

Modeling Technology Assessment via Knowledge Maps

Elan Sasson
Ben-Gurion University of the
Negev, Israel
elans@bgu.ac.il

Gilad Ravid
Ben-Gurion University of the
Negev, Israel
rgilad@bgu.ac.il

Nava Pliskin
Ben-Gurion University of the
Negev, Israel
pliskinn@bgu.ac.il

Abstract

Technology assessment (TAS) plays an important role prior to decision making about investments in existing and emerging technologies. The vast amount of data on the web has obviated the perception of using web search engine technology to look for information. However, relying on web search engines in search for relevant information to support TAS processes, decision makers face an abundance of data but are unable to screen noise or find hidden knowledge. This paper proposes a model to build knowledge-added concept map about a specific technology and the development of an underlying knowledge-mapping tool. The proposed knowledge maps are constructed on the basis of a novel method of co-word analysis based on webometric web counts. The approach is demonstrated and validated for a spectrum of information technologies. Results show that the research model assessments are highly correlated with subjective expert ($n=136$) assessment ($r > 0.91$), with inter-rater reliability scores being high as well ($ICC > 0.92$).

1. Introduction

Companies which depend heavily on technological innovation must constantly engage in a technology assessment (TAS) process when considering investments in new emerging technologies or evaluating the impact of existing technologies on the business landscape [1,2]. They are, however, due to information overload on the Internet, unable to manually process the abundance of data available about a specific technology [3, 4], making TAS a tough challenge for decision makers [5]. Information technology (IT) assessment entails an even greater challenge since new IT innovations occur at increasing speeds and with shorter life cycles [6]. A rising number of managerial analytical applications have exploited the massive amounts of available textual

documents. Some of these applications, which are of great importance to decision makers, intelligence analysts, and marketing analysts, perform text mining and co-occurrence analysis to generate concept maps [7, 8]. Generally speaking, concept maps capture concepts and concept relationships within a knowledge domain, using a two-dimensional, visually-based graphic representation of concepts and their relationships [9]. However, current research still proposes that automatically-generated concept maps, while responding to the challenge of extracting useful information for TAS purposes, leave technology assessors wondering how closely concept pairs on the map are contextually related. The main motivation of this paper is to overcome this limitation and improve concept mapping for TAS by using webometrics synthesized with co-word analysis to quantify the contextual distance between two concepts. Application of webometrics, i.e., quantitative bibliometric counts on the web (also known as hit count estimate - HCE) improves concept mapping by enhancing the relatedness proximity measure between concepts. The TAS approach in this work begins by fetching from diverse web sources a corpus of unstructured textual data about a specific technology. Then, to uncover hidden patterns in the corpus and generate a conventional concept map (co-occurrence network), information extraction (IE) is applied to the corpus to create a concept map, using a text mining (TM) technique based on natural language processing (NLP), followed by co-word analysis. The concept map, however, provides little if any confident knowledge about concept relatedness. To bridge this knowledge gap, the initial concept map is then processed further in this work into a knowledge map. The addition of contextual information upgrades the traditional concept map to a knowledge map, based upon which a technology-savvy decision maker is able to derive insights in a format of TAS propositions which were found consistent with publications by leading IT consulting firms. An automated, technology assessment knowledge (TASK) research instrument was developed on the basis of the proposed research model, one that automatically generates a knowledge

map. This study provides novel theoretical and practical contributions. From the theoretical perspective, the contribution is manifested in developing an innovative algorithmic model for upgrading a conventional concept map regarding a specific technology to a knowledge map. From the practical perspective, the contribution stems from the development of an automated research instrument capable of supporting decision makers engaged in technology assessment, and in helping gain a clear picture of knowledge about a specific technology and identifying future technological trends when they evaluate technology alternatives. Moreover, the model contribution is highlighted given the apparent advances in big data and extreme-scale analytics.

2. Background

The ability of decision makers to foresee technological advances and to assess new and current technologies is essential for anticipating future developments, understanding market position vis-a-vis their competitors, identifying upcoming innovations, and finally applying these insights to strategic business planning [10]. With an ever-quicken pace of technology development, TAS increasingly becomes a prominent task for technology and innovation management. With the advent of the web as an immense information space of diverse and often unstructured and non-standardized content formats, every decision maker turns to search engines when he/she is engaged in TAS process of a domain under exploration. However, relying on web search engines as a conceivable and major method in the information search needed to accomplish a TAS mission raises three concerns [3]. First, without information skills or a roadmap of what to look for, most people don't know how to ask for what they seek or know when it is reasonable to stop looking. Second, search result pages returned by search engines for a specific query include a vast amount of information to sort out, read and integrate. Third, there is really no metric we can use to compare the value of a 'good' search to a 'bad' one, given that relevance measurement is crucial to web search experience. Bolshakov and Gelbukh [5] acknowledge that decision makers must read and understand an enormous quantity of Internet text to make a well-informed decision. Clearly, it is beyond the ability of any person or group to comprehend large quantities of textual data without use of quantitative indicators [11]. Bibliometrics, also known as scientometrics, are methods which utilize quantitative indicators analysis and statistics to depict publication patterns within a given field or body of literature [12]. Quantitative bibliometric indicators use information,

such as word counts, date information, word co-occurrence information and citation information, to track activity in a subject area [13]. Porter and Detampel [14] highlight bibliometric analyses such as counts of publications, patents, or citations which can be used to measure and interpret scientific and technological advances. Also, they assert that a key tenet of bibliometrics are co-occurrences presented as linkage of concepts that can be detected in a specific domain and considered important in bibliometric analysis, potentially providing a powerful source of information on emerging technologies. A study by Rinia et al. [15] shows a strong correlation between assessments based on bibliometric indicators and judgments made by expert committees.

When faced with an enormous amount of information, overload makes it difficult to extract valuable insights, as expressed by Naisbitt [16] who claimed we are drowning in information but thirsty for knowledge, especially given that nearly 80% data is unstructured text. While the amount of textual data available to us is constantly increasing, the human ability to understand and process this information remains constant and limited. Given the volume and complexity of the information involved, Lee et al. [17] thus assert that manual analysis of unstructured textual data is increasingly impractical. Thus, automatic TM has the potential to give companies the competitive edge they need to survive by identifying patterns hidden inside vast collections of text data. The objective of TM is to exploit information contained in textual documents in various ways, including discovery of patterns and trends in textual data and associations among text objects (e.g., concepts) [18]. Moreover, TM involves IE, which is the task of extracting named-entities and factual assertions from text [19]. IE allows the transformation from the unstructured document space to the structured concept space, paving the way to analysis of interactions between concepts extracted from a textual corpus. There is a fairly extensive body of literature on co-word analysis [20, 21]. Feldman et al. [22] provided an early seminal work on concept co-occurrence relationships in a corpus of documents. He [23] considers co-word analysis as a powerful and proven quantitative tool for knowledge discovery in a research field. According to Rapp [24], concepts that co-occur tend to be related, demonstrating relatedness association. Therefore, co-occurring concepts have been considered as carriers of meaning across different domains in studies of science and technology, and general indicators of activity in textual document sets [25]. While maintaining essential information contained in the data, co-word analysis reduces the data into a specific visual representation based on the nature of words, revealing patterns and trends in a

specific discipline [26]. Co-word clustering is a process that begins by assessing the strength of the link value between two concepts based on their co-occurrence in a given record or document, and ends with the grouping of strongly-linked concepts into clusters. The definition used in this study for the co-occurrence measure is the Similarity Link Value (SLV), also known as Equivalence Index (E), defined by Callon et al. [26] as:

$$SLV_{ij} = \frac{C_{ij}^2}{C_i * C_j}, 0 < SLV_{ij} \leq 1, C_{ij} = C_{ji} \geq 0$$

In this definition, C_{ij} is the number of co-occurrences of terms i and j (i.e., the number of documents in which both terms co-occur), and C_i and C_j respectively count the term occurrence (i.e., the number of documents in which term appears) of term i and term j . A concept map is a common method for representing the relationships among a set of concepts, with vertices/nodes (e.g., named-entity concepts such as person, company, location) capturing concepts and edges/ links capturing the relations between concepts [28]. More specifically, a concept map is a dynamic graphical map that visually presents concepts and relevant relationship clusters, which can be portrayed as an undirected graph $G = (V, E)$ consisting of a set of vertices V and a set of edges E . Novak and Canas [29] argue that the relationships between concepts indicated by a connecting edge often represent creative leaps (i.e., meaningful learning) in the creation of new knowledge. Indeed, some sources refer to concept maps as knowledge graphs [30]. Yet, the current study differentiates between a concept map and a knowledge map, viewing the latter as upgraded concept maps improved with relatedness proximity measure. The most challenging aspect of constructing a concept map is linking the concepts into a meaningful, coherent structure that reflects understanding of a specific domain [31]. For TAS purposes, conventional concept mapping suffers from one major drawback i.e., the unreliable measure of the contextual distance between co-occurring concept pairs. This weaknesses is amplified when the textual corpus upon which the initial concept mapping is accompanied by a large amount of noise and overload of irrelevant contextual concept relationships. Indeed, a web-based corpus of textual data, such as is implemented in the current study, is often accompanied by a large amount of noise. Specifically in concept mapping, Mustafaraj et al. [32] argue that unstructured texts from the web might include errors with noise induced by the

imperfection of the concept extraction process. Information overload and noisy data may sometimes lead to imprecise outcome in co-occurrence analysis. This may result in an inaccurate or incomplete concept map where existing relations might not be discovered, discovered relations might not be the result of actual relations, or a given link might have a spurious or a missing relationship. Thus, in harnessing conventional concept maps for TAS purposes, especially ones with many nodes (i.e., concepts) and association relations (i.e., co-occurrences), there is a significant risk that information overload and noise would mislead decision makers regarding the contextual distance between co-occurring concept pairs.

3. Research Model

To overcome the drawbacks of conventional concept mapping for TAS, webometric-based co-word analysis was used for measuring relatedness proximity. Thus, adding contextual knowledge to the initial concept map, a-priori calculation of each SLV was carried out, followed by calculation of a bibliometric SLV based on webometric hit count estimates (HCEs) or web counts. Then, the two previously calculated SLV values were combined to an extended SLV value, thus following the three steps illustrated in Figure 1:

1. A-priori co-occurrence analysis yielding $aSLV_{ij}$
2. Bibliometric co-occurrence analysis yielding $bSLV_{ij}$
3. Combined co-occurrence analysis yielding $cSLV_{ij}$

In Step 1, using NLP-based TM to complete the IE task, significant semantic concepts (*named entities* such as person, company, location, product) are extracted from the time-tagged corpus of text documents (e.g., TXT files, PDF, HTML files etc.) and an a-priori $aSLV_{ij}$ co-occurrence value is calculated for each relation between Concept i and Concept j . In Step 2, the bibliometric analysis task uses the same exact concept pairs in a series of webometric queries to a web search engine, about Concept i , Concept j , and their conjunctive Concept $i + \text{Concept } j$, using the AND Boolean operator, and the bibliometric $bSLV_{ij}$ co-occurrence value is derived from the web counts of the search results retrieved for each concept pair. In Step 3, both a-priori and bibliometric SLVs i.e., $aSLV_{ij}$ and $bSLV_{ij}$ are synthesized into a combined $cSLV_{ij}$ relatedness value for each concept pair, measuring relatedness proximity for the Concept-pair i, j .

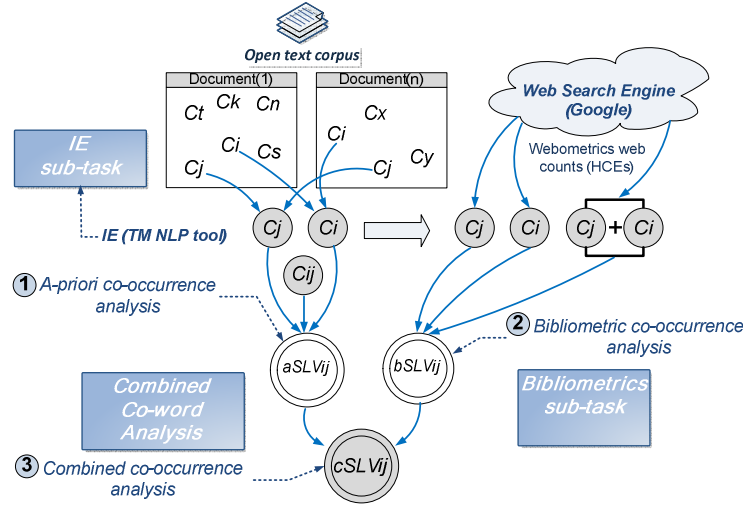


Figure 1: A conceptual workflow describing the combined co-word analysis process

Using this combined co-word analysis, weak or strong signals obtained in Step 1 are improved via weighted synthesis with the co-occurrence values obtained by applying the webometric web counts in Step 2. It is safe to assume that concept relatedness which appears in unindicative context in the a-priori co-occurrence analysis (i.e., Step 1) may appear in a very obvious context in the bibliometric co-occurrence analysis (i.e., Step 2) and vice versa. Gledson and Keane [33] consider web-searching an important part of measuring concept relatedness, as it provides up-to-date information on word co-occurrence frequencies in large available collection of documents such as the web. Similarly, Cilibrasi and Vitanyi [34] assert that relative frequencies of web pages (e.g., web counts or HCEs) containing search terms give objective information about the relationship between the terms. Finally, it seems promising to combine the two types of SLVs into one $cSLV_{ij}$ (Step 3) based on the following additive formula:

$$cSLV_{ij} = f\left(\left(\frac{aC_{ij}^2}{aC_i * aC_j}\right), \left(\frac{bC_{ij}^2}{bC_i * bC_j}\right)\right) \\ \approx \left(\frac{(aC_{ij} + bC_{ij})^2}{(aC_i + bC_i) * (aC_j + bC_j)}\right)$$

This method of improving the measurement of relatedness proximity features: (a) denoising filtering of outlier concept co-occurrences that the conventional co-word analysis process extracts from the corpus and then mistakenly presents them as significant to the decision maker, and (b) amplification filtering that enables discovery of elusive relationships not detected by the conventional co-word analysis process due to the weak signal yielded by the algorithm, and thus

overlooks hidden concept co-occurrences important to the decision maker. The result of improving the measurement of relatedness proximity is a robust knowledge-added concept map for identification and selective extraction of significant concept co-occurrences. Decreasing the number and dimensionality of extracted concept pairs and displaying only significant key ones also improves the visualization of the resulting knowledge map.

4. Methodology

A web-based research instrument was developed for mapping TAS knowledge (TASK) in order to demonstrate and validate the research model developed in the current study. The research instrument allows data collection and processing prior to knowledge mapping that yields a time-tagged textual document. Then, via an advanced interactive dashboard-oriented user interface (UI), the research instrument allows technology savvy decision makers to automatically generate a knowledge map and explore the extracted knowledge toward derivation of TAS propositions. Implementation of the TASK research instrument shown in Figure 2 followed the general CRISP-DM model. The instrument is divisible into the following six main stages and 15 tasks:

a) **Temporal GAs collection** tasks involve collecting a repository of Google Alerts (GA) email updates which include one or more URL links to domain-specific (i.e., IT topic) web documents (e.g., HTML, XML) in diverse web sites. GA is Google's content change-detection and notification service, automatically notifies subscribers when new Internet content matches a set of search terms (e.g., topic). This

non-static process of gathering links over time facilitates collection of documents into a dynamic corpus, allowing examination of concept co-word analysis not just within a given concept context but also analysis of the similarities and differences in

context relationships across different temporal segments of the corpus. Steps 1 and 2 in Figure 2 depict this stage.

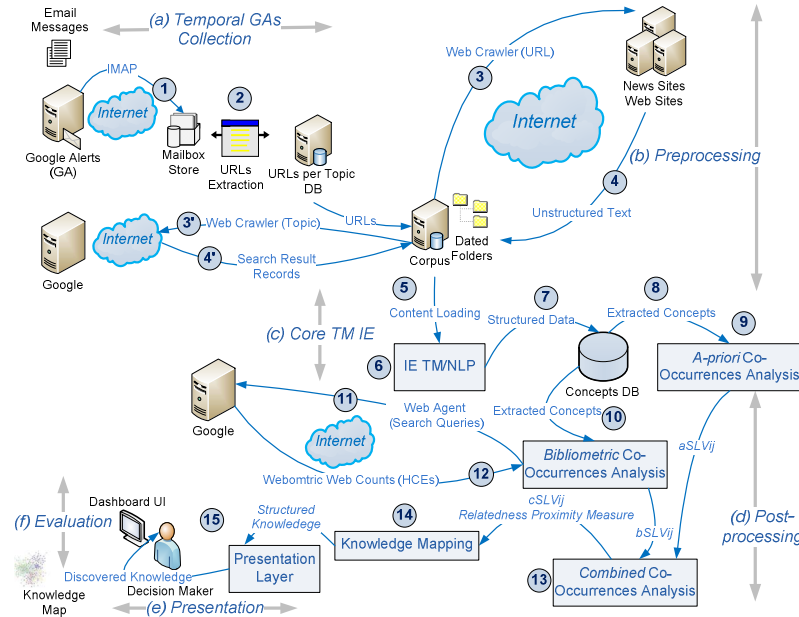


Figure 2: The stages and tasks of the TASK research instrument

b) **Preprocessing** tasks include all routines, processes, and methods required for using crawling techniques to fetch the actual HTML files. A crawler web agent is applied in order to automate the execution of the actual textual data gathering. The crawler starts from a list of URLs stored in the repository created in Stage (a), which includes all the links embedded in the GAs email messages received over time. The crawler follows all links to actually collect the required web pages and locally stores and indexes the collected textual data in a repository on a dedicated corpus server for further use and analysis. The crawling mechanism extracts only the first-level of web pages and ignores external embedded URLs other than internal links to the next page level in cases where the extracted document is made of more than one web page. Heuristically, the crawler detects specific text strings such as 'next page' or page numbering in the main paragraph section of the extracted web page to extract relevant content in next HTML pages. Moreover, a zero-level corpus is instantly built by submitting a search query to Google about a specific IT. It is necessary to build a zero-level corpus before beginning the gradual iterative process of collecting documents from the web onto the time-tagged corpus. Search engines will only return the first 1000 or so

results. The results are presented as search result records (SRRs) and thus include up to one thousands URLs which are extracted and crawled accordingly (See Steps 3' and 4' in Figure 2). The date stamps and other relevant metadata information are stored in the server. Steps 3 to 4 in Figure 2 depict this stage.

c) **Core TM and IE** NLP-based tasks are routines and processes for concept discovery in the document corpus yielded by Stage (b), which is categorized, keyword-labeled and time-stamped, toward extracting and storing for further analysis concepts and their relevant metadata (e.g., time stamp, total number of appearances, average concept distribution etc.). Steps 5 to 7 in Figure 2 depict this stage.

d) **Post-processing** analysis tasks include all procedures and methods required for conducting the relatedness proximity measurement as previously described toward knowledge mapping. Steps 8 to 14 in Figure 2 depict this stage.

e) **Presentation** tasks and browsing functionality include graphical user interface and listing capabilities. *Presentation layer* components display the knowledge map with references to co-occurrence weights calculated at each step. The browsing functionality provided by the presentation layer also includes listing

tables; graphing and GUI to dynamically present knowledge maps. Step 15 in Figure 2 depicts this stage.

f) **Evaluation** tasks are carried out by the decision maker, while valuating and interpreting the acquired results, and therefore not depicted in Figure 2. Generalization, pruning or requiring collection of additional textual data in order to enrich the corpus, may be implemented by the user.

4.1. Corpus Building and Information Extraction

Datasets used for building the time-tagged corpus were created using Google Alerts (GA) collected throughout 190 days (August 12, 2010 to February 16, 2011). Each alert is an email message in HTML format which includes aggregation of links (URLs) to the latest news articles about each technology used to demonstrate and validate research model and instrument in this study from various source types (news, web, blogs, and discussion groups' sites). In planning the corpus, the goal was to use an IT array with a spectrum of IT types in lifecycle maturation stages that are diverse enough for model demonstration and validation. Thus, Cloud Computing, which was over-hyped when corpus building commenced in August 2010 and expected to substitute Grid Computing, in addition to Business Process Management (BPM), which attracted a lot of a new attention at the time, were both included. Semantic Web, regarded as new and particularly promising, and Service Oriented Architecture (SOA), already considered then a de-facto standard on the web, were also included in the IT array. Table 1 summarizes the actual number of URLs (links) extracted for each of the five IT topics.

Table 1: Total number of links from GAs per IT topic

GA search query	# links
Cloud Computing	12,535
Grid Computing	6,470
Semantic Web	6,030
Service Oriented Architecture	5,781
Business Process Management	8,908

To accomplish the IE process, NLP-based TM analysis was applied to the TXT files in the time-tagged corpus, using IBM's SPSS/PASW Text Analytics Version13 (former SPSS Text Mining Modeler) and AlchemyAPI for rigor and robustness, although each of these tools autonomously provides all functions necessary for the IE process. Nonlinguistic

entities which include phone numbers, currencies, percentages, etc. were excluded from the extraction process. Finally, a sparse document-by-concept occurrence matrix which shows the presence of concept in a document by a 'T' (true), and its absence by a 'F' (false) was computed and uploaded to SQL database for further analysis.

The number of extracted concepts (n) posed a computational-complexity challenge of $O(n^2)$ in the co-word analysis while generating the concept-by-concept relatedness matrices. For example, a bibliometric co-occurrence analysis of 1000 concepts might require conducting 1,000,000 search queries to Google, thus posing a prohibitively unrealistic processing time considering that Google limits the number of queries to 1000 per day per IP address. To cope with the scalability challenge and based on similar studies [25, 35], 100 top concepts were used as the maximum number of concepts to be included in the computation of the concept-by-concept relatedness matrices, so that an optimized number of k concepts yields near-linear time complexity.

The extended relatedness proximity measurement was conducted disjointedly for each of the five investigated IT topics used to demonstrate and validate the research model. The relatedness proximity measurement as previously described includes the a-priori co-occurrence analysis which yields $aSLV_{ij}$ values which are then combined with the $bSLV_{ij}$ values obtained in the bibliometric co-occurrence analysis based on webometric web counts to finally produce the hybrid $cSLV_{ij}$ values.

4.2. Validation methodology and tools

Validation of the proposed model is a two-fold process. First, the relatedness proximity measurement is validated, including validation of the HCE webometric web counts used in the bibliometric co-occurrence analysis. Second, TAS propositions derived from the knowledge maps for each IT topic are validated. The relatedness proximity is validated as a targeted web-based survey aimed for domain experts such as IT practitioners and IS scholars. The web counts validation process is based on seeking a logical consistency among multiple related search queries, also known as *Metamorphic Relations*, as proposed by Zhou et al. [36]. The TAS proposition validation is accomplished by comparing derived propositions with propositions extracted from assessments reported by leading IT consulting firms such as Gartner, complemented by studies and scholar research.

Following Dochy [37], according to whom a relationship drawn between two concepts is the

smallest unit that can be used to validate a map, it is thus reasonable to validate the relatedness proximity measurement in this study via comparison to human ratings [38, 28]. The simplest way to do this is to examine the correlation between the human judgment of human raters and the calculations of the research instrument according to the research model. The typical statistical data structure is a case-by-variable structure, where cases are independent raters and the variables are the subjective ratings provided by raters.

Using this technique, each expert is asked to rate on a numerical scale the degree of relatedness between pairs of previously defined concepts. Following Reips and Funke [39], who argue in favor of implementing a visual analogue scale (VAS) as a continuous evaluation device rather than categorical scales which only reach ordinal-scale level, a VAS-based instrument was applied in this study.

In most of the 21 research studies reviewed by Ruiz-Primo and Shavelson [28], which employed expert's maps as scoring criteria, the actual number of experts varied in the range of 3 to 8. However, in the current study the number of respondents which are needed for each IT topic is set to a minimum number of 23 raters.

A complete evaluation of n concepts requires raters to compare $n * (n - 1) / 2$ pairs of concepts and is likely to place a high cognitive load on respondents, resulting in fatigue and impracticality. In the most commonly used setting, discussed in previous studies raters are provided with a set of between 10 to 15 pairs of concepts and are asked to evaluate the similarity between all possible pairs. In this study, the setting for validating relatedness proximity includes 20 pairs of concepts for each assessed IT topic. Hence, a randomly partition concept list producing fewer evaluations for each rater is generated from an integrated list of identical concepts extracted by both IE tools (SPSS/PASW and AlchemyAPI).

A survey instrument was developed and used to collect data from a target sample of experts in a survey conducted over the web. Invitation letters were sent out by e-mail to selected participants experts, inviting them to take part in the survey. Survey questionnaires were distributed internationally, targeting a database of domain experts, IT consultants, and academic researchers. The database was obtained from two major sources: LinkedIn, which is regarded as the largest web-based professional network which interconnects experts and professionals around the world, and a large list of various experts and analysts about the assessed IT topics maintained by a leading global IT consulting firm.

A total of 136 experts responded and submitted the survey.

Another aspect that must be addressed while validating the relatedness proximity measurement is the web counts or Hit Count Estimates (HCEs) used in the bibliometric co-occurrence analysis. HCEs were described by early studies as unreliable, inconsistent, and fluctuating over time [40]. However, more recent studies [36] report that existing search engines (especially Google) generally maintain a reasonable degree of self-consistency in terms of objectivity, reliability and exhibit clear evidence of good data quality for HCEs. In contrast, this study leaned toward a rigorous approach and adopted a validation process for the HCEs based on the Metamorphic Relations (MR) testing developed by Chen et al. [41]. Indeed, the actual number and % of failed MR tests observed for all five topics is a low overall total of 2.9%.

Finally, within the scope of this study, manual derivation of descriptive TAS propositions is pursued for several fundamental ITs based upon map-centric views of the knowledge maps created by the research instrument. Validation of TAS propositions is accomplished by comparing the TAS propositions derived for each IT by the technology-savvy individual to benchmark assessments in reports by leading IT consulting firms, mainly by Gartner and partially by Ovum, as well as reported by scholarly studies. Moreover, given that the data for the corpus were collected over a period of time that ended on February 2011, most of the reports and studies used for validation of the TAS propositions were explicitly chosen to include benchmark assessments from Year 2011 up to the middle of Year 2012, assuming that insights reflected in these benchmark reports and studies are relevant to a minimum of one year period. The intention behind extending the benchmarking period for more than one year beyond the end of data collection was to look beyond the scope of the current research and check whether the research model has a predictive capability in addition to its intended descriptive capability. However, the main focus here is on TAS proposition validation for assessment purposes as opposed to for prediction and forecasting purposes.

5. Results

5.1. Applying the validation instrument

A targeted web-survey was globally distributed among IT professionals, consultants and scholars to measure subjective attitudes towards the research model's relatedness proximity measurement. More specifically, a questionnaire was implemented for each IT topic to validate the $cSLV_{ij}$ values. A total number of 140 respondents for all topics were recruited. Six

questionnaires were excluded given a low rate of obtained answers. Actual numbers of valid respondents for each topic exceeded the minimum required sample size (i.e., 23). The majority of respondents (89%) have more than four years of experience. A total low-enough number of 5% of all answers collected in all questionnaires were classified as missing values.

5.2. Relatedness proximity measurement validity

To compare the $cSLV_{ij}$ results with the human rankings for validation of the relatedness proximity measurement, inter-rater reliability measures and correlation coefficient measures were statistically analyzed.

The reliability for all the expert raters averaged together is a measure of internal consistency, providing an index of homogeneity of responses based on the Intraclass Correlation Coefficient (ICC).

Table 2: ICC values

Topic	ICC
Cloud Computing	0.983
Grid Computing	0.920
Semantic Web	0.972
Service Oriented Architecture	0.978
Business Process Management	0.943

As it can be seen in Table 2, presenting obtained ICC values for each topic, the homogeneity and similarity of responses indicate high degree of inner resemblance of expert's rankings for all five topics.

A Pearson's correlation coefficient was used to compare the averaged ranking produced by the human subjects (i.e., raters) with two model-generated values: $aSLV_{ij}$ and $cSLV_{ij}$. Table 3 presents Pearson's correlations, suggesting that all measures perform well for all five topics, with high correlations for $cSLV_{ij}$ values and lower correlations for $aSLV_{ij}$ values, excluding Grid Computing.

The observed low correlation values of raters to $aSLV_{ij}$ for the IT topics (other than Grid Computing) is supportive of this research's claim that conventional co-word analysis based on a time-tagged corpus in many cases lacks the knowledge background available on the web that ought to be discovered and assimilated in the form of webometric web counts.

One argument to explain the anomaly of Grid Computing might be that Grid Computing is a mature technology. According to Google Trends, a service which shows how frequently a topic has been searched over time, the interest in the Grid Computing by the worldwide IT community was diminishing, as indicated by the on-going decrease of the *search value index* between years 2004 to 2011 from 3 to 0.4, receptively. Thus, from the perspective of relatedness proximity measurement, the research model and instrument seem more promising and valuable for an assessment of current, innovative and evolving technologies rather than mature ones, which are understandably less relevant in the TAS context.

5.3. Validation of model-based TAS propositions

Manual derivation of major propositions, based on map-centric views of the knowledge maps generated by the research instrument for each assessed IT, by either technology-savvy professionals or scholars (like the author), was pursued in the present study for four fundamental technologies: (1) Cloud computing; (2) BPM; (3) Semantic Web and (4) SOA. For the sake of brevity of this manuscript, only one snap shot of knowledge map, for SOA is presented in Figure 3 (for illustration purposes). The TAS propositions derived based on pairs of highly correlated co-occurring concepts presented on the knowledge maps were found to be compatible with the respective TAS reports by a leading IT consulting firm (i.e., Gartner) complemented with scholar assessment studies.

Table 3: Pearson correlations

Topic	Expert' Ratings vs. $cSLV_{ij}$		Expert' Ratings vs. $aSLV_{ij}$	
	Pearson's correlation coefficient	P_value*	Pearson's correlation coefficient	P_value*
Business Process Management	0.879	0.000	0.423	0.063
Cloud Computing	0.951	0.000	0.250	0.289
Grid Computing	0.939	0.000	0.763	0.000
Semantic Web	0.949	0.000	0.541	0.014
Service Oriented Architecture	0.913	0.000	0.325	0.162

*0.000→0

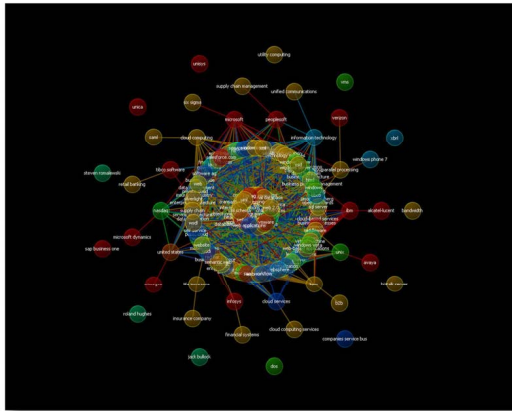


Figure 3: Knowledge map for SOA

6. Conclusion

The research instrument used to implement the knowledge-mapping research model was found valuable in assisting decision makers in assessing emerging and existing ITs but less appropriate for more mature ones. The new computed relatedness proximity measurement was found to be highly correlated with experts subjective ratings ($n = 136$): $r = 0.91$ to 0.98 . Also, high inter-rater reliability scores were found based on Intraclass Correlation Coefficient (ICC) = 0.92 to 0.94 . Moreover, TAS propositions were found to be valuable as validated by a technology-savvy decision maker. Finally, it can be concluded that the novel decision support system provided by this research model and instrument has the potential to morph into the realm of managerial decision process in the enterprise.

7. References

- [1] Henselewski, M. Smolnik, S. and Riempp, G. Evaluation of Knowledge Management Technologies for the Support of Technology Forecasting. *HICSS'06*, 2006.
- [2] Porter, AL. and others. Technology futures analysis: towards integration of the field and new methods. *Technological Forecasting and Social Change* (49) 2004.
- [3] Feldman, R. The high cost of not finding information. *Content Document and Knowledge Management KMWorld*, 2004.
- [4] Fain, D.C. and Pedersen, J.O. Sponsored search: A brief history. *Bulletin of the American Society for Information Science and Technology*, 32(2), 2006, pp. 12 -13.
- [5] Bolshakov, I.A. and Gelbukh A. *Computational Linguistics: Models, Resources, Applications*. Center for Computing Research (CIC) of the National Polytechnic Institute, the Economic Culture Fund Press, 2004.
- [6] Ashrafi, N. Xu, P. Kuilboer, J. and Koehler, W. Boosting enterprise agility via IT knowledge management capabilities. *HICSS'06*, 2006.
- [7] Su, H.N. and Lee, P.C. Mapping knowledge structure by keyword co-occurrence: a first look at journal papers in Technology Foresight. *Scientometrics*, 85(1) 2010, pp. 65-79.

- [8] Waltman, L. van Eck, N.J., and Noyons, E. A unified approach to mapping and clustering of bibliometric networks. *Journal of Informetrics*, 2010
- [9] Leake, D. Maguitman, A. and Canas, A. Assessing conceptual similarity to support concept mapping. *The Proceedings of the Fifteenth International Florida Artificial Intelligence Research Society Conference*, 2001, pp.172-186.
- [10] Halsius, F. and Lochen, C. Assessing Technological Opportunities and Threats – An introduction to Technology Forecasting. *Division of Industrial Marketing, Lulea University of Technology*, 2001.
- [11] Narin, F. Olivastro, D. and Stevens, K. A. Bibliometrics/Theory Practice and Problems. *Evaluation Review*, 18(1) 1994, pp. 65-76.
- [12] Zhu, D. Porter, A., Cunningham, S. Carlisle, J., and Nayak, A., (2004). A process for mining science & technology documents databases, illustrated for the case of “knowledge discovery and data mining. *Technology Policy & Assessment Center Georgia Institute of Technology*, 2004.
- [13] Kontostathis, A. Galitsky, L.M. Pottenger, W.M. Roy, S. and Phelps, D.J. A Survey of Emerging Trend Detection in Textual Data Mining, In: Berry, M., (ed.), *Survey of Text Mining: Clustering, Classification, and Retrieval*, Springer, 2004.
- [14] Porter, A.L. and Detampel, M.J. Technology opportunities analysis. *Technological Forecasting and Social Change*, 49(3), 1995 pp. 237-255.
- [15] Rinia, E.J. Van Leeuwen, Th.N. Van Vuren, H.G., and Van Raan, A.F.J. Comparative Analysis of a Set of Bibliometric Indicators and Central Peer Review Criteria. *Evaluation of Condensed Matter Physics in the Netherlands Research Policy*, 27, 1998, pp. 95-107.
- [16] Naisbitt, J. *Megatrends 2000*. Smithmark Publishers, New York, 1996.
- [17] Lee, S. Baker, J. Song, J., and Wetherbe, J.C. An Empirical Comparison of Four Text Mining Methods. *HICSS*, 2010, pp. 1-10.
- [18] Grobelnik, M. Mladenic, D. and Milic-Frayling, N. Text Mining as Integration of Several Related Research Areas: Report on KDD'2000 Workshop on Text Mining. *SIGKDD Explorations*, 2(2), 2000 pp. 99-102.
- [19] Wilks, Y. Information Extraction as a Core Language Technology, *Lecture Notes in Computer Science*, Springer-Verlag, 1997 pp. 1- 9.
- [20] Callon, M. Law, J. Rip, A. *Mapping the dynamics of science and technology: sociology of science in the real world*. Macmillan Press, 1986.
- [21] Courtial, J. P. A cword analysis of scientometrics. *Scientometrics*, 3, 1994, pp. 251-260.
- [22] Feldman, R. Klbgsen, W. Ben-Yehuda, Y. Kedar, G. and Reznikov, V. Pattern Based Browsing in Document Collections. *Principles of Data Mining and Knowledge Discovery: First European Symposium, PKDD'97*, 1997.
- [23] He, Q. Knowledge Discovery through Co-Word Analysis. *Library Trends*, 48, 1999, pp. 133-159.
- [24] Rapp, R. The computation of word associations: comparing syntagmatic and paradigmatic approaches. *Proceedings of the 19th international conference on Computational linguistics*, 1, 2002 pp. 1-7.
- [25] Leydesdorff, L. and Hellsten, I. Measuring the meaning of words in contexts: An automated analysis of controversies about 'Monarch butterflies,' 'Franken foods ' and 'stem cells'. *Scientometrics*, 67(2), 2006, pp. 231–258.
- [26] Ding, Y. Chowdhury, G.G. and Foo, S. Bibliometric cartography of information retrieval research by using co-word analysis. *Information Processing and Management*, 37, 2001 pp. 817-842
- [27] Callon, M. Courtial, J.P. and Laville, F. Co-word analysis as a tool for describing the network of interactions between basic and technological research: The case of polymer chemistry. *Scientometrics*, 11, 1991, pp 155 -205.
- [28] Ruiz-Primo, M.A. and Shavelson, R.J. Problems and issues in the use of concepts maps in science assessment. *Journal of Research in Science Teaching*, 33, 1996, pp. 569 – 600.
- [29] Novak, J.D. and Canas, A.J. The theory underlying concept maps and how to construct and use them. Florida Institute for Human and Machine Cognition Pensacola, 2008.
- [30] Chein, M. and Mugnier, M.L. *Graph-based knowledge representation: computational foundations of conceptual graphs*. Springer-Verlag New York, 2009.
- [31] Canas, A.J. Novak, J.D. Gonz'alez, F.M. Carvalho, M. Arguedas, M. and Cognition, M. Mining the web to suggest concepts during concept map construction. Universidad P'ublica de Navarra, 2004.
- [32] Mustafaraj, E. Hoof, M. and Freisleben, B. Mining Diagnostic Text Reports by Learning to Annotate Knowledge Roles, In: Kao, A., and Poteet, S.R., (eds.), *Natural Language Processing and Text Mining*. Springer-Verlag, New York, 2007.
- [33] Gledson, A. and Keane, J. Using web-search results to measure word-group similarity. *In the Proceedings of the 22nd International Conference on Computational Linguistics*, 1, 2008, pp. 281-288.
- [34] Cilibrasi, R.L. and Vitanyi, P.M.B. (2007). The google similarity distance. *IEEE Transactions on Knowledge and Data Engineering*, 19(3), 2007, pp. 370-383.
- [35] Hutchins, C.E. and Benham-Hutchins, M. Hiding in plain sight: criminal network analysis. *Computational & Mathematical Organization Theory*, 16(1), 2010, pp. 89–111.
- [36] Zhou, Z.Q. Zhang, S.J. Hagenbuchner, M. Tse, T.H. Kuo, F.C. and Chen, T.Y. Automated functional testing of online search services. *Software Testing, Verification and Reliability*, 2010.
- [37] Dochy, F. Assessment of domain-specific and domain-transcending prior knowledge: Entry assessment and the use of profile analysis. *Alternatives in assessment of achievements, learning processes and prior knowledge*, 1996, pp. 227-264.
- [38] McClure, J.R. Sonak, B. and Suen, H.K. Concept map assessment of classroom learning: Reliability, validity, and logistical practicality. *Journal of research in Science Teaching*, 36, 1999, pp. 475-492.
- [39] Reips, U.D. and Funke, F. Interval-level measurement with visual analogue scales in Internet-based research: VAS Generator. *Behavior Research Methods*, 40(3), 2008.
- [40] Bar Ilan, J. Search engine results over time: A case study on search engine stability. *Cybermetrics*, 3(1), 1999.
- [41] Chen, T.Y. Tse, TH. Zhou, Z.Q. Semi-proving: An integrated method for program proving, testing, and debugging. *IEEE Transactions on Software Engineering*, 2010.