

Trust and Reputation Mining in Professional Virtual Communities [★]

Florian Skopik, Hong-Linh Truong, Schahram Dustdar

Distributed Systems Group, Vienna University of Technology
Argentinierstr. 8/184-1, A-1040 Vienna, Austria
{skopik|truong|dustdar}@infosys.tuwien.ac.at

Abstract. Communication technologies, such as e-mail, instant messaging, discussion forums, blogs, and newsgroups connect people together, forming virtual communities. This concept is not only used for private purposes, but is also attracting attention in professional environments, allowing to consult a large group of experts. Due to the overwhelming size of such communities, various reputation mechanisms have been proposed supporting members with information about people's trustworthiness with respect to their contributions. However, most of today's approaches rely on manual and subjective feedback, suffering from unfair ratings, discrimination, and feedback quality variations over time.

To this end, we propose a system which determines trust relationships between community members automatically and objectively by mining communication data. In contrast to other approaches which use these data directly, e.g., by applying natural language processing on log files, we follow a new approach to make contributions visible. We perform structural analysis of discussions, examine interaction patterns between members, and infer social roles expressing motivation, openness to discussions, and willingness to share data, and therefore *trust*.

1 Introduction

The concept of virtual (or online) communities is quite common today and frequently used not only for private concerns, but also in professional working environments. Online platforms such as discussion forums, blogs, and newsgroups are regularly utilized to get introduced into new things, to find solutions for particular problems, or just to stay informed on what's up in a certain domain. Virtual communities are rapidly growing and emerging, and thus, lots of spam and dispensable comments are posted in their forums or sent via e-mail, polluting fruitful discussions. Several mechanisms have been proposed to handle this problem, such as collaborative filtering of comments and global reputation of users based on feedback mechanisms. However, because these concepts rely on manual and subjective human feedback, they suffer from several drawbacks [1], including

[★] This work is mainly supported by the European Union through the FP7-216256 project COIN.

unfair ratings, low incentives for providing feedback, and quality variations of ratings over time.

Collaborative Environments. Especially, where mentioned communication technologies are regularly embedded to connect e-professionals together, such as in professional virtual communities (PVCs), and where successful collaboration is critical for business, we identified the need for more sophisticated reputation methods. Moreover, in modern working environments, where virtual teams consisting of members from different departments or even companies work together, personally unknown to each other, various complex social factors affect the overall collaboration success. These factors can be expressed by one composite and abstract concept: *trust*. Trusted relationships between colleagues are vital to the whole collaboration and a prerequisite for successful work. A recent report about the roles of trust in today’s business world [2] discovers that besides professional skills expressed as experience, expertise and competence, soft skills, such as the willingness to exchange information, motivation and communication skills, are at least equally important. Such social skills can be discovered and evaluated in typical computer-supported discussions, common in online communities, including threaded forum discussions, instant messaging chats, and e-mail conversation.

The Autonomic Cycle. Our overall motivation for trust determination is to apply an autonomic management cycle [3] consisting of four phases (monitoring, analyzing, planning, executing), which enable the adaptation of collaboration environments and personalization of applications with respect to trust between participants. In this cycle, the collaboration behavior, such as the communication culture, the execution of tasks and the coordination of e-workers is monitored by the system and their relationships are determined by analyzing logging data and structural profiles. Depending on particular situations, different available metrics are aggregated and interpreted as trust, which enables the maintenance of a trust network. This trust network is utilized to plan further collaboration, e.g., influences work partner selection or the assignment of tasks. After that, when people perform the actual work, their collaboration is monitored by the system, which closes the loop. In this paper we focus particularly the *monitoring phase* and the *analyzing phase*.

Contributions. We show an approach for automatic inference of trust between online discussion participants. To this end, we propose a system which collects and merges data from various sources, define the notion of discussion trust, and cover related concepts, including the definition of user roles and an interaction network. The main contribution is the design of a mining algorithm, which we evaluate with a real data set.

Paper Structure. The rest of the paper is organized as follows. In Sect. 2 we consider related work. Sect. 3 is about harnessing trustworthy sources of data in PVCs. We describe trust and roles in discussions in Sect. 4, a mining algorithm using these concepts to calculate relationships based on observed communication in Sect. 5, and network-based trust inference in Sect. 6. We prove our approach with extensive experiments on real data in Sect. 7 and conclude in Sect. 8.

2 Related Work

Welser et al. [4] provides interesting studies about social roles in online discussion groups and participants' behaviors. Furthermore, Nonnecke et al. [5] and Meyer et al. [6] research the meaning of online communication and differences between traditional face-to face and threaded discussions. McLure-Wasko and Faraj [7] investigate the motivation for knowledge sharing in online platforms, and Rheingold [8] comprehensively examines the concept of virtual communities. The article [2] in The Economist draws the importance of trust in business communication, and shows that various factors which directly influence trust between individuals are based on communication.

Until now various computational trust models have been developed, as summarized by Massa in [9]. Though they are useful to deal with trust propagation, aggregation and evaluation over time, it is mostly assumed that initial trust relationships are manually determined by people. For example, the well-known TrustRank model [10] for ranking web sites with respect to their trustworthiness, needs a set of trusted web sites to be initially defined and is then able to inherit trust to further linked pages automatically.

We interpret previous communications between people as interactions and rank them according to their trustworthiness in the originating network. There are several graph based ranking algorithms, including PageRank [11], HITS [12], and Eigentrust [13]. However, in contrast to these algorithms, which operate on top of an existing network, our approach tackles the challenges beneath, i.e. creating the network based on discussion data. To this end, we develop a mining algorithm to gather individual trust relationships based on observed communications, considering detailed analysis of online discussion platforms such as Gomez et al. [14] for Slashdot.

3 Trustworthy Sources of Data

Most common online communication platforms, such as vBulletin¹, integrate reputation systems which either rank users based on simple metrics, including their posting count, or enable users to reward ('thank') others directly. In contrast to this approach, we reward the vitality of *relationships* between users and then derive the user gradings by aggregating relationship gradings. This enables us to utilize global metrics calculated from all available data, such as the overall discussion effort of a particular user with respect to the whole community, but also local metrics considering data restricted to particular users only, such as the discussion effort between two specific users. Utilizing local metrics is particularly of interest when the amount of controversial users is high [15]. With our system a user does not get one globally valid trust rank, but may be graded from each individual's view.

We developed a pluggable architecture (Fig. 1) - part of VieTE [16] - which utilizes various communication data sources through standard protocols, such as

¹ <http://www.vbulletin.com>

RSS feeds, SOAP, and e-mail. The obtained data is mapped to a generic communication schema, pre-processed and refined, and finally an interaction network, modeling the relationships between users emerging from communication, is built. Based on this network, trust between individuals and reputation from the community's view can be inferred and queried through a dedicated interface. We understand the inferred *discussion trust* to represent one dimension of general trust in PVCs, applicable in a wide range of differently organized communities. Other dimensions of trust may base on the fulfillment of service level agreements or the reliability of task execution, which are out of scope of this paper.

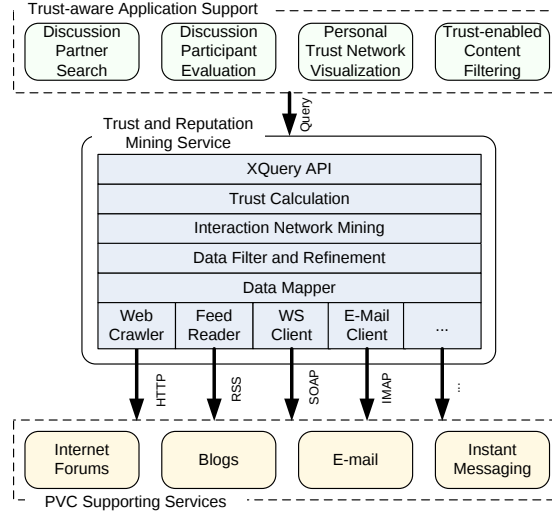


Fig. 1. Architectural overview of the trust and reputation mining service.

4 Trust and Roles in Virtual Community Discussions

In discussions we can intuitively distinguish between information providers and information consumers. Especially in online discussions we can easily track who provides information, e.g., by posting a comment in a forum or writing an e-mail to a group of people. In contrast to that, determining information consumers is tricky. We can never be sure, that people read received e-mails or new comments in a forum, even when they open forum entries in their browsers. However, if somebody replies to a particular comment, then we can certainly assume, s/he has read the message and found it worth for discussion. Thus the replier can be identified as an information consumer, but as an information provider as well.

In our approach we track exactly this discussion behavior and define, that whenever one replies to a comment of another one, an interaction between them takes place. We process these interactions and build a notion of trust on top.

We apply Mui's definition of trust [17], which states that trust is "*a subjective expectation an agent has about another's future behavior based on the*

history of their encounters". We extend this definition by the notion of context, which means trust is established by considering past interactions in particular situations as widely agreed [18–20]. In the area of online discussions, contextual information is for instance the overall discussion topic or the type of forum being used.

Particularly, in discussions it may seem intuitive, that the more comments somebody provides the more s/he can be trusted to be a good discussion partner. On the other side *lurkers* [5], referring to people just watching discussions but not actually participating, can be less trusted regarding their 'openness'. They lack the willingness to exchange information, motivation or communication skills, thus they are bad collaborators.

However, a simple comment count does not truly reflect if somebody's statements are real contributions and worth reading and discussing. Thus, we consider threaded structures as well and analyze how comments are recognized by others. To this end, we define the following novel social roles within discussion scenarios: (i) **Activator**: The role of an *Activator* reflects, that the more replies a discussion participant receives, the more one's comments seem to be worth for discussion, thus one can be trusted to have the competencies and skills to provide comments, interesting for a broad base of participants. (ii) **Driver**: The role of a *Driver* reflects, the more somebody replies to comments, the more s/he can be trusted to actively participate in a discussion, thus s/he represents a 'driver' evident for a fruitful discussion. (iii) **Affirmed Driver**: An *Affirmed Driver* is defined as a Driver whose contribution is affirmed. This is the case if there is at least one reply to a driver's comment.

According to these roles, *discussion trust* is a measure for the contribution to discussions expressing the willingness to provide information and support, but does not reflect that a particular participant offers a valid information or posts the truth. For this purpose, at least natural language processing and analyzing semantic meanings are required [21, 22], which is out of scope of this paper.

5 Discussion Mining Approach

We developed a mining algorithm to determine the contribution of people in discussions. However, in contrast to common approaches, we neither reward the participants directly (e.g., their number of provided comments), nor we utilize subjective feedback, but we mine interactions to reward particularly the relationships between each pair of discussion participants.

We make the following assumptions: (i) The notion of time can be neglected, which means our algorithms do not determine how trust relations change over time. We determine trust relations for one particular point in time, based on short history data. Temporal evaluations, e.g. by applying moving averages, temporal weighting functions or sliding windows, have to be set up on top of our approach and is out of scope of this paper. (ii) We do not apply natural language processing. Thus, we accept introducing noise and small errors by rewarding users who post useless comments (i.e., spam). In the evaluation part

we show that this is no disadvantage if we rely on larger amounts of data. We further assume that in PVCs spam occurs less frequently than in open internet forums.

5.1 Interaction Network Definition

We utilize a widely adopted graph model to reflect discussion relationships between users. However, we further incorporate context awareness in this model to allow trust determination with respect to different situations on top of the created interaction network.

This network is modeled as a directed multigraph $G = \langle V, E \rangle$ where each vertex $v, w \in V$ represents a user and the edges reflect relationships based on previous interactions between them. A relationship $e_{vw}^{Ctx} \in E$, as defined in (1), is described by various *metrics* such as the number of recent interactions, their weights and communication scores, with respect to particular situations described by context elements Ctx .

$$e_{vw}^{Ctx} = \langle v, w, Ctx, metrics[name, value] \rangle . \quad (1)$$

5.2 Discussion Mining Algorithm

We develop an algorithm which weighs the communication relations based on discussions between each pair of participants. For environments supporting threaded discussion structures as common in online forums or newsgroups, we argue that somebody who provides a comment in a discussion thread, is not only influenced by the comment s/he directly replies, but to a certain extent also by other preceding ones in the same chain of discussion. Thus, we interpret a thread to be similar to a group discussion and establish relationships between participants posting in one chain. Figure 2(a) shows a structured discussion thread where every box represents a comment provided by the annotated participant. For the highlighted comment provided by x , arrows show exemplary which interactions between participants take place. The comment provider x honors the attracting comments of w and u , and rewards the driving contribution of u , v , and y . If only affirmed drivers shall be rewarded, then the relation to y (dashed lines) is skipped, because nobody has been attracted by its comment. The weights of interactions is calculated by the interaction reward function $f_i(dt, c1, c2)$, where dt is the discussion tree, and the interaction from the author of comment $c1$ to the author of $c2$ is rewarded. We initially set $f_i(dt, c1, c2) = \frac{1}{dist(c1, c2)}$, where $dist()$ determines the distance between two comments (direct replies have $dist = 1$). However, considering further comment attributes, including time intervals between a comment and its replies or the number of replies a single comment attracts, may improve the expressiveness according to trust. All interactions between two particular participants are aggregated and directed weighted relations are created in the graph model shown in Fig 2(b).

Algorithms 1 and 2 describe formally the mode of operation. According to (1) each edge in the interaction model can have various metrics. Currently we

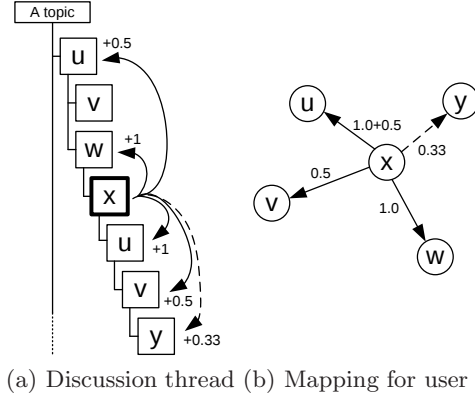


Fig. 2. Mapping from a discussion thread to the interaction network model.

apply *count*, which is the amount of interactions between two participants based to their discussion behavior, and *strength* which is the sum of the weights of all interactions between them. We utilize the function `incMetric(name, edge, value)` to increment the metric specified by `name` of the given `edge` by a certain `value`.

In Algorithm 1 relations from a comment’s provider to the providers of preceding comments are established due to their activator role. Algorithm 2 establishes relations to the providers of child comments due to driving behavior. The function `providerOf()` returns the identity of a comment provider, `parentCommentOnLevel()` determines the parent comment on the specified level ($level = dist(c1, c2)$), and `childCommentsOnLevel()` provides child comments.

Algorithm 1 function for rewarding the relations to the activators of a comment

Require: discussionThread, graphModel, comment, Ctx
commentProvider $\leftarrow$ providerOf(comment)
for level = 1 to configMaxLevelUp **do**
 parentComment $\leftarrow$ parentCommentOnLevel(comment, level)
 if $\nexists$ parentComment **or** providerOf(parentComment) = commentProvider **then**
 break
 end if
 parentCommentProvider $\leftarrow$ providerOf(parentComment)
 if $\nexists$ edge(commentProvider, parentCommentProvider, Ctx) **then**
 createEdge(commentProvider, parentCommentProvider, Ctx)
 end if
 incMetric(strength, edge(commentProvider, parentCommentProvider, Ctx), 1/level)
 incMetric(count, edge(commentProvider, parentCommentProvider, Ctx), 1)
 level $\leftarrow$ level + 1
end for **return** graphModel

Algorithms 1 and 2 are applied for every comment and reward the provider’s contribution to the overall discussion. This process can be further improved by additionally rewarding common communication patterns as well. This means, if

Algorithm 2 function for rewarding the relations to the drivers of a comment

Require: discussionThread, graphModel, comment, Ctx

```

commentProvider ← providerOf(comment)
for level = 1 to configMaxLevelDown do
  childComments ← childCommentsOnLevel(comment, level)
  if # childComments then
    break
  end if
  for all childComment ∈ childComments do
    childCommentProvider ← providerOf(childComment)
    if childCommentProvider = commentProvider then
      break
    end if
    if # edge(commentProvider, childCommentProvider, Ctx) then
      createEdge(commentProvider, childCommentProvider, Ctx)
    end if
    incMetric(strength, edge(commentProvider, childCommentProvider, Ctx), 1/level)
    incMetric(count, edge(commentProvider, childCommentProvider, Ctx), 1)
  end for
  level ← level + 1
end for return graphModel
  
```

v provides a comment replied by w , and v replies to w 's comment, then a real bidirectional communication can be observed.

6 Trust Network Model

6.1 Trust Inference

Similar to previous approaches [23, 24] trust is determined on top of the created interaction network, depending on the notions of confidence and reliability. We define that the confidence of user v in user w with respect to context Ctx can be derived from the previously described graph model by using a confidence function $c_{vw}^{Ctx} = f_c(G, v, w, Ctx)$.

Reliability, expressing the certainty of v 's confidence in w with respect to context Ctx , is determined by a reliability function $r_{c_{vw}^{Ctx}} = f_r(G, v, w, Ctx)$. The value of $r_{c_{vw}^{Ctx}} \in [0, 1]$ is basically influenced by the number and type of interactions which were used to calculate confidence, and expresses the reliability of the confidence value between totally uncertain and fully affirmed.

With the confidence of v in w and its reliability we calculate trust τ_{vw}^{Ctx} of v in w according to (2).

$$\tau_{vw}^{Ctx} = c_{vw}^{Ctx} \cdot r_{c_{vw}^{Ctx}} . \quad (2)$$

6.2 Trust Aggregation and Reputation

Aggregation of trust, often referred to as reputation, refers to (i) the combination of trust values of a group of users in one user to build a view of trust from a

community's perspective, or (ii) the combination of trust values calculated for different contexts between two users to get a notion of trust for a broader context or (iii) the combination of (i) and (ii) to get a kind of general community trust in one user.

Equation (3) is applied to determine aggregated trust T_a of a group $M = \{v_i\}$ of users in one particular user w with respect to a set of context elements $Ctxs$. The weighting factor calculated by f_a can be configured statically or obtained dynamically depending on individual properties of M 's elements, e.g., trust of long-term users have a higher impact on reputation than those of newbies.

$$T_{aMw}^{Ctxs} = \frac{\sum_{v_i \in M} \sum_{Ctx_j \in Ctxs} \tau_{v_i w}^{Ctx_j} \cdot f_a(v_i, w, Ctx_j)}{\sum_{v_i \in M} \sum_{Ctx_j \in Ctxs} f_a(v_i, w, Ctx_j)} . \quad (3)$$

6.3 Contextual Description

We distinguish two different subtypes of contextual elements: (i) *Provenance Context* describing the situation of interactions for which an edge is created, e.g., the domain of the discussion topic, or the used forum, and (ii) *Calculation Context* depicting the situation for which trust is calculated, e.g. for suggesting a discussion partner in a particular domain. Furthermore, calculation context may dynamically determine $f_i()$, $f_c()$, $f_r()$, and $f_a()$. The detailed design of the context models depend on the available information determined by the environment and area of application. We show an exemplary configuration of the trust network model in the evaluation part of this paper.

7 Evaluation

7.1 Preparing Evaluation Data

For the evaluation of our approach, we compare the output of the proposed algorithm with real users' opinions. Because our developed system is new and currently not utilized by a wide range of users, we need a dataset which offers structured discussions in various contexts and information about the real contribution of users. We fetched an appropriate dataset with the required characteristics from the famous Slashdot² community.

Slashdot is a platform which offers the ability to discuss a wide variety of topics classified in different subdomains. One nice feature is the moderation system allowing experienced users to rate the postings of other users on a scale between -1 and 5. We interpret this score as human feedback which provides information about the quality of comments and thus, when considering all posts, the average discussion quality of a person.

We developed a Web crawler to capture threaded discussions in the subdomains *Your Rights Online (yro)* and *Technology (tech)* from January 2007 to June 2008. We selected these two subdomains due to their diversity, expressing

² <http://slashdot.org>

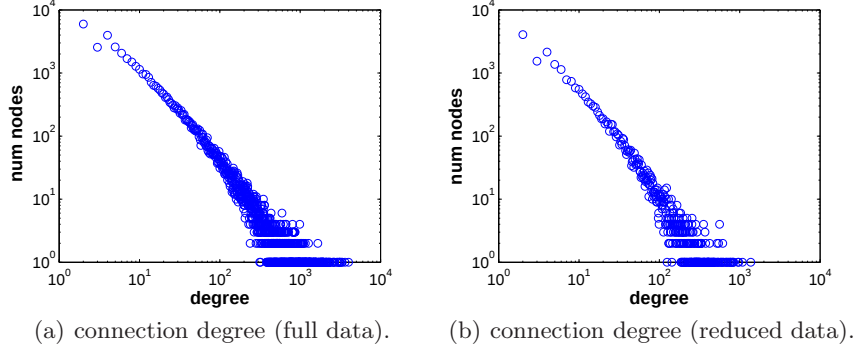


Fig. 3. Degree distribution.

different expertises of people discussing there. The subdomain in which a discussion takes place is reflected by the context of a discussion: $\text{ctx}=\{\text{yro} \mid \text{tech}\}$. Users may have established discussion relationships with respect to either *yro*, or *tech*, or both.

We have to ensure to compensate all impacts which degrade the quality of the data set and suitability for the tests. First, we remove all comments posted by anonymous users, because there is no meaningful way to map this data to particular nodes of the interaction graph model. Second, if not changed from the default settings, the Slashdot UI hides low scored comments automatically. Therefore, there is no way to distinguish if a particular comment is not replied because it is simply poor and not worth a discussion, or if it is not replied because it is hidden and thus never read. Hence, we remove low scored comments from the data set. Third, we remove all posts which potentially haven't been rated by others.

Initially the captured data set consists of 49.239 users and 669.221 comments in the given time period. After applying all steps of reduction we map the discussions to the graph model, consisting of 24.824 nodes and 343.669 edges. In the experiments we rank each user relatively to the others based on how much their discussion skills can be trusted by the rest of the community. Because our presented trust calculation method fully relies on the connectivity of a node within the graph, we have to ensure that the filtering procedures do not distort this property. Figure 3 shows the degree of connection for each node for the full data set and for the reduced one. The distribution follows a common power law function, and when applying the reduction steps, the characteristics of the user distribution and their connectivity basically do not change.

7.2 Trust Network Model Configuration

By applying the presented mapping approach we are able to grade discussion relationships between any two users v and w in the graph G with respect to the subdomain, reflected by context $\text{ctx}=\{\text{yro} \mid \text{tech}\}$.

Trust is determined by confidence and reliability as described in Sect. 6. To this end we define $f_c(G, v, w, \text{Ctx}) = \text{strength}$ to be a function which simply

returns the discussion strength from v to w in a specific subdomain. We define a notion of confidence from v in w to be fully reliable if there are at least max_{ia} interactions with respect to the same subdomain. If $f_r(G, v, w, Ctx) = \frac{count}{max_{ia}}$ is greater than 1 we set $f_r(G, v, w, Ctx) = 1$. We configure $max_{ia} = 10$ per year, which is the same amount of posts as identified in [14] to be required to calculate representative results. For trust aggregation we apply all single input trust values having the same weight $f_a(v, w, Ctx) = 1$.

For the sake of clarity we apply only the simple functions defined above, however, more complex functions can be set up, which consider similarities between subdomains, the amount of interactions compared to well-known community members or symmetry of trust relationships, just to name a few.

Furthermore, we set $configMaxLevelUp = 3$, $configMaxLevelDown = 3$ and reward bidirectional communication, i.e., post-reply-post patterns, with $bidir = 1$ extra point. By further increasing the number of levels for rewarding, the values indicating discussion strength between the users will increase as well. However, this does not highly influence the relative rankings of users.

7.3 Evaluation Approach

We evaluate our trust mining algorithm approach by comparing its results with trust values derived from the feedback of real users. We introduce the following terminology:

Link rank: The link rank of a user is calculated by our mining algorithm considering the strength of connections to others based on their nested comments within discussions. We interpret this measure as trust and argue, that it directly reflects a user’s willingness to share information and support others (driver role), and attitude to highly recognized contributions (activator role).

Score rank: The score rank of a user is calculated by averaging his/her posting scores, thus we utilize direct human feedback. We interpret the score rank as trust and argue, that users may trust posters with high average posting score more to deliver valuable contributions, than others.

Obviously both ranking methods rely on the same social properties, which reflect the value of contribution a user provides to the community.

First of all, we clarify that our proposed scoring method does not only depend on the number of posts and is completely different from simply giving reward points for every posted comment such as in common Internet forums. Figure 4 depicts the number of posts within 18 month of the top1000 linked users. However, there is a trend that frequently posting users are ranked higher, there is obviously no strong correlation between the link rank and the number of posts.

7.4 Experiments

Calculating Global Reputation. In our first experiment we determine global link ranks, built by aggregating the link strength values of all individual relations

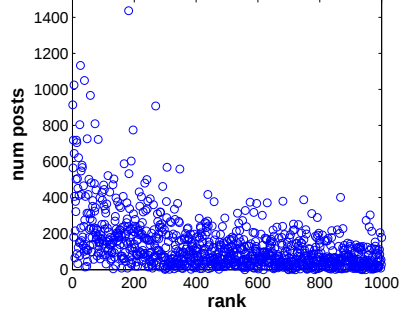


Fig. 4. Link rank compared to number of posts for top1000 linked users.

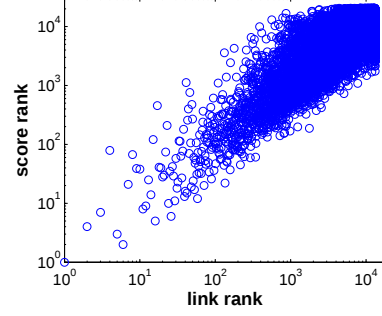


Fig. 5. Link rank compared to score rank for each user.

within the network for each user over all contexts. Besides this, we determine the global score rank as well. This means we rank each user two times: once with our algorithm based on discussion structures, and once based on humans' feedback score. For determining score ranks we degrade users' average scores by the factor $\frac{postcount}{numMinposts \cdot numMonth}$, if they posted less than $numMinposts$ posts a month to make sure that rarely posting users are not scored too high. During experiments we found out that $numMinposts = 10$ per month seems to be the value to reach the highest value for the Pearson correlation coefficient (0.77) between the results of both ranking methods for the given data set, as shown in Fig. 5.

We further calculate the Dice similarity coefficient depicted in (4), which is defined as the amount of elements included in both of two sets, in our case the sets of top scored users (TopXS) and top linked users (TopXL), where $X = \{10, 25, 50, 100, 1000\}$ determining the size of the sets.

$$s = \frac{2 \cdot |TopXS \cap TopXL|}{|TopXS| + |TopXL|} . \quad (4)$$

Table 1 shows how many percent of the top linked users and top scored users overlap after different time intervals. Obviously, the more data is used for trust calculation the more the resulting top linked users get similar to the top scored ones, which means we receive preciser results. After 18 month we finish with an overlap between 45 and 60 percent, for the top10 to top50 and approximately 65 to 70 percent for larger groups. Furthermore, we compare the amount of the top10 scored (Top10S) users who are also in the top25, top50, top100, and top1000 (TopXL) of the top linked users. The top10 scored users are the users scored best by others, and thus are most trusted to provide meaningful information. Table 1 shows that after 4 month 90 to 100 percent of the top10 scored users are included in the top50 linked users.

We conclude, that for the given data set we are able to find a similar set of users, who are trusted to post high quality comments, when ranked either by the

Table 1. Overlap similarities (OSim) of top linked and top scored users in percent.

OSim after month:	01	02	03	04	06	10	14	18
Top10 TopS10 in TopL10	10 10	30 30	30 30	30 30	40 40	50 50	60 60	50 50
Top25 TopS10 in TopL25	32 50	36 40	48 70	60 80	52 80	48 70	44 70	44 90
Top50 TopS10 in TopL50	28 50	34 60	40 80	50 90	54 100	58 90	62 100	60 100
Top100 TopS10 in TopL100	36 90	42 90	46 90	48 100	58 100	66 100	70 100	64 100
Top1000 TopS10 in TopL1000	61 100	61 100	66 100	64 100	64 100	66 100	68 100	70 100
number of users x1000	2.5	4.9	6.4	7.9	11	15	18	20

average of posting scores (scoreRank) or by the discussion structure and reply behavior (linkRank).

Enabling Context Dependent Trust Ranking. In a second experiment we consider the discussion context. Discussions in the utilized dataset take place either in subdomain **yro** or **tech**. We show that it is reasonable to calculate trust for particular situations reflected by context. We use six month of data from January 2008 to July 2008 because in this interval the amount of discussions and user distribution in both subdomains are nearly the same, thus results cannot be influenced by the number of posts. Then we rank each user two times with our algorithm, once for discussions in **yro** and once for **tech**. We rank only users with more than 10 posts, which we defined earlier as the absolute minimum for being trustworthy. There are in sum 14793 different users, where 5939 are only active in **yro** and 6288 in **tech**. Other users participate in discussions in both subdomains and thus, are ranked two times.

In Fig. 6 we compare how users are ranked with respect to both subdomains. There is an amount of approximately 40 users who are both, in the top100 wrt. **yro** and in the top100 wrt. **tech**, hence these people are highly trusted independent from the subdomain. However, there are around 60 users in the top100 of one subdomain but badly ranked in the other one, or not participating in discussions in the other subdomain at all. They are located in Fig. 6 in the top-left quadrant for **yro** and in the bottom-right for **tech** respectively.

We conclude that between the sets of top100 trusted users wrt. each subdomain there is less overlapping than diversity. These results show the usefulness of considering contextual data.

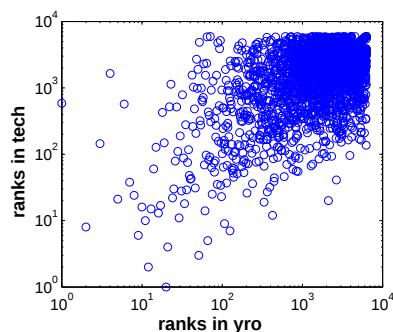


Fig. 6. Link ranks in different contexts.

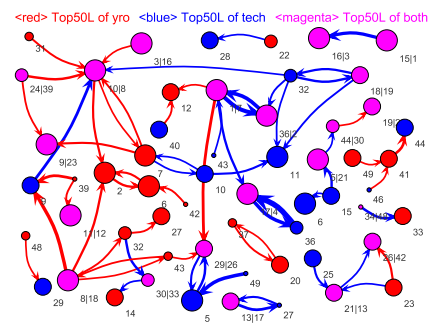


Fig. 7. Trust network (color online).

Determining Individual Trust. In contrast to reputation, which is mostly defined to be determined by the aggregated opinion of others, trust relies on personal experiences. As described in [14] in typical online communities there exist several clusters of users, which are tightly interconnected, but sparsely connected to other clusters.

Compared to most common reputation systems, which maintain only one global rank for each user from a global point of view, we are able to consider trust relations from an individual view as well. Hence, for a particular user there remain three possibilities to determine which users can be trusted: (i) trust users with highest reputation from a global view (with or without respect to context), (ii) trust the users who are directly connected strongest by utilizing local metrics (however, these users may have only an average global reputation) or (iii) combine both possibilities.

In Figure 7 we removed all connections with strength ≤ 5 , and all users who are either not in the top50L users of *yro* (red), *tech* (blue), or both (magenta), or not connected to anyone else. Therefore, the most trusted users and their strongest connections remain. The size of the circles representing users depends on their rank they received in either *yro* (red), *tech* (blue) or both (magenta), and the thickness of the lines reflect the connection strength. Obviously the trust graph splits into several only sparsely interconnected components. This justifies applying local metrics and selecting partners to trust with respect to strong personal relationships, instead of using global ranks only.

8 Conclusion and Future Work

In this paper we proposed a system for collecting communication data and performing trust determination within virtual communities. We demonstrated how our mining algorithm is able to determine trust relationships between users, after they contributed a while within the community. In the evaluation part we showed, that taking these trust relationships into account, the algorithm is able to find sets of trusted users, which are similar to sets of users top rated by humans. We further proved the usefulness of the concept of context awareness and considering local trust relationships.

In the next steps we plan to extend our framework to utilize more data sources. Especially in service-oriented collaborative working environments not only communication data, but task execution, resource utilization, and Web service invocation logs are further possible sources for better expressing the diversity of trust. We prepare our approach to be used in a project in the sector of networked enterprises to test it under real conditions and to enable research about influences of diverse interaction metrics on trust.

Furthermore, we plan to implement mechanisms to detect malicious attacks, such as artificially pushing a user's reputation rank. The evolvement of trust over time, currently neglected by our algorithm, may provide a valuable source of information about the long-term reputation of discussion participants.

References

1. Jøsang, A., Ismail, R., Boyd, C.: A survey of trust and reputation systems for online service provision. *Decision Support Systems* **43**(2) (2007) 618–644
2. The Economist: The role of trust in business collaboration. An Economist Intelligence Unit briefing paper sponsored by Cisco Systems (2008)
3. IBM: An architectural blueprint for autonomic computing. Whitepaper (2005)
4. Welser, H.T., Gleave, E., Fisher, D., Smith, M.: Visualizing the signatures of social roles in online discussion groups. *Journal of Social Structure* **8** (2007)
5. Nonnecke, B., Preece, J., Andrews, D.: What lurkers and posters think of each other. In: *HICSS*. (2004)
6. Meyer, K.A.: Face-to-face versus threaded discussions: The role of time and higher-order thinking. *Journal for Asynchronous Learning Networks* **7**(3) (2003) 55–65
7. McLure-Wasko, M., Faraj, S.: Why should i share? examining social capital and knowledge contribution in electronic networks. *MIS Quarterly* **29**(1) (2005) 35–57
8. Rheingold, H.: *The Virtual Community: Homesteading on the electronic frontier*, revised edition. The MIT Press (November 2000)
9. Massa, P.: A survey of trust use and modeling in real online systems (2007)
10. Gyngyi, Z., Garcia-Molina, H., Pedersen, J.: Combating web spam with trustrank. In: *VLDB*. (2004) 576–587
11. Page, L., Brin, S., Motwani, R., Winograd, T.: The pagerank citation ranking: Bringing order to the web. Technical report, Stanford University (1998)
12. Kleinberg, J.M.: Authoritative sources in a hyperlinked environment. *Journal of the ACM* **46**(5) (1999) 604–632
13. Kamvar, S.D., Schlosser, M.T., Garcia-Molina, H.: The eigentrust algorithm for reputation management in p2p networks. In: *WWW*. (2003) 640–651
14. Gomez, V., Kaltenbrunner, A., Lopez, V.: Statistical analysis of the social network and discussion threads in slashdot. In: *WWW, ACM* (2008) 645–654
15. Massa, P., Avesani, P.: Controversial users demand local trust metrics: An experimental study on epinions.com community. In: *AAAI*. (2005) 121–126
16. Skopik, F., Truong, H.L., Dustdar, S.: VieTE - enabling trust emergence in service-oriented collaborative environments. In: *WEBIST*. (2009) 471–478
17. Mui, L.: Computational models of trust and reputation: Agents, evolutionary games, and social networks. PhD thesis, Massachusetts Institute of Technology (December 2002)
18. Grandison, T., Sloman, M.: A survey of trust in internet applications. *IEEE Communications Surveys and Tutorials* **3**(4) (2000)
19. Marsh, S.P.: Formalising trust as a computational concept. PhD thesis, University of Stirling (April 1994)
20. McKnight, D.H., Chervany, N.L.: The meanings of trust. Technical report, University of Minnesota (1996)
21. Wanas, N.M., El-Saban, M., Ashour, H., Ammar, W.: Automatic scoring of online discussion posts. In: *WICOW, ACM* (2008) 19–26
22. Feng, D., Shaw, E., Kim, J., Hovy, E.H.: Learning to detect conversation focus of threaded discussions. In: *HLT-NAACL, The Association for Computational Linguistics* (2006)
23. Billhardt, H., Hermoso, R., Ossowski, S., Centeno, R.: Trust-based service provider selection in open environments. In: *SAC, ACM* (2007) 1375–1380
24. Huynh, T.D., Jennings, N.R., Shadbolt, N.R.: An integrated trust and reputation model for open multi-agent systems. *Autonomous Agents and Multi-Agent Systems* **13**(2) (2006) 119–154