigns. As mentioned before, the initial value of reputation score for all participants is 0.5. Category A participants who have high ToPs, produce contributions

with high ToC and hence, they get rewarded. This reward results in an increase in $trust_{RP}$ for them, which in turn, increases their reputation score. On the contrary, for category B participants with low ToPs, ToCs will also be low, and hence, they are penalised, which results in the reduction of their reputation score. To summarise, our system continually tracks the contributions made over a series of campaigns and detects participants' behaviour, which is accurately reflected in the evolution of the reputation scores.

Next, we present results for the second scenario, wherein, the behaviour of participants change for a period of time, which results in a transition from one category to another. This scenario allows us to observe the performance of the schemes in the presence of noise. For example, consider a participant P_A who is in category A, changes his behaviour for a period of time and behaves in a different manner which results in a decrease of his personal and (hence) social trust. For example P_A no longer provides timely contributions or does not care enough about the requirements of the task. This behavioural change results in a decrease in his personal trust, and consequently, others score him low with social trust. In other words, a participant may encounter a transition from category A to category B. In this scenario, we assume that 10 from 60 participants of category A transition to category B (e.g., a reduction in their personal and social factor values is created) in the period between 1000th and 4000th campaigns.

Figure 5.17 shows the reputation score of 100 participants at the end of transition period (i.e. after attending 4000 campaigns). As can be seen in this figure, the reputation of the first ten participants who encounter such a transition has a considerable decrease in comparison with others not encountering such transition. This again demonstrates the ability of our reputation scheme to adjust the reputation scores as a reflection of behavioural changes of participants.

Finally, Figure 5.18 shows the reputation score evolution of participant no. 9 encountering such transition between the 10th and 40th reputation intervals (between 1000th and 4000th campaigns). As can be observed, our proposed method shows an explicit and considerable reaction to this behavioural change, as compared with other methods. There is a decrease in reputation score due to dishonest behaviour

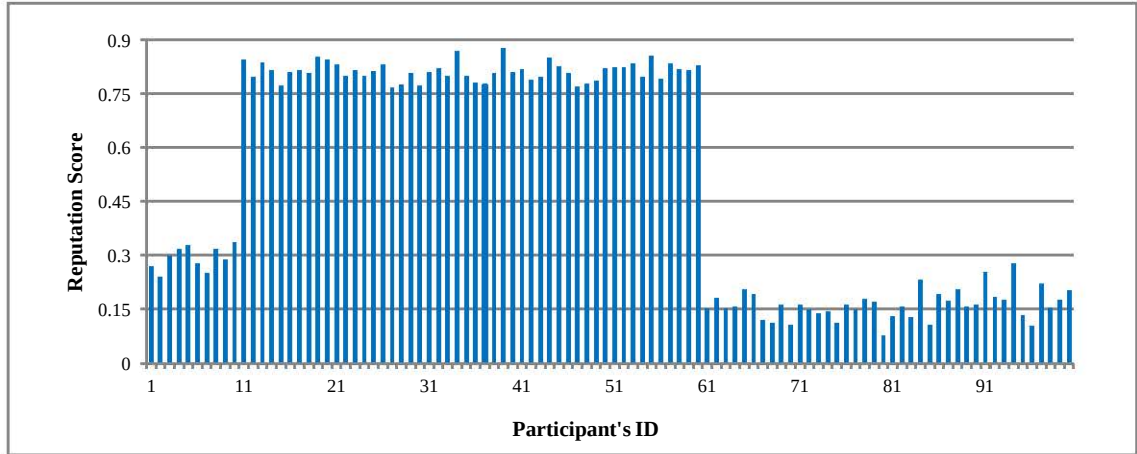


Figure 5.17: Reputation score for all 100 members at 4000th campaign, Fuzzy method, Scenario 2

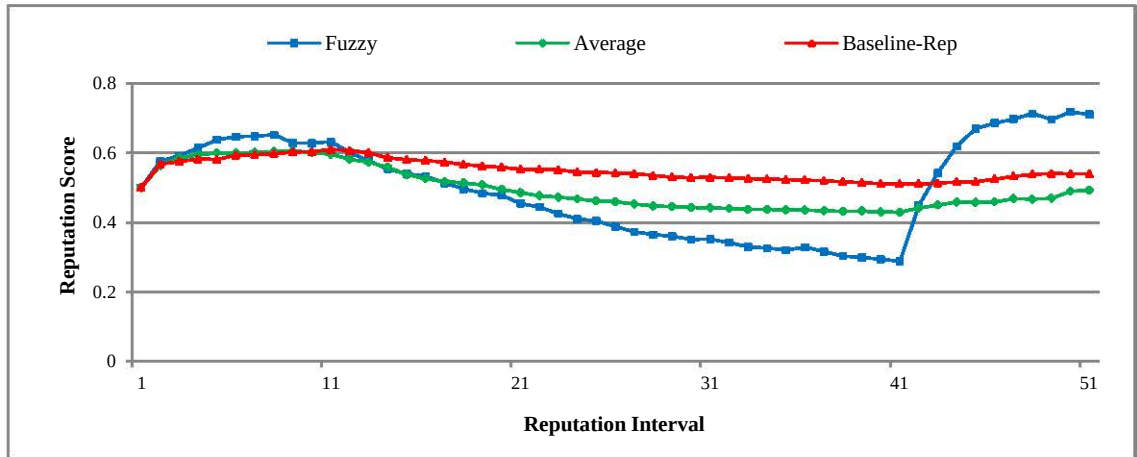


Figure 5.18: Evolution of reputation score for participant no.9 in all methods, Scenario 2

during the transition period. At the end of this transition period, the transition encountered participant resumes his normal behaviour which results in a considerable increase in his reputation score.

To summarise, in this section we aimed at investigating the performance of our proposed trust scheme. We considered two different set-ups. In the first set-up, we assumed that participants had different levels of ToP and QoC. We considered a scenario in which, the participant's contributing behaviour is changed during a limited period of time. The simulation results showed that our proposed trust scheme is able to accurately evaluate and assign a trust score to each contribution. It is

also successful in achieving high overall trust in comparison with Baseline-Rep and average methods, and revoking unreliable contributions. In the second set-up, we assumed participants had different levels of personal and social trust. The results demonstrate that the trust scheme is able to accurately detect the behaviour fluctuation of participants and show a proper reaction accordingly. It is also successful in adjusting the participants' reputation scores as a reflection of their behavioural changes.

5.5 Conclusion

In this chapter, we proposed an application agnostic trust scheme to address the issue of trust in social participatory sensing. We assumed that a sufficient number of suitable participants was recruited via the utilisation of the recruitment and the participant selection schemes, and their contributions were already at hand. Our novel trust scheme was aimed at evaluating the trustworthiness of these contributions in an accurate and comprehensive manner. Specifically, the trust assessment was carried out by evaluating the quality of the data and the trustworthiness of the participants. These two influential factors were then combined via a fuzzy inference engine to arrive at a comprehensive trust rating for each contribution. The reputation management scheme was then employed to assign a reputation score to participants by leveraging the concepts utilised in the PageRank algorithm.

We undertook extensive simulations to demonstrate the effectiveness of our trust and reputation management schemes and benchmark them against the state-of-the-art methods. The results demonstrated that by considering social relations in trust evaluation, our proposed methodologies could achieve realistic outcomes that were consistent with how human beings establish trusted social communications. We also showed that our proposed trust scheme was able to quickly adapt to rapid changes in the participant's behaviour (transitioning from high to low quality contributions) by fast and correct detection and revocation of unreliable contributions. Moreover, we found that leveraging fuzzy logic provides considerable flexibility in combining the underlying components, leading to a better assessment of the trustworthiness of

contributions. Our proposed assessment method resulted in a considerable increase in the overall trust by over 15% as compared to a method which solely associated trust based on the quality of contribution.

Chapter 6

Conclusion and Future Work

We close this dissertation with a summary of contributions made in this thesis and discuss several remaining challenges for future work.

6.1 Concluding Remarks

Encouraging well-suited individuals to contribute to participatory sensing applications is an important challenge to their success. One potential solution is to integrate online social networks with participatory sensing applications resulting in the emergence of social participatory sensing. For social participatory sensing to be a success, multiple challenges need to be addressed, namely, the sufficiency of participation, the suitability of participants and the assessment of trust. Within the scope of this thesis, we have comprehensively addressed these challenges by proposing a framework comprising three key components: (i) a recruitment scheme to address the participation sufficiency issue, (ii) a participant selection scheme to address the suitability issue and (iii) a trust scheme to address the trust assessment issue.

Chapter 3 addressed the participation sufficiency challenge by presenting a novel recruitment scheme. Our proposed recruitment scheme worked in concert with the participant selection scheme (presented in Chapter 4) to identify and recruit sufficient number of well-suited participants via the most credible paths. In particular, the proposed recruitment scheme traversed the requester’s social graph and lever-

aged multi-hop friendship relations to identify the most credible paths to well-suited participants. The credibility of a communication path is dependent on two aspects. The first aspect is the trustworthiness of the path to guarantee the integrity of the exchanged messages. The second aspect is the privacy of the communication path to provide a secure medium for the exchange of communication messages containing sensitive information. The most credible paths were then used for the recruitment of well-suited participants. In addition to credibility-based path selection, we also proposed a customised random surfer to efficiently select the communication path. Specifically, instead of investigating all the possible paths, the proposed random surfer selected the next intermediate node in the path (almost) randomly. The proposed random surfer also addressed the bootstrapping problem for new members by allowing them to be selected with a probability similar to more reputable participants. We validated our proposed scheme by performing extensive simulations and showed that our scheme successfully resolves the participant sufficiency issue by achieving an 85% participation score.

Chapter 4 addressed the challenge of identifying suitable participants. Our proposed participant selection scheme was aimed to assess the suitability of the participants who were identified by the recruitment scheme (described in Chapter 3). Specifically, the participant selection scheme first assessed the participant's initial suitability according to a set of parameters such as his reputation, expertise level and pairwise privacy. Then, an eligibility assessment technique was executed which evaluated the eligibility of the participant based on his timeliness as well as the task deadline, to ensure that he would provide timely contributions. We also proposed a collusion prevention scheme to prevent the selection of colluders as well-suited participants. The proposed scheme evaluated the possibility of collusion for each eligible participant according to a set of indicators that are related to the common approaches utilised by colluders to arrange a collusive attack. If the participant was not detected as collusive, he would then be identified as a suitable participant to contribute to the task. We demonstrated the robustness of our proposed participant selection scheme through extensive simulations and observed that our scheme was able to detect the well-suited participants with an average suitability score twice

that of comparable methods. It was also successful in accurately detecting 83% of potential opportunities for the formation of collusive groups.

Chapter 5 was intended to address the issue of trust by presenting the details of our proposed trust scheme. Our trust scheme aimed to evaluate the trustworthiness of contributions reported by the selected participants (via the co-operation of recruitment and participant selection schemes). Specifically, for each received contribution, the trust assessment scheme first evaluated the quality of the contribution and the trustworthiness of the participant, and then combined these two parameters via a fuzzy inference engine to arrive at a comprehensive trust rating for the contribution. The reputation management scheme was then executed to evaluate and assign a reputation score to participants. We conducted extensive simulations to demonstrate the effectiveness of our trust assessment and reputation management schemes, and benchmarked them against the state-of-the-art methods in use. The results demonstrated that by considering social relations in trust evaluation, our proposed scheme could achieve realistic outcomes that were consistent with the establishment of trustful social communications by individuals. Our proposed trust assessment scheme also resulted in a considerable increase in the overall trust to 0.9, which is 15% higher than what can be achieved by a method which solely associated trust based on the quality of contribution. The proposed reputation management scheme was also shown to assign the reputation scores to participants in an effective and accurate manner.

6.2 Future Directions

The contributions presented in this dissertation lay the foundations for addressing the important challenges of social participatory sensing by presenting a series of schemes and methodologies. In order to provide a comprehensive and realistic view of these solutions, several new research questions and challenges have come to light:

- **Comprehensive Privacy Preservation.** In this dissertation, we aimed to address the privacy issues in message exchange as well as participant selection

processes. However, we believe that beyond the approaches proposed in this thesis, innovative methods still need to be proposed to efficiently and comprehensively address the privacy issues in social participatory sensing. Hence, a thorough review of privacy attacks and their existing solutions, both from the perspective of information privacy as well as communication privacy, is needed to be carried out. In particular, activities such as tagging a friend in a picture taken from the task location or checking in to the sensing location may potentially expose the participant to privacy violations. Moreover, user-controlled privacy settings in social networks, while valuable in terms of flexibility, may inadvertently lead to privacy risks such as location disclosure.

- **Enhanced Collusion Detection.** We presented a collusion prevention scheme with the aim of detecting and preventing the potential collusion via a set of behavioural and content-based indicators. The assumption with these indicators is that the contributed data is in numeric format. In other words, we assumed that we are dealing with scalar sensor readings. This assumption, however, does not apply for multi-dimensional sensor readings. So, enhancing the collusion indicators in a way that they become applicable for all types of sensing data can be regarded as a desirable future work.
- **Integrating with Incentive Mechanisms.** We explained in Chapter 2 that utilising incentives mechanisms is beneficial in encouraging the users to actively participate in the sensing campaign. In this dissertation, we claimed that in social participatory sensing, the sense of community and efficacy, regarded as two effective motivations, are well satisfied. However, as a future avenue of research, social participatory sensing applications could be integrated with existing incentive mechanisms in participatory sensing to better motivate the social network members to obtain rewards by participating in tasks that originated from their friends.
- **Empirical Experiments.** As mentioned earlier, we evaluated the performance of all the contributions in this thesis using simulations. This was due to the lack of available social participatory sensing systems to provide any

possibility of evaluating the performance of our developed schemes in real-world applications. Real-world deployment of social participatory applications would allow us to investigate and refine the design of the proposed schemes and methodologies and tailor them to the real needs of users, once having tested them under real-world conditions. It would also enable us to gather insights from a technical perspective about the performance of the developed schemes and fine-tune their design to best fit the experienced conditions. While there already exists some similar work that conducted real trials on Twitter (described in Chapter 2), deploying a real-world social participatory sensing application on large-scale social networks such as Facebook could be beneficial.

Bibliography

- [1] J. A. Burke, D. Estrin, M. Hansen, A. Parker, N. Ramanathan, S. Reddy, and M. B. Srivastava, “Participatory sensing,” in *Proceedings of the 1st Workshop on World-Sensor-Web (WSW)*, pp. 1–5, 2006.
- [2] Y. F. Dong, S. S. Kanhere, C. T. Chou, and R. P. Liu, “Automatic image capturing and processing for petrolwatch,” in *Proceedings of the 17th IEEE International Conference on Networks (ICON)*, pp. 236–240, 2011.
- [3] L. Deng and L. P. Cox, “Livecompare: Grocery bargain hunting through participatory sensing,” in *Proceedings of the 10th ACM Workshop on Mobile Computing Systems and Applications (HotMobile)*, pp. 1–6, 2009.
- [4] N. Maisonneuve, M. Stevens, M. E. Niessen, and L. Steels, “Noisetube: Measuring and mapping noise pollution with mobile phones,” in *Proceedings of the 4th International Symposium on Information Technologies in Environmental Engineering (ITEE)*, pp. 215–228, 2009.
- [5] R. K. Rana, C. T. Chou, S. S. Kanhere, N. Bulusu, and W. Hu, “Ear-phone: an end-to-end participatory urban noise mapping system,” in *Proceedings of the 9th ACM/IEEE International Conference on Information Processing in Sensor Networks (IPSN)*, pp. 105–116, 2010.
- [6] I. Schweizer, R. Bärtl, A. Schulz, F. Probst, and M. Mühläuser, “Noisemap-real-time participatory noise maps,” in *Proceedings of the 2nd International Workshop on Sensing Applications on Mobile Phones (PhoneSense)*, pp. 1–5, 2011.

- [7] P. Johnson, A. Kapadia, D. Kotz, N. Triandopoulos, and N. Hanover, "People-centric urban sensing: Security challenges for the new paradigm," tech. rep., TR2007-586, Dartmouth College, Computer Science, Hanover, NH, 2007.
- [8] K. Shilton, "Four billion little brothers?: Privacy, mobile phones, and ubiquitous data collection," *Communications of the ACM*, vol. 52, no. 11, pp. 48–53, 2009.
- [9] D. Christin, A. Reinhardt, S. S. Kanhere, and M. Hollick, "A survey on privacy in mobile participatory sensing applications," *Journal of Systems and Software*, vol. 84, no. 11, pp. 1928–1946, 2011.
- [10] S. Reddy, D. Estrin, and M. Srivastava, "Recruitment framework for participatory sensing data collections," in *Pervasive Computing*, vol. 6030 of *Lecture Notes in Computer Science*, pp. 138–155, 2010.
- [11] I. Krontiris and F. C. Freiling, "Integrating people-centric sensing with social networks: A privacy research agenda," in *Proceedings of the 8th IEEE International Conference on Pervasive Computing and Communications Workshops (PERCOM)*, pp. 620–623, 2010.
- [12] I. Krontiris and F. C. Freiling, "Urban sensing through social networks: The tension between participation and privacy," in *Proceedings of the International Tyrrhenian Workshop on Digital Communications (ITWDC)*, 2010.
- [13] R. B. Hays, "The day-to-day functioning of close versus casual friendships," *Journal of Social and Personal Relationships*, vol. 6, no. 1, pp. 21–37, 1989.
- [14] J. Camenisch, S. Hohenberger, and A. Lysyanskaya, "Balancing accountability and privacy using e-cash," in *Security and Cryptography for Networks*, vol. 4116 of *Lecture Notes in Computer Science*, pp. 141–155, 2006.
- [15] K. L. Huang, S. S. Kanhere, and W. Hu, "On the need for a reputation system in mobile phone based sensing," *Ad Hoc Networks*, vol. 12, no. 0, pp. 130–149, 2014.

- [16] M. Shin, C. Cornelius, D. Peebles, A. Kapadia, D. Kotz, and N. Triandopoulos, "Anonymsense: A system for anonymous opportunistic sensing," *Pervasive and Mobile Computing*, vol. 7, no. 1, pp. 16–30, 2011.
- [17] d. m. Boyd and N. B. Ellison, "Social network sites: Definition, history, and scholarship," *Journal of Computer-Mediated Communication*, vol. 13, no. 1, pp. 210–230, 2007.
- [18] G. Pallis, D. Zeinalipour-Yazti, and M. D. Dikaiakos, "Online social networks: Status and trends," in *New Directions in Web Data Management 1*, vol. 331 of *Studies in Computational Intelligence*, pp. 213–234, 2011.
- [19] A. T. Campbell, S. B. Eisenman, N. D. Lane, E. Miluzzo, and R. A. Peterson, "People-centric urban sensing," in *Proceedings of the 2nd ACM Annual International Workshop on Wireless Internet (WICON)*, pp. 18–31, 2006.
- [20] S. Reddy, A. Parker, J. Hyman, J. Burke, D. Estrin, and M. Hansen, "Image browsing, processing, and clustering for participatory sensing: Lessons from a dietsense prototype," in *Proceedings of the 4th ACM Workshop on Embedded Networked Sensors*, pp. 13–17, 2007.
- [21] S. B. Eisenman, E. Miluzzo, N. D. Lane, R. A. Peterson, G.-S. Ahn, and A. T. Campbell, "Bikenet: A mobile sensing system for cyclist experience mapping," in *ACM Transactions on Sensor Networks (TOSN)*, vol. 6, pp. 1–39, 2009.
- [22] D. Christin, C. Roßkopf, M. Hollick, L. A. Martucci, and S. S. Kanhere, "Incognisense: An anonymity-preserving reputation framework for participatory sensing applications," in *Pervasive and Mobile Computing*, vol. 9, pp. 353–371, 2013.
- [23] A. Dua, N. Bulusu, W.-C. Feng, and W. Hu, "Towards trustworthy participatory sensing," in *Proceedings of the Usenix Workshop on Hot Topics in Security (HotSec)*, pp. 8–8, 2009.
- [24] S. Saroiu and A. Wolman, "I am a sensor, and i approve this message," in *Proceedings of the 11th ACM Workshop on Mobile Computing Systems and Applications (HotMobile)*, pp. 37–42, 2010.

- [25] “Trusted computing group.” <https://www.trustedcomputinggroup.org/home>.
- [26] S. Ganeriwal, L. K. Balzano, and M. B. Srivastava, “Reputation-based framework for high integrity sensor networks,” in *ACM Transactions on Sensor Networks (TOSN)*, vol. 4, p. 15, 2008.
- [27] A. Gupta, H. Lamba, P. Kumaraguru, and A. Joshi, “Faking sandy: Characterizing and identifying fake images on twitter during hurricane sandy,” in *Proceedings of the 22nd International Conference on World Wide Web Companion*, pp. 729–736, 2013.
- [28] A. Gupta and P. Kumaraguru, “Twitter explodes with activity in mumbai blasts! a lifeline or an unmonitored daemon in the lurking?,” tech. rep., Delhi, IIITD-TR-2011-005, 2011.
- [29] S. Papadimitriou, H. Kitagawa, P. B. Gibbons, and C. Faloutsos, “Loci: Fast outlier detection using the local correlation integral,” in *Proceedings of the 19th IEEE International Conference on Data Engineering*, pp. 315–326, 2003.
- [30] H.-P. Kriegel, P. Kröger, and A. Zimek, “Outlier detection techniques,” in *Tutorial at the 13th Pacific-Asia Conference on Knowledge Discovery and Data Mining (PAKDD)*, 2009.
- [31] K. L. Huang, S. S. Kanhere, and W. Hu, “A privacy-preserving reputation system for participatory sensing,” in *Proceedings of the 7th IEEE Conference on Local Computer Networks (LCN)*, pp. 10–18, 2012.
- [32] X. Oscar Wang, W. Cheng, P. Mohapatra, and T. Abdelzaher, “Artsense: Anonymous reputation and trust in participatory sensing,” in *Proceedings of the 32nd IEEE International Conference on Computer Communications (INFOCOM)*, pp. 2517–2525, 2013.
- [33] H. Amintoosi and S. S. Kanhere, “A trust-based recruitment framework for multi-hop social participatory sensing,” in *Proceedings of the 9th IEEE International Conference on Distributed Computing in Sensor Systems (DCOSS)*, pp. 266–273, 2013.

- [34] H. Amintoosi and S. S. Kanhere, "Privacy-aware trust-based recruitment in social participatory sensing," in *10th International Conference on Mobile and Ubiquitous Systems: Computing, Networking, and Services (MobiQuitous)*, in press.
- [35] H. Amintoosi, S. S. Kanhere, and M. Allahbakhsh, "Trust-based privacy-aware participant selection in social participatory sensing," *submitted to the Journal of Information Security and Applications*, 2014.
- [36] C. E. Shannon, "A mathematical theory of communication," in *ACM Mobile Computing and Communications Review (SIGMOBILE)*, vol. 5, pp. 3–55, 2001.
- [37] F. Spitzer, *Principles of Random Walk*. Springer, 2001.
- [38] A. Mukherjee, B. Liu, and N. Glance, "Spotting fake reviewer groups in consumer reviews," in *Proceedings of the 21st ACM International Conference on World Wide Web*, pp. 191–200, 2012.
- [39] H. Amintoosi and S. S. Kanhere, "A trust framework for social participatory sensing systems," in *Proceedings of the 9th International Conference on Mobile and Ubiquitous Systems: Computing, Networking, and Services (MobiQuitous)*, pp. 237–249, 2013.
- [40] H. Amintoosi and S. Kanhere, "A reputation framework for social participatory sensing systems," in *Mobile Networks and Applications (MONET)*, vol. 19, pp. 88–100, 2014.
- [41] H. Amintoosi, M. Allahbakhsh, S. S. Kanhere, and M. Niazi Torshiz, "Trust assessment in social participatory networks," in *Proceedings of the 3rd International eConference on Computer and Knowledge Engineering (ICCKE)*, pp. 448–453, 2013.
- [42] L. Page, S. Brin, R. Motwani, and T. Winograd, "The pagerank citation ranking: Bringing order to the web," tech. rep., Stanford Digital Library Technologies Project, 1999.

- [43] G. Grahne and J. Zhu, “Fast algorithms for frequent itemset mining using fp-trees,” *IEEE Transactions on Knowledge and Data Engineering*, vol. 17, no. 10, pp. 1347–1362, 2005.
- [44] J. Goldman, K. Shilton, J. Burke, D. Estrin, M. Hansen, N. Ramanathan, S. Reddy, V. Samanta, M. Srivastava, and R. West, “Participatory sensing: A citizen-powered approach to illuminating the patterns that shape our world,” *Foresight & Governance Project, White Paper*, 2009.
- [45] Y. Wang and D. R. Fesenmaier, “Assessing motivation of contribution in on-line communities: An empirical investigation of an online travel community,” *Electronic Markets*, vol. 13, no. 1, pp. 33–45, 2003.
- [46] P. Kollock, *Communities in Cyberspace*, ch. The Economies of Online Cooperation: Gifts and Public Goods in Cyberspace, pp. 220–242. 1999.
- [47] H. Rheingold, *The Virtual Community: Homesteading on the Electronic Frontier*. Basic Books, 1993.
- [48] M. Demirbas, M. A. Bayir, C. G. Akcora, Y. S. Yilmaz, and H. Ferhatosmanoglu, “Crowd-sourced sensing and collaboration using twitter,” in *Proceedings of the IEEE International Symposium on World of Wireless Mobile and Multimedia Networks (WoWMoM)*, pp. 1–9, 2010.
- [49] W. Khan, Y. Xiang, M. Aalsalem, and Q. Arshad, “Mobile phone sensing systems: A survey,” *IEEE Communications Surveys and Tutorials*, vol. 15, no. 1, pp. 402–427, 2013.
- [50] B. Hull, V. Bychkovsky, Y. Zhang, K. Chen, M. Goraczko, A. Miu, E. Shih, H. Balakrishnan, and S. Madden, “Cartel: A distributed mobile sensor computing system,” in *Proceedings of the 4th ACM International Conference on Embedded Networked Sensor Systems*, pp. 125–138, 2006.
- [51] M. Von Kaenel, P. Sommer, and R. Wattenhofer, “Ikarus: Large-scale participatory sensing at high altitudes,” in *Proceedings of the 12th ACM Workshop on Mobile Computing Systems and Applications*, pp. 63–68, 2011.

- [52] P. B. Crosby, *Quality is Free: The Art of Making Quality Certain*. McGraw-Hill New York, 1979.
- [53] A. J. Quinn and B. B. Bederson, “Human computation: A survey and taxonomy of a growing field,” in *Proceedings of the ACM SIGCHI Conference on Human Factors in Computing Systems*, pp. 1403–1412, 2011.
- [54] M. Allahbakhsh, B. Benatallah, A. Ignjatovic, H. R. Motahari-Nezhad, E. Bertino, and S. Dustdar, “Quality control in crowdsourcing systems: Issues and directions,” in *IEEE Internet Computing*, vol. 17, pp. 76–81, 2013.
- [55] B. Stvilia, M. B. Twidale, L. C. Smith, and L. Gasser, “Information quality work organization in wikipedia,” *Journal of the American society for information science and technology*, vol. 59, no. 6, pp. 983–1001, 2008.
- [56] L. Von Ahn and L. Dabbish, “Labeling images with a computer game,” in *Proceedings of the ACM SIGCHI Conference on Human Factors in Computing Systems*, pp. 319–326, 2004.
- [57] E. Law and L. Von Ahn, “Input-agreement: A new mechanism for collecting data using human computation games,” in *Proceedings of the ACM SIGCHI Conference on Human factors in Computing Systems*, pp. 1197–1206, 2009.
- [58] G. Paolacci, J. Chandler, and P. Ipeirotis, “Running experiments on amazon mechanical turk,” *Judgment and Decision Making*, vol. 5, no. 5, pp. 411–419, 2010.
- [59] V. S. Sheng, F. Provost, and P. G. Ipeirotis, “Get another label? improving data quality and data mining using multiple, noisy labelers,” in *Proceedings of the 14th ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*, pp. 614–622, 2008.
- [60] V. Lenders, E. Koukoumidis, P. Zhang, and M. Martonosi, “Location-based trust for mobile user-generated content: Applications, challenges and implementations,” in *Proceedings of the 9th ACM Workshop on Mobile Computing Systems and Applications*, pp. 60–64, 2008.

- [61] S. Saroiu and A. Wolman, “Enabling new mobile applications with location proofs,” in *Proceedings of the 10th ACM Workshop on Mobile Computing Systems and Applications (HotMobile)*, pp. 3:1–3:6, 2009.
- [62] P. Gilbert, J. Jung, K. Lee, H. Qin, D. Sharkey, A. Sheth, and L. P. Cox, “Youprove: Authenticity and fidelity in mobile sensing,” in *Proceedings of the 9th ACM Conference on Embedded Networked Sensor Systems*, pp. 176–189, 2011.
- [63] K. L. Huang, S. S. Kanhere, and W. Hu, “Are you contributing trustworthy data?: The case for a reputation system in participatory sensing,” in *Proceedings of the 13th ACM International Conference on Modeling, Analysis, and Simulation of Wireless and Mobile Systems*, pp. 14–22, 2010.
- [64] L. De Alfaro, A. Kulshreshtha, I. Pye, and B. T. Adler, “Reputation systems for open collaboration,” *Communications of the ACM*, vol. 54, no. 8, pp. 81–87, 2011.
- [65] A. Jøsang, R. Ismail, and C. Boyd, “A survey of trust and reputation systems for online service provision,” *Decision Support Systems*, vol. 43, no. 2, pp. 618–644, 2007.
- [66] Q. Feng, L. Liu, and Y. Dai, “Vulnerabilities and countermeasures in context-aware social rating services,” in *ACM Transactions on Internet Technology (TOIT)*, vol. 11, pp. 11:1–11:27, 2012.
- [67] M. Allahbakhsh, A. Ignjatovic, B. Benatallah, S.-M.-R. Beheshti, E. Bertino, and N. Foo, “Reputation management in crowdsourcing systems,” in *Proceedings of the 8th IEEE International Conference on Collaborative Computing: Networking, Applications and Worksharing (CollaborateCom)*, pp. 664–671, 2012.
- [68] L. Cabral and A. Hortacsu, “The dynamics of seller reputation: Evidence from ebay*,” *The Journal of Industrial Economics*, vol. 58, no. 1, pp. 54–78, 2010.

- [69] B. T. Adler and L. De Alfaro, "A content-driven reputation system for the wikipedia," in *Proceedings of the 16th International Conference on World Wide Web*, pp. 261–270, 2007.
- [70] B. T. Adler, L. De Alfaro, S. M. Mola-Velasco, P. Rosso, and A. G. West, "Wikipedia vandalism detection: Combining natural language, metadata, and reputation features," in *Computational Linguistics and Intelligent Text Processing*, vol. 6609 of *Lecture Notes in Computer Science*, pp. 277–288, 2011.
- [71] C. Selvaraj and S. Anand, "A survey on security issues of reputation management systems for peer-to-peer networks," *Computer Science Review*, vol. 6, no. 4, pp. 145–160, 2012.
- [72] S. Buchegger and J. Le Boudec, "The effect of rumor spreading in reputation systems for mobile ad-hoc networks," in *Proceedings of the Workshop on Modeling and Optimization in Mobile, Ad Hoc and Wireless Networks (WiOpt)*, 2003.
- [73] A. Whitby, A. Jøsang, and J. Indulska, "Filtering out unfair ratings in bayesian reputation systems," in *Proceedings of the 7th International Workshop on Trust in Agent Societies*, 2004.
- [74] A. Jsang and R. Ismail, "The beta reputation system," in *Proceedings of the 15th Bled Electronic Commerce Conference*, pp. 41–55, 2002.
- [75] J. F. Kenney and E. S. Keeping, *Mathematics of Statistics, part 1*. Van Nostrand, 1962.
- [76] L. Sweeney, "K-anonymity: A model for protecting privacy," *International Journal of Uncertainty, Fuzziness and Knowledge-Based Systems*, vol. 10, no. 05, pp. 557–570, 2002.
- [77] D. Chaum, "Blind signatures for untraceable payments," in *Advances in Cryptology*, pp. 199–203, 1983.

- [78] X. Wang, W. Cheng, P. Mohapatra, and T. Abdelzaher, “Enabling reputation and trust in privacy-preserving mobile sensing,” *IEEE Transactions on Mobile Computing*, vol. PP, no. 99, p. 1, 2014.
- [79] G. Chatzimilioudis, A. Konstantinidis, C. Laoudias, and D. Zeinalipour-Yazti, “Crowdsourcing with smartphones,” in *IEEE Internet Computing*, vol. 16, pp. 36–44, 2012.
- [80] J. Chandler, G. Paolacci, and P. Mueller, “Risks and rewards of crowdsourcing marketplaces,” in *Handbook of Human Computation*, pp. 377–392, 2013.
- [81] L. Von Ahn and L. Dabbish, “Designing games with a purpose,” *Communications of the ACM*, vol. 51, no. 8, pp. 58–67, 2008.
- [82] E. Peer, J. Vosgerau, and A. Acquisti, “Reputation as a sufficient condition for data quality on amazon mechanical turk,” *Behavior research methods*, pp. 1–9, 2013.
- [83] Y. Sun and Y. Liu, “Security of online reputation systems: The evolution of attacks and defenses,” *IEEE Signal Processing Magazine*, vol. 29, no. 2, pp. 87–97, 2012.
- [84] A. Alkouz, E. W. De Luca, and S. Albayrak, “Latent semantic social graph model for expert discovery in facebook,” in *Proceedings of the 11th International Conference on Innovative Internet Community Systems (IICS)*, pp. 128–138, 2011.
- [85] K. Ehrlich, C.-Y. Lin, and V. Griffiths-Fisher, “Searching for experts in the enterprise: Combining text and social network analysis,” in *Proceedings of the ACM International Conference on Supporting Groupwork*, pp. 117–126, 2007.
- [86] J. Zhang, J. Tang, and J. Li, “Expert finding in a social network,” in *Advances in Databases: Concepts, Systems and Applications*, vol. 4443 of *Lecture Notes in Computer Science*, pp. 1066–1069, 2007.

- [87] E. Smirnova, “A model for expert finding in social networks,” in *Proceedings of the 34th ACM SIGIR International Conference on Research and Development in Information Retrieval*, pp. 1191–1192, 2011.
- [88] A. Doan, R. Ramakrishnan, and A. Y. Halevy, “Crowdsourcing systems on the world-wide web,” vol. 54, pp. 86–96, 2011.
- [89] S. Reddy, K. Shilton, J. Burke, D. Estrin, M. Hansen, and M. Srivastava, “Evaluating participation and performance in participatory sensing,” in *Proceedings of the International Workshop on Urban, Community, and Social Applications of Networked Sensing Systems*, pp. 1–5, 2008.
- [90] S. Reddy, K. Shilton, J. Burke, D. Estrin, M. Hansen, and M. Srivastava, “Using context annotated mobility profiles to recruit data collectors in participatory sensing,” in *Location and Context Awareness*, vol. 5561 of *Lecture Notes in Computer Science*, pp. 52–69, 2009.
- [91] G. S. Tuncay, G. Benincasa, and A. Helmy, “Autonomous and distributed recruitment and data collection framework for opportunistic sensing,” in *ACM Mobile Computing and Communications Review (SIGMOBILE)*, vol. 16, pp. 50–53, 2013.
- [92] S. Reddy, D. Estrin, M. Hansen, and M. Srivastava, “Examining micro-payments for participatory sensing data collections,” in *Proceedings of the 12th ACM International Conference on Ubiquitous Computing*, pp. 33–36, 2010.
- [93] J.-S. Lee and B. Hoh, “Sell your experiences: A market mechanism based incentive for participatory sensing,” in *Proceedings of the IEEE International Conference on Pervasive Computing and Communications (PerCom)*, pp. 60–68, 2010.
- [94] J.-S. Lee and B. Hoh, “Dynamic pricing incentive for participatory sensing,” *Pervasive and Mobile Computing*, vol. 6, no. 6, pp. 693–708, 2010.
- [95] T. Luo and C.-K. Tham, “Fairness and social welfare in incentivizing participatory sensing,” in *Proceedings of the 9th IEEE Annual Communications*

- Society Conference on Sensor, Mesh and Ad Hoc Communications and Networks (SECON)*, pp. 425–433, 2012.
- [96] D. Yang, G. Xue, X. Fang, and J. Tang, “Crowdsourcing to smartphones: Incentive mechanism design for mobile phone sensing,” in *Proceedings of the 18th Annual International Conference on Mobile Computing and Networking*, pp. 173–184, 2012.
- [97] L. G. Jaimes, I. Vergara-Laurens, and M. A. Labrador, “A location-based incentive mechanism for participatory sensing systems with budget constraints,” in *Proceedings of the IEEE International Conference on Pervasive Computing and Communications (PerCom)*, pp. 103–108, 2012.
- [98] S. Khuller, A. Moss, and J. S. Naor, “The budgeted maximum coverage problem,” *Information Processing Letters*, vol. 70, no. 1, pp. 39–45, 1999.
- [99] D. Goldschlag, M. Reed, and P. Syverson, “Onion routing,” *Communications of the ACM*, vol. 42, no. 2, pp. 39–41, 1999.
- [100] R. Dingledine, N. Mathewson, and P. Syverson, “Tor: The second-generation onion router,” tech. rep., DTIC Document, 2004.
- [101] D. L. Chaum, “Untraceable electronic mail, return addresses, and digital pseudonyms,” *Communications of the ACM*, vol. 24, no. 2, pp. 84–90, 1981.
- [102] K. Sampigethaya and R. Poovendran, “A survey on mix networks and their secure applications,” *Proceedings of the IEEE*, vol. 94, no. 12, pp. 2142–2181, 2006.
- [103] E. M. Daly and M. Haahr, “Social network analysis for routing in disconnected delay-tolerant manets,” in *Proceedings of the 8th ACM International Symposium on Mobile Ad hoc Networking and Computing*, pp. 32–40, 2007.
- [104] A. Agarwal and S. Chakrabarti, “Learning random walks to rank nodes in graphs,” in *Proceedings of the 24th International Conference on Machine Learning*, pp. 9–16, 2007.

- [105] G. Weiss, *Aspects and Applications of the Random Walk (Random Materials & Processes)*. North-Holland, 2005.
- [106] L. Lovász, “Random walks on graphs: A survey,” vol. 2, pp. 1–46, 1993.
- [107] L. Backstrom and J. Leskovec, “Supervised random walks: Predicting and recommending links in social networks,” in *Proceedings of the 4th ACM International Conference on Web Search and Data Mining*, pp. 635–644, 2011.
- [108] I. Konstas, V. Stathopoulos, and J. M. Jose, “On social networks and collaborative recommendation,” in *Proceedings of the 32nd ACM SIGIR International Conference on Research and Development in Information Retrieval*, pp. 195–202, 2009.
- [109] P. Pons and M. Latapy, “Computing communities in large networks using random walks,” *Computer and Information Sciences (ISCIS)*, vol. 3733, pp. 284–293, 2005.
- [110] A. Sadilek, J. Krumm, and E. Horvitz, “Crowdphysics: Planned and opportunistic crowdsourcing for physical tasks,” vol. 21, pp. 125–134, 2013.
- [111] M. Huber, M. Mulazzani, and E. Weippl, “Who on earth is mr. cypher: Automated friend injection attacks on social networking sites,” in *Security and Privacy—Silver Linings in the Cloud*, vol. 330 of *IFIP Advances in Information and Communication Technology*, pp. 80–89, 2010.
- [112] S. B. Eisenman, N. D. Lane, and A. T. Campbell, “Techniques for improving opportunistic sensor networking performance,” in *Distributed Computing in Sensor Systems*, vol. 5067 of *Lecture Notes in Computer Science*, pp. 157–175, 2008.
- [113] H. Lu, N. Lane, S. Eisenman, and A. Campbell, “Bubble-sensing: A new paradigm for binding a sensing task to the physical world using mobile phones,” in *Proceedings of the International Workshop on Mobile Devices and Urban Sensing*, 2008.

- [114] H. Lu, N. D. Lane, S. B. Eisenman, and A. T. Campbell, "Bubble-sensing: Binding sensing tasks to the physical world," *Pervasive and Mobile Computing*, vol. 6, no. 1, pp. 58–71, 2010.
- [115] J. Al-Muhtadi, R. Campbell, A. Kapadia, M. D. Mickunas, and S. Yi, "Routing through the mist: Privacy preserving communication in ubiquitous computing environments," in *Proceedings of the 22nd IEEE International Conference on Distributed Computing Systems*, pp. 74–83, 2002.
- [116] A. Kapadia, D. Kotz, and N. Triandopoulos, "Opportunistic sensing: Security challenges for the new paradigm," in *Proceedings of the First IEEE International Communication Systems and Networks and Workshops (COMSNETS)*, pp. 1–10, 2009.
- [117] D. Christin, J. Guillemet, A. Reinhardt, M. Hollick, and S. S. Kanhere, "Privacy-preserving collaborative path hiding for participatory sensing applications," in *Proceedings of the 8th IEEE International Conference on Mobile Adhoc and Sensor Systems (MASS)*, pp. 341–350, 2011.
- [118] J. Ugander, B. Karrer, L. Backstrom, and C. Marlow, "The anatomy of the facebook social graph," *arXiv preprint arXiv:1111.4503*, 2011.
- [119] S.-H. Yang, B. Long, A. Smola, N. Sadagopan, Z. Zheng, and H. Zha, "Like like alike: Joint friendship and interest propagation in social networks," in *Proceedings of the 20th ACM International Conference on World Wide Web*, pp. 537–546, 2011.
- [120] L. Sweeney, "Achieving k-anonymity privacy protection using generalization and suppression," *International Journal of Uncertainty, Fuzziness and Knowledge-Based Systems*, vol. 10, no. 5, pp. 571–588, 2002.
- [121] O. Hasan, L. Brunie, and J.-M. Pierson, "Evaluation of the iterative multiplication strategy for trust propagation in pervasive environments," in *Proceedings of the ACM International Conference on Pervasive Services*, pp. 49–54, 2009.

- [122] E. Zheleva and L. Getoor, *Social Network Data Analytics*, ch. Privacy in Social Networks: A Survey, pp. 277–306. 2011.
- [123] M. H. Yaghmaee, M. Menhaj, and H. Amintoosi, “A fuzzy extension to the blue active queue management algorithm,” *Journal of Iranian Association of Electrical and Electronics Engineers (IAEEE)*, vol. 1, no. 3, pp. 3–14, 2005.
- [124] A. N. Langville and C. D. Meyer, *Google’s PageRank and Beyond: The Science of Search Engine Rankings*. Princeton University Press, 2006.
- [125] R. Levien and A. Aiken, “Attack-resistant trust metrics for public key certification,” in *Proceedings of the 7th USENIX Security Symposium*, pp. 229–242, 1998.
- [126] M. H. Yaghmaee, M. B. Menhaj, and H. Amintoosi, “Design and performance evaluation of a fuzzy based traffic conditioner for differentiated services,” in *Computer Networks*, vol. 47, pp. 847–869, 2005.
- [127] W. V. Leekwijck and E. E. Kerre, “Defuzzification: Criteria and classification,” *Fuzzy Sets and Systems*, vol. 108, no. 2, pp. 159–178, 1999.
- [128] R. Agrawal and R. Srikant, “Fast algorithms for mining association rules,” in *Proceedings of the 20th International Conference on Very Large Data Bases (VLDB)*, vol. 1215, pp. 487–499, 1994.
- [129] T. Brijs, B. Goethals, G. Swinnen, K. Vanhoof, and G. Wets, “A data mining framework for optimal product selection in retail supermarket data: the generalized profset model,” in *Proceedings of the 6th ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*, pp. 300–304, 2000.
- [130] T. Brijs, G. Swinnen, K. Vanhoof, and G. Wets, “Using association rules for product assortment decisions: A case study,” in *Proceedings of the 5th ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*, pp. 254–260, 1999.

- [131] W. Lee, S. J. Stolfo, and K. W. Mok, “Adaptive intrusion detection: A data mining approach,” *Artificial Intelligence Review*, vol. 14, no. 6, pp. 533–567, 2000.
- [132] M. Allahbakhsh, A. Ignjatovic, B. Benatallah, S.-M.-R. Beheshti, E. Bertino, and N. Foo, “Collusion detection in online rating systems,” in *Proceedings of the 15th Asia Pacific Web Conference (APWeb)*, pp. 196–207, 2013.
- [133] E.-P. Lim, V.-A. Nguyen, N. Jindal, B. Liu, and H. W. Lauw, “Detecting product review spammers using rating behaviors,” in *Proceedings of the 19th ACM International Conference on Information and Knowledge Management (CIKM)*, pp. 939–948, 2010.
- [134] G. Salton, C. Buckley, and E. A. Fox, “Automatic query formulations in information retrieval,” *Journal of the American Society for Information Science*, vol. 34, no. 4, pp. 262–280, 1983.
- [135] J. Leskovec, L. Adamic, and B. Huberman, “The dynamics of viral marketing,” *ACM Transactions on the Web (TWEB)*, vol. 1, no. 1, pp. 1–39, 2007.
- [136] Y. Emek, R. Karidi, M. Tennenholtz, and A. Zohar, “Mechanisms for multi-level marketing,” in *Proceedings of the 12th ACM Conference on Electronic Commerce*, pp. 209–218, 2011.
- [137] J. Kleinberg and P. Raghavan, “Query incentive networks,” in *Proceedings of the 46th IEEE Annual Symposium on Foundations of Computer Science (FOCS)*, pp. 132–141, 2005.
- [138] H. R. Motahari Nezhad, G. Y. Xu, and B. Benatallah, “Protocol-aware matching of web service interfaces for adapter development,” in *Proceedings of the 19th ACM International Conference on World Wide Web*, pp. 731–740, 2010.
- [139] G. Chowdhury, *Introduction to Modern Information Retrieval*. Facet publishing, 2010.

- [140] R. R. Kalidindi, K. Raju, V. Valli Kumari, and C. S. Reddy, "Trust based participant driven privacy control in participatory sensing," *International Journal of Ad Hoc, Sensor & Ubiquitous Computing*, vol. 2, no. 1, pp. 171–84, 2011.
- [141] F. Alt, A. S. Shirazi, A. Schmidt, U. Kramer, and Z. Nawaz, "Location-based crowdsourcing: Extending crowdsourcing to the real world," in *Proceedings of the 6th ACM Nordic Conference on Human-Computer Interaction: Extending Boundaries*, pp. 13–22, 2010.
- [142] G. Mesch and I. Talmud, "The quality of online and offline relationships: The role of multiplexity and duration of social relationships," *The Information Society*, vol. 22, no. 3, pp. 137–148, 2006.
- [143] Y. Altshuler, M. Fire, N. Aharony, Y. Elovici, and A. S. Pentland, "How many makes a crowd? on the evolution of learning as a factor of community coverage," in *Social Computing, Behavioral-Cultural Modeling and Prediction*, vol. 7227 of *Lecture Notes in Computer Science*, pp. 43–52, 2012.
- [144] G. M. Erickson, P. J. Currie, B. D. Inouye, and A. A. Winn, "Tyrannosaur life tables: an example of nonavian dinosaur population biology," *Science*, vol. 313, no. 5784, pp. 213–217, 2006.
- [145] A. dOnofrio, "A general framework for modeling tumor-immune system competition and immunotherapy: Mathematical analysis and biomedical inferences," *Physica D: Nonlinear Phenomena*, vol. 208, no. 3, pp. 220–235, 2005.
- [146] M. Niazi Torshiz, H. Amintoosi, and A. Movaghar, "A fuzzy energy-based extension to aodv routing," in *Proceedings of the IEEE International Symposium on Telecommunications (IST)*, pp. 371–375, 2008.
- [147] A. K. Akobeng and A. G. Thomas, "Systematic review: Tolerable amount of gluten for people with coeliac disease," *Alimentary Pharmacology & Therapeutics*, vol. 27, no. 11, pp. 1044–1052, 2008.
- [148] S. Kamvar, T. Haveliwala, C. Manning, and G. Golub, "Exploiting the block structure of the web for computing pagerank," tech. rep., Stanford University, 2003.

-
- [149] S. Shekarpour and S. Katebi, “Modeling and evaluation of trust with an extension in semantic web,” *Web Semantics: Science, Services and Agents on the World Wide Web*, vol. 8, no. 1, pp. 26–36, 2010.
- [150] K. L. Judd, “The law of large numbers with a continuum of iid random variables,” *Journal of Economic theory*, vol. 35, no. 1, pp. 19–25, 1985.