

Designing a trust evaluation model for open-knowledge communities

Xianmin Yang, Qin Qiu, Shengquan Yu and Hasan Tahir

Xianmin Yang is a lecturer at the Institute of Education at Jiangsu Normal University. His main research interests are technology-enhanced learning, mobile and ubiquitous learning, and corporate e-Learning. Qin Qiu is a project manager in China Mobile Communications Corporation. Her research interests include digital rights management and information security. Shengquan Yu is a professor at the Institute of Modern Educational Technology at Beijing Normal University. His main research interests are integration of information and communication technology and teaching, ubiquitous learning, blended learning, and development of learning management system. Hasan Tahir is a PhD candidate at the University of Illinois at Urbana-Champaign. His research interest is human resource development. Address for correspondence: Prof Shengquan Yu, Institute of Modern Educational Technology, Room 215, Yanbo Building, Beijing Normal University, 19 Xijiekouwai Street, Beijing 100875, China. Email: yusq@bnu.edu.cn

Abstract

The openness of open-knowledge communities (OKCs) leads to concerns about the knowledge quality and reliability of such communities. This confidence crisis has become a major factor limiting the healthy development of OKCs. Earlier studies on trust evaluation for Wikipedia considered disadvantages such as inadequate influencing factors and separated the treatment of trustworthiness for users and resources. A new trust evaluation model for OKCs—the two-way interactive feedback model—is developed in this study. The model has two core components: resource trustworthiness (RT) and user trustworthiness (UT). The model is based on more interaction data, considers the interrelation between RT and UT, and better represents the features of interpersonal trust in reality. Experimental simulation and trial operation for the Learning Cell System, a novel open-knowledge community developed for ubiquitous learning, show that the model accurately evaluates RT and UT in this example OKC environment.

Introduction

Recently, open knowledge communities (OKCs) have become more popular. OKCs can be used as knowledge management tools and virtual learning environments for learners (Zeng, 2011). According to the content whether it allows collaborative editing, OKCs are divided into two categories. One is represented by the Wikipedia (<http://www.wikipedia.org/>) and the Learning Cell System (LCS, <http://lcell.bnu.edu.cn>) where any valid user can create new knowledge and coedit existing knowledges. The other one is represented by Baidu Zhidao (<http://zhidao.baidu.com/>), where users have no rights to edit knowledges created by others. They can usually view the existing knowledges and post comments. In this paper, OKCs are specifically referred to the former.

OKCs have inherent advantages in attracting user participation, encouraging collaboration and promoting the sharing of knowledge. However, openness also has side effects. A major problem currently underlying OKCs is information reliability (Yang, 2012). As the most popular online encyclopedia and an excellent example of OKCs, Wikipedia encountered the crisis of confidence inevitably (Lever, 2005; Luo & Fu, 2008; Seigenthaler, 2005). Wang (2009) pointed out that the feature of completely open content editing and organization of Wikipedia resulted in considerable

Practitioner Notes

What is already known about this topic

- The confidence crisis for open-knowledge communities (OKCs) is emerging.
- Scholars begin to apply the idea of trust in social networks to solve confidence crisis in OKCs.
- The design of trust evaluation models and the application to OKCs are still at an early stage.
- Current trust evaluation models are constructed simply based on user interaction data (eg, the editing history of a paper) and to apply such trust to the evaluation of resource quality.

What this paper adds

- We developed a new trust evaluation model named two-way interactive feedback model for evaluating trust in OKCs.
- The model takes into account editing history and typical interactions in OKCs, thereby improving the modeling accuracy.
- The model considers the dynamic interrelation between user trustworthiness (UT) and resource trustworthiness (RT), and models this interrelation by iterative cross-computation.
- The model adopts useful factors in existing trust evaluation models for network communication and electronic business (eg, the time-decay effect and punishment factor), thus better representing the interpersonal trust relation in reality.

Implications for practice and/or policy

- Introducing trust mechanisms is an effective means to solve crises of confidence in OKCs. OKCs should add trust evaluation functions to help users assess the trustworthiness of knowledge and other users.
- To ensure model comprehensiveness and integrity, a successful trust evaluation model should consider various typical interaction operations under Web 2.0.
- Instead of treating RT and UT separately, an effective trust evaluation model should consider their interrelation and model them by cross-computation.

doubt as to its quality and reliability. In recent years, Wikipedia has evolved a set of collaborative mechanism during the course of its development (Aniket & Robert, 2008), including dialogue pages, historic pages, recognition of quality problems in the community and quality control of entries. Although Wikipedia has introduced new mechanisms to improve its editing process, the quality of its data entries remains seriously questioned (Dondio, Barrett, Weber & Seigneur, 2006; Wang, 2009).

Due to the vast, complex and uncontrol features of the Internet, it is rather difficult to clean up those inferior resources and malicious users in OKCs. The central problems are how to evaluate the credibility of users and the knowledge they create and how to help users identify reliable information in OKCs. Trust modeling has greater value for fixing the above problems (Lucassen & Schraagen, 2011). Some researchers have begun to study the trust evaluation models in Wikipedia (Javanmardi, Lopes & Baldi, 2010; Maniu, Abdessalem & Cautis, 2011) and in virtual learning communities (Wang & Liu, 2007). Through trust evaluation models, users can find the right and reliable learning resources and connect with the right and reliable people. It is extremely useful for improving learning experience and promoting effective learning in OKCs.

The main objective of the present study is to redesign a new trust evaluation model for OKCs. Research questions can be stated as follows: (1) What influencing factors should be taken into account while designing the trust evaluation model for OKCs? (2) How to set the computing methods for computing users' and resources' trustworthiness? (3) How to prove the validity of the trust evaluation model for OKCs?

Literature

Features of OKCs

Por (1997) described knowledge communities as self-assembled knowledge-sharing networks joined by knowledge islands. OKCs are network communities assembled by various interactions (human–human, knowledge–knowledge and human–knowledge) for the purpose of creating, spreading and sharing knowledge. Common network communities are characterized by communicating and sharing information. These communities satisfy social and daily needs of users and provide a sense of belonging (Wang & Guo, 2003). Besides the characteristics mentioned above, OKCs have some other features as follows:

User diversity and variation in knowledge quality

The openness of OKCs leads to user diversity and variation in knowledge quality (Wang, 2009). Although most users are benign, there are malignant users who spread inferior knowledge and post malicious comments to earn community scores and ranks. Similarly, although group collaboration produces high-quality knowledge, inferior and untrustworthy knowledge also arises side by side.

Multiple interaction modes

As an ecosystem, an OKC has two key species including the user and the knowledge (Yang & Yu, 2011). Interaction is an essential means for information flow in the ecosystem, and this includes interactions between knowledge (eg, a citation, link), between users and knowledge (eg, browsing, comments, subscription, editing and bookmarking) and between users (eg, collaboration, reply, invitation and sharing).

User–knowledge interactions

In an OKC, users and knowledge are interrelated and interactive. Users produce, consume and transmit knowledge. Knowledge is the essential “food” consumed by users in an effort to self-upgrade (ie, gain skills and knowledge) (Yang & Yu, 2011). The trustworthiness of a user directly affects the trustworthiness of his or her contributions (created/shared knowledge). Conversely, the trustworthiness of knowledge shared by a user also affects the trustworthiness of the user him- or herself.

Partial resemblance to real communities

As a special form of virtual communities, OKCs resemble real social communities in certain respects. For example, interpersonal trust is affected by time and social events.

Trust evaluation models

Social trust is a belief in the honesty, integrity and reliability of others (Marsh, 1994). Trust evaluation model establishes a management framework of trust relationship between entities, involving expression and measurement of trust, and comprehensive calculation of trust value (Zhou, Pan, Zhang & Guo, 2008). With the emergence of a confidence crisis for OKCs, researchers began investigating solutions from the perspective of trust.

Most recent studies on trust in OKCs have focused on Wikipedia because of its popularity. Adler *et al* (2008) developed a method to assign trust values to Wikipedia articles according to the revision history of an article and the reputation of the contributing authors. Javanmardi *et al*

(2010) designed three computational models of user reputation based on user edit patterns and statistics extracted from the entire English Wikipedia history pages. Moturu and Liu (2009) proposed a model to calculate the trustworthiness of a Wikipedia article according to the degree of dispersion of the feature values from their mean. Lucassen and Schraagen (2010) showed that textual features, references and images are key indicators of the trustworthiness of Wikipedia articles. Maniu *et al* (2011) constructed a web of trust from the interactions of Wikipedia users, and analyzed user trustworthiness (UT) and its effect on readers and article classification. Halim, Wu and Yap (2009) proposed a method to improve the trustworthiness of Wikipedia articles according to credential information provided by a third party (eg, OpenID and OAuth). In their method, a third-party signature is added to all articles edited by a user, and the trustworthiness of an article is calculated by taking into account user information, such as education level, professional expertise or affiliation. Korsgaard (2007) developed a proxy Recommender System for Wikipedia, which allows users to rate articles and thus guide other users in terms of the trustworthiness of articles and users.

In summary, the above studies present two general approaches for carrying out trust research for OKCs (as represented by Wikipedia). The first approach attempts to construct a trust model simply based on user interaction data (eg, the editing history of an article) and to apply such trust to the evaluation of article quality. The second approach calculates UT and constructs a user trust network based on user interactions or user information from other sources. Both approaches ignore the intrinsic link between UT and information trustworthiness. However, UT can be an important factor of content trustworthiness and, in turn, content trustworthiness influences UT. For example, an article on instructional design written by an educationalist is usually credible, and publishing multiple excellent articles on instructional design enhances the trustworthiness of this author. In addition, the input for trust calculation should not be limited to editing history and contribution-based user interactions. It needs to include these common and abundant interactive data under the Web 2.0 framework, such as comment, subscription, bookmark, invitation and so on.

Since Marsh (1994) first introduced trust in social networks to computer networks, trust evaluation models have been widely studied and used in network communication (Denko, Sun & Woungang, 2008; Tian, Zou, Wang & Cheng, 2008; Yu, Zhang & Zhong, 2009) and electronic business (Jones & Leonard, 2008; Li & Wang, 2011; Wang, Xie & Zhang, 2010). Compared with current models for OKCs, these models consider the effects of time decay on trust evaluation, and some have introduced punishment factors to differentiate the effects of positive versus negative interactions (Liu, Yau, Peng & Yin, 2008; Wang *et al*, 2010). These design details suit the features of social trust in reality and provide a valuable reference for developing appropriate trust evaluation models for OKCs.

Overall, the design of trust evaluation models for OKCs are still at an infancy stage. There are three major drawbacks: (1) ignoring the intrinsic link between UT and information trustworthiness, (2) excluding some key interactions (eg, comment, subscription, bookmark and invitation) that affect trust calculation and (3) unrepresenting the interpersonal trust relation in reality (eg, the time-decay effect and punishment factor).

Trust evaluation framework

For an OKC, the objects to be evaluated include its users and knowledge. Accordingly, their credibilities should be evaluated separately. Knowledge may appear in different forms, such as entries in Wikipedia, items in Hudong and Knol pages in GoogleKnol. Despite their various forms, digital resources act as the carrier of knowledge in all OKCs. Therefore, in the present work, knowledge and a “learning resource” are considered synonyms; similarly, knowledge trustworthiness evaluation purports to evaluate the trustworthiness of learning resources.

Dong (2010) proposed qualitative method and quantitative method to design trust evaluation models in Grid Service. According to the quantitative method, four steps should be finished in order: (1) analyze factors that influence trust computing, (2) set basic assumptions for constructing model, (3) build trust evaluation model and (4) develop trust computing methods guided by the model. Because of its clear flow and operability, we will use this quantitative method to design the trust evaluation model for OKCs.

Analyze influencing factors

There are lots of factors we should consider carefully for constructing the trust evaluation model for OKCs. However, which factors actually affect the trustworthiness is the first question to figure out. By extending factors properly in current trust evaluation models for Wikipedia (Adler *et al.*, 2008; Korsgaard, 2007; Maniu *et al.*, 2011), the influencing factors of resource trustworthiness (RT) and UT are identified. Now we will discuss factors affecting RT and UT.

Factors affecting RT

RT can be evaluated in two ways. First, the systems provide trust evaluation functions allowing users to score RT directly, which is called direct evaluation. Second, RT can also be evaluated through other rich user–resource interaction data indirectly, which is called indirect evaluation.

Direct evaluation. There is no universal system for the visible evaluation of RT. OKCs currently use different trustworthiness indicators according to their individual features and requirements. Wikipedia now assesses articles in terms of reliability, objectivity, completeness and writing formality. The LCS assesses content in terms of accuracy, objectivity, completeness, citation formality and updating timeliness. Baidu Baike and Hudong Baike grade articles with up to five stars and invite readers to vote to what extent “This is helpful.”

Indirect evaluation. Invisible evaluation relies on records of user–resource interactions, such as collaborative editing, subscribing, bookmarking, browsing and citing. Obviously, different OKCs support different interaction modes. To some extent, user–resource interaction reflects user perception of RT. For example, an increase in user subscription of resource A suggests its attractiveness and trustworthiness.

Factors affecting UT

The trustworthiness of an OKC user is determined by the average trustworthiness of the resources that he or she created and also by his or her interactions with other users. Different interaction modes reflect invisible evaluations between users. Common factors of UT include the following.

Trustworthiness of user contributions. The trustworthiness of resources created by a user affects his or her own trustworthiness. For instance, if user A has created high-quality and reliable resources, his or her trustworthiness is increased.

Number of invitations/cancellations for collaboration. An invitation from user A to user B can be considered a positive vote for B by A. Conversely, cancelling an invitation is deemed a negative vote. If many users have invited user B to collaborate on resource editing, then user B has high trustworthiness.

Number of additions/cancellations as a friend. User A adding User B as a friend is considered a positive vote for B by A. Conversely, cancelling friendship is regarded as a negative vote. If many users have added user B as a friend, user B has high trustworthiness.

Number of accepted/rejected content revisions. The acceptance of content editing made by user A is a positive vote for A, and the rejection of editing is a negative vote. A greater probability of editing acceptance relates to higher trustworthiness.

Set assumptions

Inspired by the research findings of trust evaluation models in network communication and electronic business (Denko *et al*, 2008; Jones & Leonard, 2008; Wang *et al*, 2010), the following assumptions are determined to guide the model construction for OKCs. One of the most important principles is to reflect the trust relationship in the real society as much as possible.

Time-decay effect

It is assumed that trust decays with time (Li & Wang, 2011). The effects of user–resource and user–user interactions on trust are both time-dependent and time-limited. The effects of an interaction operation on trust decay with time. Thus, compared with earlier interactions, recent interactions have greater effects on trust.

Differential effect

It is assumed that interactions between one object (resource/user) and different users are associated with different effects on the trustworthiness of that object. Interactions with highly trusted users contribute better to the trustworthiness of that object and vice versa.

Participant size effect

Evaluations made by a large number of participants are assumed to be reliable. If a large number of users voted for a resource (in visible evaluation), the voting result is considered an accurate indication of the trustworthiness of that resource. Conversely, if few users voted for that resource, the result is regarded as uncertain or unreliable.

Two-way interactive effects

If a resource is cited, recommended, subscribed or bookmarked many times, it is considered well received and trusted by users. Similarly, if a user has made many accepted revisions, has created many credible resources, or has been invited to collaborate and been added as a friend many times, then he or she is considered well accepted by other users and his or her operations are thus believed to be more credible.

Build trust evaluation model

Employing the above influencing factors, analyses and assumptions, an OKC-oriented trust evaluation model (see Figure 1)—a two-way interactive feedback model (TIFM) is constructed.

The TIFM includes two core interactive components: UT and RT. In Figure 1, information boxes on both sides explain the factors affecting the two components, and the oval in the center lists the four assumptions behind the evaluation. As opposed to trust between peer nodes in P2P networks, trust is defined as global trust in this TIFM. The trustworthiness of a source represents the overall trustworthiness evaluation made by all community users for this resource. Likewise,

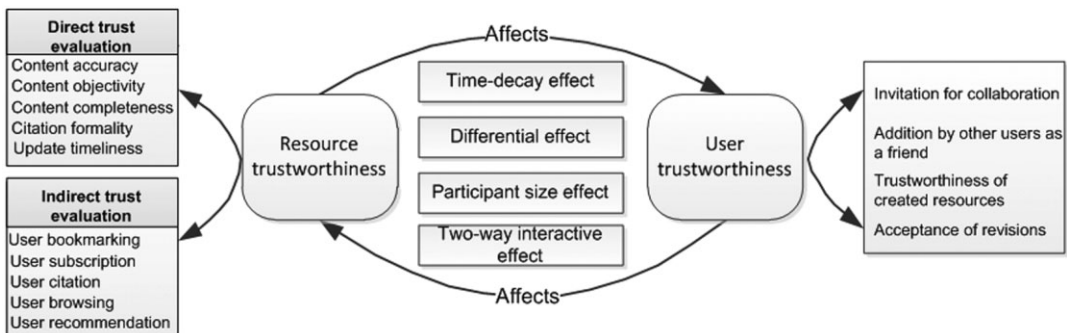


Figure 1: Framework of the two-way interactive feedback model as a trust evaluation model

the trustworthiness of a user means the overall trustworthiness evaluation made by all other users for that particular user.

Develop computation methods

In this part, the second research question is answered by developing the computation methods of RT and UT under the guidance of TIFM.

Trustworthiness: definition and calculation

Trustworthiness can be expressed by discrete values or continuous numbers. Discrete representation of trustworthiness resembles the features of human perception but lacks computability. Continuous representation, on the other hand, is amenable to modeling and computation, but it is not straightforward for users to rate UT/RT. In the TIFM, trustworthiness is expressed as continuous real numbers in the range (0, 1) and mapped into five trustworthiness ranks (full, strong, medium, weak, very weak) to provide users with a straightforward assessment of UT/RT.

RT

RT includes direct RT (DRT) and indirect RT (IRT). DRT is calculated from direct (ie, visible) user evaluations, and IRT is calculated from records of human–resource interactions. Shortly after creation of a resource, only a small number of users are expected to have posted direct evaluations of this resource. Therefore, at that stage, the contribution (or weight) of DRT to RT should be low. The weight (w) is a function that increases with the number of direct trust evaluations as the independent variable. It serves to dynamically adjust the relative importance of direct trust evaluations in RT calculation. With an increasing number of DRT evaluations, w increases and DRT accounts for an increasing proportion of RT. Equation 1 is used to calculate the resource trustworthiness by combining DRT and IRT.

$$RT = w \times DRT + (1 - w) \times IRT. \quad (1)$$

Furthermore, DRT can be expressed by an indicating factor set denoted by DRT indicating factors (DIFs) = ($dif_1, dif_2, dif_3 \dots dif_n$), where n is the total number of indicators and dif_i represents the i^{th} indicator. Individual OKCs focus on different aspects, and their DIFs are accordingly individualized. Even for a specific OKC, the DIF can be dynamically adjusted when required. Here, we give a relatively general DIF for RT: DIF = (content accuracy, content objectivity, content completeness, citation formality, updating timeliness).

Correspondingly, the weight set for this DIF can be given as $DW = (DW_1, DW_2, DW_3 \dots DW_n)$, where DW_i represents the weight of the i^{th} indicator among all indicators (ie, $\sum DW_i = 1$). The purpose for creating Equation 2 is to provide a formula to calculate DRT, which is part of Equation 1.

$$DRT = \frac{\sum_{j=1}^n \left(UT_j \times \sum_{i=1}^{|DIF|} (f(u_j, dif_i) \times DW_i) \times \partial^{t-t_j} \right)}{m \times itemScore}, \quad (2)$$

where m is the total number of users that posted a direct evaluation of the RT, U_j represents the j^{th} user that posted a direct evaluation, UT_j represents the trustworthiness of that user (j^{th}), $|DIF|$ is the number of indicators, and $f(u_j, dif_i)$ represents the score (ie, evaluation) that the j^{th} user gave to the i^{th} indicator. $itemScore$ represents the full mark for an indicator (eg, $itemScore = 5$ on a 5-point rating scale). Moreover, Equation 2 contains an item ∂^{t-t_j} representing the effect of time, where $\partial \in (0, 1)$, t is time in the trust calculation, and t_j is the time when the j^{th} user posted his or her comment (evaluation). The time difference for this item is given in months.

IRT is calculated from records of user–resource interactions (excluding the visible evaluations considered above). Similarly, IRT is expressed by another set of indicating factors denoted by IRT indicating factors (IIF) = ($iif_1, iif_2, iif_3 \dots iif_n$), where n is the total number of user–resource

Table 1: Interaction operations that affect indirect resource trustworthiness

Interaction mode	Positive interaction	Negative interaction
Recommend	Recommend a resource (thumb up)	Disrecommend a resource (thumb down)
Subscribe	Subscribe to a resource	Cancel subscription to a resource
Bookmark	Bookmark a resource	Cancel bookmarking a resource
Cite	Cite a resource	Cancel citation to a resource
Browse	Browse a resource	—

interaction modes (eg, recommend, subscribe, bookmark, browse and cite), iif_i represents the i_{th} mode (eg, iif_1 : user recommended this resource, iif_2 : user subscribed to this resource). Furthermore, each mode supports positive and negative interactions. Positive interactions are user operations that enhance the trustworthiness of a resource, such as making a subscription. Negative interactions are those reducing trustworthiness, such as cancelling a subscription. Table 1 summarizes common positive and negative interactions in OKCs.

The weight set for user–resource interactions can be written as $IW = (IW_1, IW_2, IW_3 \dots IW_n)$, where IW_i represents the weight of the i^{th} interaction mode among all indicators (ie, $\sum IW_i = 1$).

IRT is modeled according to number accumulation and subsequent normalization (to confine the result within $[0, 1]$). The purpose for creating Equation 3 is to provide a formula to calculate IRT before normalization. IRT before normalization (IRT_BN) represents the value of IRT before normalization.

$$IRT_BN = \sum_{i=1}^n (UT_i \times \alpha \times IW_j \times \partial^{t-t_i}), \quad (3)$$

$$\alpha = \begin{cases} 1 & \text{if this is a positive interaction} \\ -(1 + \rho) & \text{if this is a negative interaction} \end{cases}$$

where n is the total number of times that community users interacted with a resource, UT_i is the trustworthiness of the user that committed the i^{th} interaction operation and α is a regulatory factor [$\alpha = 1$ for positive interactions, and $\alpha = -(1 + \rho)$ for negative interactions; ρ is a punishment factor, $\rho \in (0, 1]$], and IW_j is the weight assigned to the i^{th} interaction operation according to its mode. The punishment factor is introduced to amplify the consequence of negative interactions because their effects are more intense than counterpart positive interactions in real social networks. Additionally, similar to user–user interaction (Equation 2), user–resource interaction also features a time-decay effect; correspondingly a time-decaying item, ∂^{t-t_i} , is incorporated into Equation 3.

The IRT derived from Equation 3 is transformed to a number within $(0, 1)$ by segment mapping according to Equation 4, which is part of Equation 1. The segment setup and the mapped values can be dynamically adjusted according to actual requirements.

$$IRT = \begin{cases} 0.0 & IRT_BN < a \\ 0.2 & a \leq IRT_BN < b \\ 0.4 & b \leq IRT_BN < c \\ 0.6 & c \leq IRT_BN < d \\ 0.8 & d \leq IRT_BN < e \\ 1.0 & e \leq IRT_BN \end{cases} \quad (4)$$

UT

UT is described by a quadruple $UT = \{UT_{res}, UT_{col}, UT_{fri}, UT_{rev}\}$, where UT_{res} is the trustworthiness component for a user calculated from the resources that he or she created, UT_{col} is a trustworthi-

ness component calculated from his or her interaction with other uses, UT_{fri} is a component calculated from friendship relations between this user and other community members, and UT_{rev} is a component calculated from his or her editing history in the community. Equation 5 is used to calculate the UT by combining UT_{res} , UT_{col} , UT_{fri} and UT_{rev} .

$$UT = UW_1 \times UT_{res} + UW_2 \times UT_{col} + UW_3 \times UT_{fri} + UW_4 \times UT_{rev}. \quad (5)$$

The relative importance of the four components is described by a user weight set: $UW = (UW_1, UW_2, UW_3, UW_4)$ ($\sum UW_i = 1$). UT_{res} is simply defined as the average RT for all resources that this user created. As part of Equation 5, Equation 6 is used to calculate UT_{res} .

$$UT_{res} = \frac{\sum_{i=1}^n RT(r_i)}{n} \quad (6)$$

where n is the total number of resources he or she has created, and $RT(r_i)$ is the RT for the i^{th} resource created by this user (calculated using Equation 1).

UT_{col} is modeled, similarly, according to the accumulation number followed by normalization. As part of Equation 5, Equation 7 is used to calculate UT_{col} .

$$UT_{col} = \text{normalize_utcol} \left(\sum_{i=1}^{|invited_col_log(u)|} (UT_i \times \alpha \times \partial^{t-t_i}) \right). \quad (7)$$

Normalization can be accomplished by segment mapping, and the exact `normalize_utcol` algorithm (Equation 7) can be decided and adjusted empirically. $|invited_col_log(u)|$ in Equation 7 represents the total number of invitations/cancellations for collaboration for this user. α is a regulatory factor [$\alpha = 1$ for an invitation to collaborate (positive action), and $\alpha = -(1 + \rho)$ for cancellation of such an invitation (negative action); ρ is a punishment factor, $\rho \in (0, 1)$]. Similarly, ∂^{t-t_i} represents the time-decay of the effect of an action on UT_{col} .

UT_{fri} is also modeled according to the accumulation number plus normalization. As part of Equation 5, Equation 8 is used to calculate UT_{fri} . Again, normalization can be done by segment mapping, and the actual `normalize_utfri` algorithm (Equation 7) can be decided empirically. `add_fri_log(u)` in Equation 8 represents the total number of times that this user is added/cancelled as a friend. α is a regulatory factor [$\alpha = 1$ for being added as friend (positive action), and $\alpha = -(1 + \rho)$ for the cancellation of friendship (negative action); ρ is a punishment factor, $\rho \in (0, 1)$]. ∂^{t-t_i} represents the time-decay of the effect of an action on UT_{fri} .

$$UT_{fri} = \text{normalize_utfri} \left(\sum_{i=1}^{|added_fri_log(u)|} (UT_i \times \alpha \times \partial^{t-t_i}) \right) \quad (8)$$

UT_{rev} is defined as the probability of a user's revisions being accepted. As part of Equation 5, Equation 9 is used to calculate UT_{rev} . In this equation, `rev_accept(u)` is the number of accepted revisions contributed by this user, and `rev_total(u)` the total number of revisions made by this user.

$$UT_{rev} = \frac{\text{rev_accept}(u)}{\text{rev_total}(u)} \quad (9)$$

Weight function setting

The weight (w) of DRT is an increasing function and, ideally, should reach a plateau (ie, 1) following a Sigmoid Curve. That is, w should initially increase slowly when there are only a small number of direct evaluations (denoted by n), then increase sharply with n , and finally

experience slow growth as it approaches the value of 1. Correspondingly, the growth rate (dw/dn) of w increases to a peak value and decreases thereafter. In our implementation, we let the growth rate increase linearly with n with a slope of a ($a > 0$) until $n = m$ (where m is a positive integer, corresponding to n at which the growth rate peaks). When $n > m$, we let the growth rate decrease linearly with a slope of $-a$. When $n > 2m$, we let the growth rate equal zero. Additionally, we let $w = 0$ when $n = 0$, and $w = 1$ when $n > 2m$. Thus, we can write

$$\frac{dw}{dn} = \begin{cases} an, & 1 \leq n \leq m; \\ -an + 2am, & m < n \leq 2m; \\ 0, & n > 2m. \end{cases}$$

Solving this differential equation, we have

$$w = \begin{cases} \frac{n^2}{2m^2}, & 1 \leq n \leq m; \\ -\frac{n^2}{2m^2} + \frac{2n}{m} - 1, & m < n \leq 2m; \\ 1, & n > 2m. \end{cases} \quad (10)$$

The actual value of m can be decided according to requirements and implantation strategies for different OKCs.

Weight setting

The TIFM has several weight sets for trustworthiness indicators, and the assignment of weight values affect the accuracy and validity of model. Here, we determine these weight values employing an analytic hierarchy process (AHP). An AHP is an effective technique for solving complex problems involving subjective judgment (Saaty, 1990). Moreover, the AHP can rationally determine the weight assigned to each criterion and has thus been used in many studies requiring weight decisions (Lin, 2010; Zhao & Qi, 2008; Zheng, Wang, Guo & Yang, 2010).

For the TIFM, the weight sets that should be decided are DW for DRT evaluation ($DW = [DW_1, DW_2, DW_3, DW_4, DW_5]$), IW for IRT evaluation ($IW = [IW_1, IW_2, IW_3, IW_4, IW_5]$) and UW for UT evaluation ($UW = [UW_1, UW_2, UW_3, UW_4]$). DW_1 – DW_5 are weights for the content accuracy, content objectivity, content completeness, citation formality and the timeliness of updates, respectively. IW_1 – IW_5 are weights for resource recommendation, subscription, bookmarking, browning and citation by another user, respectively. Moreover, UW_1 – UW_5 are weights for the trustworthiness of a user contribution (resources created by the user, termed *resource creation*), user–user collaboration (simply referred to as *collaboration*), friendship relation (simply referred to as *friendship*) and the history of content revision (referred to as *revision history*), respectively.

In model implementation, we used an AHP package (Yaahp 0.5.2, Foreology Software Ltd., Beijing, China) to determine the above weights. Briefly, eight educationalist were asked to rate the relative importance of the weight items. The results were transformed to the scales suggested by Saaty (2008) to construct judgment matrices. The matrices were analyzed in a consistency test, and the weights were calculated by taking the normalized column average. Tables 2–4 show matrices and the weights determined employing this method.

Table 2 shows that the judgment matrix for DW gives a consistency ratio (CR) of less than 0.1, indicating satisfactory consistency. The weight values were rounded to two decimal points; thus $DW = (0.50, 0.30, 0.11, 0.04, 0.05)$.

Table 3 shows that the judgment matrix for DW also gives a CR of less than 0.1, indicating satisfactory consistency. The weight values were rounded to two decimal points; thus $IW = (0.28, 0.18, 0.44, 0.03, 0.07)$.

Table 2: Judgment matrix for DW and consistency ratios

Indicator	Content accuracy	Content objectivity	Content completeness	Citation formality	Update timeliness	Weight
Content accuracy	1	3	6	9	7	0.5014
Content objectivity	1/3	1	5	7	8	0.3043
Content completeness	1/6	1/5	1	4	5	0.1132
Citation formality	1/9	1/7	1/4	1	1/2	0.0354
Update timeliness	1/7	1/8	1/5	2	1	0.0457

Note: Consistency index (CI) = 0.0942; consistency ratio (CR = CI/RI) = 0.0942/1.12 = 0.0841; RI (random index): refer to the table for the average consistency index for 1st–10th-order matrices.

Table 3: Judgment matrix for IW and consistency ratios

Indicator	Recommendation	Subscription	Bookmarking	Browsing	Citation	Weight
Recommendation	1	2	1/2	7	5	0.2795
Subscription	1/2	1	1/3	6	4	0.1811
Bookmarking	2	3	1	8	7	0.4394
Browsing	1/7	1/6	1/8	1	1/4	0.0325
Citation	1/5	1/4	1/7	4	1	0.0674

Note: consistency index (CI) = 0.0513; consistency ratio (CR = CI/RI) = 0.0513/1.12 = 0.0458; RI (random index): refer to the table for the average consistency index for 1st–10th-order matrices.

Table 4: Judgment matrix for UW and consistency ratios

Indicator	Resource creation	Collaboration	Friendship	Revision history	Weight
Resource creation	1	3	6	1	0.3919
Collaboration	1/3	1	5	1/3	0.1643
Friendship	1/6	1/5	1	1/6	0.0519
Revision history	1	3	6	1	0.3919

Note: consistency index (CI) = 0.0397; consistency ratio (CR = CI/RI) = 0.0397/0.90 = 0.0441; RI (random index): refer to the table for the average consistency index for 1st–10th-order matrices.

The judgment matrix for DW (Table 4) gives a CR of less than 0.1, indicating satisfactory consistency. The weight values were rounded to two decimal points; thus UW = (0.39, 0.16, 0.05, 0.39). UW was finely adjusted to (0.39, 0.16, 0.06, 0.39) to satisfy $\sum W_i = 1$.

Solving “the chicken or the egg” problem through cross-computation

Obviously, the TIFM involves an interdependence: the calculation of RT requires UT and vice versa. This poses the dilemma of “Which comes first—the chicken or the egg?” Here, we solve this problem by iterative cross-computation/approximation. The basic idea is to calculate the trustworthiness values for all users and resources in a given community by iterative computation, until the differentials between adjacent computation results (for all UT and RT items) become smaller than preset maximum errors (ie, stability is reached). Figure 2 explains the procedures. First, data are input into the model and the maximum error was set as <0.1, and the iteration counter (n) is initialized at zero. The trustworthiness for all users [$UTS_0 \in (0, 1)$] and resources [$RTS_0 \in (0, 1)$] is initialized (see below). The next UT (UTS_{n+1}) and RT (RTS_{n+1}) are then calculated iteratively using Equations 1 and 4. The differential UT and RT between adjacent iteration steps

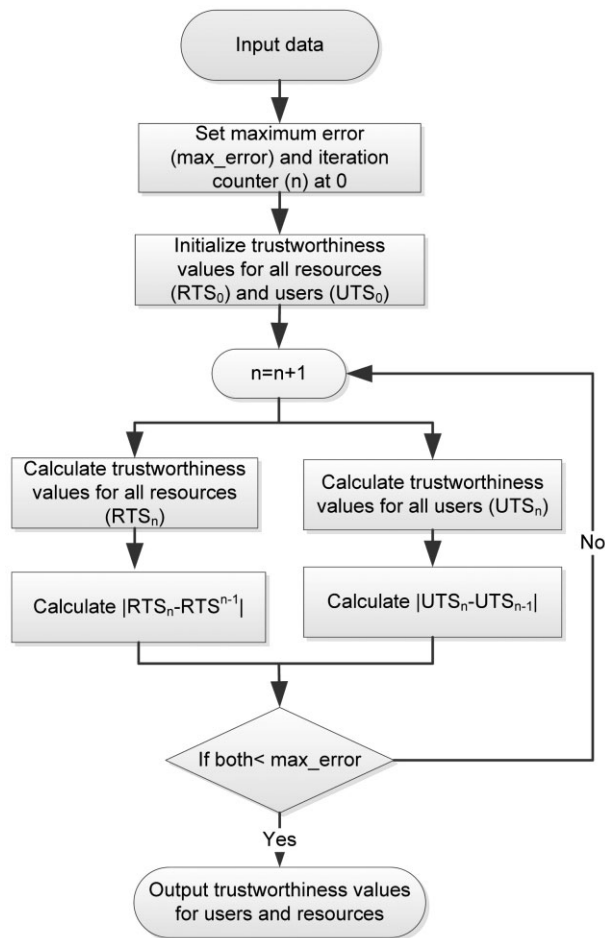


Figure 2: Flowchart showing procedures (iterative cross-computation) to solve “the chicken or the egg” dilemma in the TIFM

are calculated and compared with `max_error`. If all differentials are less than `max_error`, the iteration ends and the UTS and RT values are output. Otherwise, the iteration continues.

Validation of trust evaluation model

Methodology

In order to answer the third research question proposed in the introduction, the methods of experimental simulation and trial operation were adopted to check the validity of TIFM.

The method of experimental simulation is widely used in the area of computer science and mathematics (Shuang, 2011; Wu, Zhang & Xu, 1999). Researchers often use this method to validate the effectiveness of algorithms and formulas. In this study, two experimental simulations were performed.

Experimental simulation 1 (ES1) was done to examine the validity of the method of solving “the chicken or the egg” problem through cross-computation.

Experimental simulation 2 (ES2) was done to check whether the weight function (Equation 10) followed a Sigmoid Curve in the process of trust growth.

Additionally, in order to check the effect of TIFM in practical application, an actual TIFM system was realized for the LCS and trialed for 6 months. The trustworthiness values generated by the TIFM at the end of the trial were compared with counterpart values given by experts to check the validity of this model. Furthermore, a network interview was conducted for gathering the feedback information of TIFM from practical users.

Experimental simulation

ES1: Solving the “chicken or the egg” interdependence

The procedure of simulation 1 was as follows:

1. Create a resource group (res_num = 1000) and a user group (user_num = 100);
2. Initialize RT (init_rt = 0.1), UT (init_ut = 0.1) and maximum error (max_error = 0.0001);
3. Randomly configure all interaction records (eg, resource subscription, user being invited to collaborate);
4. Apply the iteration program until convergence (all differentials are less than max_error);
5. Save RT set (RTS) and UT set (UTS).
6. Vary the parameters (eg, user_num, res_num, init_rt, init_ut and max_error) and repeat steps 3–5;
7. According to different values of UTS and RTS, examine simulation results.

A Java program was created to realize the above procedure, and run on a personal computer (Intel Core i5 CPU 2.4 GHz, Intel Corporation, Santa Clara, CA, USA; 32-bit OS, 2.00-GB RAM) to examine the performance of this method under different conditions.

Simulation 1: trust calculation under constant interaction records but different initial RT and UT values

Interaction data obtained in Step 3 were stored, and the program was run under three init_rt and init_ut configurations (Group 1: init_rt = init_ut = 0.01; Group 2: init_rt = 0.01, init_ut = 0.1; Group 3: init_rt = 0.1, init_ut = 0.01). From Figures 3 and 4, we can see three curves representing three different groups actually overlap together. It showed that regardless of the setup of initial values, TIFM gave the same UT and RT values as long as the interaction data were kept constant.

Simulation 2: iteration number and computation overhead under constant resource/user numbers and interaction records but different maximum error values

The resource and user numbers were kept constant (res_num = 1000, user_num = 100). The initial RT and UT were also kept constant (init_rt = 0.01, init_ut = 0.1). The program was run

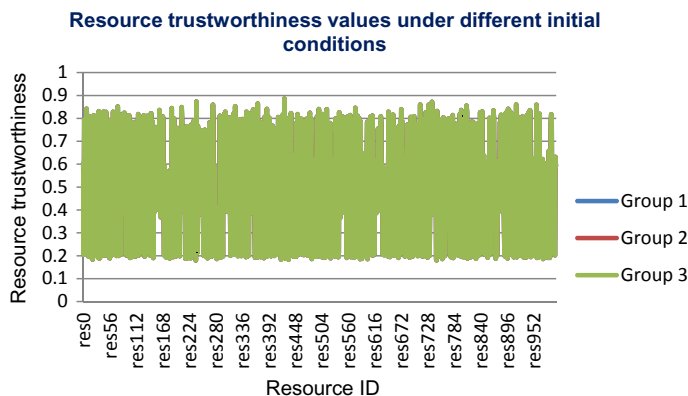


Figure 3: Final RT values calculated under different initial RT and UT conditions

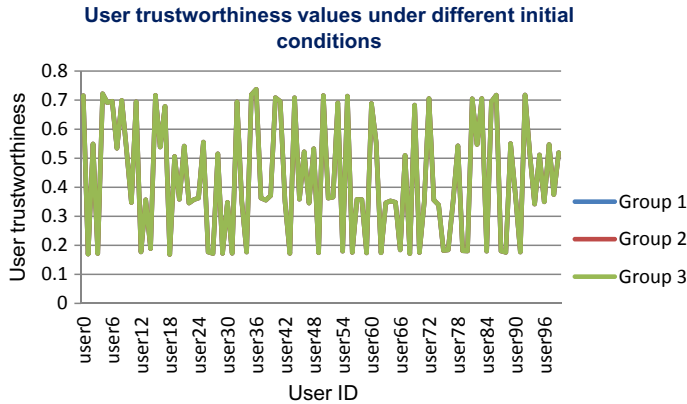


Figure 4: Final UT values calculated under different initial RT and UT conditions

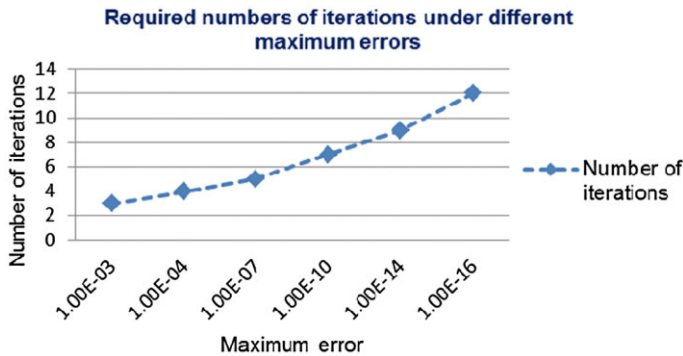


Figure 5: Required number of iterations under different maximum-error configurations

under six max_error configurations (1×10^{-3} , 1×10^{-4} , 1×10^{-7} , 1×10^{-10} , 1×10^{-14} and 1×10^{-16}). Experiments indicated max_error to be an important parameter affecting the required iteration number (see Figure 5). With an increase in precision (eg, smaller max_error), the number of required iterations increased and the computation overhead was proportional to the number of iterations (see Figure 6).

Simulation 3: Iteration numbers under different initial RT/UT values but otherwise constant configurations (resource/user numbers, interaction records and max-error)

The program was run under 50 groups of different RT/UT values, with other parameters kept constant (res_num = 1000, user_num = 100 and max_error = 0.0000000001). Experiments indicated that the iteration number and overhead were stable regardless of the setup of initial RT/UT values (see Figure 7).

ES2: validating the w function

As described in the section of weight function setting, ideally w should increase following a Sigmoid Curve. In our implementation, m was set at 20, and Equation 10 was realized using Java. As expected, the resulting w function (see Figure 8) increased with n (number of comments) following a Sigmoid Curve and reached 1 when $n > 2m$.

We also analyzed the relation between resource trustworthiness indices (RT, DRT and IRT) and n . It was found that when n was small, RT and IRT were essentially the same (see Figure 9),

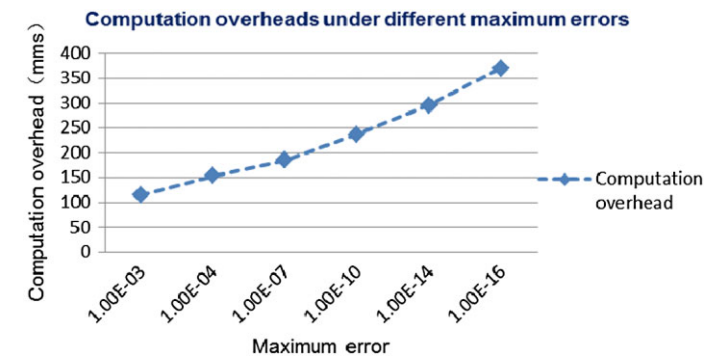


Figure 6: Computation overheads under different maximum-error configurations

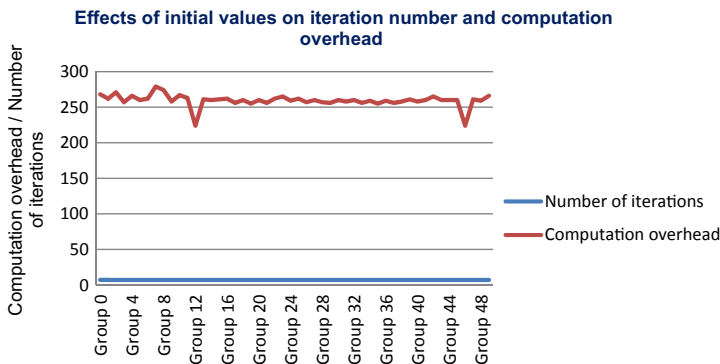


Figure 7: Required numbers of iterations and overheads for different initial RT/UT values

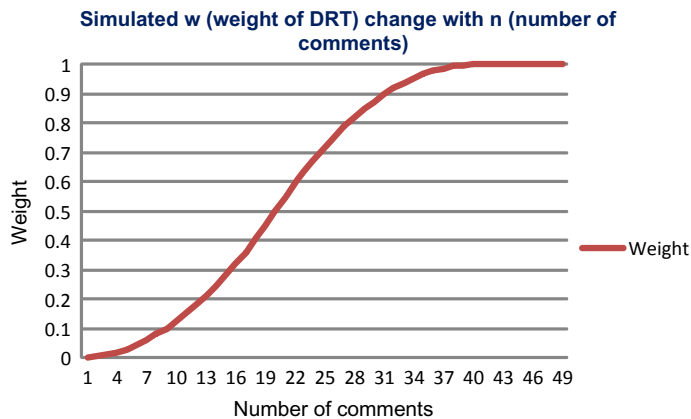


Figure 8: Simulated w (weight of DRT) change with n (number of comments)

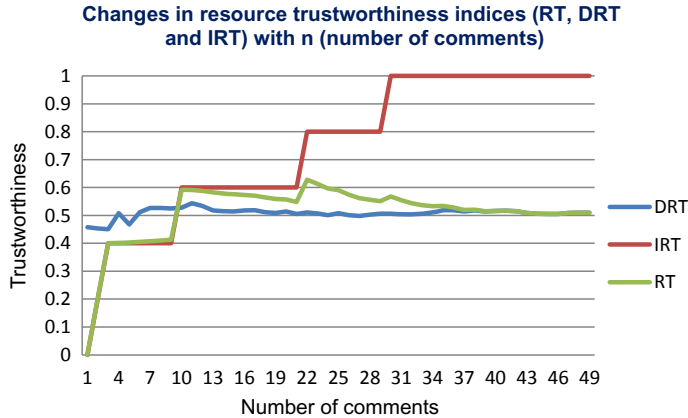


Figure 9: Changes in resource trustworthiness indices (RT, DRT and IRT) with n (number of comments)

indicating RT was dominated by DRT. With an increase in n , RT approached DRT. When $n > 2$ m, RT overlapped with DRT, indicating RT was dependent entirely on DRT under such a condition.

Trial operation—validation for LCS

Experiment design

The TIFM was applied to the LCS platform. The trustworthiness values generated by the TIFM were compared with counterpart values given by experts to check the validity of our model. The comparison experiment was conducted according to the steps outlined in Figure 10.

Thirty learning cells were randomly selected in an LCS environment employing subject of education technology. First, the trustworthiness values of these learning cells and creators were calculated using the TIFM and expressed as two sets: RTS_1 (resource trustworthiness set 1) and UTS_1 (user trustworthiness set 1). RTS_1 and UTS_1 were then transformed to five trustworthiness ranks. Five education technology experts who have actively participated in the LCS (according to login records) were then asked to rate the trustworthiness of the selected learning cells and their creators using a 5-point Likert scale. The expert results were averaged for each user/learning cell and expressed as two sets: RTS_2 (resource trustworthiness set 2) and UTS_2 (user trustworthiness set 2). RTS_1 – RTS_2 and UTS_1 – UTS_2 were analyzed in a kappa consistency test.

RT evaluation was done by 5-point grading based on direct trustworthiness comment tools currently available in the LCS. The evaluation considered five factors: content accuracy, content objectivity, content completeness, citation formality and update timeliness. For UT evaluation, user information (resource creation, invitation for collaboration, invitation as a friend and acceptance of content revision) was retrieved from the LCS for expert review. Excel questionnaire spreadsheets were prepared and e-mailed to experts, who were invited to rate the selected users and resources. In order to ensure the credibility, the unified guide information was provided to help experts evaluate RT and UT. Furthermore, another five sample resources and five sample users were provided for experts' trial evaluations. The results were returned by e-mail.

Besides the above comparison experiment, in order to validate the effect of TIFM further in judging the credibilities of resources and users in the LCS, a network interview was conducted. Thirty most active users registered in the LCS were randomly selected as the interviewees, whose login numbers were over 50 times in the last month. The interview outline including three questions (see Table 5) was sent to the interviewees by e-mail. Moreover, all the interview results were collected by e-mail.

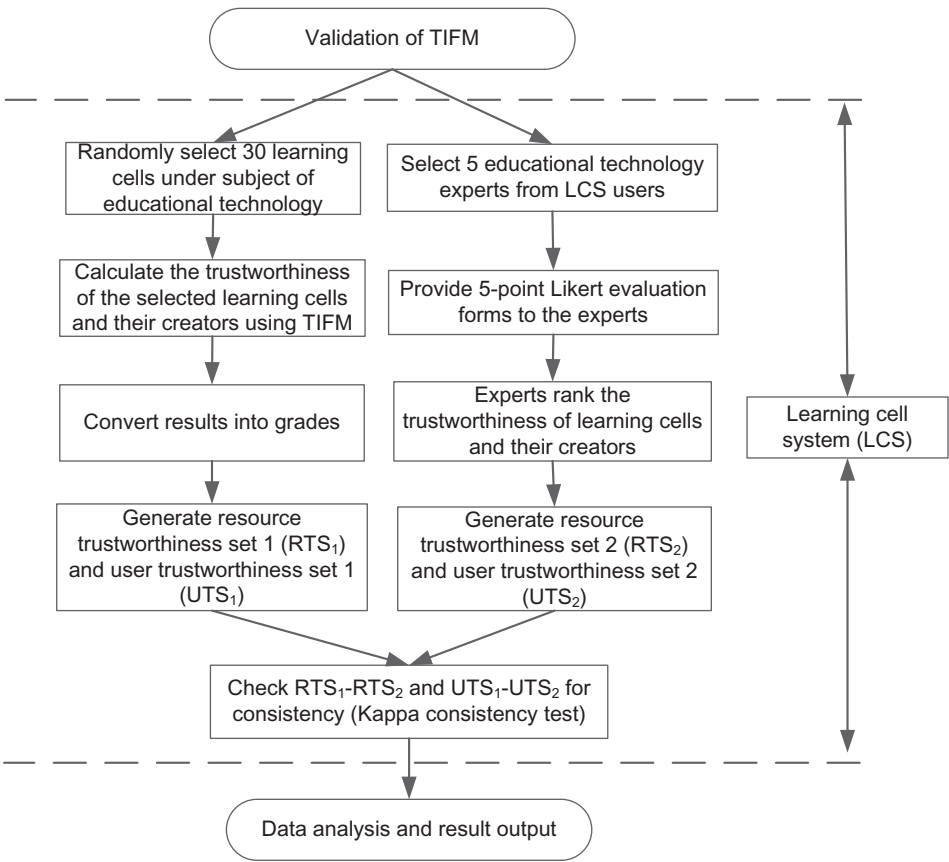


Figure 10: Flowchart showing the design of the experiment to validate the TIFM for the LCS

Table 5: The interview outline

ID	Question
1	What do you think of the trust evaluation function provided in the LCS?
2	How do you use the trust values of resources and users in the LCS?
3	Any suggestions for improving the trust evaluation function in the LCS?

Result of comparisons

Expert evaluations and TIFM results for RT and UT were analyzed with SPSS13.0 (SPSS, Chicago, IL, USA), as shown in Tables 6 and 7.

Table 8 shows the acceptable ranges for kappa. Kappa consistency tests showed that the expert evaluations and TIFM results were significantly consistent for both RT (kappa = 0.676 > 0.61; 95% confidence interval, $\alpha < 0.05$) and UT (kappa = 0.682 > 0.61; 95% confidence interval, $\alpha < 0.05$), demonstrating our evaluation method to be reliable.

Analyses produced kappa > 0.61 for both UT and RT evaluations, showing our TIFM developed here to be a reliable trust evaluation model for an OKC. However, future adjustment and improvement are needed according to actual issues during LCS operation.

In this study, the TIFM gave trustworthiness values as real numbers within (0, 1), while users rated UT and RT on a 5-point scale. Thus, conversion of continuous trustworthiness values to

Table 6: Kappa consistency test results for RT

	Value	Asymptotic standard error*	Approximate t†	Approximate significance
Measure of agreement kappa n of valid cases	0.676 30	0.105	6.308	0.000

*Not assuming the null hypothesis.

†Using the asymptotic standard error assuming the null hypothesis.

Table 7: Kappa consistency test results for UT

	Value	Asymptotic standard error*	Approximate t†	Approximate significance
Measure of agreement kappa n of valid cases	0.682 16	0.135	5.328	0.000

*Not assuming the null hypothesis.

†Using the asymptotic standard error assuming the null hypothesis.

Table 8: Acceptable ranges for Kappa

Kappa	Consistency	Kappa	Consistency
<0.00	Poor	0.41~0.60	Moderate
0.00~0.20	Very weak	0.61~0.80	Significant
0.21~0.40	Weak	0.81~1.00	Optimum

discrete values is a key factor affecting the operation of the TIFM, and different conversion functions obviously would give different kappa values. We conducted preliminary research to search for the optimum conversion function. Ten graduate students studying education technology were asked to evaluate the trustworthiness of 16 users and 30 learning cells in the LCS. Parameters in the conversion function were adjusted according to the evaluations made by the graduate students. The final functions (convert_r for RT conversion and convert_u for UT conversion) are shown below.

$$\text{convert}_r(\text{RT}) = \begin{cases} 1 & \text{RT} < 0.10 \\ 2 & 0.10 \leq \text{RT} < 0.30 \\ 3 & 0.30 \leq \text{RT} < 0.50 \\ 4 & 0.50 \leq \text{RT} < 0.95 \\ 5 & 0.95 \leq \text{RT} \leq 1.00 \end{cases}$$

$$\text{convert}_u(\text{UT}) = \begin{cases} 1 & \text{UT} < 0.05 \\ 2 & 0.05 \leq \text{UT} < 0.10 \\ 3 & 0.10 \leq \text{UT} < 0.35 \\ 4 & 0.35 \leq \text{UT} < 0.70 \\ 5 & 0.70 \leq \text{UT} \leq 1.00 \end{cases}$$

Although initial results are promising, we realize that the LCS has been running for a relatively short period, and interaction data are increasing in volume and complexity. Therefore, the

conversion functions will be adjusted and improved in the future. Moreover, future studies should consider a potential attack on trustworthiness to enhance the robustness of the TIFM.

Result of interviews

Among 30 randomly selected registered users in the LCS, 25 were interviewed with a response rate of 83.3%. In question 1, 20 users mentioned the trust evaluation function was a distinguishing feature, and approved its usefulness. In question 2, most users mentioned that they mainly used the trust values to help judging resources' qualities and selecting high credible resources. Some interviewees said they would like to subscribe and collect high credible resources. In addition, half of them said they used to see trust values of users while inviting friends or accepting invitations. In question 3, the interviewees mainly gave two suggestions. First, the trust values of resources and users should be placed in a more prominent position. Second, it is better to reduce five trustworthiness ranks (full, strong, medium, weak and very weak) to three trustworthiness ranks (strong, medium, weak) for enhancing the ease of use.

According to the interview data, it can be found that TIFM plays an important and desired role in judging credibilities of resources and users in the LCS. The trust evaluation function developed based on TIFM is very useful to help users selecting high quality resources and highly credible users as friends. It will expediate the elimination of inferior resources and improve the whole quality of a knowledge ecosystem.

Discussion

Effectiveness and advantages of TIFM

On the whole, the result of the trial run for the LCS shows that the TIFM can evaluate UT and RT accurately in OKCs. It also can be found that the application effects of TIFM are also approved and praised by practical users through analyzing interview data. The result of ES1 indicates that the cross-computation is a useful way to solve "the chicken or the egg" problem. This way is also successfully applied in the PageRank algorithm (Page, Brin, Motwani & Winograd, 1999) for Google. Additionally, by means of ES2 the weight function in Equation 1 is proved to be applicable to the TIFM with a desired changing curve.

Compared with current trust evaluation models (Adler *et al.*, 2008; Javanmardi *et al.*, 2010; Moturu & Liu, 2009) for OKCs, the TIFM has three advantages. First, it takes into account editing history and typical interactions in OKCs, thereby improving the modeling accuracy. Second, it considers the dynamic interrelation between UT and RT, and models this interrelation by iterative cross-computation. Third, the model adopts useful factors in existing trust evaluation models for network communication and electronic business. The integration of the time-decay effect and punishment factor can make trust evaluation models represent the interpersonal trust relationship more accurate in reality (Liu *et al.*, 2008; Wang *et al.*, 2010).

Cold start problem and its coping strategy

TIFM works by relying on the interaction data. For a new registered user or a new created resource, no interaction history could be identified for computing trust. Then the problem of cold start appears. Cold start is a common phenomenon that happens in the area of personalized resources recommendation (Drachslar, Hummel & Koper, 2008; Luo, Wang, Du, Liu & He, 2007). Cold start means one computing method loses efficacy without initial data set. Currently, in order to solve the cold start problem encountered by TIFM, any new user or new resource will be assigned zero as the initial trust value. Meanwhile, the initial trust value of zero will be recognized as unknown trust level in OKCs. Along with the interaction data growing, TIFM can utilize these data to compute trust values of resources and users more precisely. The method mentioned above is just a preliminary treatment to the cold start problem in OKCs. Cold start has become a hot research topic in recommendation of personalized resources, and scholars (Guo & Deng, 2008; Li

& Liang, 2012; Sun, He & Zhang, 2012) have done a lot of research on this topic. Therefore, we should study the problem further by referencing their works.

Conclusion

In this work, a new OKC-oriented trust evaluation model named TIFM is proposed. Three research questions are examined: (1) What influencing factors should take in account while designing the trust evaluation model for OKCs? (2) How to set the computing methods for computing users' and resources' trustworthiness? (3) How to prove the validity of the trust evaluation model for OKCs? Results of experimental simulations and a trial run for the LCS confirms the effectiveness of this model.

Those who are interested in online knowledge management and responsible for knowledge dissemination in organizations, like the chief knowledge officer, the knowledge manager and the architect of knowledge management system, will find TIFM very useful in determining the credibility of information. TIFM would be valuable for the design and development of OKCs. By integrating TIFM into OKCs, the quality of knowledge ecosystems provided by OKCs would be improved. Users will have better experiences with a high trust level towards OKCs, which will promote the spreading of credible knowledge inside and outside the organizations.

Some implications for developing OKCs especially in designing trust evaluation models have also been found:

1. Keep its integrity. To ensure model comprehensiveness and integrity, a successful trust evaluation model should consider various typical interaction operations under Web 2.0. In addition, direct evaluation and indirect evaluation should both be taken into account.
2. Keep it consistent with real world as far as possible. Trust in the virtual world has a lot of similarities to that in the real world (Jones & Leonard, 2008; Wang *et al.*, 2010). So the design of the trust evaluation model for OKCs should consider time-decay effect and punishment factor.
3. Keep RT and UT correlated. Instead of treating RT and UT separately, an effective trust evaluation model should consider their interrelation and model them by cross-computation.

Until July 2012, the LCS had been working for 1 year. The generated data size was relatively small with 11 000 resource entities and 7000 registered users. Therefore, long-term tracking and analyses are required to further assess the performance of the TIFM. Moreover, attack-resistant mechanisms are unavailable in the current model and need to be developed in future studies. In addition, though the TIFM works well on judging resources and uses' credibilities, the process is a little complex and slow. Next, the major principles in this work will be developed into a qualitative heuristic. We also plan to design another trust evaluation model and computation methods using the extrapolation of principles, and then conducting a comparative study with results of this research.

Acknowledgements

This work was supported by the National Social Science Foundation of China for the Youth Scholars of Education (project number: CCA130134). Authors would like to thank Dr Fengguang Jiang and Dr MinChen for their help during the study. [Correction added on 7 March 2014, after first online publication: The Acknowledgement section was corrected.]

References

- Adler, B. T., Chatterjee, K., de Alfaro, L., Faella, M., Pye, I. & Raman, V. (2008). Assigning trust to Wikipedia content. In *Proceedings of the 4th International Symposium on Wikis*, September, Porto, Portugal.
- Aniket, K. & Robert, E. K. (2008). Harnessing the wisdom of crowds in Wikipedia, quality through coordination. *Proceedings of the 2008 ACM conference on Computer supported cooperative work*, November 08-12, 2008, San Diego, CA, USA.

- Denko, M. K., Sun, T. & Woungang, I. (2008). Trust management in ubiquitous computing: a Bayesian approach. *Computer Communications*, 34, 3, 398–406.
- Dondio, P., Barrett, S., Weber, S. & Seigneur, J. M. (2006). Extracting trust from domain analysis: a case study on the Wikipedia Project. *Lecture Notes in Computer Science*, 4158, 362–373.
- Dong, X. H. (2010). Research on the trust mechanism in Grid Service. Ph.D. thesis, Department of Computer Science, Chongqing University.
- Drachsler, H., Hummel, H. G. K. & Koper, R. (2008). Personal recommender systems for learners in lifelong learning networks: the requirements, techniques and model. *International Journal of Learning Technology*, 3, 4, 404–423.
- Guo, Y. H. & Deng, G. R. (2008). Hybrid recommendation algorithm of item cold-start in collaborative filtering system. *Computer Engineering*, 34, 23, 11–13.
- Halim, F., Wu, Y. & Yap, R. (2009). Wiki credibility enhancement. In *Proceedings of the 5th International Symposium on Wikis and Open Collaboration (WikiSym '09)*. ACM, New York, NY, USA.
- Javanmardi, S., Lopes, C. & Baldi, P. (2010). Modeling user reputation in wikis. *Statistical Analysis and Data Mining*, 3, 2, 126–139.
- Jones, K. & Leonard, N. K. (2008). Trust in consumer-to-consumer electronic commerce. *Information & Management*, 45, 2, 88–95.
- Korsgaard, T. R. (2007). Improving Trust in the Wikipedia. Master's thesis, Department of Informatics and Mathematical Modelling, Technical University of Denmark.
- Lever, R. (2005). Wikipedia faces crisis. *ABC Science*, 2005-12-12.
- Li, C. & Liang, C. Y. (2012). Cold-start eliminating method of collaborative filtering based on n-sequence access analytic logic. *Systems Engineering—Theory & Practice*, 32, 7, 1537–1545.
- Li, Z. Y. & Wang, R. C. (2011). Dynamic secure trust management model for P2P e-commerce environments. *Journal on Communications*, 32, 3, 50–59.
- Lin, H. (2010). An application of fuzzy AHP for evaluating course website quality. *Computers & Education*, 54, 4, 877–888.
- Liu, Z., Yau, S. S., Peng, D. & Yin, Y. (2008). A Flexible trust model for distributed service infrastructures. In *11th IEEE Symposium on Object Oriented Real-Time Distributed Computing (ISORC)* (pp. 108–115).
- Lucassen, T. & Schraagen, J. M. (2010). Trust in wikipedia: how users trust information from an unknown source. In *Proceedings of the 4th workshop on Information credibility (WICOW '10)*. ACM, New York, NY, USA (pp. 19–26).
- Lucassen, T. & Schraagen, J. M. (2011). Evaluating WikiTrust: a trust support tool for Wikipedia. *First Monday [Online]*, 16, 5, 2 May 2011. Retrieved November 3, 2011, from <http://firstmonday.org/htbin/cgiwrap/bin/ojs/index.php/fm/article/viewArticle/3070/2952>
- Luo, X. J., Wang, T. C., Du, X. Y., Liu, H. Y. & He, J. (2007). Recommendation based on category-a method to solve cold-start problem in collaborative filtering. *Journal of Computer Research and Development*, 44, z3, 290–295.
- Luo, Z. C. & Fu, Z. Z. (2008). The impact analyze of some external forces on the wikipedia order process. *Document, Informaiton & Knowledge*, 3, 28–33.
- Maniu, S., Abdessalem, T. & Cautis, B. (2011). Casting a web of trust over Wikipedia: an interaction-based approach. In *Proceedings of the 20th international conference companion on World wide web (WWW '11)*. ACM, New York, NY, USA (pp. 87–88).
- Marsh, S. (1994). Formalising trust as a computational concept. Ph.D. thesis, Department of Mathematics and Computer Science, University of Stirling.
- Moturu, S. T. & Liu, H. (2009). Evaluating the trustworthiness of Wikipedia articles through quality and credibility. In *Proceedings of the 5th International Symposium on Wikis and Open Collaboration*, Orlando, Florida.
- Page, L., Brin, S., Motwani, R. & Winograd, T. (1999). The PageRank citation ranking: bringing order to the web. Technical Report. Stanford InfoLab.
- Por, G. (1997). Designing knowledge ecosystems for communities of practice. In *Advancing Organizational Capability Via Knowledge Management*. Los Angeles, September 29–30.
- Saaty, T. L. (1990). How to make a decision: the analytic hierarchy process. *European Journal of Operational Research*, 48, 1, 9–26.
- Saaty, T. L. (2008). Relative measurement and its generalization in decision making: why pairwise comparisons are central in mathematics for the measurement of intangible factors—the analytic hierarchy/network Process. *Review of the Royal Spanish Academy of Sciences, Series A, Mathematics*, 102, 2, 251–318.
- Seigenthaler, J. (2005). A false Wikipedia 'Biography'. *USA Today*, 2005-11-29.
- Shuang, Y. (2011). Combinatorial Proofs of q—analogues of two Summation Formular. *Journal of Inner Mongolia University for the Nationalities(Natural Sciences)*, 26, 6, 633–636.

- Sun, D. T., He, T. & Zhang, F. H. (2012). Survey of cold-start problem in collaborative filtering recommender system. *Computer and Modernization*, 5, 59–63.
- Tian, C. Q., Zou, S. H., Wang, W. D. & Cheng, S. D. (2008). A new trust model based on recommendation evidence for P2P networks. *Chinese Journal of Computers*, 31, 2, 271–281.
- Wang, D. D. (2009). Quality assurance system under the mode of wikipedia self-organization. *Library Theory and Practice*, 8, 21–24.
- Wang, H. & Guo, Y. J. (2003). Community in cyberspace and its features in intercourse. *Journal of Beijing University of Posts and Telecommunications (Social Sciences Edition)*, 5, 4, 19–21.
- Wang, K. W., Xie, F. D. & Zhang, Y. (2010). P2P e-commerce model based on punishment mechanism. *Computer Engineering*, 36, 12, 265–268.
- Wang, S. & Liu, Q. (2007). Study of trust mechanism in virtual learning community. *Distance Education Journal*, 3, 12–15.
- Wu, Q. H., Zhang, J. H. & Xu, X. H. (1999). An ant colony algorithm with mutation features. *Journal of Computer Research and Development*, 36, 10, 12–16.
- Yang, X. M. (2012). The research of resource evolution in the ubiquitous learning environment. Ph.D. thesis, School of Educational Technology, Beijing Normal University.
- Yang, X. M. & Yu, S. Q. (2011). The construction of ubiquitous learning ecosystem and analysis of some key issues. *Journal of Nanjing Normal University*, 11, 1–5.
- Yu, H. T., Zhang, J. H. & Zhong, Z. (2009). Trust model in P2P network based on time factor and punitive measures. *Computer Engineering and Applications*, 45, 20, 115–117.
- Zeng, C. X. (2011). Community-based knowledge of Web2.0 application framework. *Journal of Chongqing College of Electronic Engineering*, 20, 3, 152–154.
- Zhao, L. M. & Qi, L. Z. (2008). The network courses' evaluation based on AHP. *Distance Education Journal*, 2, 65–67.
- Zheng, K. F., Wang, X. J., Guo, S. Z. & Yang, Y. X. (2010). Evaluation model on service quality of information system based on user's appraisal. *Journal of Beijing University of Posts and Telecommunications*, 33, 1, 84–88.
- Zhou, G. Q., Pan, F. R., Zhang, W. F. & Guo, J. (2008). Research progress of trust evaluation model. *Wuhan University Journal of Natural Sciences*, 13, 4, 391–395.