

SCIENTIFIC PROGRAMME

SIMULATION IN OR AND KNOWLEDGE MANAGEMENT

A BAYESIAN APPROACH FOR KNOWLEDGE MANAGEMENT EVALUATION

Katia Passerini
New Jersey Institute of Technology
University Heights
Newark, NJ 07109, USA
pkatia@njit.edu

Kemal Cakici
School of Business and Public Management
George Washington University
Washington, DC 20052, USA
cakici@gwu.edu

KEYWORDS

Statistical Analysis, Probabilistic Models, AI in simulation, Decision Models, Government, Policy Making

ABSTRACT

The authors review the World Bank's Knowledge Assessment Methodology (KAM) - which defines a country's knowledge index based on a set of 66 variables - and apply Bayesian methods to determine the causality relationships among the variables as well as an optimized model for knowledge management. The analysis draws on the researchers' model simulation of causal relationships and contrasts them against uses of country level data and indexes. The goal is to establish a causal model for national knowledge management which will provide insights also at the micro-economic level of the individual firm.

OVERVIEW

There are numerous examples of worldwide knowledge-based initiatives which are conducted not only at the micro-level of the firm, but also at the macro-level of a nation [Macdonald, 1999; World Development Report, 1998/99; Dahlman and Andersson, 2000; Dahlman and Aubert, 2001; World Bank Institute, 2001]. Knowledge management evaluation efforts and implementation initiatives have displayed similarities both at the level of national alliances (trade, global business / political linkages) [Arthur, 1996], and at the level of firms' strategic partnerships [Choi, 1997]. These similarities drive the authors' approach to defining a knowledge management network model that draws upon the international models established by organizations such as the World Bank [World Bank Institute, 2002] and the National Research Council in Washington DC [National Research Council, 1996].

KNOWLEDGE ASSESSMENT MODELS

The Knowledge Assessment Methodology (KAM) established by the World Bank [World Bank Institute, 2002] is a process to evaluate a nation's potential to participate in the knowledge economy and reap the benefits of knowledge management. The KAM uses 'knowledge assessment scorecards' to evaluate how a country compares to other nations, particularly in its use of knowledge for its overall social and economic development. The KAM consists of 51 variables (**Figure**

1) plus 15 variables (**Figure 2**), for a total of 66 variables in the model.

The 66 variables are grouped by macro-areas (performance indicators, and four key areas of knowledge-development which include economic incentives, institutional regime, innovation system, information infrastructure) and play a different role in fostering knowledge potential. This research uses Bayesian simulation techniques to define the mutual role of the variables, as well as their optimal combination. Opponents of Bayesian techniques argue that these techniques do not generate optimal results. However, Jouffa (2002), Munteanu and Bendou (2001), Jouffa and Munteanu (2001) present results that using SopLEQ learning method that utilizes the equivalence properties of Bayesian networks and total characterization of the data. Indeed generate results.

The analysis first focuses on a subset of variables for the assessment of the variables' contribution to knowledge creation (the 'knowledge assessment scorecards'). The scorecards consist of a subset of the crucial variables, normalized to represent key performance measures based on the above-mentioned macro-areas. The variables and sources of the most recent implementation of the KAM exercise (2002) are illustrated in **Figure 2**.

In the 2002 World Bank's KAM implementation, the scorecards are used to compare 98 countries, which include most of the developed economies and about sixty developing economies. The results are keyed in radar graphics that allow, for example, direct identification of the countries relative positioning compared to the G7 countries. The described framework represents a structured approach for a quantitative overview of the countries knowledge potential but it does not clearly map the role played by the different variables nor defines their optimal probability distributions. Therefore, a Bayesian network and decision graph is used against the results of the most recent international assessments (KAM-2002) to clearly map the direct and indirect relationships of the variables in the achievement of the country knowledge potential.

BAYESIAN NETWORKS AND DECISION GRAPHS

In this paper, Bayesian networks and decision graphs are used to map the relationships between the KAM variables, the knowledge drivers, and the resulting country knowledge potential (dependent variable). Bayesian networks are selected as the method of discovery as they

provide a compact and easy-to-use representation of the probabilistic information embedded in the data. The network structure is an effective way to communicate dependencies among the variables. Bayesian networks effectively discover the set of interactions responsible for the observed data. Learning in Bayesian networks can be

either automatically from data sources and/or manual modeling by experts. Once this is done, probability distribution of each variable can be updated given the states of other variables (Bayesia S.A., 2004). Therefore, they define the most probable model given the data.

Figure 1
51 Variables of the KAM Model

Performance Indicators
1. Gender development index 1999 (Human Development Report, UNDP, 2001)
2. Poverty index 1999 (Human Development Report, UNDP, 2001)
3. Composite ICRG risk rating 2000 (World Development Indicators, 2001)
4. Unemployment rate, % of total labor force 1996-98 (World Development Indicators, 2001)
5. Productivity growth (% change of GDP x person employed) 2000 (IMD World Competitiveness Yearbook, 2001)
Economic Incentives
6. Overall central government budget deficit as % of GDP, 1998 (World Development Indicators, 2001)
7. Trade as % of GDP, 1999 (World Development Indicators, 2001)
8. Intellectual Property is well protected (WEF Global Competitiveness Report, 2000)
9. Soundness of banks (WEF Global Competitiveness Report, 2000)
10. Adequate regulations & supervision of financial institutions (IMD World Competitiveness Yearbook, 2001)
11. Local competition (WEF Global Competitiveness Report, 2000)
12. Protection of property rights (WEF Global Competitiveness Report, 2000)
Institutional Regime
13. Regulatory framework (WBI, 1999)
14. Government Effectiveness WBI, 1999)
15. Voice and accountability (WBI, 1999)
16. Political stability (WBI, 1999)
17. Press freedom 2001 (Freedom House, 2001)
Innovation System
18. Technology Assessment Index (Human Development Report, UNDP, 2001)
19. Royalty and license fees payments (millions) (1999) (World Development Indicators, 2001)
20. Scientists and engineers in R&D per million 1987-97 (World Development Indicators, 2001)
21. Research collaboration between companies and universities (WEF Global Competitiveness Report, 2000)
22. Entrepreneurship among managers (IMD World Competitiveness Yearbook, 2001)
23. Easy to start a new business (WEF Global Competitiveness Report, 2000)
24. Availability of Venture capital (WEF Global Competitiveness Report, 2000)
25. Number of technical papers per million people 1997 (World Development Indicators, 2001)
26. Patent Applications granted by the USPTO (per million pop.) 2000 (USPTO)
27. Private sector spending on R&D (WEF Global Competitiveness Report, 2000)
Human Resources
28. Primary Pupil-teacher ratio, pupils per teacher, 1998 (2001 SIMA database)
29. Life expectancy at birth, years, 1999 (2001 SIMA database)
30. Management/worker relations (WEF Global Competitiveness Report, 2000)
31. Flexibility of people to adapt to new challenges (IMD World Competitiveness Yearbook, 2001)
32. Public spending on education as % of GDP 1999 (World Development Indicators, 2001)
33. Professional and technical workers as % of the labor force 1999 (ILO, 2000)
34. 8th grade achievement in mathematics (TIMMS 1999)
35. 8th grade achievement in science (TIMMS 1999)
36. National culture is open to foreign influence (IMD World Competitiveness Yearbook, 2001)
37. Companies invest heavily to attract, motivate and retain staff (WEF Global Competitiveness Report, 2000)
38. Mgt education locally available in 1 st class business schools (WEF Global Competitiveness Report, 2000)
39. Well educated people do not emigrate abroad (IMD World Competitiveness Yearbook, 2001)
40. University education meets the needs of a competitive economy (IMD World Competitiveness Yearbook, 2001)
Information Infrastructure
41. Telephones per 1,000 people, 1999 (International Telecommunication Union, 2000)
42. Mobile phones per 1,000 people, 1999 (International Telecommunication Union, 2000)
43. TV Sets per 1,000 people, 1999 (World Development Indicators, 2001)
44. Radios per 1,000 people, 1999 (World Development Indicators, 2001)
45. Daily newspapers per 1,000 people, 1996 (World Development Indicators, 2001)
46. Investment in telecoms as % of GDP 1998 (IMD World Competitiveness Yearbook, 2001)
47. Rating of computer processing power as % total world-MIPS 1998 (IMD World Competitiveness Ybook, 2001)
48. International telecoms: cost of call to US in \$ per 3 minutes, 1999 (World Development Indicators, 2001)
49. Information Society Index 2000 (IDC 2000)
50. % of Companies that use the Internet for electronic commerce (WEF Global Competitiveness Report, 2000)
51. ICT Expenditures as a % of GDP 1999 (World Development Indicators, 2001)

Source: World Bank Institute, 2002

APPLICATION OF A BAYESIAN APPROACH FOR KNOWLEDGE MANAGEMENT

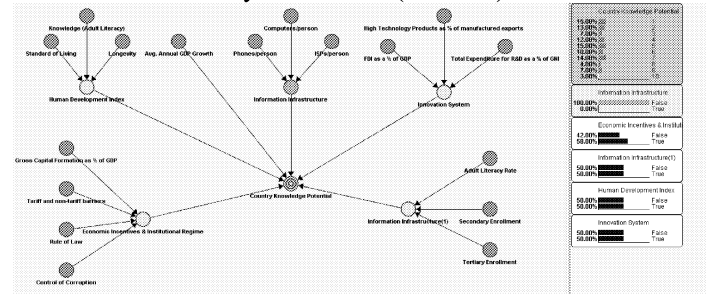
There are two key steps in this process:

- 1) To consider that all possible alternative models constitute a set of mutually exclusive explanations for the data, and add a probability distribution over them. If little is known about the model, it is assumed that all models are equally likely (as presented later in **Figure 3**);
- 2) To update the uniform probability on the basis the observed data to obtain the posterior probability of the model given the data. A Scoring Metric, a measure to assess the fitness of a model to the current data and a, search procedure, exploiting the set of possible models are used for this task. As a result, the marginal probability distribution of a chance node can be determined taking into account the knowledge coded by the network (either through database and/or experts) and possible observations (instantiations) made on other nodes.

After defining a Bayesian network, either by manually or by an automated discovery process from data (as defined in Step 1), model can be used to reason about new problems, for prediction, diagnosis, and classification. One of the most useful properties of Bayesian network is the ability to propagate evidence irrespective to the position of a node in the network, contrary to standard Data Mining methods. From parametric learning such as the estimation of the conditional probabilities to various types of structural learning algorithms make any type of data mining tasks possible including but not restricted to: “unsupervised learning for discovering the whole of the probabilistic relations, unsupervised learning for the search of new concepts and supervised learning for the characterization of a particular variable.” (Bayesia S.A., 2004)

Figure 3 shows a simple Bayesian network that models the 15 scorecard variables presented in Figure 2. The relationship between the areas of knowledge development and these 15 variables are assumed to be causal (represented by arcs or arrows), and variables are represented as nodes (circles). Using the available data from each variable respective source (such as those listed in Figure 2), probability distributions for each variable can be assigned and used to improve the Bayesian model. This process will determine the best model and will reveal the most accurate causal relationships with and between the variables and the macro-development areas. Additional data obtained by different simulation methods can be applied to the model to assess the predictive capabilities of the final model for KAM.

Figure 3
Bayesian Model (Baseline)

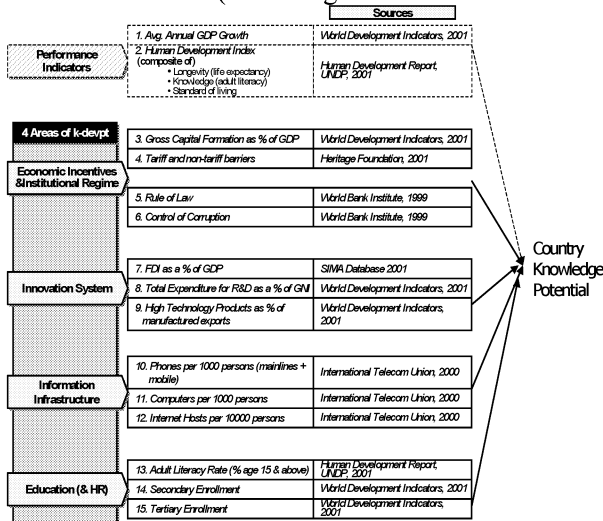


In Figure 3, the Country Knowledge Potential is identified by the Target Node, shown with circle rings. Knowing any of the variables current value, a new set of probability distributions can be calculated for all the child nodes. The model in Figure 3 is based on dummy data. Plugging real country data in the model will determine the causal links and dependencies. The power of this simulation model is that of assessing the impact on all knowledge activities caused by country level economical, social, and political change (expressed through sets of probability distributions). This model can be used to determine the impact of a change in a variable caused by a political decision, a cultural norm shift, an economical policy decision, a trade law and educational system change on the Country's Knowledge Potential. Extending this methodology to incorporate the time-effect on such data using Dynamic Bayesian models may reveal hidden economical, cultural and political variables that play role in Country's Knowledge Potential.

ADDITIONAL RESEARCH EFFORTS

In addition to defining the causal relationships based on real country data, further research applications will include the verification and refinement of the framework resulting from the simulation for application in different areas, and particularly in industry settings. The primary goal is to identify whether the derived model displays firm level and industry trends similar to the country level mapping. Future analysis will include expanding the methodology to

Figure 2
15 KAM Variables (Knowledge Assessment Scorecards)



Adapted From: World Bank Institute, 2002

different industries and developing a common template for Bayesian analysis and assessment across companies. The differences in the variables of interest and the resulting model will be compared and contrasted.

Reconstruction & Development Publisher. Washington D.C. ISBN 0 8213 4107

REFERENCES

- Arthur, W. B. (1996). "Increasing Returns and the New World of Business". *Harvard Business Review*. July/August, 100-108.
- Bayesia S.A (2004). "BayesiaLab 3.0 Online Documentation." France.
- Choi, C. J., and Lee, S. H. (1997). "A Knowledge-Based View of Cooperative Inter-organizational Relationships." In *Cooperative Strategies: European perspectives*, 33-58.
- Dahlman Carl and Andersson T. (2000). "Korea and the Knowledge Based Economy: Making the Transition." World Bank Institute and OECD Publishers. Washington, D.C.
- Dahlman Carl and Aubert J. (2001). China and the Knowledge Economy: Seizing the 21st Century. WBI Development Studies. The International Bank for Reconstruction
- Human Development Report (2001). "Making new technologies work for human development. United Nations Development Programme (UNDP)." Oxford University Press. New York. ISBN 0 19 521835 3
- Jensen, F. V. (2001). "Bayesian networks and decision graphs." New York: Springer.
- Jouffe, L. (2002). "Nouvelle classe de méthodes d'apprentissage de réseaux bayésiens." Journées francophones d'Extraction et de Gestion des Connaissances (EGC). Montpellier, January 2002
- Jouffe, L. and Munteanu, P. (2001). "New Search Strategies for Learning Bayesian Networks." Proceedings of Tenth International Symposium on Applied Stochastic Models, Data Analysis. June 2001
- Macdonald Lord (1999). "Scotland: Towards the Knowledge Economy. The Report Of The Knowledge Economy Taskforce." The Scottish Office Publisher. ISBN 0 7480 8220 4
- Munteanu P., Bendou P. "The EQ Framework for Learning Equivalence Classes of Bayesian Networks." First IEEE International Conference on Data Mining (IEEE ICDM). San José. November 2001.
- National Research Council (1996). "Prospectus for National Knowledge Assessment. Committee on Knowledge Assessment: Office of International Affairs. National Academy Press. Washington, D.C.
- World Bank Institute (2001). Monitoring the Republic of Korea's Implementation of the Strategy for the Transition to a Knowledge-Based Economy. Knowledge Networks and Distance Learning. World Bank Institute Publisher. Washington, D.C.
- World Bank Institute (2002). [On-line]. Knowledge for Development. URL: <http://www1.worldbank.org/gdln/kam.htm> [2003, February 1].
- World Development Indicators (2001). The World Bank Publisher. Washington, D.C. ISBN 0 8213 4898 1
- World Development Report. (1998/99). "Knowledge for Development." The International Bank for

A Business Simulation and Knowledge Management System for Complex Adaptive Systems: “Causalysator”

Carl Henning Reschke
University of Witten/Herdecke and University of Cologne
chreschke@yahoo.com

Abstract

The goal of this article is to describe a prospective system for dealing with complex systems.

The system that is proposed here is a business intelligence system that allows to deal with change in complex adaptive economic systems. Such a system is likely to increase users' awareness of his surroundings and thus to have reflexive repercussions in social systems. The main challenges are relate to endogenous change, emergence and reflexivity, while technically necessary elements are available albeit in dispersed form. The technical challenge therefore should lie in the integration of these elements into one system.

1. Introduction

Predicting the future and the course of developments of social systems has always been an ardent interest of humans. Examples range from the Greek Delphi oracle to utopian conceptions, in idealistic form such as Thomas Morus' "Utopia" or pessimistic form such as Orwell's "1984", and to literary visions such as Herrmann Hesse's novel "Das Glasperlenspiel" and Isaac Asimov's science fictions series "Foundation Trilogy". While the goal of predicting (and controlling as in the latter two examples) the future will remain ultimately elusive, steps towards such a vision are possible. The vision is ultimately elusive since the effects and variety of human reaction to others' action and conditions cannot be controlled. This is reflected in the unintended consequences and reflexivity concepts of sociology. But the usual patterns of human and social behavior in history and constraints of

human behavior can be discerned and accounted for.

Goertzel (2001) has proposed wide ranging changes in the web in the next decades towards the creation of intelligent web applications based on complex adaptive systems methodology. This requires the design and development of systems that reflect mental structures of humans. One example is the system for organizational and business intelligence described here. It is based on an evolutionary view of social systems and behavioral perspectives.

The knowledge management module is based on the concept of mental representations, reflecting the influence of perceptions of users on their interpretation of reality and their subsequent actions. In organizations, these perceptions are usually shared and trigger similar mental model – action complexes in group members that have been termed routines (Nelson and Winter 1982). One effect of such an explicit modeling will be an opportunity to reflect on those mental complexes. This will likely lead to optimization of mental representation and action.

2 Scientific Perspectives

2.1 Social Sciences

There is a longstanding debate about the evolutionary character and modeling of economic systems (Veblen 1898, Alchian 1950, Nelson and Winter 1982 to name some of the most important ones). Although there are a number of unsolved problems, the advantage of an evolutionary approach is that it a) allows to deal with endogenous change in such systems and that it b) nicely meshes with a complex systems view (see Reschke 2001). Evolutionary economics is a school of economics, located between neoclassical

economics and sociology, that focuses on the real behavior of human actors as well as a (co-) evolutionary view of economic systems. It allows to integrate organizational, industrial and system-wide levels.

Changes in the system become easier to discern, fluctuations can be separated from shifts in development processes, if the nature of change processes is studied. At this point the subject area of (evolutionary) economics, sociology, history and future studies begin to mesh: Flechtheim argues that Comte's sociology as well as List's economics were explicitly directed at studying the future (Flechtheim 1971). In a similar vein, Schiller (1789) argued that the study of history can be understood as enterprise to understand historical development processes with the aim of influencing developments. In response to Schiller and Kant, Hegel developed his theory of the "Weltgeist" and the dialectic process of history. This was related to the idea of a "qualitative mathematics" to deal with change. Evolutionary economics tries to cope with the problem of endogenous change, unfortunately with a focus on models of random change. Systemic evolutionary perspectives may have something better to offer here (Reschke 2004). But the identification of change that is "system breaking", will likely always remain problematic.

2.2 Complexity, Emergence and Structural Breaks

We all (hopefully) know that extrapolation does not always work and at some point turns into a danger. Therefore a methodology to deal with qualitative novelty and structural breaks is required. Also the system described here will not be able to predict how such changes will develop ultimately, but it may well increase the awareness that such a transformation is underway and may help to identify the trends that will turn out to be important.

Goertzel (1992) has proposed a measure to deal with emergence by comparing the

patterns in entities. This may help in detecting emergence and structural breaks ex post, and possibly even early on, but does not help in predicting them ex ante. I have suggested to use a framework based on multidimensional vectors for analyzing and modeling emergent events (see Reschke 2004). But the endogenous generation of change seems difficult and predicting system breaks troublesome. Catastrophe theory, developed by Thom and Zeeman on the basis of topology, seems to offer a better method for the prediction of system breaks part. Realistically modeling the effects of combinatory emergence seems doubtful given limitations of computer systems (Taylor 1998). The usual approach taken is to work with random generation of new features, since we do not know how to deal with the surprise element in emergence. While Geisendorf (2003) suggests using genetic algorithms to deal with the simulation of novelty, I think the best approach currently possible seems the use of two other methods of evolutionary computation: combining the representation of evolutionary algorithms with hierarchical recombination of program elements used in evolutionary programming.

2.3 Psychology: Personal Constructs

One perspective from psychology that matches in an evolutionary characterization of social processes is personal construct theory (Kelly 1955, Addams-Webber 1979, Bannister 1985). It posits that humans perceive and structure their world mentally along dichotomous conceptual poles. These concepts can be seen as elements in cognitive maps describing the structure and content of actors' worldviews. This perspective manages to show the existence of a binary, hierarchical ordering of human concepts. It focuses largely on changes in this ordering and self-perception.

Researchers can construct binary orderings of finite constructs by applying Boolean set theory to cognitive maps. This method,

called repertory grid, is based on the analysis of correlations between the elements of cognitive maps and some other phenomenon. It is possible to relate the binary orderings of human concepts to observable attributes of real world phenomena such as artifacts and to actions as well as conceptual structures. This approach runs also under the heading of cognitive architecture. The cognitive architecture of a concept is thus a similar dichotomous splitting of the attributes perceived relevant by actors leading to a hierarchical tree structure of entities. This suggests the application of various methods developed in computer science, among others the use of tree-based algorithms, and evolutionary computation (Banzhaf 1998 et al).

3 System structure

3.1 Knowledge Management Module

The system consists of two major modules, one for knowledge management and one for simulation. The knowledge management module is based on personal construct methodology. This allows to construct cognitive maps of the representations of reality in individuals and teams as well as binary trees of such representations where need be. The simulation engine is used to explore the interaction between elements in the knowledge base.

The basis of the system is a nested object-oriented freeform database system with meta description and linking capabilities. This allows the user to make arbitrary entries and to link knowledge “pieces”. Entity properties can be described with a meta-language such as XML. This information can be used by the simulation engine. Objects can be nested into each other. Knowledge pieces are objects with defined properties, which can be associated with behaviors or extended through a simple scripting language by the user.

Principally, the information entry is easy: if there is a new piece of knowledge to be entered, the users can create a new object or change an existing one. This means,

they define its properties respectively enter the knowledge pieces in windows. Relations between entities can be described by RDF schemas.

3.2 User Interface and Visual Structuring

Knowledge pieces are described textually and possibly graphically in individual windows. Meta information, like dates, keywords, assumptions and hypotheses, links to other pieces of knowledge, project etc. are attached to or displayed in the window belonging to a knowledge piece. Objects / windows describing knowledge pieces can be nested hierarchically.

Knowledge pieces can be represented by graphical symbols and the structure displayed in maps of the structure of knowledge systems. This structure can be changed by drag and drop operations changing the respective properties or structures in the underlying database.

3.3 Simulation Module

The simulation module can execute program modules linked to pieces of knowledge and exchange information between program modules. In this way hypotheses about relationships and effects of interactions between entities described in the KM module can be tested. This can be done on the level of subsystems as well as the whole system.

The simulation module is to use the non-standardized information that is supplied by the KM-module. This will at first only be possible by humans intervening and writing the required non-standard simulation code. The simulation is also to use the meta-information that is supplied by XML description of entities and RDF schemas. This part of the simulation can be automated. If all information relevant for simulation is contained in meta-descriptions (XML, RDF) and simulation modules exist that describe the behavior of entities based on these descriptions these simulations can be automated fully.

3.4 Simulation Process

The simulation engine is to run along a process that is described in differential equations that follow:

A population of entities O (products, technologies, and organizations) is defined as a set of objects with certain technical and social characteristics c with a value x :

$$(1) \quad p_i := \{O_i \{c_{xi}\}\}$$

Mental representations of users define the boundaries between objects with different characteristics. Fuzzy definitions reflecting inexact perception of actors are possible.

The development of a population over time depends on the number of individuals n_i in the preceding period and the statistical distribution D of relevant characteristics cx :

(2)

$$p_i^{t+1}(D_{\mu,\sigma}^{t+1}(c_{xi}^{t+1}), n_i^{t+1}) = p_i^t(D_{\mu,\sigma}^t(c_{xi}^t), n_i^t) + \Delta_i^{t \rightarrow t+1}(D, c, n)$$

,

with $\Delta_i^{t \rightarrow t+1}(D, c, n)$ as change operator, which affects the distribution, characteristics and population size.

I have not specified the change operator explicitly here, since the final answer requires empirical research on the characteristics of change and emergence discussed above. The beauty of the system is that this knowledge can be built with the knowledge base on events captured in the system. Partly, we can of course also rely on the results of the concerned disciplines. The precise form and effects of the change operator D depends on the rules of perception and human actions, the laws of emergence, the technical opportunities and constraints operating in a society at time t , and the attention problems receive.

3.5 A Case Example

While the tool may be used for other areas of complex systems that require the build-up of knowledge bases from diverse sources and testing of hypotheses on the basis of this information, we focus on a business application here:

A manager is interested in the potential of a new market and collects bits and pieces of information from diverse sources. While the process of search alone creates an intuitive idea of the market situation and competitive opportunities as well as threats, the tool requires him to make his assumptions and mental structures explicit in ordering that information and enables him to discuss them with colleagues. The manager tries to bring the collected pieces of information into one or maybe several alternative coherent structures and creates a visual map of relationships between the pieces of information. He is able to develop a notion of what information is missing and which hypotheses about the market need to be tested, to define project tasks and integrate these with project management tools. He is able to create an outline of relevant issues for reports (or presentations) and “compile” initial versions of documents containing knowledge pieces. Simulation: At first the equations presented above will not be much more than a mental guideline, a methodological tool, to consider relevant issues in the search process, but with the development of the system, the properties of the knowledge pieces should inform the simulation engine so as to be able to develop and execute a simulation of relevant events.

4 Strategic Implications

As not all users will have the time respectively understanding of the system’s workings, it will very likely be necessary to have a team of specialists that go over the entries of the “normal” employees and define entity properties properly as well as structure entities / object structures. This part may also be partly taken over by other employees checking on the validity of entries of colleagues harnessing the power of distributed knowledge creation. In the end, this may lead to the creation of a dedicated business intelligence unit in companies.

Further implications for strategy are possibly wide-ranging. Organizations can become easier aware of their perceptions of their environment and make adaptations easier, where necessary.

Inertia in management decisions despite mounting change pressure is an old problem in management science. Examples can be observed today in the context of companies founded on new technologies threatening the position of established players. While the cause is to be sought in the domains of psychology and social behavior, the solution might be an intelligently applied system such as the one described here. Organizational change becomes easier when differing representations of reality between team members and organization parts can be exposed. Likewise, it is principally possible to expose the differences in mental representations between competing organizations. This allows for instance to identify the degree of competitive threat from such competitors and may make it easier to identify where they are headed.

The simulation part allows to make “predictions” on future courses of actions under several assumptions. Initially, these simulations will not be realistic since relevant information is missing, but with time going by, a concerted effort should be able to build an electronic brain of an organization. If this “brain” contains the relevant information, it should be able to help an organization over the loss of employees and come up with predictions no individual would have developed.

5 Conclusion

The described system can be implemented in several stages focusing first on the knowledge management part and later on the simulation module. Elements of the system can be taken from the shelf but need to be integrated. Visual elements will improve adoption by users and increase productivity relative to textually oriented implementations.

It is crucial to inform possible early implementers about the time frame

required to construct a fully functional system. Developing the system module-wise would ease the task of development as well as the strain on implementers. The issues of emergence, change operators and reflexive effects due to the interaction of actors’ mental representations need further research.

Acknowledgements: I want to thank Arne Seib and Eike Kock for critical comments on the technical side of the ideas presented here and Wolfgang Streeck for making me aware of the social engineering issues in my thinking.

6 Literature

- Alchian, A. (1950): Uncertainty, Evolution, and Economic Theory, *Journal of Political Economy*, 58, 211-221.
- Addams-Webber, J.R. (1979): Personal Construct Theory. Concepts and Applications, John Wiley, Cichester etc.
- Bannister, D. (1985): Issues and Approaches in Personal Construct Theory, Academic Press, London etc.
- Banzhaf, W., Nordin, P. Keller, R.E., Francone, F.D. (1998): Genetic Programming. An Introduction. Morgan Kaufman, San Francisco, and dpunkt, Heidelberg.
- Geisendorf, S. (2003): Der ökonomische Evolutionsbegriff: Selbstorganisation und Selektion, paper for Buchenbach V workshop, to appear in revised form in: Buchenbach V, Metropolis, forthcoming.
- Goertzel, B. (1992): Self-Organizing Evolution, *Journal of Social and Evolutionary Systems*, 15,7-54.
- Schiller, F. (1789): Was heißt und zu welchem Ende studiert man Universalgeschichte? Akademische Antrittsrede am 26.5.1789 in Jena, appeared in Deutscher Merkur, November 1789, available via Project Gutenberg.
- Kelly, G.A. (1955): The Psychology of Personal Constructs, 2 vols., Norton, New York.
- Nelson, R.R. and Winter, S.G. (1982): An Evolutionary Theory of Economic Change, Harvard University Press, Cambridge, Ma.

Reschke, C.H. (2004): Economic Evolution, Behavior and Complexity, Paper for the AKSOE 2004 Workshop, March 8-12, Regensburg

Reschke, C.H. (2001): Evolutionary Perspectives on Simulations of Social Systems, *JASSS*, 4/4/8, <http://jasss.soc.surrey.ac.uk/4/4/8.html>

Taylor, T. (1998): Nidus Design Document. Departmental Working Paper No. 269. Department of Artificial Intelligence, University of Edinburgh, June 1998.

Veblen, T. (1898): Why is Economics not an Evolutionary Science, *Quarterly Journal of Economics*, 12, 373-397.

Bio: Carl Henning Reschke has studied business administration and economics at the Universities of Passau, Maastricht, California at Santa Cruz, Strasbourg (Institut Beta, ULP). He has worked as a researcher at MERIT and Content Manager/Business Analyst for an online community. He is taking courses in sociology and natural sciences at the University of Cologne and is currently completing a dissertation on 'Evolutionary Processes in Economics. The Example of the Life-Sciences' at the University of Witten/Herdecke, Germany.

TESTING METHODS AND ALGORITHMS FOR THE NEXT GENERATION OF KNOWLEDGE MANAGEMENT SYSTEMS

Joris Vertommen
Bert Vandermeulen
Joost Duflou
Centre for Industrial Management
Katholieke Universiteit Leuven
Celestijnenlaan 300A
B-3001 Heverlee, Belgium

Dries Van Dromme
Bart De Moor
SCD-ESAT
Katholieke Universiteit Leuven
Kasteelpark Arenberg 10
B-3001 Heverlee, Belgium

KEYWORDS

Knowledge management, KMS, clustering, user profiles.

ABSTRACT

This paper describes techniques which can be applied in advanced knowledge management systems (KMS) and reports results of two experiments which illustrate the applicability of these methods in real business environments. The techniques are mathematical and originate from the data mining and text mining fields of research, with an added management perspective. The incorporation of user profiles to characterise employees using the system is seen as another typical feature of a next generation of KMS, opening perspectives for better human resource management oriented applications. The need for such systems is motivated and a short description of the used data mining techniques is provided.

The experiments reported in this article demonstrate that there is a good correlation between the clustering results from the developed algorithms and document classifications made by human interpreters. Also, reduction of vector space dimensionality to decrease computation times on high document volumes has shown no problematic loss of classification quality.

INTRODUCTION

Need for New KM Tools

The fast penetration of digital documents, data carriers and web based communication in today's business has led to new needs concerning the management of these volatile data. Especially the fact that this increasing volume of electronic media also contains a substantial part of the company's valuable knowledge assets has enforced management's interest in good knowledge management tools.

The McKnow project (<http://www.mcknow.com>) aims at developing advanced algorithms and methods to help dealing with this changed situation in a more efficient way. At the same time, it combines the use of user profiles with a robust approach starting from unstructured data.

The Unconditional Automated Approach

Contrary to approaches in which documents are classified according to predefined taxonomies and enhanced with metadata, the unconditional approach to data and knowledge management does not impose preconditions about the way the information should initially be organised. This has the advantage that it is much less time-consuming because there is no need to build manmade taxonomies or to review and organise the load of digital media.

The following typical example illustrates the usefulness of this approach. When employees leave the company, they often leave a hard drive full of legacy documents which are still valuable, and may even contain the answer to a new employee's questions. The problem is that these media are often hard to access and recycle by other people than the author, because the document owner structured them in a way which seemed most logical to himself, but possibly not to other staff members.

The proposed automatic approach tackles this problem by processing the collection of documents fast and accurately without the need for human review or intervention. Useful data are identified, extracted and transferred to the main knowledge repository, where they stay available to other employees.

Applications in Business

Data mining techniques for automated knowledge identification and extraction are tailored for the World Wide Web, where they are applied to HTML-files from web sites (Ackerman et al. 1997). However, extending the use of similar techniques to documents in a business environment requires a number of adaptations:

- Flexible extraction: the content can be found in a number of different file types.
- Distributed environment: the files are spread over a number of different physical locations, hard drives or network spaces.
- User profiles: the fact that all users of the knowledge system belong to a limited and well-known group of people (the employees of the company) can be used to improve system performance. Integrating a profile for each user

creates new opportunities in querying and business analysis. It will not only be possible to link documents, but also to match users with documents and users with other users.

- Security and privacy issues are encountered. Companies do not want to open their databases to third parties, and employees do not like the feeling of being spied on.

DATA MINING TECHNIQUES

Vector Representation

Each document or user is represented by a vector, containing stemmed terms and the weights associated with each stem. These vectors (or profiles), the generation of which is briefly explained in the next two subsections, reside as knowledge entities in an n -dimensional space where n is the number of unique stemmed terms occurring in the total volume of textual data. With m the total number of documents and users, an $n \times m$ matrix allows to represent the knowledge carriers. Typically, such a matrix is very sparse (containing a large number of zeros).

The drawback of such an approach is that it introduces high dimensionality, making it computationally expensive and thus requiring a hardware configuration that is dedicated to this task. Improved algorithms, classification and dimensionality reduction help to overcome this problem.

Cosine proximity of vectors can be interpreted as semantic similarity between the original content of the documents or the user profiles, and forms the basis for clustering.

The cosine proximity between two vectors r and s is calculated as in (1)

$$\cos \theta = \frac{r \cdot s}{\|r\| \|s\|} \quad (1)$$

with θ the angle between vectors r and s and $r \cdot s$ the standard vector dot product, defined as in (2).

$$\sum_{i=1}^n r_i s_i \quad (2)$$

The norm $\|r\|$ is defined as in (3).

$$\sqrt{r \cdot r} \quad (3)$$

This is illustrated in the experiments reported in the next section. Using Euclidean distance in a normalised space will lead to the same result as calculating cosines.

By defining different vector spaces, multiple language environments can be covered.

Document Profile Generation

The profile vectors are derived from textual content. First, the language of the document is identified, and then the right stemming algorithm is triggered (Porter 1980, Kraaij and Pohlmann 1997). The stems of the words are counted (indexing) and weighed by a TF-IDF scheme (Salton and McGill 1986). This row of weighed stems forms the document profile vector.

User Profile Generation

The user profiles are generated (Hermans et al. 2003) in a similar way to the document profiles. Currently, each user is represented by a single vector, but it is possible to use a more refined profile consisting of several vectors (Çetintemel et al. 2001).

One approach for initialising a profile is to collect documents of interest to the user, which can be characterised by means of a document profile as described above. The user profile is then constructed as a combination of the document profiles, respecting the TF-IDF weights in the latter and, if applicable, a ranking of importance between these documents.

Clustering

Hierarchical clustering is used to increase the speed and accuracy of searching. For this purpose document and user profiles are clustered in a pre-processing step. At first, the search algorithm will allocate a higher priority to the cluster of vectors that is closest to the profile of the user who launches a query.

EXPERIMENTS AND RESULTS

Enterprise Context Testing

The presented approach was tested with a dataset collected in an international, R&D-oriented company.

Objective

The objective of this test was to prove that the clustering algorithm would group together people working in the same subdivisions of the company, thereby acknowledging a significant similarity between their profiles and a difference with those from employees working in a different subdivision.

Dataset

A set of documents was collected from each of 10 employees at the company. The document-sets varied from 8 to 34 in size. Each of the documents was assigned a weight of importance by the corresponding employee. Besides documents, function descriptions of the test subjects were also collected, from which a part of the company structure could be reconstructed. More precisely, a distinction was made between the R&D department and the IT department.

Experimental outline

In a first stage of this experiment, each of the users was assigned a number of profiles, each based on a different amount of documents. A pairwise comparison (by calculating the amount of shared terms) between the profiles in this collection was made per user. From the result of this test it could be concluded that no more than 8 documents were needed to establish a stable profile for a user in this test set.

In a second stage, the user profiles were fed to the hierarchical clustering algorithm, to see if the constructed clusters would reflect the company structure.

Analysis of results

The results obtained from the clustering algorithm are shown in figure 1.

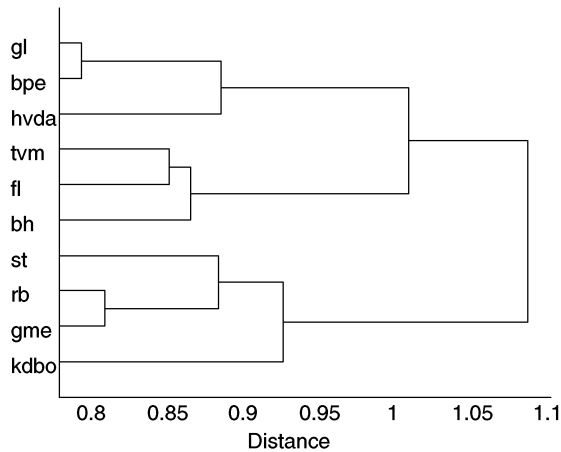


Figure 1: Dendrogram showing the Cluster-evolution

This dendrogram shows the evolution of the clusters being formed, starting with 10 clusters each containing one profile (at 0.79 or less on the horizontal axis). The abbreviations in front of each branch are the names of the employees in the dataset. One by one the profiles are grouped together in clusters, resulting in one big cluster at x-value 1.09. The abscissa value is a measure for the distance between clusters.

It is clear how, halfway the clustering-process, three distinct clusters are formed (shown in table 1).

Table 1: The 3 Main Clusters

Cluster	Members
A	gl, bpe, hvda
B	tvn, fl, bh
C	st, rb, gme, kdbo

When comparing the formed clusters with the company structure, we see that the algorithm has grouped together people who can be expected to have the same interests and/or competences;

Cluster A contains employees who work in the Research and Technology Development division. Cluster C, on the other hand, contains employees of the IT division. Cluster B required some more thorough analysis since the clustered employees are affiliated with different functional units of the company. When examining the data more in detail it became clear that these were higher-ranked employees and that they were clustered together on the fact that they had more general management interests than specialized technological skills.

This structure is reflected when investigating the most important stems on which profiles are grouped together (See table 2).

Table 2: Dominant Stems per Cluster

Cluster	Stems
A	fast, damp, modal, figur, structural, experimenta, sine, vibraat, ...
B	effici, compani, relationship, organisaa, longer, sale, memo,...
C	dataserv, intern, studi, extern, emb, client, xml, securiti, web, intranet

Summary

Constructing user profiles based on user-related documents and a TF-IDF weighing-scheme was found to be an effective way of representing a user. This conclusion stems from experiments on the use of the developed clustering algorithm to separate related users from non-related users using the constructed profiles. The experiment reported above forms illustration of this observation.

Newspaper Articles

A second series of experiments was performed on a database of Dutch language newspaper articles. For a number of articles, the author (or KMS user) was known, so these files were used to build the user profiles.

Objective

The objective was to test the statement that a user profile based on a limited number of documents is an effective representation of that person's interests and competences.

Dataset

The dataset consisted of 1776 newspaper articles (of which the authors were unknown) and 39 articles belonging to 13 known authors, all gathered in April 2003. The authors have different specialisations, as listed in table 3.

Table 3: Newspaper Authors (KMS users) and their Topics

# id	Name (abbr.)	Main topic
1	We.Ma.	Weather forecasting
2	To.Ys.	General domestic news
3	Pe.Va.	Culture (music)
4	Pa.De.	Economics
5	Mi.Do.	International news
6	Ha.Se.	Sports (soccer)
7	Gu.Te.	News & politics
8	Gu.Fr.	General domestic news
9	Fr.Co.	Sports (soccer)
10	Bo.Va.	News & politics
11	Be.Bu.	International news
12	Ba.Do.	Politics
13	Ba.Br.	Politics

Experimental outline

For each author, a vector was created using the textual input from the articles corresponding to those authors, resulting in 13 user profile vectors. For each of the 1776 anonymous articles, a document vector was created.

The vector space containing all user and document profiles (total of 1789 vectors) was normalised and subjected to hierarchical clustering. Results are visualised by a dendrogram as shown in figure 2. The stemmed terms in front of each branch are the words with the highest overall weight in that cluster, accompanied by the cluster id. The abscis value is a measure for the distance between clusters.

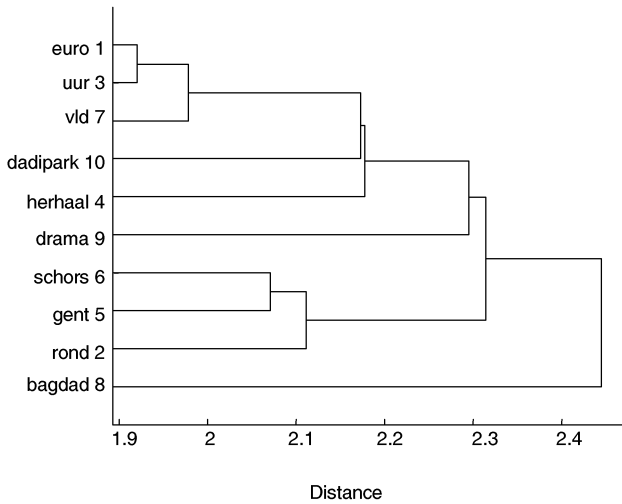


Figure 2: Dendrogram after Clustering into 10 Clusters

The vector space had dimensions 1789×39363 and was processed unmodified in a first stage. Secondly, a reduced space (1789×1000) was created. The differences will be discussed below.

Analysis of results

The coherence of the clusters formed by the algorithm was investigated and it was checked whether the classification can be considered logical from an external observer's point of view.

Looking at the documents in the different clusters, it can be observed that one large cluster and several smaller clusters were identified. The distribution of the documents over ten different clusters is shown in figure 3.

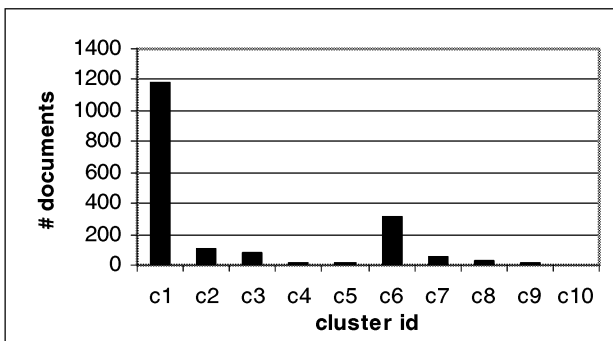


Figure 3: Distribution of Documents over 10 Clusters

This phenomenon was merely due to the nature of the dataset, which contained a vast number of often very small news items on various topics. The small clusters are homogeneous and have clearly identifiable topics such as sports (soccer and cycling), movies, or the Iraqi crisis.

Table 4 shows the topics of the clusters and the distribution of the authors over these clusters. The third column contains the total number of items (documents and users) in the cluster, the fourth shows which users are in each cluster.

Table 4: Cluster Topics

Cluster # id	Main topic of cluster	# items	Users in cluster
1	Various news items	1184	1, 2, 3, 4, 5, 8, 11
2	Sports (cycling)	99	-
3	TV programs	80	-
4	Elections	16	-
5	Listings	12	-
6	Sports (soccer)	305	6, 9
7	Politics	50	7, 10, 12, 13
8	Iraqi crisis	24	-
9	Movies	13	-
10	Local news	6	-

This result is satisfying because authors with specific interests are clustered together with documents on the same topic. For example, both authors 6 and 9 have been clustered into the "soccer" cluster, which is where they should be. Also, authors 7, 10, 12 and 13 are grouped together in the cluster about politics, which clearly matches their writings.

An overview of the processing times for the different tasks is given in Table 5. It can be seen that pairwise distance calculation is the bottleneck task. The tasks were completed on an Intel® Pentium® III, 1200 MHz with 256 MB RAM, running Windows® XP Professional.

In a second stage of the experiment, the number of dimensions of the vector space was reduced from 39363 to only 1000 by means of singular value decomposition (SVD) (Berry and Browne 1999, Golub and Van Loan 1989). This resulted in a decrease of needed computation time by a factor 211 for the pairwise distance calculation on this dataset.

When comparing the clustering results for this reduced space with the original results, we can only notice a small number of documents (<5%) shifting to other clusters, sometimes even improving the cluster quality. The main structure of the dendrogram remained unchanged.

Table 5: Processing times for the different tasks

Task	% of total time
Text extraction from original files	0,26%
Indexing and stemming terms	0,40%
TF-IDF weighting	0,04%
Normalisation of vectors	0,50%
Pairwise distance calculation	97,95%
Clustering	0,85%

Summary

The representation of users and documents in a vector space can lead to coherent clusters, based on human semantic interpretation of the texts. This way, users of the KMS can be clustered in a group of documents which are closely related to their area of interest or expertise.

Reduction of dimensionality by SVD of the document-term vector space does not seem to have a negative impact on clustering, which means it is safe to apply this technique to reduce calculation time.

CONCLUSIONS AND PROSPECT

From the experiments, it can be concluded that the developed profiling and clustering techniques allow to group or differentiate different users according to their fields of interest or expertise. This feature is useful when looking for experts in a specific field of knowledge, when composing teams, or for human resource management in general.

Also, documents can be processed to obtain clusters on the same topic. Moreover, grouping users and documents together creates the possibility to provide users with documents that are of interest to them, which is applicable in an information push system. The experiments indicate that it is possible to achieve an acceptable quality for this task.

Calculation times increase dramatically when the number of knowledge items in the vector space grows, but dimensionality reduction can help to tackle this problem, since it does not seem to have a serious influence on clustering quality.

Further research is oriented towards automatic updating of user profiles based on explicit and/or implicit system feedback. This is expected to result in a dynamic user characterisation that automatically adjusts to evolving user competences and domains of interest.

ACKNOWLEDGEMENTS

The authors would like to recognise the financial support from IWT-Vlaanderen (Instituut voor de Aanmoediging van Innovatie door Wetenschap en Technologie).

REFERENCES

- Ackerman, M.; Billsus, D.; Gaffney, S.; Hettich, S.; Khoo, G.; Kim, D.J.; Klefsstad, R.; Lowe, C.; Ludeman, A.; Muramatsu, J.; Omori, K.; Pazzani, M.J.; Semler, D.; Starr, B.; and Yap, P. 1997. "Learning Probabilistic User Profiles Applications for Finding Interesting Web Sites, Notifying Users of Relevant Changes to Web Pages, and Locating Grant Opportunities". In *AI Magazine* 18 (2), pp. 47-56.
- Berry, M.W. and Browne, M. 1999 *Understanding Search Engines: Mathematical Modeling and Text Retrieval*. SIAM, Philadelphia, pp. 53-59.
- Çetintemel, U.; Franklin, M. J.; and Giles, C. L. 2001 "Self-Adaptive User Profiles for Large-Scale Data Delivery". In *Proceedings of the 16th International Conference on Data Engineering*, San Diego, California, pp. 622-633.
- Golub, G.H., and Van Loan, C.F. 1989 *Matrix Computations*, 2nd ed. Baltimore, Johns Hopkins University Press.
- Hermans, K.; Duflou, J.; and De Moor, B. 2003 "Automated User Profile Generation For Knowledge Management Systems". In *Proceedings of the Knowledge Management Aston Conference*, Birmingham, UK, pp. 223-233.
- Kraaij, W. and Pohlmann, R. 1994 "Porter's stemming algorithm for Dutch". In *Informatiewetenschap 1994: Wetenschappelijke bijdragen aan de derde STINFON Conferentie*, Noordman, L.G.M. and De Vroomen, W.A.M. (Eds). pp. 167-180.
- Porter, M.F. 1980 "An algorithm for suffix stripping". *Program* 14(3) pp. 130-137.
- Salton G. and McGill M.J. 1986 *Introduction to Modern Information Retrieval*. McGraw-Hill, Inc., New York, NY.

AUTHOR BIOGRAPHY

JORIS VERTOMMEN holds degrees in Computer Sciences and Industrial Management, both obtained at the KULeuven, Belgium. His current research focusses on automated and user oriented methods and algorithms for knowledge management. More specifically, he is performing research on the construction and maintenance of user profiles for knowledge management systems.

joris.vertommen@cib.kuleuven.ac.be

BERT VANDERMEULEN holds degrees in Applied Biological Sciences and Industrial Management, both obtained at the KULeuven, Belgium. His current research focusses on automated and user oriented methods and algorithms for knowledge management.

bert.vandermeulen@cib.kuleuven.ac.be

WORLD INNOVATION NETWORKS: PROJECT MANAGING COLLABORATIVE NEW PRODUCT DEVELOPMENT PROJECTS ACROSS TIME-ZONES

Olaf Diegel
Institute of technology and Engineering
Massey University
Auckland, New Zealand
O.Diegel@massey.ac.nz

KEYWORDS

Innovation Management, Product Development, Project Management, Work Breakdown Structure

ABSTRACT

High-tech products that come to market six months late but on budget will earn 33% less profit over five years. In contrast, coming out on time and 50% over budget cuts profit by only 4% (McKinsey, 1989). If a company can develop its products on budget, but in shorter than expected time frames, they develop an immense commercial advantage and increased flexibility.

This paper describes the concept of World Innovation Networks, in which new product development occurs across time-zones, so that at the end of its working day, the first design team hands the project over to the next design team in an adjacent time-zone, and so on. This twenty-four hour a day design gives the advantages of much shortened total product development cycles, and increases the innovative potential of the new product.

The Design Breakdown Structure is proposed as the tool that forms the basis of effective project management and communication for these collaborative projects.

INTRODUCTION

In a perfect world we would be able to take any product development project that is not restrained by physical location and have three collaborative design teams located in different time-zones develop the product non-stop 24 hours a day. The cycle would begin with the first team working on the project for the eight working hours of their day, and at the end of the day, the project would be handed over to the next team in an adjacent time-zone that is just beginning its day. The second team would then hand it to the third team who, in turn would hand it back to the first team at the end of their day.

Again in a perfect world, the product would be developed in one third of the time required by a single design team.

Reality is, unfortunately, not quite so easy. The time currently required for one team to clearly and accurately communicate current project status is such that the second

team may take much of their day trying to gain an understanding of the project status before they even get started on real design work. The second design team must not only be able to understand the physical project status, but must also gain an understanding of design intent, technical problems, ideas, thought processes and constraints faced by the previous team so that they can continue along the same path. In fact, they might well decide to take a different path, but they should only do so consciously after they have all the information and understanding required to make an informed decision.

Managing New Product Development (NPD) is about managing innovation. Project managers and design teams therefore require a new paradigm in management tools that allow them to effectively monitor and control their projects, be they collocated or not.

Unless one can first overcome the communication and management difficulties encountered by such collaborative teams, the total project time is unlikely to be much reduced.

If, however, one can develop effective communication and project management protocols, one can then develop a tool that can be used to effectively manage such cross-time-zone projects within a reduced time-span.

This paper expands the concept of the Design Breakdown Structure (Diegel, 2003) to serve as an effective starting point for developing such a project management tool.

Before one can begin to truly manage innovation, one must first gain an understanding of the cognitive processes involved in innovation

A PERSPECTIVE VIEW ON INNOVATION

The following may almost seem like a contradiction in terms: Though the end product of innovation is about the new, the process of innovation is about the old.

In the words of Gerald F Smith, "We cannot imagine anything whatsoever, only things constructed out of existing knowledge. I cannot imagine a colour beyond my visual experience" (Smith, 1998).

Innovation is about finding new uses for things we already know. We cannot think of anything, new or old, that is not already stored in the knowledge banks of our minds. If we recall something from long-term memory without attempting to alter it, then it is just that: a memory. But if we take that data we have in memory, and use it for a different purpose than it was originally intended or if we take two bits of data and combine them to form something new and useful, then that is innovation. The process of innovation is therefore about taking existing knowledge and data and converting it into something novel and useful. (Barnett 1953, Amabile 1983).

So, if innovation is about using and adapting things we already know, then why is it so difficult to control? Is it completely random, or do we have some amount of mental control over the process? To try and understand these questions, we need to briefly delve into some background cognitive theory about how our brain thinks and resolves problems.

Mainstream cognitive theory behind how our thought processes function is relatively straightforward: We receive environmental inputs that are stored and processed in short-term memory. We are limited to around seven simultaneous external stimuli, and only a very limited amount of data can be stored in short-term memory for very a limited time (around 7 items for between 5 and 20 seconds). Short-term memory includes a scanning process that continuously scans the data stored in this buffer and determines its usefulness, and then either refreshes it or discards it. This combination forms the working memory, similar to the combination of RAM (short-term memory) and software (scanning process) in a computer. The results of working memory can be used to recall other data from long-term memory as, and when, required. (Blumenthal, 1977, Ellis, 1978).

Long-term memory could be compared to a hierarchical semantic database in which extra time is required (approximately 0.75msec) for each extra level of the database you descend to (Collins & Quillian, 1969). The links that join the nodes of data in long-term memory can weaken or get corrupted or broken with time, thus explaining memory loss or incorrect memory retrieval. Each data node can be stored as a mental representation of a single bit of data, as an image or, in the case of often used bits of data, as 'chunks' of data in which an entire concept is retrievable in a single operation (Birch & Clegg, 1996, Collins & Quillian, 1969).

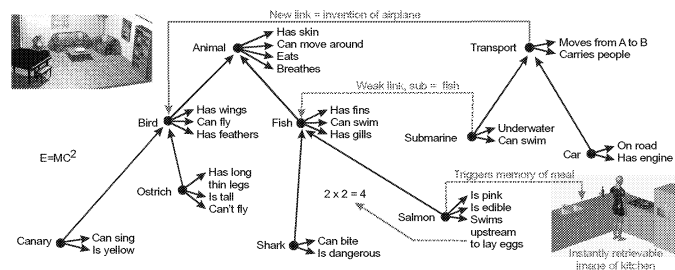


Fig 1: Long term memory hierarchical semantic database

When short-term memory attempts to retrieve data from long-term memory, it can follow any of a great number of data link paths to reach the data it wants to retrieve. It is our ability to create new links and choose which existing links to follow through our long-term memory database which allows us to be innovative or not.

If we follow a path of existing links, a memory is recalled. If we create a new link between one bit of data and another that was previously unrelated, or if we follow a link that would not, under normal circumstances, be associated with the current data in short-term memory, then that is a new thought.

One of the principles of Gestalt theory (A theory in which human beings are viewed as open systems in active interaction with their environment) is the law of Prägnanz, which states that psychological organisation will take the route requiring the least effort or energy in the achievement of the spatial and temporal stability of experience (Wertheimer, 1924). To put it crudely, it implies that the mind is lazy and will look for the simplest and most obvious links in long-term memory first. What this also implies is that thinking innovatively requires directed effort and a conscious decision to expend the extra energy required (Blumenthal, 1977).

THE MATHEMATICS OF INNOVATION

It is of interest to do some simple mathematical calculations in order to try and grasp the level of innovation that our minds are capable of. Based on the basic numbers suggested by mainstream cognitive theory we know that short term memory is capable of containing seven items at any one moment in time, and that scanning long-term memory requires approximately 75msecs per node traversed [Blumenthal, 1977, Ellis, 1978, Birch & Clegg, 1996, Collins, & Quillian, 1969, Atkinson & Shiffrin, 1971].

Based on the assumption that innovation is in fact based on our ability to recombine or find new uses for knowledge we already possess, one can do a calculation of possible combinations.

A combination of n taken m at a time is defined as a selection of m out of the n items without regard to the order. The total number of all the possible combinations is denoted as:

$$C_n^m = \frac{n!}{m!(n-m)!} = \binom{n}{m} = \binom{n}{n-m} \quad \text{where } n \geq m$$

Repeating and summing this equation for the seven items in short-term memory in groups ranging from one to seven tells us that at any one instant we have the potential for 127 new combinations in memory, which also means that at any one moment in time we have the potential for 127 innovations.

This assumes however that we think in combinations only, and not permutations, which is not necessarily the case. In other words, is our innovative outcome identical if we consider the concept of an apple combined with an orange, as opposed to an orange and an apple?

If one considers that we may have different innovative outcomes based on permutations of the data in short term memory, then the total potential for innovations can once again be calculated:

A permutation of n taken m at a time is defined as an ordered selection of m out of the n items. The total number of all the possible permutations is denoted as:

$$P_m^n = n(n-1)(n-2)\dots(n-m+1) \quad \text{where } n \geq m$$

Once again using the seven items in short-term memory, this calculation suddenly expands the number of innovative outcomes from 127 to 5913.

Taking these numbers a step further to calculate our innovative potential over a period of one second we find that over a period of 1 second, we have the potential to traverse 1333 data items in long-term memory (based on 75msecs per node) and process them in short-term memory. Repeating the combination and permutation calculations based on this number, we can now calculate that, over a period of one second, we have the potential for 16,291 innovations based on combinations, or 7,882,209 based on permutations.

Now, whether we take the lower or the higher of these two numbers, we can see that we have a tremendous potential for innovation at any one time. This then raises the question as to why we therefore have so much difficulty in thinking innovatively.

FUNCTIONAL FIXEDNESS, ONE OF THE ROADBLOCKS TO INNOVATION

One of the main mental barriers that prevent us from easily forming the new ideas with existing data is functional fixedness (Duncker, 1945). Functional fixedness is the effect of only being able to see something for what we have traditionally been taught it is by our education, environment, culture, etc.

A commonly used example of functional fixedness is Maier's two-string problem (Maier, 1931). In this problem, the subject is in a room with two strings tied to the ceiling. Both strings are of equal length. The objective is to tie the ends of the two strings together. The problem is that while the strings are long enough to be tied together they are short enough that one is unable to just take hold of one string, walk over to the other string, and tie them together. Scattered around the room there are a number of objects. These objects include a plate, some books, a chair, a pair of pliers, an extension cord, and a book of matches.

To resolve a problem, the real source of the problem must first be located. In the above example, the fundamental source of the problem can be viewed as one of the following:

1. The string is too short
2. My arms are too short
3. The end of one string won't stay anchored in place while I get the other
4. The string won't come to me

Depending on what objects are located around the room, any one of these problem sources can be resolved by, for example, using an object such as the extension cord to lengthen one of the strings, or using an object (such as a chair, for example) to lengthen one's arms, etc. However, if the only object in the room were a pair of pliers, then the solutions become much more limited, as this object cannot be used to resolve all the possible problem sources.

About 60% of the participants in Maier's study failed to find a solution within a 10-minute time limit.

These participants saw the pliers only as the traditional tool they are, not recognizing that the pliers could be used as a pendulum bob, swinging at the end of one of the two strings, thus resolving the "string won't come to me" problem source.

Most of us have difficulty in seeing the pair of pliers in the above example, as anything other than a tool as that is what we have always been taught they are. Through force of habit, we are fixated by the fact that the object's function is that of a pair of pliers. If we can overcome this fixedness, then we can see that they could have many other uses. The pair of pliers could be used as a weight (paperweight, pendulum weight, weapon, fishing sinker, etc.), an electricity conductor (emergency fuse, car jump start kit, etc.), and so on.

The reason that overcoming functional fixedness is so important is that, as innovation consists in finding new uses for knowledge we already have, we need to try to get past the barrier that a particular bit of knowledge only has the use it was originally intended for.

The Creativity Breakdown Structure and Design Breakdown Structure techniques proposed in this paper help to overcome functional fixedness by breaking the object or concept one is looking at into its fundamental problem areas.

THE CREATIVITY BREAKDOWN STRUCTURE

The Creativity Breakdown Structure (CBS) technique revolves around the following: Developing an innovative idea always consists in finding a novel use for an existing object or concept.

As an example, if one were looking for novel uses for a brick, one might come up with several initial ideas (door-stop, paper weight, etc.). Our functional fixedness about

what a brick is will have a tendency to limit the alternative uses we can find for the brick. However, if one breaks the brick down into the fundamental properties that make it up (such as weight, rectangular, heavy, porous, does not conduct electricity, rough, small enough to be picked up in one hand, holds heat, etc.), one can usually come up with many more ideas for its use, than when one is thinking of just the brick as a whole.

When problem solving, which is a large component of the product development process, one is in effect attempting to use the above technique, but in reverse. One has already identified the properties that one requires, and one is looking for the “brick” or other solution that will fulfil those properties in a novel way.

The Creativity Breakdown Structure helps us to graphically break down each problem to overcome into its fundamental problem areas (such as theory, knowledge, energy source, timing, cost, equipment, materials, components, mechanical, etc.), so that we can identify what the properties that the potential solution will have to possess, and thus reverse-engineer it.

It should be noted that breaking things down into sub-units, or decomposition as it is often referred to in product development, is nothing new. However, though of some help, just breaking a product down into subassemblies is of limited use, particularly if done in a non-graphical format such as a list. It is the graphical and structural nature of the Work Breakdown Structure (WBS), Design Breakdown Structure (DBS) and Creativity Breakdown Structure (CBS) that make them powerful tools, as one can almost instantly see the relationships and hierarchy between the various elements that make them up. They also act as a common communication tools so that, when talking about the idea generation process, all team members are talking the same language.

PROJECT MANAGEMENT CURRENT PRACTICE

Within the field of modern project management, a Project is broadly defined as an activity of finite duration with the ultimate delivery of a defined goal (Wideman, 2002).

The development of new products fits this definition of a project perfectly. Therefore one should, in theory, be able to apply standard project management practices to product development projects and thus be able to effectively manage them, and deliver better products that delight the customer and are on time and on budget.

Though current project management practices do undoubtedly help in managing certain aspects of product development projects, they are also somewhat lacking in their ability to help manage the many innovative factors often required by product development projects. This effect is magnified when one examines product development projects that are constrained by tight deadlines as many often are.

A large number of product development and project management models, as shown in table 1, have been developed by a great many people from a variety of disciplines.

They all contain a number of distinct phases or stages that can be generically diagrammed as a combined project management / product development model somewhat along the following lines:

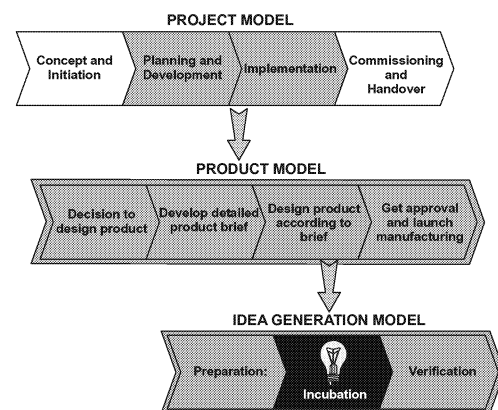


Fig 2: Models for Project, Product and Idea Generation

Each of these phases is further broken down as required and each phase has a reasonably well-defined number of tools, techniques and methods developed to help control it.

Thought formation	Scientific method	Research	Discovery & invention	Invention	Design & Manufacture		Management
Wallace, 1926	Encyclopaedia Britannica	Walkup, 1958	Zwicky, 1969	Rossmann, 1964	Azimow, 1962	Von Fange, 1959	Army, 1964
Preparation	Observation and experiment	Problem formulation, fact finding analysis	Problem concisely formulated	Observation of a need or difficulty. Analysis of need, Survey available information	Analyse problem situation	Investigate direction	Problem identification, Problem research, Problem definition
			Localise/ analyse important parameters			Establish measures	
Incubation, Illumination	Analysis and synthesis, Hypothesis inference	Incubation, Decision	Build morphological box containing all potential solutions, Evaluate solutions	Formation of all objective solutions, Analysis of solutions, birth of new ideas	Synthesise solutions, Evaluate and decide	Develop methods	Idea hopper, Idea filter
Verification	Comparison and analogy	Action	Optimum solutions selected and applied	Experimentation test, selection, perfection of final embodiment	Revise Implement	Complete solution	Idea tester Formulate plan for change, Co-ordinate action
						Convince all	

Table 1: Comparison of Product Development Processes

It is these methods and techniques that allow designers and project manager to control their projects.

Within these models, one area that is invariably under-defined in terms of methods and techniques for control is that of the Idea Generation stage, which is almost always described as a "black-box" process (Incubation and Illumination are other words often used to describe this process) in which various inputs are entered, and over an undefined period of time, something magically occurs and a good workable idea is generated as an output. (Amabile, 1996, Kmetovicz, 1992, Rosenthal, 1992, Reilly, 1999, etc.)

What this means from a project management point of view is that the manager has an extremely comprehensive toolbox of methods and techniques to deal with all but the most vital area of his project, that of generating the ideas that will make his product a success, and tell him what the very project he is working on actually involves.

The project management toolbox does include a tool that is vital to the management of the project called the Work Breakdown Structure (WBS) that gives details of the product or deliverable to be made. What is lacking is a tool to allow us to effectively manage and control the construction of the Work breakdown Structure.

THE WORK BREAKDOWN STRUCTURE

As defined in PMBOK (1987), A Work Breakdown Structure is a task-oriented 'family tree' of activities that organizes, defines and graphically displays the total work to be accomplished in order to achieve the final objectives of a project. Each descending level represents an increasingly detailed definition of the project objective. It is a system for subdividing a project into manageable work packages, components or elements to provide a common framework for scope/cost/schedule communications, allocation of responsibility, monitoring and management.

The WBS consists in breaking down the project into the various work packages that make it up. From these work packages, schedules, task lists, resource allocations, etc. can be derived thus allowing the project manager to effectively manage and control the project.

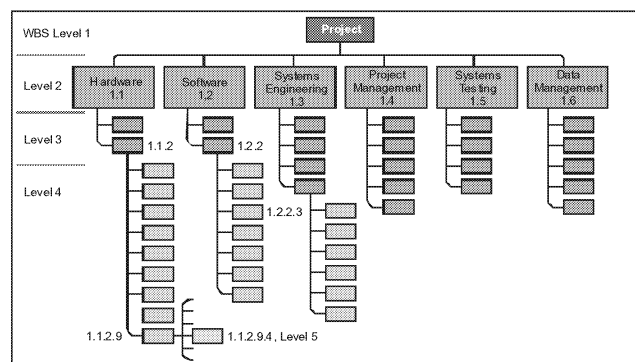


Fig 3: Sample of traditional Work Breakdown Structure

Thought the WBS is invaluable for helping to manage projects that are relatively well defined (in other words, something for which you already have a good set of plans), it does not cater well for areas of unknown. How do you break down something that you do not yet know the shape of, or how you are going to achieve? The catch with the WBS is that it can only be generated once we know details of the product that will be produced from it. Once the project manager has a Work Breakdown Structure he can relatively easily manage the project from that point forward.

A better set of tools is therefore required to help manage the early conceptual stages of NPD projects before a meaningful WBS can be generated.

THE DESIGN BREAKDOWN STRUCTURE

By adding a further refinement to the WBS, which we will call a Design Breakdown Structure (DBS), we will have a tool that will permit us to deal with those areas of unknown, and to construct the foundations of the product that will eventually be transformed into the WBS.

The DBS consists in breaking down areas and components into the fundamental problem areas that need to be resolved in order to achieve those goals set out by the component. Much like the WBS, it is a graphical family tree showing all the problems that need to be overcome, potential solutions to these problems, and how they are related. The DBS could also contain several CBSs for individual problems that may need to be resolved.

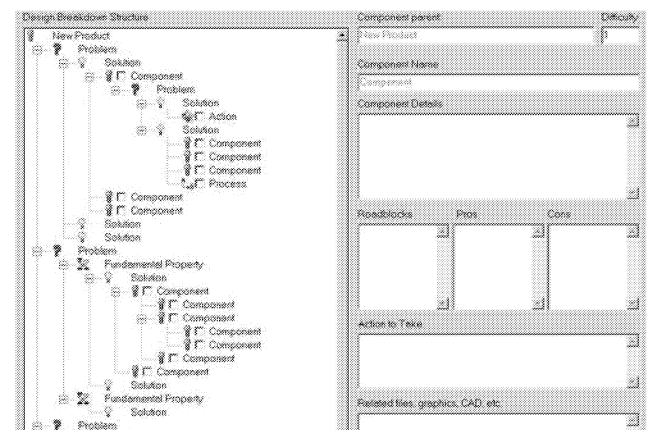


Fig 4: The Design Breakdown Structure

Figure 4 shows a much simplified diagram of a DBS showing only the components, problems, and solutions. Other elements that would normally be included (if relevant) in the course of constructing a DBS are:

- **Actions:** Any action needing to be taken before a problem can be solved. If further research and information is required before a problem can be solved, then this becomes a task that the project manager and team need to know of so that it can be effectively managed.

- New Processes: Certain solutions may require fundamentally new processes within the organisation. The adoption of a nuclear battery, for example, may require the organisation to change the way they operate, and an awareness of this therefore needs to be taken into account.
- TRIZ elements: TRIZ (the Theory of Inventive Problem Solving) is a useful problem solving technique that isolates the technical contradictions within a problem and then use one or more of 40 inventive principles to resolve them. It complements the Creativity Breakdown Structure method well, as contradictions can often be derived from the Creativity Breakdown elements.

As each problem area is identified, alternative solutions can then be listed below their respective problem areas. Each alternative can then be further developed in a manner similar to the conventional WBS, showing the deliverables that that solution may contain. This would generally however only be done in order to allow the comparison of alternative solutions. For the comparison of potential solutions to be effective, one also needs to develop a good framework of comparison criteria, together with a project by project weighting system.

Whereas a WBS concentrates on the 'WHAT' outputs of the project, the DBS puts emphasis on the 'HOW' questions of the project.

In essence the Design Breakdown Structure attempts to map the “black-box” process of incubation, thus allowing the project management a certain amount of control over this process.

The DBS is a structured, hierarchical, graphical

documented team process where the entire design team, including the project manager, works through the fundamental issues they must overcome to reach their preset goal. It is this well-structured documentation that acts a road map of the project, and spreads an awareness of both the issues that each individual designer faces, and the interconnectedness of all the issues. It forces the team to work as a coordinated, well-knit unit rather than as a team of isolated individuals. This road map is essential towards allowing the project manager to effectively manage the project.

The DBS also helps to relatively easily isolate the true constraints of the project as well as isolate those areas that require high levels of creative thinking and those that require merely basic knowledge or mechanical design.

In a recent series of tests, conducted at Massey University, the Design Breakdown Structure was implemented in the form of a software package called InnovationWorks (www.massey.ac.nz/~odiegel/win/downloads.html), and used to manage two major Mechatronics research projects. The first project, for the design of an Autonomous Guided Vehicle (AGV) showed design time reductions of over 30% as well as the development of several new and innovative methods of constructing omni-directional wheels. The second project, for the development of an Internet controlled robotic lawn mower, showed time reduction of over 40%.

It should be noted that it is difficult to categorically state that development time have been reduced when such a project has never been undertaken in the past. The time estimates above therefore rely on the expertise of people (Professors and Senior lecturers) who have the appropriate knowledge to estimate how long the project should have taken with conventional approaches.

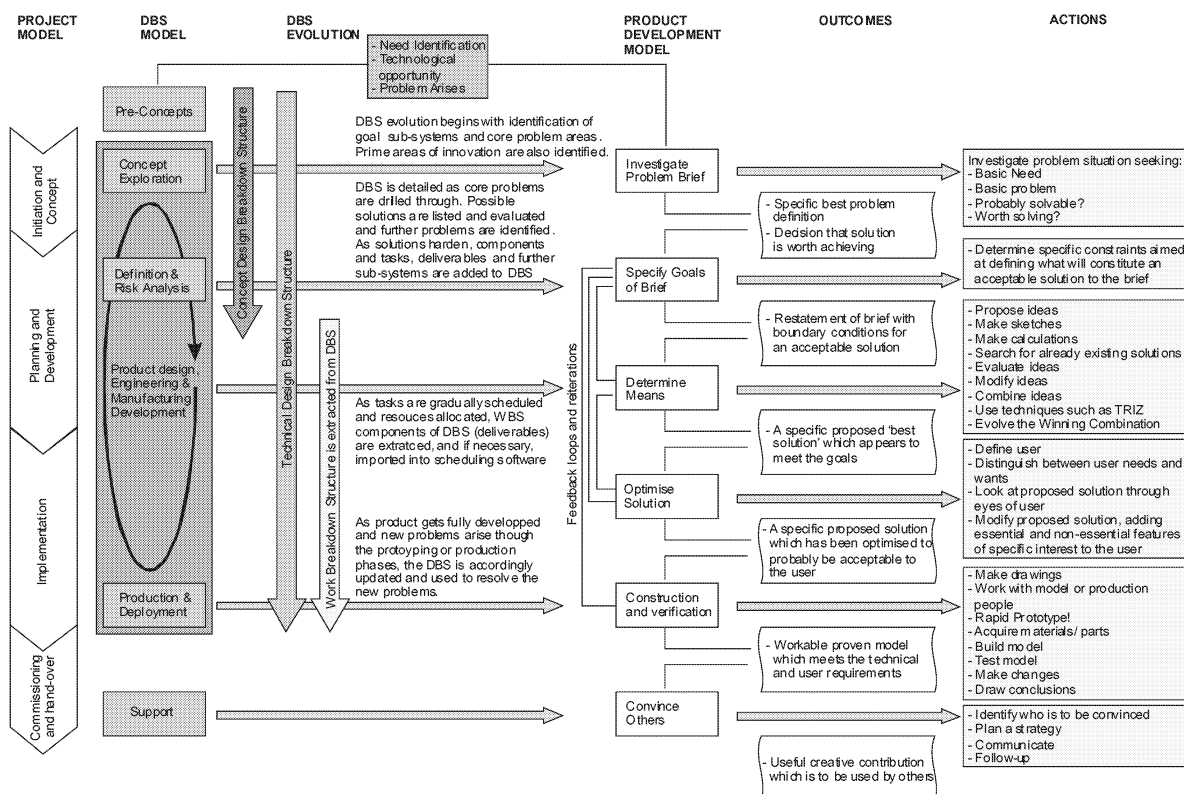


Fig 5: The Design Breakdown Structure Evolution

APPLYING THE DESIGN BREAKDOWN STRUCTURE TO NON-COLOCATED COLLABORATIVE TEAMS

The principals of World Innovation Networks are relatively simple. Several product development teams, located in different countries, work in collaboration on a single project. The idea is to achieve 24 hour a day product development by having each team pass the project on to the next time zone at the end of their day.

There are several technical problems involved in the concept, including:

- File and document transfers
- Computer Aided Design file sharing
- Video & Telephone conferencing
- Live project monitoring
- Centralized data repositories
- Knowledge databases

These technical problems are, however, relatively easily overcome through the use of technologies such as broadband Internet connections, common project file web servers, and advanced on-line 3D collaboration tools.

The main point of importance in this area is that the technology should not hinder the flow of communication. This, of course, may have implications in that the format of communication may need to be adapted to suit currently available technologies. An example of this adaptation may be to, for example, use .ISF format files for 3D data exchange as they only require a few hundred Kb compared to the more conventional CAD formats which often require many Mb. These tools and technologies are improving rapidly, and will make this type of collaboration easier and easier.

The area that poses the most difficulty is in the communication of ideas, and intangible factors such as design intent, potential roadblocks (be they mental or physical), and overall thought processes. This is an area where the highly structured by flexible Design Breakdown Structure can come into play to help alleviate these difficulties.

Using a Design Breakdown Structure has benefits to both the project manager and to the design team.

It allows the Project Manager to:

- Better understand needs or obstacles the designers may come across.
- Understand, Manage and control the process that occurs in the “black box” prior to the generation of a WBS
- Act as a better facilitator in getting the designers the resources or information that they need to overcome the problems.
- Allocate his design resources more accurately, as he can match the designer with the expertise best suited to solving each problem area.

- More accurately monitor the project, as the sections of the design have been broken into smaller more easily measurable tasks.

For the design team it:

- Crystallises all the problems that the designer may need to overcome and shows the relationships between all the elements of the product
- Helps to overcome the functional fixedness that often acts as a barrier to innovative solutions
- Gives them an idea of which areas will require the most creative thought
- Gives the designer has a distinct number of paths to follow
- Gives the designer the ability to give a more accurate estimate of the time it will take to come up with a solution
- Allows design teams to work together as a unit, where one persons' ideas generates new ideas in others

As the DBS must be well documented to be truly effective, the documentation (particularly if computerised) acts as a knowledge database for future projects.

For the DBS to be effective it must be used as a live communication tool, and much like any other project status type tool, it must be kept up to date and grow with the project.

A DBS would normally be used as an integral part of a WBS. A DBS can be used either for the product as a whole, or for individual elements of the product WBS. DBS elements are often drawn directly from the elements of the WBS and, having gone through the DBS process, are then reintegrated to the WBS. Once the product design has been finalised, if a traditional Work Breakdown Structure is required or an existing WBS needs to be updated, it can easily be generated from the Design Breakdown Structure, by simply removing all the ‘non-deliverable’ elements. The DBS can then be integrated to form part of the WBS dictionary, where each DBS element can be included with it’s corresponding WBS element, thus showing the process that was gone through to reach the final solution and strengthening the relationship between the various WBS elements.

CONCLUSIONS

Ultimately, the objective of any research into product development must be to improve the development process in terms of speed and product outcome quality (cost generally being a factor of these). Using Design Breakdown Structures as an extension of existing project management and design management techniques can be a useful aid in shortening the design time required by innovative product development projects. It is a tool that is of great help to overcoming the functional fixedness that often acts as a barrier to efficient innovative product development.

It is of benefit to the design team in aiding them to comprehend all the issues involved in the product design, and in helping them to communicate design intent and ideas. It is of benefit to the Project manager in that it gives him the information he requires to fulfil some of his roles as facilitator, scheduler, prodger, and negotiator, etc.

One of the main keys to innovation is the ability to convert existing knowledge, ideas and concepts into new ones. The Creativity Breakdown Structure component of the Design Breakdown Structure facilitates the designer team's ability to break a problem down into its fundamental building blocks, thus allowing them to more easily overcome the functional fixedness which prevents them from generating new ideas out of existing knowledge.

For non-located product development teams, all working on the same project, effective communication is the key difficulty that must be overcome in order for the team to work efficiently, and not have to waste time in trying to understand the previous teams, design intent and thought processes. The Design Breakdown Structure acts as the starting point for a core communication tool in effectively communicating design intent, problems encountered, potential solutions, thought processes, etc. It allows the early conceptual stages of NPD projects to be managed and controlled.

The Design Breakdown Structure is a tool that can be extremely effective when used in close conjunction with the traditional Work Breakdown Structure. It must be remembered, however, that the Design Breakdown Structure is not suggested as a replacement for the traditional Work Breakdown Structure, but rather as an extension of it.

The Design Breakdown Structure, when used in the context of World innovation Networks, gives the design team the common understanding of the project that they need in order to reduce the total development time.

REFERENCES

- Amabile, Teresa M., *The social psychology of creativity*, Springer-Verlag, New York, 1983
- Amabile, Teresa M., *Growing up creative*, Creative Education Foundation, 1050 Union Rd., Buffalo, NY 14224
- Asimow, M., *Introduction to Design*, Prentice-Hall Inc, 1962, pp42
- Barnett, H.G., *Innovation The basis of cultural Change*, McGraw-Hill Book Company Inc, New York, 1953
- Birch, P., & Clegg, B., *Imagination Engineering, The toolkit for business creativity*, Pitman publishing, London, 1996, p3-25
- Blumenthal, Arthur L, *The Process of Cognition*, Prentice Hall Inc, New Jersey, 1977, p19-33
- Collins, A.M., & Quillian, M.R., *Retrieval time from semantic memory*, Journal of verbal learning and verbal behaviour, 8, pp240-247, 1969

Diegel, O., *Design Breakdown Structures: a tool for innovative targeted problem solving in collaborative product development teams*, PhD Thesis, Massey University, New Zealand 2003

Dunker, K. (1945). *On problem solving*. Psychological Monographs, 58, Whole No. 270.

Ellis, H. C., *Fundamentals of Human Learning, memory and Cognition*, 2nd Edition, Wm.C.Brown Company, Iowa, 1978, p86-98

Encyclopaedia Britannica, "Scientific Method", Vol 20, Encyclopaedia Britannica Inc, 1949

Kmetovicz, R. E., *New Product Development Design and Analysis*, Wiley-Interscience Publication, John Wiley & sons inc. New York, 1992

Maier, N. R. F. (1931). *Reasoning in humans II: The solution of a problem and its appearance in consciousness*. Journal of Comparative Psychology, 12, 181-194.

McKinsey and Company, Fortune Magazine, Feb 13, 1989

PMBOK, Project Management Body of Knowledge, Project Management Institute, 1987

Reilly, N., *The Team based product development guidebook*, ASQ Quality Press, Milwaukee Wisconsin, 1999

Rosenthal, S., *Effective product design and development, How to cut lead time and increase customer satisfaction*, Business One Irwin, Homewood, Illinois 60430, 1992

Rossmann, J., *Industrial Creativity: The psychology of the Inventor*, University books, 1964, pp288

Smith, Gerald F., *Quality Problem Solving*, ASQ Quality Press, Milwaukee, Wisconsin, 1998

US Department of Defence Handbook, *Work breakdown Structure*, MIL-HDBK-881, 1998

Von Fange, E., *Professional Creativity*, Prentice-Hall Inc, 1959

Walkup, L., *Individual Creativity in Research*, Battelle Technical review, Battelle Memorial Institute, Columbus, Ohio, August 1958

Wallace, G., *The art of thought*, New York, NY, Harcourt, Brace, 1926

Wertheimer, M., *Über Gestalttheorie*, the Kant Society, Berlin, 1924, translated by Willis D. Ellis published in his "Source Book of Gestalt Psychology," New York: Harcourt, Brace and Co, 1938.

Wideman, M., Combined from the various definitions at: *Wideman Comparative Glossary of Project Management Terms 2.0*, <http://www.maxwideman.com/pmglossary/index.htm>, Jan 2000

Workbook for military creative problem solving, US Army management school, 1964

Zwicky, F., *Discovery, Invention, Research*, First American Edition, The Mac Millan Co, 1969, pp107

Introduction of evolutionary computation into a shared process : product design and project management

*Claude BARON, Member, IEEE, **Daniel ESTEVE, *Citlali GUTIERREZ

*LESIA, INSA, 135 av. de Rangueil, 31077 Toulouse cedex 04, France

**LAAS, CNRS, 7 avenue du colonel Roche, 31077 Toulouse cedex 04, France

e-mail : claud.baron@insa-tlse.fr, daniel.esteve@laas.fr, gutierre@insa-tlse.fr

Abstract— This paper explores the interest and the possibility to join system design and project management methods and tools. Our motivation is to prevent the obvious incompatibilities between technical objectives and socio-economical requirements in the enterprise. What we recommend is to work on a generic unique model based on the classical top down design steps, to which costs models and non-functional requirements are associated. Project management thus appears as an activity of diagnosis and optimisation, allowing to choose certain realisations between the different possible scenarios and to optimise the management by an allocation of tolerances, which is calculated for each supplier on the base of a global objective. This analysis concludes on the interest of two complementary tools : the evolutionary algorithms to arbitrate the scenarios, and the Monte-Carlo methods for the allocation of tolerances.

Index Terms—evolutionary computing, selection and optimization techniques, modeling, system design, project management.

I. INTRODUCTION

The accelerated development of technologies offers a wide range of materials, components, production modes ... to the engineer. This is useful to design and sell products which life time is short : either products are destined to be consumables, either they become old-fashioned in front of more innovative products. Manufacturers, in this very concurrent international context, must be very reactive to their client needs and very efficient to shorten the time to market.

Once the product specifications are established, manufacturers have to simultaneously master both the product design methodology and project management. These dual of this shared problem are separately conducted and disposed of separated tools: CAD tools, economic planning tools, financial tools.... This practice presents risk of incoherence and lengthening delays for at least two reasons:

- the innovation process is not correlated enough to economical requirements, and enough introduced into the company life,
- project management is conducted on an insufficient knowledge of technical difficulties.

Obviously, these risks could be reduced if all technical, administrative and financial decisions relied on a shared model between partners.

In this perspective, our proposition relies on five items :

- this unique model must rely on a detailed representation of tasks and technological steps induced by the product design,
- the tasks and methodological steps must be detailed according to a refinement process till the practical designation of supplies and suppliers,
- the results of the previous steps must conduct to the development of possible scenarios : planning and scheduling,
- these scenarios must be improved by these pieces of information in order to facilitate their selection :
 - time constraints : delays, tolerances, ...
 - supply constraints : at least two suppliers,
 - direct costs : manpower, investments and common expenses,
 - economical environment : origins and financing conditions of the project,
 - other performance constraints : security, reliability, quality ...
- choices and decisions must be based on chosen criteria, useful to manipulate the database previously defined for optimization procedures at the project beginning and whenever necessary [1].

The shared modeling we recommend does not reconsider project management nor product design methods themselves. Our goal is rather to favor their cooperation. It suggests the development of preliminary design step for which we use new tools like Hiles [2] and the optimization , to which are associated both functional and non-functional product requirements. This paper treats the selection and optimization tools that this shared approach suggests.

II. THE SHARED MODELING

At the beginning of the project, one can imagine an approach in which the project manager and the design engineer jointly define a global architecture for the project process as a whole: technological, financial ... choices. This step consists in precisely estimating, scheduling, and anticipating the best general organization for the project. This basic general organization might subsequently become an initial generic project shared model, so far as the major steps usually followed during a project, according to project management or technical design tools, are not very different at a high level. Moreover, this architecture might be obtained in an assisted way , by the expression of precedence relations between steps, the choice of technological criteria

related to the product to design, the financial constraints, the fixed delivery delays, ...

If privileging the technical approach, we are convinced that the shared model could apply the top down refinement system design steps (specifications, preliminary design, virtual prototyping, production...) as useful steps to project management. This first description level can be considered as a generic model expressing technical tasks as well as management steps. It must be detailed and temporally improved regarding the following technical points on the one hand :

- specification : functional and procedural specifications, performance objectives, environment constraints, security directives,
- preliminary design : functional representation, verification, derivation of constraints and directives towards structural or temporal characteristics,
- virtual prototyping : replacement of functional components by real components, introduction of new constraints relative to these components, of production constraints,
- production : introduction of real technological constraints and performances, necessary adjustments on components supplies, validation criteria,

and, on the other hand, regarding non technical objectives, i.e. relative to the project management: costs, market, supply constraints, supplying delays, quality, certification, any types of risks...

The intersection between the tasks' content and the technical and non-technical objectives lead to establish several scenarios which are coherent with general specifications of products. These different scenarios can thus be optimized, hierarchically arranged and selected on the basis of complex compromise criteria associating technical and non-technical considerations.

Of course, with such specifications, no design nor project management operational tool already exists. The schema shown in figure 1 tries to illustrate contents and data structures. At each step, non-functional data modules are associated to functional data ones. For example, we distinguished non-functional technical requirements, costs and execution time requirements, ... The functional approach can lead to several architectures which are compatible with the specifications.

Thus, a shared model can be reached by two different ways :

- from the point of view of traditional project management tools, such as planning and evaluation technico-economics tools,
- from the point of view of technical design tools, such as C.A.D. tools.

In the first case, data characterizing technical constraints will be associated to planning tasks. In the second case, costs and time constraints will be associated to design steps [3].

However, a shared model imposes that planned tasks and design steps should be completely equivalent. Of course, considering the numerous parameters taken into account in this decomposition, several competing scenarios are possible.

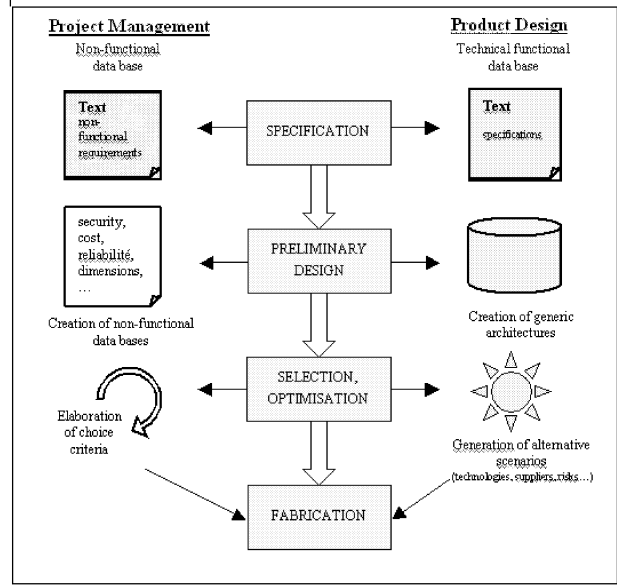


Fig.1 Shared approach : project management and product design

III.GENERATION OF MULTIPLE SCENARIOS

Scenarios are deduced from the options attached to each step- they correspond to global solutions respecting both technical specifications and strategic project requirements.

Of course, the number of scenarios increases as requirements and specifications are relaxed... This relaxing of requirements corresponds to a risk level that we judge acceptable on the basis of both strategic and technical plans.

Scenarios are thus deduced from initial options that only the project management team can determine under the form of a systematic questionnaire. The latter can be a technical one, for instance :

- Are there other product architectures, other types of supply, predictable evolution of technology, ...?
- Which technical risks are associated to each option? Which solutions can be found to reduce these risks?

This questionnaire can also include administrative and financial points:

- Are allocated means sufficient?
- Are deadlines compatible with commercial and financial ambitions?

The criteria used to validate an option are simply the tests of whether this option verifies the technical specifications and the non-functional requirements. Acceptable scenarios will result from a compromise between compatible options, as illustrated on figure 2.

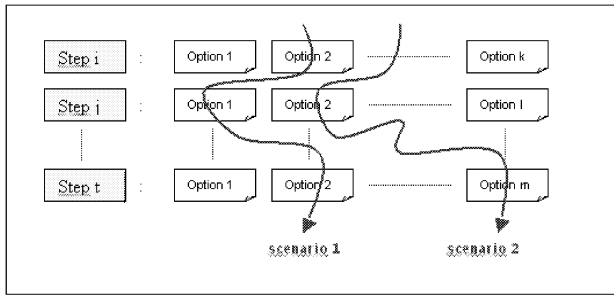


Fig.2 Generation of scenarios from compatible options

This is a first way to generate multiple scenarios; they can be classified according to criteria based on non-functional aspects, using several algorithms that will be exposed in section 4.

Other generation modes can be found into architectural variants that can be imagined by the designer to answer non-functional requirements such as tolerance allocation, reliability, functional risks, and other “project” risks... A special accent can be put on risk analysis during this preliminary phase [4].

Options are stored in a database. The initial database, obtained through the questionnaire and corresponding to the respective experiences of project managers and design engineers, will be progressively improved as projects are conducted.

Then, during the project, regular adjustments will be necessary to take into account events that have occurred which present new risks: longer delays, supplier bankruptcy, new security requirements on the product, insufficient performance... These adjustments will be easier to make if we make good use of the previously mentioned database; however, if it was not detailed enough to obtain one or several satisfactory solutions, new options could be added. This demonstrates the importance of the preliminary risk analysis and of the exhaustiveness of the questionnaire.

The selection of new options, and thus of alternative scenarios, will be done according to a new choice of functional (new technical performances) or non-functional (restricted budget) criteria. The question will thus be how to generate a set of possible solutions from the current state of the project that integrate new constraints; these solutions must also be close enough to the initial one in order to induce a minimum of perturbations into the different aspects of the project (financial, human, technical, ...).

These solutions are acceptable from only one point of view of their conformance to the modified project architecture. They will be submitted to the decision-maker and he will select a set of solutions that best satisfy a multi-criteria compromise (for example, global low costs and delays, but man-power increase).

IV. DIFFERENT APPROACHES FOR THE SELECTION AND OPTIMISATION OF MULTIPLE SCENARIOS

The process previously described offers the designer such elements as:

- A unique description: project tasks and steps, different options by tasks, multiple scenarios that conform to specifications and formulated requirements.
- These scenarios can be classified on the base of :

- technical optimization criteria of potential performances by the examination of precise technological questions,
- more complex optimization criteria dealing with technico-economical compromises (quality and cost for example),
- economical profitability criteria by the anticipation of production and industrial exploitation phases...

Optimization will lead to different hierarchies of scenarios that the decision-maker will arbitrate.

We will only discuss here the selection procedures of scenarios.

4.1. Monte-Carlo methods applied to the allocation of tolerance

This chapter focuses on the decisions made inside a scenario. When market and requirements analyses have defined the scenario and the objectives of a project, the goal of the project management is to strictly answer the deduced requirements for this project. Objectives can be varying: cost objective, time objective, performances objectives... What is considered here is that reaching the global objective results from actions on intermediate influent variables: product global cost depends of each component cost, global performance of each component performance.

Having a global adapted model is essential to appreciate the influence of each parameter on the global objective. This is conducted with a sensibility analysis to parameter variations. The analysis indicates to the project manager which are the sensitive parameters to examine but does not guide him with the decision strategy to adopt. For that, a model describing the consequences of gaps from the objective must be included. It is the proposition of Taguchi [5] to introduce a loss function when the objective is not exactly reached. Our hypothesis is that this idea is interesting whatever the objective is, either a technical performance or a socio-economical question. If you do not reach the goal, the client and the whole society will have to assume the consequences of these gaps. If you surpass the goal, the manufacturer will have to support the consequences.

This synthetic and attractive approach [6,7] motivates us to:

- represent in terms of costs all the consequences of the product requirements,
- introduce a generalized notion of tolerance which will express that each requirement of precision has a cost...

To decide thus consists, for the project manager, in calculating the allocations of tolerances for suppliers and partners. There are calculated with the global predictive model of the product, or the system to design, and a strategic vision which will appear with the loss function associated to gaps relative to nominal values initially fixed by the retained scenario.

The function to minimize corresponds to the sum of costs related to the tolerance requirement on each component parameters and on an estimation of costs related to gaps in the results relative to the fixed objective. The Tabuchi proposition [8] is to estimate this complex dependence for each component. The statistical approach is best appropriated. But this approach is costly in terms of processing time because it requires the exploration of the whole research space with Monte-Carlo draws [9]. Hopefully, some optimization modes can be used in order to reduce this time.

In the context of the optimization of multiple parameters, Monte-Carlo methods are interesting when the number of parameters is reduced to two or four thanks to their simplicity. If the number of parameters is higher, Monte-Carlo methods can be used in a restricted research space to evaluate the sensitivity of a technological device to technological uncertainties. However, other methods, such as neural networks, can be used: the simplified model is established by a direct identification of inputs-outputs. It is then exploited to allow a more rapid minimization, directly processed in terms of standard deviation [10].

4.2. Evolutionary algorithms for the scenarios selection

Evolutionary Algorithms (EA) can be used to select particular scenarios among multiple scenarios. Their principles are inspired from the "Intelligence" of Nature, that can be defined in such a way: "the capability of a system to adapt its behavior to meet its goals in a range of environments" [11]. If no general proof exists of the EA efficiency, it is easy to notice that the selection mechanism is quite efficient a posteriori.

1) General principles of evolutionary algorithms

Three types of EA have been separately developed in the sixties: genetic algorithms, evolution strategies, and evolutionary programming. Initially different, they now constitute convergent techniques [12] and are known under the term of Evolutionary Computation.

Among the EA previously mentioned, genetic algorithms (GA) seem to offer a good compromise between power, generality and ease of programming. They are inspired by the Neo-Darwinism movement- they are based on natural selection mechanisms. Indeed, they use the selection of best adapted individuals and the principles of genetic inheritance propagation. Intuitively, one can associate the problem to a given environment and the solutions to individuals evolving in this environment. At each generation, best adapted individual are selected. After a certain number of generations, the remaining individuals are particularly adapted to the given environment. In this way, one can obtain solutions that are very close to the optimal solution.

The applications of GA are numerous : optimization of difficult numerical functions, image processing, design optimization [13], industrial system control [14], neural network learning [15], etc. GA are used at every step in research, development and production for optimization or selection questions such as the problem that we are concerned with here.

2) Application of genetic algorithms for the selection of scenarios

The generation of scenarios is processed from the different options of figure 2 ; the generated scenarios are already validated and optimized from a functional point of view. Here is how the genetic mechanisms proceeds.

A task, as defined on figure 1, is defined with three main parameters, cost, duration and prerequisites, and some additive informative categories.

A scenario is build as a combination of chosen options (an array) at each step. Options are also stocked into arrays at each step. A scenario contains the following pieces of information: total cost, total duration, fitness and some additive informative elements. The cost and duration parameters are calculated with simple addition operations. An initial population (an array) of scenarios is then randomly or quasi-randomly generated with all their non-functional characteristics. To it are associated three parameters : best individual fitness, best scenario and average fitness. At the beginning of the project, one must fix the objectives in terms of cost and delays, generate the different options.

The genetic engine then makes this population evolve in order to obtain either the best valid and optimal scenario, or a set of optimized scenarios. The evolution of different scenarios is shown on figure 3; one can see that, in order to simultaneously make selection and optimization of scenarios, the classical scheme has been improved with a step of technical validation for candidates: before optimization, they are evaluated according to technical criteria, simulations of performances for instance. Scenarios are then evaluated according to criteria related to the project management domain.

A selection of individuals is then made among the population of candidates in order to favor "good" individuals according to the selected evaluation criteria; however, as a certain diversity has to be respected into the population, a few individuals, less adapted, must survive too. What has been chosen for the moment is to apply the roulette principle.

Selected individuals are then crossed and mutated in different percentages, often empirically determined, in order to constitute the next population. For the moment, only a single point crossover is implemented, the objective being to validate the principle of use of the genetic algorithm.

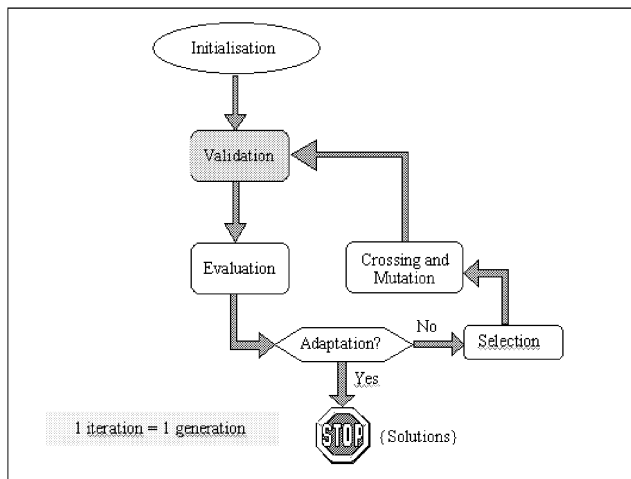


Fig.3 Genetic algorithm functioning

The algorithm proceeds in this way until the solution(s) is obtained, which means that one or more scenario(s) that is functionally satisfactory and corresponds to other technico-economical non-functional constraints have finally been obtained. This stop criterion could be improved further in the future version of the tool... When no optimal solution is reached, the algorithm can help the decision-maker to select the best approximated compromise...

V.CONCLUSION

Nowadays, project management is basically funded onto tasks scheduling and resources (human and financial) management considerations. It supervises product design tasks in the way that the decisions made determine the allocation of resources. This situation is not totally satisfactory because it induces misunderstandings, as the project manager can be very far from technical requirements, and reciprocally the product designer can be unaware of financial constraints.

The main contribution of this paper is to submit a first exploration of an organization more closely associating project management and product design. This proposition consists of three recommendations :

- First: a shared model to describe technical design tasks and project management steps similarly at a high level. Our hypothesis is that this model is founded on top down systems design steps: specification, preliminary design, virtual prototyping, optimization, material prototyping, ...
- Second: proceed to a model exploration by associating options at the task level, and by generating multiple scenarios at the project level. The generation of options and scenarios must be systematically based on possible technological variants, on risks analyses and on financial and administrative variants.
- Third: process optimization treatments in order to select the most effective scenarios. These treatments must be activated at the beginning of the project, then regularly during the project. According to technical or financial criteria, choices could highlight some incompatibilities that decision-makers will have to arbitrate.

In this paper, we focused on the development of the database and on the optimization tools that can be envisaged.

Considering the database development, supposing that the shared model is obtained on the base of the technical tasks decomposition, we suppose that each task be systematically documented with non-functional data related to project management : costs, manpower, deadlines, marketing requirements, strategic constraints, ... Data will consist of, with the product functional model, the shared model project management – product design, on which the optimization methods will rely.

Considering optimization methods, we illustrated their fundamentals on two points: the use of genetic algorithms for the selection of scenarios and the use of Monte-Carlo methods to tackle allocation of tolerance questions, those tolerances being either technical or non-technical. The Monte-Carlo draws are used to explore a system behavior around the nominal values of the model parameters. On the basis of data thus obtained, an optimization of the tolerance allocation is processed on the parameters considered as pertinent that best correspond to the allocation required for the global system. The paper showed that computational techniques can be employed to reduce computation delays which are often long as far as Monte-Carlo statistic evaluation are concerned...

VI.REFERENCES

- [1]Maders H., Gauthier E., "Conduire un projet d'organisation (Guide méthodologique) ", *Les Editions d'Organisation*, 1998.
- [2]Hamon J.C., Esteve D, Pampagnin P., "HiLeS Designer: A tool for systems design". *International Symposium Convergence 03: Aeronautics, Automotive & Space*, Paris, décembre 2003.
- [3]Provost H., "La Conduite de Projet de la conception à l'exploitation des réalisations industrielles", *éditions TECHNIP*, 1994.
- [4]Guiochet J., Baron C., "UML based FMECA in risk analysis", *European Simulation and Modelling Conference (ESMc'2003)*, Napoli, Italy, oct. 2003
- [5]Pillet M., "Les plans d'expériences par la méthode TAGUCHI", *Les Editions d'Organisation*, Paris, 1997.
- [6]Steele G., Byers S., Young D., Moore R., "An Analysis of Injection Molding by Taguchi Methods", *Proceedings of ANTEC '88 Conference*, 1988.
- [7]Warner J.C., O'Connor J., "Molding Process is Improved by Using the Taguchi Method", *Modern Plastics* :65-68, 1989.
- [8]Govaerts B., « La planification expérimentale vue par Taguchi », pp.1-16., <http://www.stat.ucl.ac.be/cours/stat2520/transparent/taguchi.pdf>, Institut de Statistique, Université Catholique de Louvain.
- [9]Al-Mohammed M., "Conception des systèmes électronique : Les étapes d'optimisation et d'allocation des tolérances", thèse, INP, Toulouse, 2003.
- [10]Goldberg D. E, *Genetic Algorithms in Search, Optimisation, and Machine Learning*, Addison-Wesley, Reading (MA), 1989.
- [11]Fogel J. L, *Intelligence Through Simulated Evolution*, New York, N.Y. ISBN-0-471-33250-X, Wiley-Interscience, 1999.
- [12]Bäck T., Hammel U., Schwefel H.P., "Evolutionary Computation: Comments on the History and Current State", *IEEE Transactions on Evolutionary Computation*, vol. 1, n°1, p. 3-17, april 1997.
- [13]Baron C., Geffroy J.-C., Zamilpa C., "Identification of Evolutionary Sequential Systems. Genetic Simulation Experiments", *Communications in Numerical Methods in Engineering*, pp. 631-637, vol. 17, John Wiley Editor, july 2001
- [14]Beasley D., Bull D.R., Martin R.R., "An Overview of Genetic Algorithms : Part 1, Fundamentals", *University Computing*, vol. 15, n°2, p. 58-59, 1993.
- [15]Renders J., "Algorithmes génétiques et réseaux de neurones", *Hermès Science Publications*, 1995.

RISK ANALYSIS MANAGEMENT

INDUSTRIAL RISKS MANAGEMENT USING AGENTS AND DYNAMIC CASE-BASED REASONING

Gaële Simon
Hadhoum Boukachour
Laboratoire d'Informatique du Havre,
25, rue Ph. Lebon
76058 Le Havre Cedex, France
{gsimon,Hadhoum.Boukachour}@iut.univ-lehavre.fr

KEYWORDS

Industrial risks, crisis scenarios, multi-agent systems, Case-Based Reasoning.

ABSTRACT

This article presents the architecture of a multiagent system using dynamic Case-Based Reasoning for the evaluation of potentially risky situations in an industrial park. This architecture allows a real-time comparison of the observed situation with past situations describing crises stored in a case base. It implements a dynamic multiagent Case-Based Reasoning. Having presented the goals of the system, we make a comparison of the proposed approach with CBR techniques for dynamic situations. We present then in detail the behaviour of the agents used in the architecture.

CONTEXT: INDUSTRIAL RISKS MANAGEMENT

Numerous applications of computer systems must allow to represent an evolving situation in order to be able to analyse it. This problem exists in different application domains such as road traffic, epidemics anticipation, meteorology, risks management, etc. We are particularly interested in preventive monitoring systems (Boukachour et al. 2003b). The need of this kind of systems has appeared during multiple interviews we had with industrialists (eg. TotalFinaElf), managers and firemen in the town of Le Havre (Le Havre is a town with a huge industrial park including numerous industries classified as Seveso sites). The goal of such a system is to allow a monitoring of a potentially risky situation in an industrial park. In this context, information about the current situation comes from either managers who can interact with the system or data bases which can be connected to sensors (weather data for instance). More precisely, the system must provide a real-time representation of the situation and help to anticipate potential anomalies or accidents in order to allow managers to take decisions which could avoid them. That's why the system must also provide information allowing to choose the best decisions

and actions. The anticipation consists in evaluating in real-time the risk associated to the current situation.

This anticipation can be performed by comparing the observed situation with crises that already occurred (for example, the catastrophe of AZF firm in Toulouse). Indeed, experts have many descriptions of scenarios of crises resulting from experience feedback. A scenario generally contains a description of the course of a past situation which has yielded to a crisis, that is to say a situation where anomalies have occurred. Such a description is based on values of a set of parameters which are evaluated, by experts, to be the most relevant in that case. As a consequence, different sets of parameters can be used in different scenarios. A scenario also contains a description of decisions and actions which have been used to help to solve the crisis. In real life, during the analysis of a new situation, the decision-making process often relies on a comparison between the current course of the situation and some of these scenarios. The goal is to determine if some of them are similar enough to the situation in order to be able to reuse and adapt for the current case the decisions taken for past situations. Unfortunately, taking into account the quick evolution of the situation, the large number of parameters used to describe it and the significant number of scenarios to reuse, the experts are unable to completely carry out the reasoning. As a consequence, the system we propose is supposed to help experts to achieve this task.

ARCHITECTURE : CONSTRAINTS AND CHOICES

The observed situation generally contains a large number of dynamic parameters, that is to say parameters whose value change over time. Systems allowing the management of such situations must be dynamic in order to be able to handle these evolutions. As a consequence, to design these systems, a flexible and adaptive architecture is needed. This led us to choose a multi-agent architecture (Jennings et al. 1998). We are thus interested in the development of multi-agent systems dedicated to the modelization and the evolution evaluation of dynamic

situations.

Such a system must not only represent the observed situation, but also has to allow its evaluation. As explained in the previous section, this can be achieved using previous situations whose consequences are known. So, a reasoning based on analogy can be used relying on the following hypothesis: if a A situation looks like a B situation, the consequences of the A situation ought to be similar to those of the B situation. To perform such a reasoning, we must elaborate:

- a multi-agent CBR. The representation of the current situation is, in our context, based on a set of agents.
- a dynamic CBR. The target case of the CBR process is an evolving situation, so the CBR has to take this evolution into account incrementally. In other words, when the situation changes, it must not be considered as a new target case.

Using such a dynamic multi-agent CBR, the aim of the system is to select as soon as possible the cases of the base which seem to be the most similar to the current situation in order to be able to anticipate its consequences. Of course, this selection must be adapted to the evolution of the situation over time. Indeed, new information on the situation can eventually modify the set of cases which have been selected during previous steps.

For instance, in the case of industrial risks management, the base might contain the three following cases (arrows represent temporal ordering):

1. gas leak in a tanker → explosion near the tanker → toxic cloud
2. liquid leak → flood → ground water contamination
3. storm → strong rain → flood

At the beginning, without information, the system considers all the cases of the base as potential cases for selection. If the current situation first contains a piece of information dealing with a leak, cases 1 and 2 must be selected by the system. If a new piece of information specifies the kind of leak (gas), then the system must select only the first case which will be used, without new information, to evaluate the consequences of the current situation. On the other hand, if the second piece of information deals with a storm, according to the case base, the system cannot yet produce any result. The three cases of the base must remain selected by the system.

The goal of the following paragraphs is, at first, to highlight the specificities of a CBR adapted to multi-agents systems and to dynamic situations. The second one is to propose an architecture suitable to various domains, especially to industrial risks management, allowing to implement such a reasoning.

CASE-BASED REASONING FOR DYNAMIC SITUATIONS

CBR (Kolodner 1993) is used in many application domains. Among them, we distinguish applications based on CBR dealing with dynamic situations. Their specificity is to handle target cases with temporal characteristics (called situations) and to perform a search of the best case of the base using a similarity between histories (often called *chronics*). The reasoning cycle is the same as in standard CBR. The main difference is the way target and source cases are represented (by chronics) which implies major modifications in the similarity calculus between cases. Indeed, this calculus must take time into account. Many kinds of systems are based on these techniques. A survey of these systems can be found in (Rougegrez 1994). Some examples are REBECAS (Rougegrez 1995), ICONS (Schmidt et al. 1998), SINS (Ram and Santamaria 1997) and (Bull et al. 1997).

Our problem is, *a priori*, close to CBR for dynamic situations. Indeed, we also deal with dynamic situations, represented by target cases with temporal properties. Our goal is also to anticipate the evolution of the situation represented by the target case. However, the CBR which we propose has the following specificities:

- It is a multiagent one.
- It is a dynamic one. This is its major specificity because it modifies the cycle of the reasoning itself. Indeed, the CBR we propose, reuses, at time $t+1$, data and results obtained at time t . In other words, the cycle of CBR in our approach is a continuous cycle. Consequently, the elaboration of the target case and the recall step are made in real time and in an incremental way, that is to say updated according to the evolution of the current situation (the target case). This fundamental aspect for our problem is not present in most systems using CBR for dynamic situations.
- In CBR for dynamic situations, an history is used for the representation of the source and target cases. In our context, the source case does not necessarily contain a real history. For example, in our application in the risks management domain, cases represent past situations associated to different kinds of industrial accidents. They are only described by parameters which are considered to be directly linked to the causes of the accidents and to the context in which these accidents occurred. So, the result may sometimes be a kind of abstraction of the real history. Furthermore, from a case to another, the relevant parameters may not be the same. As a consequence, there is not a fixed distinction between source and target parameters. The problem of cases representation will not be detailed in this article. More information can be found in (Boukachour et al. 2003a) and (Simon et al. 2002).
- As for CBR for dynamic situations, in the source cases, the limit between the past part and the future part of the

history is variable. But it does not vary in the same way. In the CBR for dynamic situations, for a given source case, this limit can change according to the target case. In our approach, this limit can change over time for the same target case, according to its evolution.

In the following paragraph, a multi-agent architecture allowing to implement such a dynamic CBR is presented.

OUR PROPOSAL : PRINCIPLES AND IMPLEMENTATION

Principles: The Different Kinds of Agents Used

First of all, the observed situation is modelled by a set of "informational" agents. They are a generalisation of "aspectual agents" proposed by Durand in (Durand 1999). This set of agents receives information about the situation which are sent to the system by actors or by distributed data bases. Each informational agent is supposed to represent one of these pieces of information which is called "item". For example, if the situation is focused on a gas leak, one of the items can be the kind of gas. In this article, we do not make any assumption about the way these items are modelled inside the agents (for more details, see (Boukachour et al. 2003a)). The main advantage to use agents to represent information about the current situation is that it allows to obtain a flexible representation which can be easily adapted as the situation evolves.

Each informational agent must provide numerical measures of its evolution over time. More precisely, these measures must allow to evaluate the level of reinforcement of the agent inside the organisation it belongs to. Indeed, it is supposed that the more an agent is reinforced, the more its item must be taken into account in the evaluation of the situation. This reinforcement must be based on a similarity measure between items which can use semantic, temporal and spatial aspects (Boukachour et al. 2003a). These mechanisms allow to take into account the fact that, for example, a piece of information introduced very early in the system can turn out to be non relevant later. In this article, we do not make any assumption about the kind of measures which are used to evaluate the reinforcement level of an agent.

Each agent must also provide a temporal validity measure allowing to evaluate the "freshness" of the piece of information associated to its item. For example, a piece of information about the meteorology has not the same lifetime as a piece of information about the traffic state in a town.

Our architecture is based on a multi-agent architecture as proposed by Marcenac in (Marcenac and Giroux 1998). This kind of architecture uses several hierarchical agent

layers, a layer of the level n having a view on the layer of the level $n-1$. Our system uses three different layers :

- the lowest one : it contains agents allowing to model the current state of the situation, that is to say informational agents,
- the intermediate one : it contains synthesis agents used to analyse the previous layer,
- the highest one : it contains prediction agents which must provide information about the potential evolution of the situation using dynamic CBR techniques.

In order to be able to use CBR techniques, the system must contain cases describing past situations. Such cases are called "scenarios". They must allow to characterize, for each past situation, the set of decisive factors which seem to be related to the way the situation went on. For example, in the case of an explosion caused by a gas leak, any fact associated to the leak is a decisive factor. On the contrary, facts associated to the road traffic volume the same day are not decisive factors. As a consequence, each scenario contains a list of items associated to the decisive factors of the past situation. This list can, eventually, be organized temporally. These factors can be determined using experience feedback provided by domain experts.

A prediction agent is associated to each scenario stored in the system. Its goal is to compare the course of the current situation, represented by the informational agents, with the one described in its scenario. This comparison, which must be made in real time, consists in determining if the factors which seem to be important in the current situation are similar to the decisive factors of the situation described in the scenario. In order to do that, a prediction agent must know the factors, that is to say items, which are considered to be the most representative of the current state of the situation. Calculating these factors is the job of the synthesis agents of the system which are described in the next paragraph.

Synthesis agents

The goal of these agents is to provide a synthetic view of the global behaviour of the informational agents layer in order to facilitate the comparison with past situations stored in the scenarios base.

More precisely, the goal is to classify informational agents into groups of agents having similar measures values. These groups can be representative of important aspects of the current situation which will be used by prediction agents to manage the comparison with past situations.

The goal of synthesis agents is to dynamically build these groups called clusters. In (Coma et al. 2003), we propose a dynamic method for agents clustering. Each cluster is modified over time according to informational agents evolution.

Prediction agents

The goal of the prediction agents layer is to provide a continuous recall process of cases of the base, unlike the one used in CBR for dynamic situations described before. Notice that, for the moment, the adaptation step will be done by domain experts which will be in charge to evaluate if the non matched part of the recognised scenario can be used for the current situation. Indeed, the continuous evolution of the analysed situation may decrease the relevance of the adaptation process result. That's why, after having discussed with experts in industrial risks, it has been chosen to provide elements about the main similarities and differences between the scenario and the current situation which can be used by them in order to manage their own adaptation.

A prediction agent has to continuously evaluate the similarity degree between the current situation and the past situation described in the scenario it is associated to. As a consequence, there is as much prediction agents as scenarios in the base of the system. To evaluate this similarity, prediction agents use the result provided by synthesis agents.

The behaviour of a prediction agent is specified by a macro-automaton described in figure 1. This macro-automaton is made of three macro-states: initialisation, supervision and warning. These three states correspond to different behaviours of the agent which are specified by three automata.

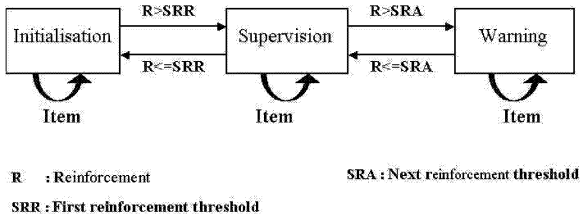


Figure 1: Macro-automaton of a prediction agent

Transitions between macro-states are labelled with reinforcement thresholds. Indeed, prediction agents will be more or less reinforced according to the degree of similarity between their scenario and the current situation. If a new piece of information introduced in the system seems to be similar to an item of a scenario, the prediction agent associated to this scenario will be reinforced. As a consequence, at a given moment, prediction agents with the higher reinforcement values provide the set of scenarios which seem to be the most similar to the current situation.

The ultimate goal of this whole process is to allow to detect similar scenarios as soon as possible while avoiding to provide to users scenarios which would be similar

during a too short time. This process must also update the set of scenarios proposed to the user if the course of the analysed situation changes significantly.

The behaviour of the agent inside each three macro-states is detailed in the following paragraphs.

Initialisation state

A prediction agent, as informational agents, receives all items introduced in the system. Each received item is compared with the items of its scenario. This comparison is based on the similarity measure between items used by informational agents. For example, this measure must allow to match the two following items : "a green vehicle" and "an apple green car". That way, if an item RI, received by a prediction agent PA, is evaluated as being similar to an item of its scenario, and if RI has not been processed before by PA, PA is then reinforced. The reinforcement is proportional to the degree of similarity between the two items. If the scenario contains a temporally sorted list of items, the prediction agent uses an additional criterion to reinforce itself. Indeed, in that case, RI must match with the first item of the list which has not yet been matched. When the prediction agent has been enough reinforced in this state so that its reinforcement becomes equal to the first reinforcement threshold of the automaton, its behaviour is modified according to the supervision state.

Supervision state

In this state, the agent continues to receive and process new items introduced in the system as in the previous state. This can eventually reinforce it again. In parallel to this task, the agent must perform a second one called the supervision task. It consists in evaluating if the items which have yet been used to reinforce the agent are reliable. In other words, it must verify if the informational agents containing these items are important ones in terms of reinforcement.

In order to achieve this task, the prediction agent uses the informational agents groups (clusters) provided by synthesis agents. First of all, for each cluster, it computes the mean of reinforcement levels (MRL) of the agents belonging to the cluster. MRL is used to sort the set of clusters. This task allows the prediction agent to associate to each matched item of its scenario the level of the cluster to which the informational agent containing this item or a similar one belongs. This scheduling task is performed continuously by the prediction agent in order to take into account the clusters evolution.

Each time the clusters levels are updated, the prediction agent compares them to the previous ones. For example, let consider an item I of its scenario. Let suppose that the level of its corresponding cluster C is 1. That means that C contains the most important informational agents. This

implies that I is, at this moment, a relevant item for the description of the current state of the situation. Let now suppose that, after the next evaluation of the levels of the clusters, the level of C becomes 15. This means that the informational agent containing I is less important in this informational layer. This is interpreted by the prediction agent as a loss of relevance of I in the description of the analysed situation. As a consequence, as this item has previously been used in the comparison with its scenario, the prediction agent weakens. On the contrary, if the level of C has increased, the prediction agent is reinforced. This mechanism allows to take into account the fact that some items introduced in the system can finally be non relevant for the situation description. As for the previous state, if, during this process, the reinforcement of the prediction agent reaches the next threshold, its behaviour is modified according to the warning state.

Warning state

In this state, the behaviour of the agent is similar to the previous state's one. The main difference is that, when the agent is in the warning state, its scenario is considered to be very similar to the situation. As a consequence, it can be directly provided to the users. The presentation to the user must be updated in real-time, emphasizing the matched items, the non-matched ones and the opposite ones. This may allow the user to understand quickly what distinguishes the scenario and the situation in spite of the observed similarities.

Notice that, by this mechanism, several scenarios can be provided to the user at the same time. This behaviour is coherent because it's unlikely that the system contains at a moment all the needed information about the analysed situation in order to be able to select only one scenario.

CONCLUSION

In this article, we have presented a multiagent architecture allowing the implementation of a dynamic CBR for the evaluation of the potential evolution of an observed situation. This architecture relies on three agents layers. The lower layer allows to build a representation of the target case, i.e. the current situation. The second layer allows to implement a dynamic elaboration of the target case. Finally, the upper layer implements a dynamic process of source cases recall allowing the search for past situations similar to the current one.

The main prospect for this work is to validate the architecture which is being implemented using MadKit development platform (Gutknecht 2003). This represents a long-term work because it is necessary to test and validate the behaviour of each of the three layers before to be able to evaluate the behaviour of the whole system. These tests will also allow to determine the best measures of the reinforcement of informational agents which is an

essential problem. That will also allow to determine rules for establishing the thresholds of the behaviour automaton of prediction agents. An other prospect is to model real scenarios of past crises as cases in the base of the system. We are currently doing this work on scenarios provided by TotalFinaElf. A longer-term prospect concerns the study of the integration of an adaptation step into the process performed by prediction agents.

REFERENCES

- Boukachour, H. T. Galinho, P. Person and F. Serin. 2003a. "Towards an Architecture for the Representation of Dynamic Situations". ICAI, Las Vegas, USA.
- Boukachour, H. T. Galinho, P. Person, F. Serin, G. Simon, M. Coletta and D.Fournier. 2003b. "Vers une Architecture Multi-agents pour la Représentation et l'Evaluation de Situations Dynamiques ». CCECE 2003 - *Canadian Conference on Electrical and Computer Engineering, Towards a Caring and Humane Technology*, IEEE Canada Montreal.
- Bull, M. G. Kundt and L. Gier. 1997. "Case-based risk detection and forecasting". In a geographic-medical system, German workshop on CBR, page 59-64, in Bergman and M. Wilke editors.
- Coma, R. G. Simon, and M. Coletta. 2003. "A multi-agent architecture for agents clustering". *Agent Based Simulation. ABS'2003*, Montpellier.
- Durand, S. 1999. « Représentation des points de vues multiples dans une situation d'urgence : une modélisation par organisations d'agents ». PhD thesis, Université du Havre.
- Gutknecht, O. 2003. Madkit <http://www.madkit.org>
- Jennings, N.R. M. Wooldridge. and K. Sycara. 1998. " A roadmap of agent research and development". *Autonomous Agent and Multi-Agent Systems*, 1(5):7-38.
- Kolodner, J. 1993. "Case-based reasoning". San Mateo, CA : Morgan Kaufman.
- Marcenac., P. S. Giroux. 1998. "GEAMAS : A Generic Architecture for Agent-Oriented Simulations of Complex Processes". *International Journal of Applied Intelligence, Neural Networks, and Complex Problem-Solving Technologies*, Kluwer Academic Publishers, Vol. 8, N°3, pp. 247-267.
- Ram, A. and J.C Santamaria. 1997. "Continuous case-based reasoning". *Artificial intelligence*, pages 25-27.
- Rougegrez, S. 1994. « Prédiction de processus à partir de comportement observés. le système rebecas ». Thèse de doctorat de l'université Paris 6, rapport technique TH94/05.
- Rougegrez, S. 1995. "Prediction without Modelling : A Demonstration of the Use of Case-based Reasoning for the Prediction of process Behaviour". In *Proceedings of Expert Systems 95*, page 109-120, A. Macintosh AIAI and C. Cooper editors.
- Schmidt, R. B. Pollwein, and L. Gierl. 1998. "Medical multiparametric time course prognoses applied to kidney function assessments". In *International on medical informatics*.
- Simon, G. H. Boukachour, and M. Coletta. 2002. « Vers une architecture multi-agents pour la modélisation et l'évaluation de situations dynamiques ». Technical report, LIH, Université du Havre.

EUROPEAN CALL OPTIONS UNDER UNCERTAINTY

Yuji Yoshida

Faculty of Economics and Business Administration
the University of Kitakyushu

4-2-1 Kitagata, Kokuraminami, Kitakyushu 802-8577, Japan

E-mail: yoshida@kitakyu-u.ac.jp

KEYWORDS

Uncertainty modeling, European call options, valuation of price, fuzzyrelated methodologies.

ABSTRACT

A model of European options with uncertainty of both randomness and fuzziness in output is presented, by introducing fuzzy logic to the stochastic financial model. The randomness and fuzziness in the systems are evaluated by both probabilistic expectation and fuzzy expectation, taking account of seller's/buyer's subjective judgment. Prices of European call/put options with uncertainty are given and their valuation and properties are discussed under a reasonable assumption. This paper demonstrates Black-Scholes formula to give rational expected price of the European options and buyer's/seller's permissible range of expected prices. The meaning and properties of rational expected prices are discussed in a numerical example. The hedging strategies are also considered for marketability of the European options for portfolio selection.

INTRODUCTION

Uncertainty in financial modeling

The theory of option pricing has many applications in finance. European option has been applied and improved by Black-Scholes stochastic model, which is developed on the basis of a log-normal stochastic differential equation in books (Elliott and Kopp 1999; Karatzas and Shreve 1998; Ross 1999) and so on. However, when we sell or buy stocks, there sometimes exists a difference between the actual price and the theoretical value which derived from Black-Scholes method. The losses/errors become bigger between the decision maker's expected price and the actual price if the market are changing rapidly. One of the reason of the losses/errors is from the uncertainty of volatility. Actually, it is difficult to identify the exact value of present volatility.

Mathematical modeling of stochastic systems in decision-making has many applications to engineering, economics, etc.. In general, when we apply stochastic model to actual data, it is important to estimate uncertain factors in dynamical movement. On the option pricing, it is not easy to explain the losses/errors in rapidly changing systems by only probabilistic methods because there exist uncertain factors like volatil-

ity which arise from a difficulty to identify the actual values strictly. Probability theory is constructed on randomness whether something occurs or not in future. Therefore, to estimate uncertain factors like volatility which come from a lack of knowledge regarding the present stock market, we need a method like fuzzy logic to deal with vagueness. Fuzzy methodology plays important role in actual application related human judgment and it is found in many articles, for example a book/journal (Klir and Yuan 1995; Yao and Su 2000).

In this paper, probability is applied as the uncertainty such that something occurs or not with probability, and fuzziness is applied as the uncertainty such that we cannot specify the exact values because of a lack of knowledge regarding the present stock market. This paper discusses a model on European options with uncertainty of both randomness and fuzziness in output, which is a reasonable and natural extension of the original log-normal stochastic processes in Black-Scholes method, by introducing fuzzy logic for uncertainty of identification to the stochastic model. We discuss the stochastic process with uncertainty (randomness and fuzziness) from the viewpoint of fuzzy expectation, taking account of seller's/buyer's subjective judgment in journals (Yoshida 1996; Yoshida 1997; Yoshida 2003). We present confidence intervals of European option prices in the Black-Scholes pricing model with uncertainty (randomness and fuzziness) to secure the expected prices under seller's/buyer's fuzzy goal which is considered as a utility function.

Fuzzy methodology for uncertainty

In order to describe stochastic systems with fuzziness, we need to extend real-valued random variables in the classical probability theory to *fuzzy random variables*, which are random variables with fuzzy number values. In the next section, we introduce a *fuzzy stochastic process* by fuzzy random variables to define prices in European options with uncertainty, and the prices are called *fuzzy prices* in this paper. In the fuzzy stochastic process, the randomness and fuzziness are evaluated by both probabilistic expectation and fuzzy expectation defined by a possibility measure from the viewpoint of journals (Yoshida 1996; Yoshida 1997). Further, this paper gives prices in European call/put options with uncertainty and we discuss their valuation and properties under a reasonable assumption. We give an explicit formula for the fuzzy prices in European options, and we consider rational expected price of

the European options and buyer's/seller's permissible range of expected prices. The meaning and properties of rational expected prices are discussed in a numerical example. In the last section, we consider hedging strategies for marketability of the European options.

Bond price processes and stock price processes

We describe notations regarding bond price processes and stock price processes. We consider European options in a finance model where there is no arbitrage opportunities (refer to books (Elliott and Kopp 1999; Karatzas and Shreve 1998)). Let $(\Omega, \mathcal{M}, P)$ be a probability space, where $\mathcal{M}$ is a σ -field of Ω and P is a non-atomic probability measure. $\mathbb{R}$ denotes the set of all real numbers. Let μ be the *appreciation rate* and let σ be the *volatility* ($\mu \in \mathbb{R}, \sigma > 0$). Let $\{B_t\}_{t \geq 0}$ be a standard Brownian motion on $(\Omega, \mathcal{M}, P)$. $\{\mathcal{M}_t\}_{t \geq 0}$ denotes a family of nondecreasing right-continuous complete sub- σ -fields of $\mathcal{M}$ such that $\mathcal{M}_t$ is generated by B_s ($0 \leq s \leq t$). We consider two assets, a *bond price* $\{R_t\}_{t \geq 0}$ and a *stock price* $\{S_t\}_{t \geq 0}$, where the bond price process $\{R_t\}_{t \geq 0}$ is riskless and the stock price process $\{S_t\}_{t \geq 0}$ is risky. Let r ($r \geq 0$) be the instantaneous interest rate, i.e. interest factor, on a bond. Let a bond price process $\{R_t\}_{t \geq 0}$ is given by $R_t = e^{rt}$ for $t \geq 0$. Let a stock price process $\{S_t\}_{t \geq 0}$ satisfy the log-normal stochastic differential equation: S_0 is a positive constant, and

$$dS_t = \mu S_t dt + \sigma S_t dB_t, \quad t \geq 0. \quad (1)$$

It is known in a book (Elliott and Kopp 1999) that there exists an equivalent probability measure Q such that $\{S_t/R_t\}_{t \geq 0}$ is a martingale under Q , by setting $dQ/dP|_{\mathcal{M}_t} = \exp((r - \mu)/\sigma)B_t - \frac{1}{2}((r - \mu)/\sigma)^2 t$, $t \geq 0$. Under Q , $W_t := B_t - ((r - \mu)/\sigma)t$ is a standard Brownian motion and it holds that $dS_t = rS_t dt + \sigma S_t dW_t$. By Ito's formula, the stock price S_t is represented by

$$S_t = S_0 \exp\left((r - \frac{\sigma^2}{2})t + \sigma W_t\right), \quad t \geq 0. \quad (2)$$

In this paper, we present option models where a stock price process S_t takes fuzzy values using fuzzy random variables, which are introduced in the next section.

FUZZY STOCHASTIC PROCESSES

Fuzzy random variables

Fuzzy random variables, which take values in fuzzy numbers, were first studied by a journal (Kwakernaak 1978), and have been studied by many authors. It is known that the fuzzy random variable is one of the successful hybrid notions of randomness and fuzziness. First we introduce fuzzy numbers. A fuzzy number is denoted by its membership function $\tilde{a} : \mathbb{R} \mapsto [0, 1]$ which is normal, upper-semicontinuous, fuzzy convex and has a compact support. Refer to a journal (Zadeh 1965) regarding fuzzy set theory. In this paper, we identify fuzzy numbers with its corresponding membership functions. $\mathcal{R}$ denotes the set of all fuzzy numbers. The α -cut of a fuzzy

number $\tilde{a} (\in \mathcal{R})$ is given by $\tilde{a}_\alpha := \{x \in \mathbb{R} \mid \tilde{a}(x) \geq \alpha\}$ ($\alpha \in (0, 1]$) and $\tilde{a}_0 := \text{cl}\{x \in \mathbb{R} \mid \tilde{a}(x) > 0\}$, where cl denotes the closure of an interval. We write the closed intervals as $\tilde{a}_\alpha := [\tilde{a}_\alpha^-, \tilde{a}_\alpha^+]$ for $\alpha \in [0, 1]$. A fuzzy-number-valued measurable map $\tilde{X} : \Omega \mapsto \mathcal{R}$ is called a fuzzy random variable. Next we need to introduce expectations of fuzzy random variables in order to describe fuzzy-valued European option models in the next section. Let $\tilde{X}$ be an integrally bounded fuzzy random variable. The expectation $E(\tilde{X})$ of the fuzzy random variable $\tilde{X}$ is defined by a fuzzy number

$$E(\tilde{X})(x) := \sup_{\alpha \in [0, 1]} \min\{\alpha, 1_{E(\tilde{X})_\alpha}(x)\}, \quad x \in \mathbb{R}, \quad (3)$$

where $E(\tilde{X})_\alpha := [\int_\Omega \tilde{X}_\alpha^-(\omega) dP(\omega), \int_\Omega \tilde{X}_\alpha^+(\omega) dP(\omega)]$, $\alpha \in [0, 1]$.

Now, we consider a continuous-time fuzzy stochastic process by fuzzy random variables. Let $\{\tilde{X}_t\}_{t \geq 0}$ be a family of integrally bounded fuzzy random variables. We assume that the map $t \mapsto \tilde{X}_t(\omega) (\in \mathcal{R})$ is continuous on $[0, \infty)$ for almost all $\omega \in \Omega$. $\{\mathcal{M}_t\}_{t \geq 0}$ is a family of nondecreasing sub- σ -fields of $\mathcal{M}$ which is right continuous, and fuzzy random variables $\tilde{X}_t$ are $\mathcal{M}_t$ -adapted. We call $(\tilde{X}_t, \mathcal{M}_t)_{t \geq 0}$ a fuzzy stochastic process.

Fuzzy expectations

We introduce a valuation method of fuzzy prices, taking into account of decision maker's subjective judgment. Give a *fuzzy goal* by a fuzzy set $\varphi : [0, \infty) \mapsto [0, 1]$ which is a continuous and increasing function with $\varphi(0) = 0$ and $\lim_{x \rightarrow \infty} \varphi(x) = 1$. Then we note that the α -cut is $\varphi_\alpha = [\varphi_\alpha^-, \infty)$ for $\alpha \in (0, 1)$. For an exercise time T and call/put options with fuzzy values $\tilde{X}_T = \tilde{C}_T$ or $\tilde{X}_T = \tilde{P}_T$, which will be given in the next section, we define a *fuzzy expectation* of the fuzzy numbers $E(\tilde{X}_T)$ by

$$\begin{aligned} \tilde{E}(E(\tilde{X}_T)) &:= \int_{[0, \infty)} E(\tilde{X}_T)(x) d\tilde{m}(x) \\ &= \sup_{x \in [0, \infty)} \min\{E(\tilde{X}_T)(x), \varphi(x)\}, \end{aligned} \quad (4) \quad (5)$$

where $\tilde{m}$ is the possibility measure generated by the density φ and $\int d\tilde{m}$ denotes Sugeno integral, which is introduced by a doctoral thesis (Sugeno 1974). The fuzzy number $E(\tilde{X}_T)$ means a fuzzy price, and the fuzzy expectation given by Equation (4) implies the degree of buyer's/seller's satisfaction regarding fuzzy prices $E(\tilde{X}_T)$. Then the fuzzy goal $\varphi(x)$ means a kind of utility function for expected prices x in Equation (5), and it represents a buyer's/seller's subjective judgment from the idea of a journal (Bellman and Zadeh 1970). Hence, a real number $x^* (\in [0, \infty))$ is called a *rational expected price* if it attains the supremum of the fuzzy expectation given by Equation (4), i.e.

$$\tilde{E}(\tilde{V}) = \sup_{x \in [0, \infty)} \min\{\tilde{V}(x), \varphi(x)\} = \min\{\tilde{V}(x^*), \varphi(x^*)\}, \quad (6)$$

where $\tilde{V} := E(\tilde{X}_T)$ is a *fuzzy price of European options*.

EUROPEAN OPTIONS UNDER UNCERTAINTY

Fuzzy prices

In this section, we introduce European option with fuzzy prices and we discuss their properties. Let $\{a_t\}_{t \geq 0}$ be an $\mathcal{M}_t$ -adapted stochastic process such that the map $t \mapsto a_t(\omega)$ is continuous on $[0, \infty)$ and $0 < a_t(\omega) \leq S_t(\omega)$ for almost all $\omega \in \Omega$. Let T ($T > 0$) be an *exercise time* and let K ($K > 0$) be a *strike price*. Define a fuzzy stochastic process $\{\tilde{S}_t\}_{t \geq 0}$ by

$$\tilde{S}_t(\omega)(x) := L((x - S_t(\omega))/a_t(\omega)) \quad (7)$$

for $t \geq 0$, $\omega \in \Omega$ and $x \in \mathbb{R}$, where $L(x) := \max\{1 - |x|, 0\}$ ($x \in \mathbb{R}$) is the triangle-type shape function shown in Fig.1 and $\{S_t\}_{t \geq 0}$ is defined by Equation (1). In this paper, we call $\{\tilde{S}_t\}_{t \geq 0}$ a *fuzzy stock price process*. Hence, $a_t(\omega)$ is a spread of triangular fuzzy numbers $\tilde{S}_t(\omega)$ and corresponds to the amount of fuzziness in the process. The fuzziness in the process increases as $a_t(\omega)$ becomes bigger, and $a_t(\omega)$ should be an increasing function of the stock price $S_t(\omega)$ since the fuzziness in the process depends on the volatility σ and stock price $S_t(\omega)$ in Equation (1) (see Assumption S in this section). The α -cuts of Equation (7) are

$$\tilde{S}_{t,\alpha}(\omega) := [\tilde{S}_{t,\alpha}^-(\omega), \tilde{S}_{t,\alpha}^+(\omega)], \quad (8)$$

where $\tilde{S}_{t,\alpha}^\pm(\omega) := S_t(\omega) \pm (1 - \alpha)a_t(\omega)$, $\omega \in \Omega$.

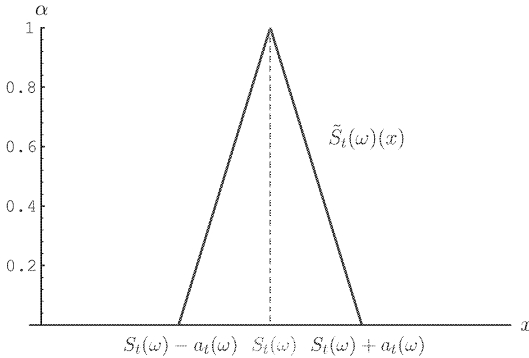


Figure 1: Fuzzy random variable $\tilde{S}_t(\omega)(x)$.

Call option and put option

We define fuzzy stochastic processes of European call/put options by $\{\tilde{C}_t\}_{t \geq 0}$ and $\{\tilde{P}_t\}_{t \geq 0}$:

$$\tilde{C}_t(\omega) := e^{-rt}(\tilde{S}_t(\omega) - 1_{\{K\}}) \vee 1_{\{0\}}, \quad (9)$$

$$\tilde{P}_t(\omega) := e^{-rt}(1_{\{K\}} - \tilde{S}_t(\omega)) \vee 1_{\{0\}} \quad (10)$$

for $t \geq 0$ and $\omega \in \Omega$, where $\vee$ means the maximum in the sense of fuzzy numbers, and $1_{\{K\}}$ and $1_{\{0\}}$ denote the crisp numbers K and zero respectively. Refer to a book (Klir and Yuan 1995) regarding basic calculus of fuzzy numbers. We evaluate these fuzzy stochastic processes by the expectations introduced in the pervious section. Then, the *fuzzy price processes of European call/put options* are given as follows:

$$\tilde{V}^{\tilde{C}}(y, t) := e^{rt} E(\tilde{C}_T | S_t = y) \quad (11)$$

$$= E(e^{-r(T-t)}(\tilde{S}_T - 1_{\{K\}}) \vee 1_{\{0\}} | S_t = y); \quad (12)$$

$$\tilde{V}^{\tilde{P}}(y, t) := e^{rt} E(\tilde{P}_T | S_t = y) \quad (13)$$

$$= E(e^{-r(T-t)}(1_{\{K\}} - \tilde{S}_T) \vee 1_{\{0\}} | S_t = y) \quad (14)$$

for an initial stock price y ($y > 0$) and $t \in [0, T]$, where $E(\cdot)$ denotes expectation with respect to the equivalent martingale measure Q . Their α -cuts are

$$\tilde{V}_\alpha^{\tilde{C}, \pm}(y, t) = E(e^{-r(T-t)} \max\{\tilde{S}_{T,\alpha}^\pm - K, 0\} | S_t = y); \quad (15)$$

$$\tilde{V}_\alpha^{\tilde{P}, \pm}(y, t) = E(e^{-r(T-t)} \max\{K - \tilde{S}_{T,\alpha}^\mp, 0\} | S_t = y). \quad (16)$$

Assumption derived from uncertainty

Now we introduce a reasonable assumption. We can develop the method in this paper without the following Assumption S and triangle-type shape functions for the fuzzy stock price given by Equation (6) (see journals (Yoshida 2002; Yoshida et al. 2000) for the discrete-time case), however this paper adopts them for the numerical computation which is important for its application.

Assumption S. The stochastic process $\{a_t\}_{t \geq 0}$ is represented by $a_t(\omega) := cS_t(\omega)$, $t \geq 0$, $\omega \in \Omega$, where c is a constant satisfying $0 < c < 1$.

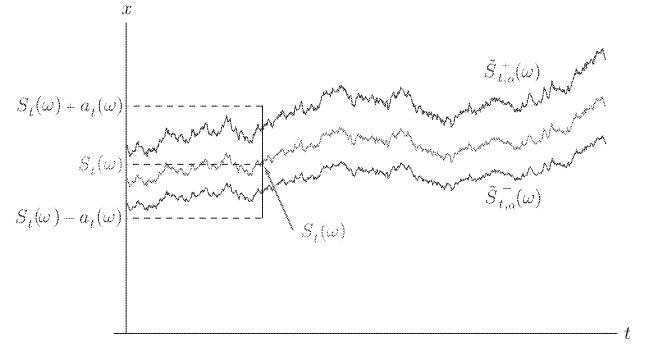


Figure 2: Fuzzy stock price and Assumption S.

Since Equation (1) can be written as

$$d \log S_t = \mu dt + \sigma dB_t, \quad t \geq 0, \quad (17)$$

one of the most difficulties is estimation of the volatility σ in actual cases, which we find in a book (Ross 1999). Therefore, Assumption S is reasonable since $a_t(\omega)$ corresponds to a size of fuzziness which is shown in Figs.1 and 2 and so it is reasonable that $a_t(\omega)$ should depend on the fuzziness of the volatility σ and the stock price $S_t(\omega)$ of the term $\sigma S_t(\omega)$ in Equation (1). In this model, we represent by c the fuzziness of the volatility σ , and we call c a *fuzzy factor* of the process. From now on, we suppose that Assumption S holds.

Black-Scholes formula under uncertainty

Here, we obtain formulae to calculate fuzzy price in European options.

Black-Scholes formula for fuzzy prices. Suppose that Assumption S holds. Let $\alpha \in [0, 1]$. Let an initial stock price $y(= S_0 > 0)$.

- (i) The rational fuzzy price of European call option is given by

$$\tilde{V}_{\alpha}^{\tilde{C},\pm}(y, 0) = b^{\pm}(\alpha)y\Phi(z_1) - Ke^{-rT}\Phi(z_2), \quad (18)$$

where $b^{\pm}(\alpha) := 1 \pm (1 - \alpha)c$ ($\alpha \in [0, 1]$), $\Phi(z) = (1/\sqrt{2\pi}) \int_{-\infty}^z e^{-w^2/2} dw$ ($z \in \mathbb{R}$) is the standard normal distribution function, and z_1 and z_2 are given by

$$z_1 := \frac{\log b^{\pm}(\alpha) + \log(y/K) + T(r + \sigma^2/2)}{\sigma\sqrt{T}}; \quad (19)$$

$$z_2 := \frac{\log b^{\pm}(\alpha) + \log(y/K) + T(r - \sigma^2/2)}{\sigma\sqrt{T}}. \quad (20)$$

- (ii) The rational fuzzy price of European put option is given by the following call-put parity:

$$\tilde{V}_{\alpha}^{\tilde{P},\pm}(y, 0) = \tilde{V}_{\alpha}^{\tilde{C},\mp}(y, 0) - b^{\mp}(\alpha)y + Ke^{-rT}. \quad (21)$$

In the next section, we consider the expected prices of European options based on this formula.

THE EXPECTED PRICES OF EUROPEAN OPTIONS

Expected prices of European options

Fix an initial stock price $y(> 0)$. In this section, we discuss the expected price, which is introduced in the previous section, of European call/put options $\tilde{V} = \tilde{V}^{\tilde{C}}(y, 0)$ or $\tilde{V} = \tilde{V}^{\tilde{P}}(y, 0)$. Define a grade $\alpha^{\tilde{C},+}$ satisfying $\varphi_{\alpha^{\tilde{C},+}}^{-} = \tilde{V}_{\alpha^{\tilde{C},+}}^{\tilde{C},+}(y, 0)$.

Expected prices of European call option. For the fuzzy expectation in Equation (5) generated by possibility measures, the following (i) and (ii) hold.

- (i) The grade of the fuzzy expectation of European call option price $\tilde{V}^{\tilde{C}}$ is given by

$$\alpha^{\tilde{C},+} = \tilde{E}(\tilde{V}^{\tilde{C}}(y, 0)) = \tilde{E}(E(\tilde{C}_T) | S_0 = y). \quad (22)$$

- (ii) Further, the rational expected price of European call option is given by

$$x^{\tilde{C},+} = \varphi_{\alpha^{\tilde{C},+}}^{-}. \quad (23)$$

Since the fuzzy expectation is defined by possibility measures in Equation (5), Equation (23) gives an upper bound on rational expected prices of European call option. Therefore, similarly we can define another grade, which gives a lower bound on rational expected prices of European call option as follows:

$$x^{\tilde{C},-} = \varphi_{\alpha^{\tilde{C},-}}^{-}, \quad (24)$$

where $\alpha^{\tilde{C},-}$ is defined by $\varphi_{\alpha^{\tilde{C},-}}^{-} = \tilde{V}_{\alpha^{\tilde{C},-}}^{\tilde{C},-}(y, 0)$. Hence, from Equations (23) and (24), we can easily check the interval $[x^{\tilde{C},-}, x^{\tilde{C},+}]$ is written as

$$[x^{\tilde{C},-}, x^{\tilde{C},+}] = \{x \in \mathbb{R} \mid \tilde{V}^{\tilde{C}}(y, 0)(x) \geq \varphi(x)\}, \quad (25)$$

which is the range of prices x such that the reliability degree of the optimal expected price, $\tilde{V}^{\tilde{C}}(y, 0)(x)$, is greater

than the degree of buyer's satisfaction, $\varphi(x)$ (refer to Fig.3). Therefore, $[x^{\tilde{C},-}, x^{\tilde{C},+}]$ means *buyer's permissible range of expected prices under his fuzzy goal φ* .

Numerical method of European options under uncertainty

Now we show a numerical procedure for the expected prices of European options under uncertainty through a simple example.

An example. Consider a fuzzy goal

$$\varphi(x) = \begin{cases} 1 - e^{-2x}, & x \geq 0 \\ 0, & x < 0. \end{cases} \quad (26)$$

Then the inverse function is $\varphi_{\alpha}^{-} = -2^{-1} \log(1 - \alpha)$, $\alpha \in (0, 1)$. Put an exercise time $T = 0.5$, a volatility $\sigma = 0.25$, an interest factor $r = 0.05$, a fuzzy factor $c = 0.05$, an initial stock price $y = 25$ and a strike price $K = 30$. From Equations (18), (23) and (24), we can easily calculate that the grades of the fuzzy expectation of the fuzzy price are

$$\alpha^{\tilde{C},-} \approx 0.55909 \quad \text{and} \quad \alpha^{\tilde{C},+} \approx 0.690012. \quad (27)$$

The grades (27) mean the degrees of buyer's satisfaction in pricing. From Equations (23) and (24), the corresponding permissible range of rational expected prices in European call option under his fuzzy goal φ is

$$[x^{\tilde{C},-}, x^{\tilde{C},+}] \approx [0.409458, 0.585611]. \quad (28)$$

We can calculate the crisp case, which is also obtained when we let the fuzzy factor $c \rightarrow 0$. Then the expected prices in European call option is 0.504167, which is surely included in the interval (28). The interval (28) gives the buyer a confidence interval of expected prices in European call option under uncertainty, i.e. randomness and fuzziness as Fig.3 shows.

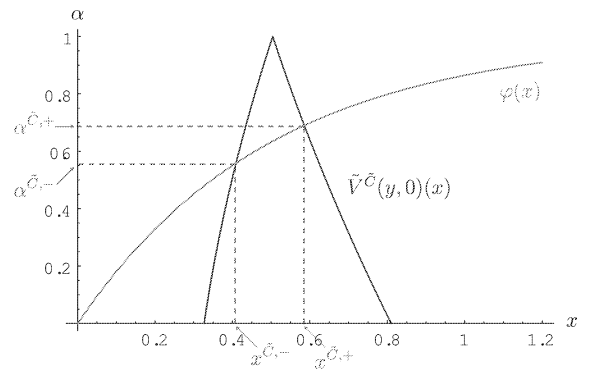


Figure 3: Fuzzy reward $\tilde{V}^{\tilde{C}}(y, 0)(x)$ and fuzzy goal $\varphi(x)$.

Next, we consider another fuzzy goal

$$\varphi(x) = \begin{cases} 1 - e^{-0.5x}, & x \geq 0 \\ 0, & x < 0. \end{cases} \quad (29)$$

Put an exercise time $T = 0.5$, a volatility $\sigma = 0.25$, an interest factor $r = 0.05$, a fuzzy factor $c = 0.05$, an initial stock price $y = 30$ and a strike price $K = 35$. Similarly, in European put option, we can easily calculate the grades, the degree of

seller's satisfaction, and the corresponding permissible range of rational expected prices is as follows:

$$\alpha^{\tilde{P},-} \approx 0.903016 \quad \text{and} \quad \alpha^{\tilde{P},+} \approx 0.911651; \quad (30)$$

$$[x^{\tilde{P},-}, x^{\tilde{P},+}] \approx [4.66641, 4.85291]. \quad (31)$$

We can calculate the crisp case similary to the case of European call option. Then the expected prices in European put option is 4.76346, which is included in the interval (31). The interval (31) gives the seller a confidence interval of expected prices in European put option under uncertainty as Fig.4 shows.

Buyer/seller should take into account of the permissible range of rational expected prices under their fuzzy goal φ .

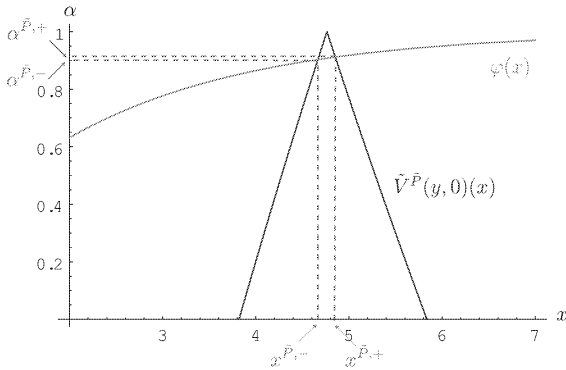


Figure 4: Fuzzy reward $\tilde{V}^{\tilde{P}}(y, 0)(x)$ and fuzzy goal $\varphi(x)$.

HEDGING STRATEGIES

Hedging strategies for European call option

Finally, we deal with hedging strategies in European call option. Fix any $\alpha \in [0, 1]$. A hedging strategy is an $\mathcal{M}_t$ -predictable process $\{(\pi_t^{0,\pm}, \pi_t^{1,\pm})\}_{t \geq 0}$ with values in $\mathbb{R} \times \mathbb{R}$, where $\pi_t^{0,\pm}$ means the amount of the bond and $\pi_t^{1,\pm}$ means the amount of the stock at time t , and it satisfies

$$V_{t,\alpha}^{\pm} = \pi_t^{0,\pm} R_t + \pi_t^{1,\pm} \tilde{S}_{t,\alpha}^{\pm}, \quad t \geq 0, \quad (32)$$

where $V_{t,\alpha}^{\pm} := e^{rt} E(\tilde{C}_{T,\alpha}^{\pm} | \mathcal{M}_t)$ is called a *wealth process*. A hedging strategy $\{(\pi_t^{0,\pm}, \pi_t^{1,\pm})\}_{t \geq 0}$ is called *self-financing* if $d\pi_t^{0,\pm} R_t + d\pi_t^{1,\pm} \tilde{S}_{t,\alpha}^{\pm} = 0$, $t \geq 0$. Then, we obtain the following results.

Hedging strategies. The minimal hedging strategy $\{(\pi_t^{0,\pm}, \pi_t^{1,\pm})\}_{t \in [0, T]}$ for the fuzzy price of European call option is given by $\pi_t^{0,\pm} = \Phi(\log b^{\pm}(\alpha) + Z_t^{0,\pm})$ and $\pi_t^{1,\pm} = -e^{-rT} K \Phi(\log b^{\pm}(\alpha) + Z_t^{1,\pm})$ for $t < T$, where

$$Z_t^{0,\pm} := \frac{\log(S_t/K) + (T-t)(r + \sigma^2/2)}{\sigma\sqrt{T-t}}, \quad (33)$$

$$Z_t^{1,\pm} := \frac{\log(S_t/K) + (T-t)(r - \sigma^2/2)}{\sigma\sqrt{T-t}}. \quad (34)$$

The corresponding wealth process is

$$V_{t,\alpha}^{\pm} = \pi_t^{0,\pm} R_t + \pi_t^{1,\pm} \tilde{S}_{t,\alpha}^{\pm}. \quad (35)$$

REFERENCES

- R.E.Bellman and L.A.Zadeh. 1970. "Decision-making in a fuzzy environment," *Management Sci. Ser. B*, 17, 141-164.
- R.J.Elliott and P.E.Kopp. 1999. *Mathematics of Financial Markets* Springer, New York.
- I.Karatzas and S.E.Shreve. 1998. *Methods of Mathematical Finance* Springer, New York.
- G.J.Klir and B.Yuan. 1995. *Fuzzy Sets and Fuzzy Logic: Theory and Applications* Prentice-Hall, London.
- H.Kwakernaak. 1978. "Fuzzy random variables. Part I : Definitions and theorem," *Information Sciences* 15, 1-29.
- S.M.Ross. 1999. *An Introduction to Mathematical Finance* Cambridge Univ. Press, Cambridge.
- M.Sugeno. 1974. "Theory of fuzzy integrals and its applications," *Doctoral Thesis, Tokyo Institute of Technology*.
- Jing-Shing Yao and Jin-Shieh Su. 2000. "Fuzzy inventory with backorder for fuzzy total demand based on interval-valued fuzzy set," *Europ. J. Oper. Res.* 124, 390-408.
- Y.Yoshida. 1996. "An optimal stopping problem in dynamic fuzzy systems with fuzzy rewards," *Computers Math. Appl.* 32, 17-28.
- Y.Yoshida. 1997. "The optimal stopped fuzzy rewards in some continuous-time dynamic fuzzy systems," *Math. and Comp. Modelling* 26, 53-66.
- Y.Yoshida. 2002. "Optimal stopping models in a stochastic and fuzzy environment," *Information Sciences* 145, 221-229.
- Y.Yoshida. 2003. "A Discrete-Time European Options Model under Uncertainty in Financial Engineering". In *Multi-Objective Programming and Goal-Programming*, T.Tanino, T.Tanaka and M.Inuiguchi (eds.) Springer, Heidelberg, Chapter III-16, 415-420.
- Y.Yoshida, M.Yasuda, J.Nakagami and M.Kurano. 2000. "Optimal stopping problems in a stochastic and fuzzy system," *J. Math. Anal. and Appl.* 246, 135-149.
- L.A.Zadeh. 1965. "Fuzzy sets," *Inform. and Control* 8, 338-353.

AUTHOR BIOGRAPHY

YUJI YOSHIDA obtained the degree of his Dr. of Science in mathematics from Kyushu University, and he was an Associate Professor of Chiba University (1989 - 1993) and the University of Kitakyushu (1993- 1996). Now he is a Professor of the University of Kitakyushu. He is an editor of a book, *Dynamical Aspects in Fuzzy Decision Making* (Physica-Verlag, Springer, Heidelberg, 2001). He is an author or a coauthor of chapters in 5 books and many research papers including 47 papers are published in international academic journals: Fuzzy Sets and Systems, J of Math. Anal. and Appl., Information Sciences, Europ. J. of Oper. Res., etc.. He is also an editor of International Journal of Hybrid Intelligent Systems.

His current research interests are mainly decision making in fuzzy/stochastic systems and mathematical modeling of economics/management sciences.

SIMULATION AND DATAMINING

Episode Detection with Vector Space Model in Agent Behavior Sequences of MMOGs

Ji-Young Ho and Ruck Thawonmas

Intelligent Computer Entertainment Laboratory

Department of Computer Science, Ritsumeikan University

Kusatsu, Shiga 525-8577, Japan

E-mail: ruck@cs.ritsumei.ac.jp

KEYWORDS

Episode detection, Behavior sequences, MMOG

ABSTRACT

Generally action sequences of players in MMOGs (Massively Multiplayer Online Games) are lengthy and often contain consecutively meaningless actions. In our previous work, we have suggested a preprocessing algorithm using items instead of actions to improve classification ability, and validated the results. In fact, the algorithm has limited applicability, because it is derived by a priori based on our observation of the simulated MMOG data. In this paper, we propose a more universal and effective algorithm to transform action sequences to episode sequences. We show that episode sequences are more informative and improve the classification ability over both action sequences and item sequences.

INTRODUCTION

Recently the market size of MMOGs (Massively Multiplayer Online Games) has become significantly large with the growth of entertainment industry. Thus it is very important to grasp the players' needs and to satisfy them through furnishing appropriate contents for each player or each specific group of players. This leads to a new research field called Game Mining, proposed originally by Tveit et al. at NTNU (<http://abiody.com/gamemining/>) and later but independently by the authors' group. In Game Mining, data mining techniques are exploited to improve quality of MMOGs in various aspects such as contents, designs, stories, costs and so forth.

In MMOGs, four player types are typically identified by their characteristics, namely, achievers, explorers, socialisers, and killers in [1]. Achievers set their main goal to gather points or to raise levels while explorers want to find out interesting things about the virtual game world and then to expose them. Socialisers are interested in

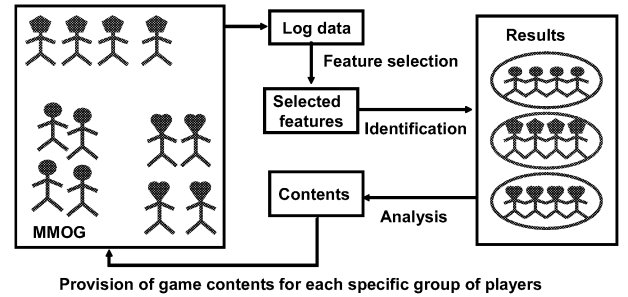


Figure 1: Typical implementation of Game Mining

relationships among players, and killers just want to kill players and monsters with the tools provided by games.

Following the above categorization, a typical implementation of Game Mining can be depicted in Figure 1. Namely, players are identified the types with appropriate selected features from the log data and are provided contents according to their favorites. Through this, the players should enjoy the games with more amusement and the contents providers should make more profits.

It has already been reported in [2] that, for classification of player types based on their specific characteristics, item sequences perform better than action sequences of player agents in a simulated MMOG. Item sequences discard meaningless portions in action sequences thus the resulting data are more compact and informative. For example, consider two different sequences, "Walk, Walk, Attack, Attack, Attack, Walk, PickKey" of an agent of type 1 and "Walk, Attack, Walk, Attack, Walk, Attack, Walk, PickKey" of an agent of type 2. Although the agent of type 1 fought with one monster and the agent of type 2 fought with three monsters, the input features for classification are same in terms of the frequency of performing each action. The proposed preprocessing algorithm in [2] makes those action sequences to so-called item sequences as follows: "Monster, Key" for the agent of type 1 and "Monster, Monster, Monster, Key" for the agent of type 2.

Namely, it generates new sequences based on our observation that information on items acquired might be better than that on actions performed for the MMOGs of interest. The applicability of the algorithm in [2] is, however, limited. If there are a large amount of item types and action types, it will be difficult to generate a suitable algorithm that incorporates human knowledge based on observation of game data.

In this paper, we propose a more universal and effective algorithm for detection of episodes from sequences of players' actions in MMOGs. We show that episode sequences are more informative and improve the classification ability over both action sequences and item sequences.

EPISODE DETECTION WITH VECTOR SPACE MODEL

In general, action sequences are lengthy and have many consecutively uninformative actions such as "Walk". We propose below an algorithm to detect episodes and to transform a long action sequence to a shorter meaningful episode sequence. The algorithm is divided into 2 phases, generation of action vectors with training data and transformation of sequences in which a partial action sequence is transformed to the most representative episode. The algorithm for generation of action vectors is a modified version of the algorithm for generation of term vectors in Information Retrieval [3]. The major modification is that of controlling the importance of each action with its number of occurrence. In this study, an action performed many times compared with other actions is considered less important. Thus we decrease the importance of such an action by dividing each component of its action vector by the rate of occurrence.

Preliminaries

In this paper, an action is a type of behaviors that an agent performs in the game and an episode is a partial collection of different actions that occur together. In the research field of sequence analysis, episodes have been used to represent sequence data and find general rules with consensus patterns [4] and frequent episodes [5]. Both of them are defined by partial ordered sequences. However, they neglect partial sequences that rarely occur though such sequences might contain useful information.

We represent episodes by combinations of actions and thus given n action types, $2^n - 1$ episode candidates are generated. However, we consider only existing episodes in the training data, which greatly reduces the number of episodes to be considered. The number of existing

episodes is notated by m . More definitions of terms used in the algorithm are given as follows.

An episode j , notated by e_j , is associated with an $m \times 1$ unit episode vector $\mathbf{b}_j$ that is pair-wise orthogonal with other episode vectors. A divided sequence k , notated by d_k , is a partial subsequence of an action sequence divided by the window of width w , the only one parameter the user has to decide. A threshold ρ is defined as the window width divided by the number of existing episodes, i.e., $\rho = \frac{w}{m}$. Finally, o_i is a parameter to control the weight of action i and is defined by the product of the window width and the rate of occurrence. Its formula is given in (3).

Episode Detection Algorithm

- **Step 1 (Begin of Action Vector Generation)**
- Divide all training action sequences into l divided sequences with the window width w .
- **Step 2** - Find m episodes from d_k , for $k = 1, 2, \dots, l$, and generate unit episode vectors $\mathbf{b}_j$ for $j = 1, 2, \dots, m$.
- **Step 3** - Calculate the $l \times (n+1)$ frequency matrix $\mathbf{F}$ that represents the occurrence number of each action in each divided sequence, where the last element of each row is the corresponding episode type.
- **Step 4 (End of Action Vector Generation)**
- Generate $m \times 1$ action vectors $\mathbf{a}_i$ using episode vectors $\mathbf{b}_j$ and frequency matrix $\mathbf{F}$. Equation (1) shows the formula to calculate action vectors.

$$\mathbf{a}_i = \frac{\sum_{j=1}^m c_{ij} \cdot \mathbf{b}_j}{o_i \sum_{k=1}^l f_{ki}} \quad (1)$$

$$c_{ij} = \sum_{k=1, \text{ when } f_{k(n+1)} = e_j}^l f_{ki} \quad (2)$$

$$o_i = \begin{cases} w \frac{\sum_{k=1}^l f_{ki}}{\sum_{k=1}^l \sum_{i=1}^n f_{ki}} & \text{if } w \frac{\sum_{k=1}^l f_{ki}}{\sum_{k=1}^l \sum_{i=1}^n f_{ki}} > 1 \\ 1 & \text{otherwise} \end{cases} \quad (3)$$

- **Step 5 (Begin of Sequence Transformation)**
- Sum up the action vector for each action in the action sequence of each agent until the largest component of the resulting vector is larger than the threshold ρ .
- **Step 6** - Represent the current partial action sequence with the episode corresponding to the largest component.

Table 1: Example action sequences

Id	Sequence	Type
agent 1	"ccccaaccccaacbbcbbc"	Type 1
agent 2	"cbbcbccccaaccccaac"	Type 1
agent 3	"caacaacccbbcccbbc"	Type 2
agent 4	"cccbbcccbcaacaac"	Type 2

- **Step 7 (End of Sequence Transformation)** - Repeat from Step 5 for the remaining actions in the action sequence of interest.

The above algorithm is illustrated by the following example. Table 1 shows four simple action sequences of two different types of agents composed of three action types, i.e., a , b , and c . In this example, a , b , and c stand for "attack A-type monster", "attack B-type monster", and "walk", respectively. Agents of type 1 prefer to attack B-type monsters and agents of type 2 prefer A-type monsters. Thus if agents of type 1 find a B-type monster, they attack the monster with no hesitation. However, they only attack an A-type monster when there is no other choice. As a result, before attacking an A-type monster, agents of type 1 usually roam around. This trend applies to agents of type 2, but vice versa. The window width w is set to 4 in this example.

At Step 1, each action sequence is scanned from left to right and action by action with the window width of 4. For agent 1, 16 divided sequences are generated as follows: $d_1="ccc"$, $d_2="ccca"$, $d_3="ccaa"$, $d_4="caac"$, $d_5="aacc"$, $d_6="accc"$, $d_7="cccc"$, $d_8="ccca"$, ..., $d_{15}="bcb"$, $d_{16}="bbc"$.

As mentioned earlier, $2^3 - 1$ episode candidates exist with 3 type of actions, i.e., $\{abc, \bar{a}bc, a\bar{b}c, ab\bar{c}, \bar{a}\bar{b}c, a\bar{b}\bar{c}, \bar{a}b\bar{c}\}$. The algorithm, however, considers only 4 episodes found at Step 2 in the scanned 64 divided sequences. The found episodes are $\{e_1 = \bar{a}\bar{b}c, e_2 = a\bar{b}c, e_3 = abc, \text{ and } e_4 = \bar{a}bc\}$. The algorithm thus generate 4 episode vectors represented by the following set of orthogonal basis vectors:

$$\mathbf{b}_1 = (1, 0, 0, 0)^T, \mathbf{b}_2 = (0, 1, 0, 0)^T,$$

$$\mathbf{b}_3 = (0, 0, 1, 0)^T, \mathbf{b}_4 = (0, 0, 0, 1)^T$$

, where T represents the transpose operation.

At Step 3, the frequency matrix $\mathbf{F}$ is Calculated and the result is shown in Table 2. Table 3 shows the results of c_{ij} calculated at Step 4.

Action vectors $\mathbf{a}_1$, $\mathbf{a}_2$ and $\mathbf{a}_3$, corresponding to action types a , b and c , respectively, are generated based on elements in Tables 2 and 3 as follows:

Table 2: Frequency matrix $\mathbf{F}$

$k \setminus i$	1(a)	2(b)	3(c)	4(episode type)
1(d_1)	0	0	4	e_1
2(d_2)	1	0	3	e_2
3(d_3)	2	0	2	e_2
$\vdots$	$\vdots$	$\vdots$	$\vdots$	$\vdots$
15(d_{15})	0	3	1	e_4
16(d_{16})	0	2	2	e_4
17(d_{17})	0	2	2	e_4
18(d_{18})	0	3	1	e_4
$\vdots$	$\vdots$	$\vdots$	$\vdots$	$\vdots$
63(d_{63})	3	0	1	e_2
64(d_{64})	2	0	2	e_2

Table 3: Results for each c_{ij}

$i \setminus j$	1(e_1)	2(e_2)	3(e_3)	4(e_4)
1(a)	0	49	6	0
2(b)	0	0	6	49
3(c)	32	55	4	55

$$\mathbf{a}_1 = \frac{0 \cdot \mathbf{b}_1 + 49 \cdot \mathbf{b}_2 + 6 \cdot \mathbf{b}_3 + 0 \cdot \mathbf{b}_4}{1 \times 55}$$

$$= 0.8909\mathbf{b}_2 + 0.1091\mathbf{b}_3$$

$$\mathbf{a}_2 = \frac{0 \cdot \mathbf{b}_1 + 0 \cdot \mathbf{b}_2 + 6 \cdot \mathbf{b}_3 + 49 \cdot \mathbf{b}_4}{1 \times 55}$$

$$= 0.1091\mathbf{b}_3 + 0.8909\mathbf{b}_4$$

$$\mathbf{a}_3 = \frac{32 \cdot \mathbf{b}_1 + 55 \cdot \mathbf{b}_2 + 4 \cdot \mathbf{b}_3 + 55 \cdot \mathbf{b}_4}{4 \frac{146}{256} \times 146}$$

$$= 0.0961\mathbf{b}_1 + 0.1651\mathbf{b}_2 + 0.0120\mathbf{b}_3 + 0.1651\mathbf{b}_4$$

At Step 5, transformation of a given action sequence to an episode sequence is conducted. Table 4 shows the transformation process for the action sequence of agent 1. The process proceeds from top to bottom accumulating each action vector and at the same time checking the threshold, $\rho = \frac{4}{4} = 1$ in this example. Underlined elements represent the element vectors whose value exceeds the threshold. Eventually, the sequence "ccccaaccccaacbbcbbc" is transformed to " $e_2e_2e_2e_4e_4e_4$ ". More precisely, "cccca", "ac", "ccca", "ac" are transformed to e_2 and "bb", "cb", "bc" to e_4 .

Table 4: Transformation process for the action sequence of agent 1

action	Accumulated vector
c	$0.0961b_1 + 0.1651b_2 + 0.0120b_3 + 0.1651b_4$
c	$0.1922b_1 + 0.3302b_2 + 0.0240b_3 + 0.3302b_4$
c	$0.2883b_1 + 0.4953b_2 + 0.0360b_3 + 0.4953b_4$
c	$0.3844b_1 + 0.6604b_2 + 0.0480b_3 + 0.6604b_4$
a	$0.3844b_1 + 1.5513b_2 + 0.1571b_3 + 0.6604b_4$
a	$0b_1 + 0.8909b_2 + 0.1091b_3 + 0b_4$
c	$0.0961b_1 + 1.0560b_2 + 0.1211b_3 + 0.1651b_4$
c	$0.0961b_1 + 0.1651b_2 + 0.0120b_3 + 0.1651b_4$
c	$0.1922b_1 + 0.3302b_2 + 0.0240b_3 + 0.3302b_4$
c	$0.2883b_1 + 0.4953b_2 + 0.0360b_3 + 0.4953b_4$
a	$0.2883b_1 + 1.3862b_2 + 0.1451b_3 + 0.4953b_4$
a	$0b_1 + 0.8909b_2 + 0.1091b_3 + 0b_4$
c	$0.0961b_1 + 1.0560b_2 + 0.1211b_3 + 0.1651b_4$
b	$0b_1 + 0_2 + 0.1091b_3 + 0.8909b_4$
b	$0b_1 + 0b_2 + 0.2182b_3 + 1.7818b_4$
c	$0.0961b_1 + 0.1651b_2 + 0.0120b_3 + 0.1651b_4$
b	$0.0961b_1 + 0.1651b_2 + 0.1211b_3 + 1.0560b_4$
b	$0b_1 + 0_2 + 0.1091b_3 + 0.8909b_4$
c	$0.0961b_1 + 0.1651b_2 + 0.1211b_3 + 1.0560b_4$

Comparisons of action, item, and episode sequences in Table 5 show the predominance of episode sequences in terms of the discrepancy of the input features, derived by the procedure in the next section, for the two types of agents.

PLAYER DATA ACQUISITION AND PRE-PROCESSING

MMOG Simulator and Player Modeling

To acquire player data in MMOGs, we use Zereal [6] that is a Python-based multiple agent simulation system running on a PC cluster system. Zereal can simulate multiple game worlds simultaneously, each world with multiple numbers of agents modeling player characters and monsters, and with various kinds of items. A game world is executed on a node of the PC cluster system. The special node called master node sends log data to ZerealViewer that is a visualization client originally developed by our group to manage data and observe the game situation. Figure 2 shows the architecture of the system used in this study.

In the version of Zereal we used, three type of player agents, i.e., Killer, MarkovKiller, and PlanAgent, are provided with seven common actions, Walk, Attack,

Table 5: Comparisons of action, item, and episode sequences

Id	Sequence	Input Features
agent 1	"ccccaaccccaacbbcbbc"	(1.00, 1.00, 1.00)
agent 2	"cbbcbccccaaccccaac"	(1.00, 1.00, 1.00)
agent 3	"caacaaccccbcccbbc"	(1.00, 1.00, 1.00)
agent 4	"ccccbccccbbcaacaac"	(1.00, 1.00, 1.00)

(a) Resulting action sequences

Id	Sequence	Input Features
agent 1	"AABB"	(1.00, 1.00)
agent 2	"BBAA"	(1.00, 1.00)
agent 3	"AABB"	(1.00, 1.00)
agent 4	"BBAA"	(1.00, 1.00)

(b) Resulting item sequences

Id	Sequence	Input Features
agent 1	"e ₂ e ₂ e ₂ e ₂ e ₄ e ₄ e ₄ "	(1.00, 0.75)
agent 2	"e ₄ e ₄ e ₄ e ₂ e ₂ e ₂ e ₂ "	(1.00, 0.75)
agent 3	"e ₂ e ₂ e ₂ e ₄ e ₄ e ₄ e ₄ "	(0.75, 1.00)
agent 4	"e ₄ e ₄ e ₄ e ₄ e ₂ e ₂ e ₂ "	(0.75, 1.00)

(c) Resulting episode sequences

PickFood, PickPotion, PickKey, LeaveWorld, and EnterWorld, as well as five common items, Monster, Food, Potion, Key, and Door. Figure 3 shows a screenshot of ZerealViewer when one game world is being simulated.

Each type of player agents has different main characteristics described as follows:

- **Killer** puts the highest priority on killing monsters.
- **Markov Killer** gets as many items as possible to be stronger. Player agents of this type also kill monsters, but attack monsters according to the corresponding state-transitional probability.
- **Plan Agent** finds a key and leaves the current game world.

Killer, **Markov Killer** and **Plan Agent** correspond to, to some extent, "killer", "achiever", and "explorer" in [1] respectively.

Input Feature Generation

For performance comparisons, three types of sequences are generated by different algorithms. Action sequences are generated from log data by extraction only action information. Items sequences are generated by the following algorithm proposed in [2]:

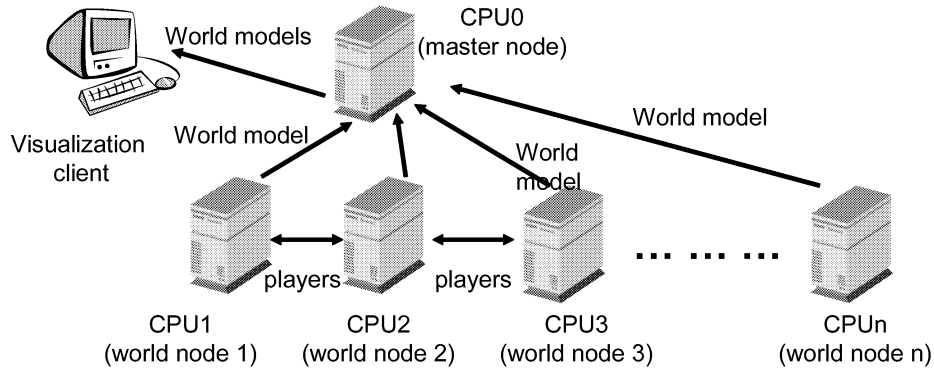


Figure 2: Architecture of the MMOG simulation system

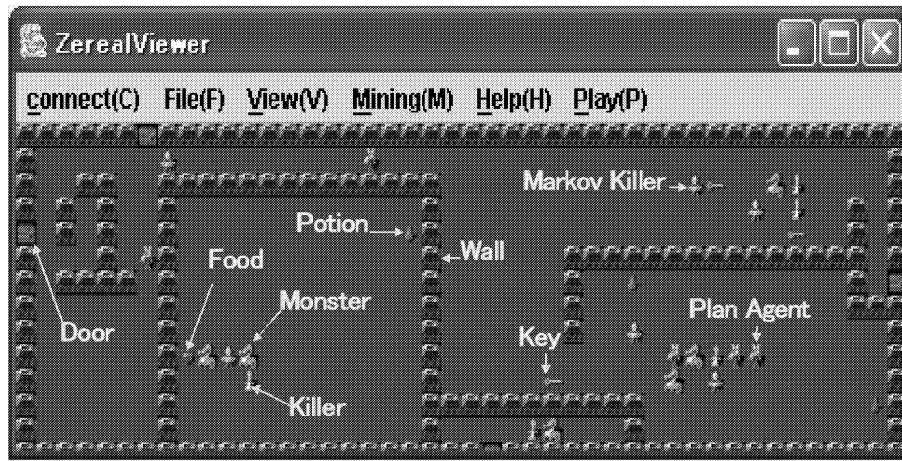


Figure 3: Screenshot of ZerealViewer with one game world

- **For Monster items**, if a player agent attacks a particular monster, add one Monster item to the item sequence of that player agent. If the player agent attacks the same monster many times, only one Monster item is added.
- **For Food, Potion, and Key items**, if a player agent picks Food, Potion, or Key, add one Food, Potion, or Key item to the item sequence of that player agent, respectively.
- **For Door items**, if a player agent leaves the world through a door, add one Door item to the item sequence of that player agent.

Episode sequences are generated by the algorithm explained in the previous section.

Figure 4 shows typical acquired log data. Figures 5, 6, and 7 show, respectively, typical action, item, and episode sequences. We apply the following algorithm

proposed in [7] to these different types of sequences to generate input features for a classifier discussed in the next section.

- **Step I** - For each player agent, sum up the total number of each action that the player agent performed.
- **Step II** - For each player agent, divide the result of each action in Step I by the total number of actions that the player agent performed.
- **Step III** - For each player agent, divide the result of each action in Step II by that of the agent who most frequently performed the action.

Feature-selection algorithms for item sequences and episode sequences are the same as the one above, except that action and performed are replaced by item and acquired respectively for the former, and action by episode for the latter.

```

10||1||2003-12-8: 19:5:50||1000450||walk||{(11,79)}||(10,80)||Killer1
10||1||2003-12-8: 19:5:50||1000451||attack||{(99,30)}||(98,29)||Killer1
10||1||2003-12-8: 19:5:50||1000452||attack||{(60,89)}||(60,90)||Killer1
10||1||2003-12-8: 19:5:50||1000453||walk||{(100,60)}||(99,61)||Killer1
11||1||2003-12-8: 19:5:53||1000055||walk||{(137,32)}||(136,33)||MarkovKiller1
11||1||2003-12-8: 19:5:53||1000056||walk||{(105,127)}||(106,128)||MarkovKiller1
11||1||2003-12-8: 19:5:53||1000057||walk||{(124,67)}||(125,68)||MarkovKiller1
11||1||2003-12-8: 19:5:53||1000058||pickup||{(75,46)}||(74,46)||MarkovKiller1
:
82||1||2003-12-8: 19:8:45||1000239||walk||{(93,133)}||(93,132)||Monster
82||1||2003-12-8: 19:8:45||1000242||attack||{(74,82)}||(75,82)||Monster
82||1||2003-12-8: 19:8:45||1000244||walk||{(111,93)}||(110,92)||Monster
82||1||2003-12-8: 19:8:45||1000251||attack||{(4,115)}||(3,115)||Monster
82||1||2003-12-8: 19:8:45||1000259||walk||{(113,21)}||(114,20)||PlanAgent1
82||1||2003-12-8: 19:8:45||1000260||walk||{(54,6)}||(54,7)||PlanAgent1
82||1||2003-12-8: 19:8:45||1000271||walk||{(57,4)}||(57,3)||PlanAgent1
82||1||2003-12-8: 19:8:45||1000285||walk||{(99,33)}||(100,32)||PlanAgent1

```

Figure 4: Typical log data

```

:
:
:
Killer || 1000459 || waaaaaaaaaaaaawwwaaaaaaaaaaaaaaaaaaaaawwwwwwwwwf
wwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwww
Killer || 1000460 || wwwwwawwwwwwwwwwwwwwwwwwwwwwwwwawwwwwwwwwww
wwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwww
wwwwwaaaaaaaaawwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwww
MarkovKiller || 1000059 || wwwwwwwwwwwwwfwwwwwwwwwwwwwwwwwwwwwwfwwww
wwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwww
wwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwfww
wwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwww
MarkovKiller || 1000060 || wwwwwwwwwfwwwwwwwwwwwwwwwwwwwwwwwwwwwwww
wwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwww
wwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwww
wwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwww
wwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwww
PlanAgent || 1000259 || wwwkwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwww
wwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwww
wwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwww
wwwpwwwwr
PlanAgent || 1000260 || wwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwwww
wwwwwwwwwwww
:
:
:

```

Figure 5: Typical action sequence data

EXPERIMENTS

In the experiments, we use as a classifier our originally developed adaptive memory-based reasoning (AMBR), a variant of memory-based reasoning (MBR) [8]. Given an unknown data to classify, MBR performs majority voting of the labels (player types in our case) among the k nearest neighbors in training data set, where the parameter k has to be decided by the user. On the contrary, AMBR is MBR with k initially set to 1; When ties in the voting occur, it increments k accordingly until ties are broken. Detailed procedure of AMBR is described below:

- **Step I** - For each known data in the training data set, compute the distance from the given unknown data. Mark each data type with a flag.
- **Step II** - Initialize k to 1.
- **Step III** - For all data of the marked data types, find the k nearest neighbors for the unknown data, then count the number of each marked data type.

```

:
:
:
Killer || 1000459 || mffr
Killer || 1000460 || mmmmmmmmmmr
MarkovKiller || 1000059 || fpffmr
MarkovKiller || 1000060 || fffk
PlanAgent || 1000259 || kkkkkkkfkppr
PlanAgent || 1000260 || kpr
:
:
:

```

Figure 6: Typical item sequence data

```

:
:
:
Killer || 1000459 || e1 e1 e1 e1 e1 e1 e1 e1 e1 e1 e1 e1 e1 e5 e2 e2 e5
Killer || 1000460 || e1 e1 e1 e2 e1 e1 e1 e1 e1 e1 e1 e1 e1 e1 e1 e1
MarkovKiller || 1000059 || e5 e10 e5 e2 e2 e5 e2 e5 e2 e1 e1
MarkovKiller || 1000060 || e5 e2 e5 e5 e2 e2 e2 e2 e2 e2 e2 e2 e2 e2 e2
PlanAgent || 1000259 || e8 e2 e2 e2 e8 e5 e2 e10
PlanAgent || 1000260 || e2 e10
:
:
:

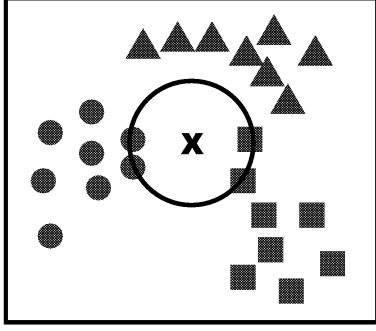
```

Figure 7: Typical episode sequence data

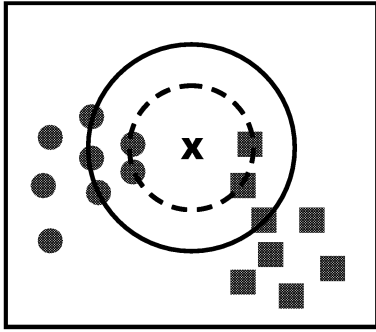
- **Step IV** - If one marked data type has the highest majority, the unknown data is predicted to be that type. Otherwise (if a tie occurs), after unmarking all data types that are not in the tie, set k to 1 added by the number of the data of the marked data types in the tie, then restart from Step III.

Figure 8 depicts a concept of AMBR with 3 types of data represented by circle, triangle and square, respectively. In this figure, only the data of the marked data types are made visible. To predict the type of unknown data represented by cross, first the procedure attempts to find the 1 nearest neighbor (Fig. 8.a), but a tie occurs with two circles and two squares. According to the procedure, after unmarking the triangle type that is not in the tie, k is increased to 5 (Fig. 8.b) by which five circles and three squares are found in the next step. Finally the unknown data is predicted as circle.

Any ideal classifier should be able to correctly classify unknown data not seen in training data set. This ability is called generalization ability. To approximate the generalization ability of the above AMBR, we use leave-one-out method discussed in [9]. In the leave-one-out method, supposing that the total number of training data is M , first, data number 1 is used for testing and the other data including data number 1 are used for training the classifier of interest. Next, data number 2 is used for testing and the other data are used



(a) $k = 1$



(b) $k = 5$

Figure 8: Concept of Adaptive Memory Based Reasoning

for training the classifier. The process is iterated in total M times. In the end, the averaged recognition rate for the test data is computed, and used to indicate the generalization ability of the classifier.

For experiments, log data were generated by running 10 independent Zereal games with 500 simulation-time steps. In each game, we simulated 100 player agents of each type, 100 monsters, and 100 items for each of the other game objects. For the generated log data, we conducted the feature selection algorithms discussed in the previous sections, and obtained input features to AMBR for each sequence type, i.e., action, item, and episode. For the last type, they were generated with the window width $w = 10$. From Table 6, the average length of episode sequences is comparable to that of item sequences.

Table 7 shows the generalization ability of AMBR with action sequences, item sequences, and episode sequences. Based on these results, we performed hypothesis tests for the equality or homogeneity of variances using F-test with 95% confidence and for the mean difference between paired data sets using T-test with 99%

Table 6: Average length for each sequence type

Game	Action	Item	Episode
1	165	9	12
2	182	10	10
3	176	9	10
4	189	10	9
5	186	10	59
6	178	10	13
7	181	10	14
8	154	9	8
9	175	10	58
10	186	10	45
Mean	172.2	9.7	23.8

confidence. Table 8 and Table 9 show the results of F-test and T-test for each paired case, respectively; An upper 1-tail test is used for F-test, and a lower 1-tail test is used for T-test. In the F-test for each paired case, the null hypothesis $\mathbf{H}_0$ is that the variances of two data sets are equal and the alternate hypothesis $\mathbf{H}_1$ is that the variance of the first data set is larger than that of the second data set. Since all P-values of F-test are larger than 5%, all null hypotheses are not rejected. We, thus, can say that the variances of all three data sets are equal with 95% confidence. Regarding T-test, the null hypothesis $\mathbf{H}_0$ is that the means of the two data sets are equal, and the alternate hypothesis $\mathbf{H}_1$ is that the mean of the second group is larger than that of the first group. As all P-values are lower than 1%, all null hypotheses are rejected. We can say that the mean of the second data set in all paired cases is larger than the mean of the first data set and that the mean difference of each paired case is statistically significant with 99% confidence. Consequently, the performance of episode sequences is better than those of action sequences and item sequences.

CONCLUSIONS AND FUTURE WORKS

In this paper, we proposed an episode detection algorithm based on vector space model for players' action sequences in MMOGs. An episode is a combination of actions that exists in the training data. In addition, our algorithm applies a weighting mechanism based on the occurrence frequency of each action to decrease the importance of meaningless actions. In the experiments, we showed that the proposed algorithm transforms action sequences to episode sequences successfully and the resulting episode sequences are better than other kinds of sequences in classification of player types in the simu-

Table 7: Comparisons of the generalization ability of AMBR with each sequence type

Game	Action	Item	Episode
1	0.913333	0.930070	0.930380
2	0.896667	0.913979	0.934426
3	0.913333	0.940351	0.958580
4	0.890000	0.911565	0.938889
5	0.910000	0.917526	0.923333
6	0.896667	0.915493	0.920904
7	0.919732	0.923611	0.948718
8	0.853333	0.847222	0.905882
9	0.913333	0.930070	0.936667
10	0.910000	0.923875	0.932660
Mean	0.9016398	0.9153762	0.9330439

Table 8: Results of F-test

$$\mathbf{H}_0 : \sigma_1^2 = \sigma_2^2 \text{ vs } \mathbf{H}_1 : \sigma_1^2 > \sigma_2^2$$

	Action-Episode	Item-Episode	Action-Item
F-value	1.76	3.03	0.58
P-value	20.71%	5.70%	21.43%

Table 9: Results of T-test

$$\mathbf{H}_0 : \mu_1 = \mu_2 \text{ vs } \mathbf{H}_1 : \mu_1 < \mu_2$$

	Action-Episode	Item-Episode	Action-Item
T-value	-7.19	-3.27	-4.54
P-value	0.00%	0.49%	0.07%

lated MMOG.

In the future works, we'll conduct experiments with a larger number of intelligent agents that have more complex behaviors and detailed characteristics. Moreover we plan to apply the algorithm to different kinds of behavioral sequence data such as real MMOG log data, online transaction data, human behavior data in ubiquitous computing environment, and so forth.

ACKNOWLEDGEMENTS

The first author is now being awarded a scholarship by the Ministry of Education, Culture, Sports, Sciences and Technology, Japan. This work has been supported in

part by the Ritsumeikan University's **Kyoto Art and Entertainment Innovation Research**, a project of the 21st Century Center of Excellence Program funded by the Japan Society for Promotion of Science.

References

- [1] Bartle, R., "Hearts, Clubs, Diamonds, Spades: Players Who Suit MUDs", *The Journal of Virtual Environments*, 1(1), Apr., 2003,
- [2] Ho, J.Y., Matsumoto, Y., and Thawonmas, R., "MMOG Player Identification: A Step toward CRM of MMOGs," *Proc. the 6th Pacific Rim International Workshop on Multi-Agents (PRIMA2003)*, Seoul, Korea, pp. 81-92, Nov., 2003.
- [3] Wong, S.K.M., Ziarko, W., Raghavan, V.V., and Wong, P.C.N., "On Modeling of Information Retrieval Concepts in Vector Spaces", *ACM Transactions on Database Systems*, 12(2), pp. 299-321, Jun., 1987.
- [4] Kim, H.C.(Monica), Pei, J., Wang, W., and Duncan, D., "ApproxMAP: Approximate Mining of Consensus Sequential Patterns." *Proc. the 2003 SIAM International Conference on Data Mining (SIAM DM 2003)*, San Francisco, CA, May, 2003.
- [5] Mannila, H., Toivonen, H., and Verkamo, A., "Discovery of Frequent Episodes in Event Sequences", *The Journal of Data Mining and Knowledge Discovery*, 1(3), pp.259-289, Nov. 1997.
- [6] Tveit, A., Rein, Y., Jrgen, V.I. and Matskin, M., "Scalable Agent-Based Simulation of Players in Massively Multiplayer Online Games.", *Proc. the 8th Scandinavian Conference on Artificial Intelligence (SCAI2003)*, Bergen, Norway, Nov., 2003.
- [7] Thawonmas, R., Ho, J.Y., and Matsumoto, Y., "Identification of Player Types in Massively Multiplayer Online Games," *Proc. the 34th Annual conference of International Simulation And Gaming Association (ISAGA2003)*, Chiba, Japan, pp. 893-900, Aug., 2003.
- [8] Michael, J.A.B. and Linoff, G., *Data Mining Techniques-For Marketing, Sales, and Customer Support*, John Wiley & Sons, Inc., N.Y., 1997.
- [9] Weiss, S.M. and Kulikowski, C.A., *Computer Systems That Learn*, Morgan Kaufmann Publishers, San Mateo, CA, 1991.

Mapping System Theory Problems to the Field of Knowledge Discovery in Databases

Van Welden D.

Ghent University, Applied Mathematics, Biometrics and Process Control
Coupure Links 653, B-9000 Gent
Belgium

ABSTRACT

Some General System Theory (GST) problems in the domain of induction can be transferred to the domain of Knowledge Discovery in Databases (KDD) where they can more easily be tackled. The corresponding mapping delegates not only the problem solving to the people working in KDD, but also it enriches the latter domain with some additional techniques in the application of data mining techniques to dynamic systems. This paper gives a very brief overview of the mapping and shows some commonality between GST and KDD.

INTRODUCTION

Machine Learning (ML), statistics, and KDD have in common a strong link : they all acknowledge the importance of induction as a normal way of thinking. Induction, however, is also known in the field of system theory, where it deals more specifically with (often/mostly linear) time-varying systems, Ljung [1987].

Induction systems for non-linear systems were made in GST and put in a framework by Klir [1985]. However, the implementation of these concepts showed to be far from trivial. Problems were tackled by taking an exhaustive approach at first, but this limited the application to simple real-world applications. The use of heuristics improved upon that by circumventing the combinatorial combinations problem for more complex cases. However, at that time KDD was emerging and new techniques became available for the analysis of large amounts of data.

In this paper, it is shown how some induction problems from the field of GST can be mapped into the field of KDD resulting in a new way of analysing dynamic real-world systems.

SHORT HISTORY AND TERMINOLOGY

A tool that performs induction on dynamic systems can be found under the name of SAPS (System Approach Problem Solver). It is a pattern recognition approach for system identification, based on GSPS. It searches a kind of “mental” model for dynamic directed (i.e., there is a distinction between input variables and output variables) black box systems with an underlying mainly deterministic relation-

ship between outputs and inputs. Initially, an exhaustive search of states was executed, but it soon turned out that this was not feasible for many real-world systems. The use of heuristic search techniques (hill-climbing) made it possible to tackle more realistic problems, [Van Welden 92]. Further extending this concept led eventually to a mapping of the initial transformed state variables to a tool called CART (Classification and Regression Trees). The latter tool is situated in the domain of KDD, which makes it possible to undertake induction of real-world dynamic systems. Or, put it in a more fashionable way: to data mine dynamic systems.

A key concept to this endeavour is a general idea of a model. A nice definition can be found in Minsky [1965]: “An object 'A' is a model of an object 'B' for an observer, if the observer can use 'A' to answer questions that interest him about 'B' “.

The GST point of view is interdisciplinary and has brought a new perspective, a new way of doing science. It provides an organized body of knowledge for dealing with systems in general, in which there are basic concepts, a number of principles, some rigorous enough to be considered “laws,” and, a general framework for theory construction. In this framework, isomorphic schemes play an important role. Hence, GST touches virtually all traditional disciplines, from mathematics, technology and biology to philosophy and the social sciences and includes AI, neural networks, dynamical systems, chaos, and complex adaptive systems.

KNOWLEDGE DISCOVERY IN DATABASES

Storage has become much cheaper the last decade and databases are larger than ever before due to the pervasive digitalisation of (for example: customer) data. Processing power has increased tremendously and the emergence of new induction techniques opened a new field in which the analysis of large numbers of data has become feasible. This new field, called KDD, is a data exploration methodology that is defined to be the non-trivial extraction of implicit, previously unknown, potentially useful, ‘relatively simple’, and not predefined information from large databases.

KDD has the possibility in it to give a good return on investment. Correspondingly, it is employed in finance:

e.g., fraud detection; stock market prediction; credit assessment, in CRM (Customer Relation Management) where one tries to see how to improve the targeting efficiency, in quality control, in medicine: e.g., effect of drugs, diagnosing, hospital cost analysis, in astronomy (cataloguing), molecular biology (finding patterns in molecular structures), text mining, web mining, etc.

KDD starts with the *goal definition*, which must be formalised and made executable so that it can be related to relevant data, which are hopefully present in the database. The *data pre-processing* step prepares and reshapes the data for subsequent processing. It involves data and attribute focusing, data cleaning, data projection, and data augmentation. The *data mining* step induces the model. It consists of a model specification, model fitting, model evaluation, and model refinement. *Consolidation of the newly found knowledge* and output generation consists of interpreting and documenting the found patterns, and if more model types were used, comparing them. Any conflicts (if any) that may arise with previous knowledge in the knowledge base must be resolved. A priori knowledge can be used in any step.

KNOWLEDGE- OR DATA-BASED APPROACH?

The 2 approaches from the title are depicted in Table 1.

	knowledge based approach	data based approach
synonyms	modelling	system identification
	top-down modelling	bottom-up approach
reasoning	deduction	induction
does what?	encodes the (inner) structure of the system	encodes the behaviour of the system (via experimental data)
problem type	analysis	synthesis

Table 1 : The two main approaches to model construction

A bottom-up approach tries to infer the structural information from experimental data and to come up with a usable model under a given experimental frame. This approach may generate an infinite number of models satisfying the observed input-output relationships. So, there is no straightforward procedure for determining the structure of a model. A set of guiding principles and quantitative procedures for inferring structure parts from data sets is needed (inductive bias, Michie [1994]).

KDD can be used for light gray systems (even white systems for fault diagnosis), but its major field of application is towards black box systems. This affects the kind of

validity that prevails. Validity concerns the process of determining whether or not the model is an acceptable description of reality (model-reality relation). Replicative (the easiest thing to achieve) and predictive validity is important in both GST and KDD, while structural validity is more an issue in GST.

LIFE CYCLE OF GST AND KDD

A merged life cycle of GST and KDD is represented in Figure 1. In this life cycle, three major conceptual blocks are to be distinguished:

- A knowledge or model base
- The business part
- The modelling construction part

The Knowledge Base

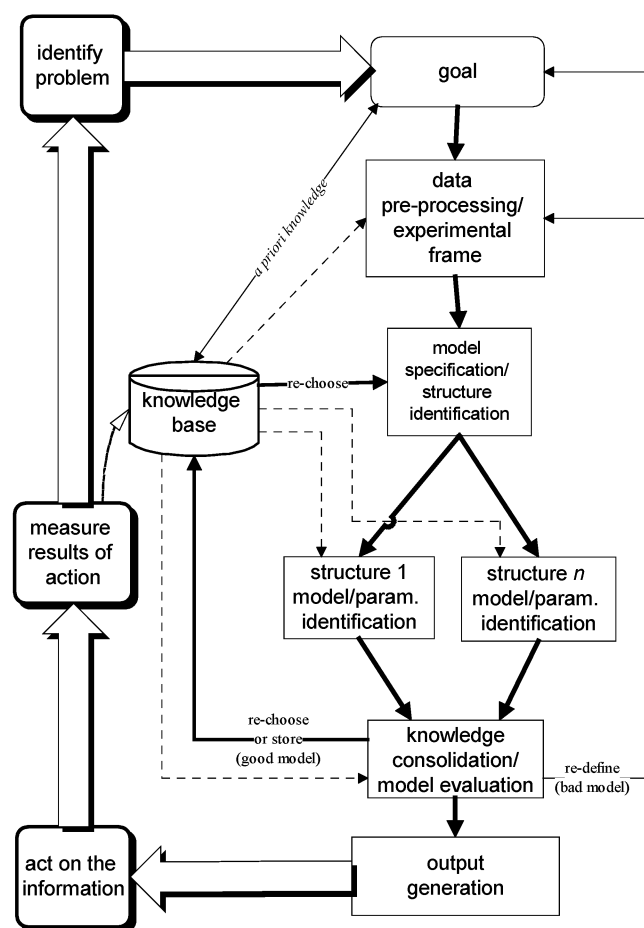


Figure 1 : General modelling and KDD

Scheme's for modeling and simulation that take into account the re-use of models are integrated in a virtuous cycle of data mining [Berry 1997], which shows its complete enterprise approach to KDD. Figure 1 combines these two in an unifying attempt.

Models of standard components should be saved in libraries (model bases) so reuse is promoted. A model can

have any number of components defining the interface and the behaviour. The description of the actual behaviour of a model will be called the models' realization.

The knowledge/model base supports model reuse, the models have to be expressed in a language which admits structural decomposition with well defined interfaces between submodels. Apart from rules, the KB may also store entire models or sub-models. Also, model inheritance and model decomposition are 2 orthogonal concepts for enhancing model storage efficiency.

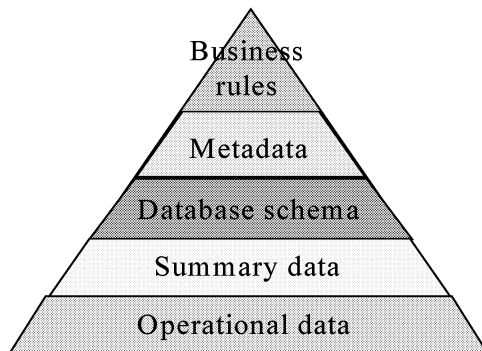


Figure 2 : The use of meta data in data warehousing (from [Han 1999])

The concept of meta-data exists in both GST and KDD, but it seems meta-data has more semantic richness in KDD.

A first aspect of this richness can be illustrated by looking at the data types in a data warehouse, which are shown in Figure 2. Operational data is the data itself. In data warehousing, this also means where it comes from, when it was stored, etc. Summary data gives summaries of the data, so it is a kind of meta-data already. The database scheme gives the physical layout of the data. The meta-data level itself is the logical model. The meta-data repository also encompasses business terms and definitions, ownership of data, and charging policies. For example, operational meta-data concerns:

- data lineage: history of migrated data and sequence of transformations applied
- currency of data: active, archived, purged
- monitoring information: warehouse usage statistics, error reports, audit trails

A second aspect of the richness of the term meta-data in KDD is found in the data mining approach. Here, meta-data says something about the measurement scale; whether variables are derived from others (data projection); if there are bounds on data; if there are structural zero's; whether there are distributional assumptions made (parametric versus non-parametric), etc. Meta-data is thus much related to the use of a-priori knowledge in data-pre-processing (structural zero's, derived variables), and in model specification and selection (which data mining method to use). Remark that the distinction between data and meta-data may be blurred.

The Business Part

The 'management' part of Figure 1 (identify problem, measure results of action and act on the information) is very general. Thus, it is valid for any modelling attempt, be it via GST or via KDD. It is a process that applies knowledge that is gained from increased understanding of customers, markets, products, and competitors to internal (continuously evolving) processes. It is not the aim of this paper to elaborate on this part.

The Modelling Construction Part

Of course, these concepts will have an impact on the life cycle of modeling. One distinguishes

- the goal formulation in both GST and KDD,
- the data pre-processing versus experimental frame definition,
- the model specification versus structure identification (use of shallow versus deep models)
- data mining or parameter estimation
- knowledge consolidation and model evaluation

Goal Formulation

All problem statements start with the identification of the problem. They try to set a goal in order to solve the problem at hand. To be able to meet a goal, it has to be formalised so that a rigorous approach to problem solving can be obtained.

With regard to goal formulation, similarities are present: e.g., gaining insight (in KDD) corresponds with understanding (in GST). Both domains stress the principle that a model should not be more complicated than absolutely necessary (Occam's razor). GST puts some more emphasis on a correct classification and prediction. A speedy or less costly classification and prediction is less the issue in GST (except perhaps in the domain of control theory), but of more importance in KDD. Hence, the 'Quick decision' problem type, where accuracy may be less important than comprehensibility, is less stressed in classical modelling. A similar argument applies for the cost of a model.

Data Pre-Processing Versus Experimental Frame Definition

An experimental frame isolates specific input/output behaviour both of the real system and its model, [Elzas 1984]. In KDD, this corresponds with the data pre-processing step. Attribute focusing is in fact nothing more than defining what are supposed to be relevant variables for the system under investigation. This is consistent with the notion of an experimental frame and with what Klir calls an observation channel. It also fits in the viewpoint, taken in systematic modelling, which states that a model is conceived as a collection of variables and relations among them, [Ören 1984]. The collection of relevant variables is

Model specification	Rules	Differential equations	Tree classifiers	Neural networks	Time series
Model structure	Conjunctive, disjunctive	Linear, non-linear, time-invariant, ...	Classification, regression, survival ...	Kohonen, back-propagation, ...	AR, ARMAX, MA, ...
Model complexity	Rule size	Order	Number of nodes	Number of neurones	Order

Table 2 : Model specification, structure and complexity

called attribute focusing, while identifying the relations is part of the data-mining step. Data focusing or sampling data is less used in systems theory, because the nature of data is often different than for KDD. In KDD, one usually starts with a collection of static independent records. Taking a relevant subset poses no major problems. However, in general systems theory, taking a subset of time-depending data is usually not done: data records do depend on each other. This does not mean that GST has no way of dealing with an abundance of data records. A kind of data reduction can be done in GST via the determination of the Nyquist frequency and via filtering techniques. Still, it illustrates the static(KDD) and dynamic(GST) modelling aspects of the respective domains.

Model Specification Versus Structure Identification

The data-mining step in KDD is much related to modelling in system theory. Model specification can be compared with model structure identification. It involves deciding what type of model to use (in system theory: Bond graphs, Petri nets, block diagrams, ... in data mining: classification trees, hierarchical clustering, linear regression, neural networks, etc.). The model specification step is straightforward applicable to both domains and can be considered at different abstraction levels. On a very abstract level this involves choosing if one wants to use Bond graphs, Petri nets, block diagram, rules, neural nets, etc. (for systems theory) or trees, clustering, rules, neural nets, regression models, etc. (for KDD). On a more concrete level, one has to further specify the model. Examples are: linear, logistic, or non-linear regression (KDD), linear or non-linear models (GST), state-space or transfer functions in block diagrams (GST), kind of neural net (both domains), kind of tree (decision, regression, ...) (KDD), order of a differential equation (GST), type of clustering (KDD), etc. In both domains, the goal has a large impact on the used model types: when comprehensibility is more important certain representations may be more preferred than others. For example in GST, a block diagram may be more comprehensible than a Bond graph (it is also relative to the field of expertise). A tree structure is more comprehensible than a neural network (KDD). Rules are more comprehensible than some other model types (both domains), and neural networks are usually the least comprehensible (both domains). The issue about comprehensibility has a lot to do with the greyness of the model: black box models are always less comprehensible than white box models.

The amalgamation of terminology from GST and KDD may shed new light on the terms 'model specification', 'model structure', and 'model complexity'. Model specification has a lot to do with the goal setting, while the model structure is more determined by the experimental frame. Model complexity is related to validation and parameter estimation.

Table 2 shows the terms in relation to each other and it gives an idea of the degree of abstraction involved.

Data Mining Or Parameter Estimation

Model fitting involves parameter estimation (identification) in modelling. It is also called model calibration, [Elzas 1984]. This reduces to estimating parameters or coefficients in differential equations (GST), parameters in state-space models (GST), regression models (KDD), etc.

Model validation consists of comparing the behavioural data of the system under investigation and the calibrated (fitted) model. Usually, a train-and test method is used. A rule of thumb is that 2/3 is used for fitting the model (training) and 1/3 for validation (testing). Cross-validation is commonly used in KDD, while it is not so popular in GST. Even more specifically, bootstrapping is known too in KDD, but almost unknown in GST. Replicatively validity is known in GST: it consists of fitting the model on the training set. In KDD, this is better known as internal estimates (resubstitution estimates, [Breiman 1984]). Predictively validity is done on a test set; it gives true estimates.

Knowledge Consolidation And Model Evaluation

KDD uses an interesting function for *evaluating* a model. The evaluation can be based on more than just accuracy performance. For example, KDD can take into account what is economically interesting (cost of model), or it can take the faster model with regard to prediction. This may prove an important point for the modelling society when they want to evaluate their models in an economic context (cost of modelling), or when speed is of the utmost importance. KDD provides a more general framework for dealing with these situations.

Model refinement is equally applied in GST and KDD; when a model does not validate well, another model (struc-

ture) is chosen. When this fails too, one can go one step further back and redefine the experimental frame or even adjust the goal. The refinement of an existing model is a major issue in modelling.

From Figure 1, it can be seen that many models may be used in parallel. The models can be of a different specification (e.g., neural nets and genetic algorithms), and they can be situated on different epistemological levels (e.g., rules versus decomposed models). In the latter case, one speaks of shallow versus deep models (in GST).

Models are not only evaluated with regard to an interesting function, or validated with regard to a goal setting, but in KDD, model *specifications/paradigms* are also compared with each other. This belongs to the knowledge consolidation step. Here, model evaluation is on another epistemological level than in the data-mining step. Models are not only compared on a test set to validate the parameter estimation, but they are compared on yet another (independent) test (or evaluation) set to evaluate the chosen model specification.

In both GST and KDD, the consolidation with regard to storing the found knowledge is present. In GST, the newly found model is stored in a model base (called modelling in the large, while in KDD this is not stated so explicitly).

Finally, both GST and KDD acknowledge the necessity of feedback from steps that appear later in the cycle to steps that appear earlier. This is indicated in Figure 1. Therefore, the steps in the life cycle are more intertwined than one should expect at first sight.

In GST, Klir showed how one could map time information in the state domain, but this generates a lot of extra data. However, the main advantage is that the resulting data set is now static in structure and hence is very well suited for DM techniques. Elaboration in detail in this paper would lead too far, but it should be clear now that both domains have their own focus and can complement each other in their problem solving approach.

CONCLUSION

GST focuses more on accuracy of a model, sometimes speed and less the cost of a model. It does not lay emphasis on data reduction as KDD does. Modelling of dynamical systems is the primary focus of research. Decomposition of models is common place.

KDD puts more emphasis on static systems and focuses more on model comparison in the large (via a third test set), on data warehousing and relies on a richer semantic meta-data structure. One has more experience with very large databases and with high dimensionality problems. Data reduction is commonly used. The interestingness function is more general than the evaluation functions used in GST.

Hence, the complementary aspect of both domains makes it fruitful to have a good cross-fertilization.

REFERENCES

- Berry M. J. A., Linoff G. [1997], *Data Mining Techniques For Marketing, Sales and Customer Support*, Wiley Computer Publishing, 1997.
- Breiman L., Friedman J. H., Olshen R. A., Stone C. J. [1984], *Classification and Regression Trees*. Chapman & Hall, 1984.
- Elzas M.S. [1984], "*System Paradigms as Reality Mappings*", *Simulation and Model-Based Methodologies: An Integrative View*, ed. Ören T.I. et al. NATO ASI Series F: Computer and System Sciences, vol. 10, Springer Verlag, p. 41- 68, 1984.
- Han J. [1999], Personal communication, (from a lecture for the IBM chair in Antwerp 1999).
- Klir G.J. [1985], *Architecture of Systems Problem Solving*. Plenum Press, 1985.
- L Ljung [1987], *System Identification - Theory for the User*, Prentice-Hall, Englewood Cliffs, N J, 1987.
- Minsky M.L. [1965], "*Matter, Mind, and Models*", *Proceedings of the IFIP Congress*, 1, Spartan Books, p. 45-49, 1965
- Michie D., Spiegelhalter D.J., Taylor C.C. [1994], *Machine Learning, Neural and Statistical Classification*. Ellis Horwood, 1994.
- Ören T.I. [1984], "*Model-Based Activities: A Paradigm Shift*", *Simulation and Model-Based Methodologies: An Integrative View*, ed. Ören T.I. et al., NATO ASI Series F: Computer and System Sciences, vol. 10, Springer Verlag, p. 3-40, 1984.
- Van Welden D., Vansteenkiste G.C. [1992], "*A Mixed Deductive-Inductive Approach to Model Recognition*", *Proceedings of the 1992 European Simulation Multiconference*, York, UK, June 1-3, p. 112-118, 1992.

BUSINESS GAMING AND SIMULATION

An Information-Rich, Virtual Trading Environment

Les Green
Faculty of Information Technology
University of Technology, Sydney
PO Box 123, NSW 2007,
Australia
lesgreen@it.uts.edu.au

KEYWORDS

E-Market, Synthetic Environment, Business Process

ABSTRACT

The E-Markets group at the University of Technology, Sydney has been undertaking research on the evolution of trading environments. The main focus is on mechanisms for innovation. Work to date includes: simulation experiments to investigate the relative impact and cost of the fundamental evolutionary mechanisms, the construction of an e-Market kernel embedded in a virtual world, the construction of smart mining bots to extract and condense information from the market context, and a process management system. In this project every transaction is managed as a business process. Future work focuses on understanding the evolution of business networks in the virtual marketplace. This paper presents the current state-of-art project in information rich, virtual trading environments.

INTRODUCTION

The overall goal of this work is to derive fundamental insight into how the next generation of e-market will evolve. The work is mapping out and exploring future generation trading environments. e-Marketplaces will be complete and immersive. **Complete** in that all transactions, including requests for information, will be handled within the e-marketplace and so it contains all information sources. **Immersive** in that e-marketplaces are 'virtual realities' in which players interact as lifelike avatars.

Another key aim is to understand the evolution of business networks that are established in a virtual marketplace. Three fundamental evolutionary mechanisms are innovation, imitation and improvement of existing procedures. The main focus is on mechanisms for innovation. The perturbation of market equilibrium through entrepreneurial action is the essence of market evolution by innovation. Entrepreneurship relies both on intuition and information discovery. The term 'entrepreneur' is used

here in its technical sense (Watkins 2002). Smart systems that deliver timely information to support the market evolutionary process in a virtual marketplace are described here. The deep research goals that motivate this work are to understand the processes that will drive both market evolution and the evolution of business relationships in a future generation trading environment.

Work to date includes: simulation experiments to investigate the relative impact and cost of these three fundamental evolutionary mechanisms, the construction of an e-Market kernel, the construction of smart mining bots to extract and condense information from the market context, and a process management system that handles the transactions in such a place. The goal of the bots constructed to date is to identify timely information for traders in an e-market. The traders are the buyers and sellers.

The e-market was designed by Professor John Debenham and Professor Simeon Simoff. It is a virtual marketplace that is embedded in virtual worlds technology. The actors' assistant is being constructed in the Faculty of Information Technology at the University of Technology, Sydney. It is funded by an Australian Research Council three-year Discovery Grant, and various other grants. Our e-marketplace consists of a market kernel embedded in an information-rich environment. It is implemented within virtual reality space.

The smart information system addresses the problem of identifying timely information for e-markets with their rapid, pervasive and massive flows of data. This information is distilled from individual signals in the markets themselves and from signals observed on the unreliable, information-overloaded Internet. Distributed, concurrent, time-constrained data mining methods are managed using business process management technology to extract timely, reliable information from this unreliable environment.

One feature of the whole project is that every transaction is treated as a business process and is managed by a process management system. In other words, the process management system makes the whole thing work. These transactions include simple

market transactions such as “buy” and “sell” as well as transactions that assist the actors in buying and selling. The process management system is based on a robust multiagent architecture and has been constructed in Jack. The use of multiagent systems is justified first by the distributed nature of e-business, and second by the critical nature of the transactions involved.

The overall goal of this work is to investigate the evolution of e-markets. We look at some factors that cause changes in electronic markets. A market is said to be in *equilibrium* if there is no opportunity for arbitrage—ie. no opportunity for no-risk or low-risk profit. Real markets are seldom in equilibrium. Market *evolution* occurs when a player identifies an innovative opportunity for arbitrage and takes advantage of it possibly by introducing a novel form of transaction. At a deeper level, business networks that support market activity also evolve and are the key focus of this investigation.

A trading agent strives to make informed decisions in an information-rich environment. Its design was provoked by the observation that agents are not always utility optimizers, and by the quotation: “Good negotiators, therefore, undertake integrated processes of knowledge acquisition combining sources of knowledge obtained at and away from the negotiation table. They learn in order to plan and plan in order to learn” (Watkins 2002). The agent attempts to fuse the negotiation with the information generated both by and because of it. It reacts to information derived from its opponent and from the environment, and proactively seeks missing information that may be of value.

Business networks are the result of the *interaction* between the players (traders) in the electronic market environment. These interactions are facilitated by synthetic characters presented in more detail in the next section.

SYNTHETIC CHARACTERS

To represent autonomous agents in the information rich environment, anonymous social agents have been embedded in the UTS eMarket environment. Avatars represent these agents visually in the eMarket space. They provide casual, but context-related conversation with anybody who chooses to engage with them. This feature markedly distinguishes this work from typical role-based chat agents in virtual worlds. The latter usually are used as “helpers” in navigation, or for making simple announcements. They are well informed about recent events, and so “a quick chat” can prove valuable. A demonstration that contains one of these synthetic characters is available at: <http://research.it.uts.edu.au/emarkets/> — first click on “virtual worlds” under “themes and technologies” and then click on “here”. To run the demonstration you will need to install the Adobe Atmosphere plugin.

The following issues have to be addressed in the design of any synthetic character:

- its appearance

- its mannerisms
- its sphere of knowledge
- its interaction modes
- its interaction sequences

and how all of these fit together.

Appearance is a key issue in initiating contact but is not addressed here. Mannerisms are equally important—a researcher associated with the UTS group is investigating the impact of facial expressions on interaction. This is a major issue that is beyond the scope of this discussion.

An agent’s sphere of knowledge is crucial to the value that may be derived by interacting with it (Barthelemy et al. 2004). We have developed extensive machinery, based on unstructured data mining techniques, that provides our agent with a dynamic information base of current and breaking news across a wide range of financial issues. For example, background news is extracted from on-line editions of the Australian Financial Review.

Our agent’s interaction modes are presently restricted to passive, one-to-one interaction. They are passive in that these agents do not initiate interaction. We have yet to deal effectively with the problem that occurs when a third party attempts to “barge in” on an interaction.

The interaction sequences are triggered by another agent’s utterance. The machinery that manages this is designed to work convincingly for interactions of up to around five exchanges only. Our agents are designed to be “strangers” and no more—their role is simply to “toss in potentially valuable gossip”.

The aspects of design discussed above cannot be treated in isolation. If the avatars are to be “believable” then their design must have a unifying conceptual basis. This is provided here by the agent’s *character*. The first decision that is made when an agent is created is to select its character using a semi-random process that ensures that multiple instances of the agent in close virtual proximity have identifiably different characters.

The dimensions of character that we have selected are intended specifically for a finance-based environment. They are:

- Politesse
- Dynamism
- Optimism
- Self-confidence

The meaning of each of these dimensions is (moderately) independent of the others:

- Politesse means the use of polite words, phrases and forms
- Dynamism is the tendency to react rapidly, succinctly and vigorously
- Optimism here means a tendency to use up-beat phrases and the tendency to not use negations
- Self-confidence here means the tendency to respond with declarative statements rather than tentative propositions or questions

The selection of the character of an agent determines: its appearance (ie: which avatar is chosen for it) and the style of its dialogue. Future plans will

address the avatars mannerisms. The selection of agent's character provides an underlying unifying framework for how the agent appears and behaves, and ensures that multiple instances of the agent appear to be different.

The selection on an agent's character does not alone determine its behavior. Each agent's behavior is further determined by its moods that vary slowly but constantly. This is intended to ensure that repeated interactions with the same agent have some degree of novelty. The dimensions of moods that we have identified are:

- Happiness
- Sympathy
- Angry

Synthetic characters provide only one of the components of the immersive trading environment. In the next section we present the negotiation support.

AUTOMATED NEGOTIATION

We illustrate the automated negotiation approach on the example of negotiation between two agents in an information-rich environment. That is, the agents have access to general information that may be of use to them. They also exchange information as part of the negotiation process. The two agents are called "me" and my opponent ω . The environment here is the Internet, in particular the World Wide Web, from which information is extracted on demand using special purpose 'bots'. So my agent, "me", may receive information either from ω or from one of these bots. In a 'real life' negotiation, the sort of information that is tabled during a negotiation includes statements such as "this is the last bottle available", "you won't get a better price than this", and so on. To avoid the issue of natural language understanding, and other more general semantic issues, the interface between each of the negotiating agents and these information sources is represented using the language of first-order, typed predicate logic, and a set of pre-agreed, pre-specified predicates. All information passed to "me" is expressed in this way.

As well as being unambiguous, the use of first-order typed predicate logic has the advantage of admitting metrics that describe, for example, how "close" two statements are. These metrics are useful in enabling the agent to manage the information extracted from the environment in a strategic way.

Negotiation between two trading agents is a two-stage process. First, the agents exchange offers whilst acquiring and exchanging information. Second, they attempt to reach a mutually satisfactory agreement in the light of that information. This process aims to reach informed decisions in eMarket bargaining by integrating the exchange of offers and the acquisition and exchange of information drawn from the environment.

Negotiation then proceeds by a loose alternating offers protocol that is intended to converge when the agents believe that they are fully informed.

Each agent exchange offers at each discrete time. The agents enter into a commitment if one of them accepts a standing offer. The protocol has three stages:

1. Initial offers from both agents;
2. A sequence of alternating offers, and
3. An agent quits and walks away from the negotiation.

The negotiation ceases either in the second or final round if one of the agents accepts a standing offer or in the final round when one agent quits and the negotiation breaks down.

PROCESS MANAGEMENT

To successfully manage a large number of negotiation threads, every market transaction is managed as a business process. To achieve this, suitable process management machinery is required. To investigate what is 'suitable' the essential features of these transactions are related to two classes of process that are at the 'high end' of process management feasibility. The two classes are goal-driven processes and knowledge-driven processes. The term "business process management" is generally used to refer to the simpler class of workflow processes (Fischer 2000), although there notable exceptions using multiagent systems (Jennings et al. 2000) (Huhns and Singh 1998).

Goal-Driven Processes

A *goal-driven process* has a process goal, and achievement of that goal signals the termination of the process. The process goal may have various decompositions into possibly conditional sequences of sub-goals such that these sub-goals are associated with (atomic) activities and so with atomic tasks. Some of these sequences of tasks may work better than others, and there may be no way of knowing which is which (van der Aalst and van Hee 2001). A task for an activity may fail outright, may fail to achieve its goal, or may violate process constraints (Debenham 2004). In other words, a central issue in managing goal-driven processes is the management of task failure. Hybrid multiagent architectures whose deliberative reasoning mechanism is based on "succeed/fail/abort plans" (Rao and Georgeff 1995) are well suited to the management of goal-driven processes.



Figure 1. A simplified view of goal-driven process management.

Knowledge-Driven Processes

A second class of process, whose management has received little attention, is called knowledge-driven processes (Debenham 2000). *Process knowledge* is all the knowledge that is relevant to a process instance. It includes common-sense knowledge, knowledge that was available when an instance is created, and knowledge acquired during the time that the instance exists. A *knowledge-driven process* may have a process goal, but the goal may be vague and may mutate. In so far as the process goal gives direction to goal-driven—and activity-driven—processes, the process knowledge gives direction to knowledge-driven processes. The body of process knowledge is typically large and continually growing and so knowledge driven processes are seldom considered as candidates for process management. They are typically supported, rather than managed, by CSCW systems (van der Aalst and van Hee 2001). But, even complex knowledge-driven processes are “not all bad”—they typically have goal-driven sub-processes which may be handled as described above. *Knowledge-base processes* are a special type of knowledge-driven process for which the process knowledge *can* be represented and accessed by a process management system. This proves to be a useful concept for managing e-marketplace transactions.

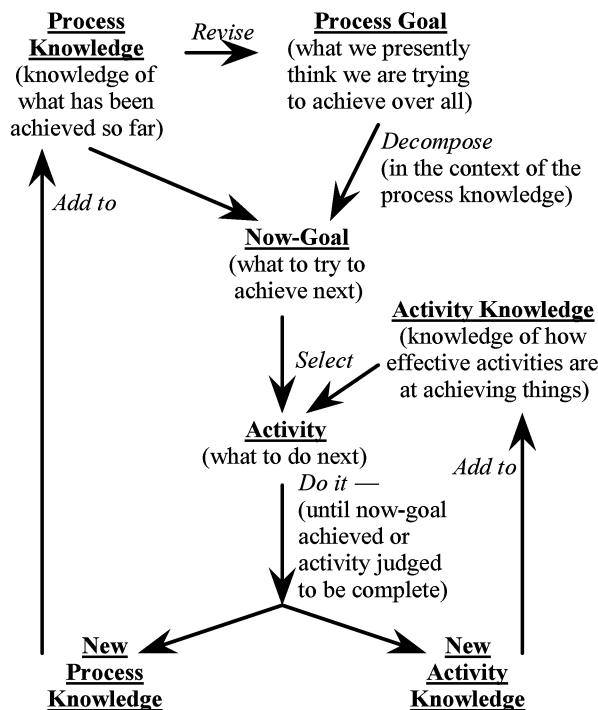


Figure 2. A simplified view of knowledge-driven process management.

The management of goal-driven and knowledge driven processes are radically different. Goal-driven processes may be managed by a goal/plan decomposition process (see Fig. 1), and knowledge-driven processes are

managed by continually reviewing the growing corpus of process knowledge—this is illustrated in Fig. 2. That Figure is deceptively simple in that the business of managing the process knowledge and of revising the process-goal and next-goal in the light of that growing body of knowledge is far from trivial in even simple examples.

This approach allows us to incorporate and synchronise the work of the main players – the trading agents in the virtual trading environment with the requests and provision of the necessary market information, performed by data and text mining agents.

DATA AND TEXT MINING

Data and text mining agents are responsible for extracting, verifying, combining, filtering and condensing information drawn from the trading environment and from the Internet. This information is then fed to the players in the e-marketplace. A data stream can include both structured (numerical and symbolic data) and unstructured (text, images and video data). Specialised data mining ‘bots’ (software agents), each of which applies a specific data mining method, are developed for extracting patterns from these data streams.

The agents in this project support data mining as a concurrent process, which allows us to identify potentially useful patterns within a large variety of rapidly-changing signals. These technologies include mining text data from news feeds as well as text and numerical data from financial reports and web pages. Text mining bots identify particular topics and trace those topics in the news stream. An initial prototype of a multi-agent system for identifying relevant events in news feeds and tracking those events from news sources on the Internet is fully operational. Potential future work includes the development of image and video mining technologies which complement the text mining tools, allowing us to extract context information from additional (alternative) media sources.

This work does not, for example, attempt to use gathered data to predict whether the US dollar will rise or fall against the UK pound. What it does attempt to do is predict the value that an actor in the system will place on the data (Han et al. 1999). So the feedback here is provided by the ‘bots’ in the form of a rating of the material found. A five point scale runs from ‘totally useless’ to ‘very useful’ to classify this data.

As mentioned in section 1, the environment is complete, i.e. all transactions are handled within. Information about these transactions is collected in the application level logs. This information is the source for a further type of data mining – Network Mining, which aims at discovery and analysis of business networks and their evolution.

NETWORK MINING

Uncovering the patterning of people's interaction has been the subject of social network analysis (Wasserman and Faust 1994). Relational data mining

(Dzeroski and Lavrac 2001) extends these techniques beyond the comparative static methods typical of existing analysis systems (e.g. UCINET, Inflow) to the examination of large continuous flow data sets representing links between a variety of entities (e.g. friendship ties, links between places, things, ideas, data, web pages, events). Existing data mining techniques that are currently considered include a number of probabilistic relational techniques (Kersting and De Raedt 2001).

The work involves the analysis of activities and the diffusion of information in networks of virtual communities leading to the emergence of new business models and networks (Hagel and Armstrong 1997). Relevant background to this project are recent works on identifying the “network value” of a person (node) within a social network structure (e.g. (Domingos and Richardson 2001)), based on person interactions and information diffusion, and research on the role and importance of social capital and network competence in business performance and innovation (Ritter et al. 2002). Our work uses the continual, extensive and complex data flows from the virtual market to detect, measure and track static and dynamic patterns of change in social, business and knowledge networks and relate this to innovation events and the evolution of the players’ strategies.

CONCLUSIONS AND FUTURE RESEARCH

E-markets in their current form must evolve into Virtual Trading Environments where well informed decisions are made from an information rich, immersive environment. In order for these decisions to be made, information must be gathered from the environment as well as other actors and extracted from the interactions which happen between them. Data and text mining as well as network mining techniques have been applied in this area.

In the discussed work, all transactions are handled by a process management system where both knowledge based and goal based processes are discussed. The environment has been expressed by a virtual world where avatars represent the actors involved and other contributing characters such as informative ‘bots’.

The E-Markets group at the University of Technology, Sydney has developed an appropriate framework to handle the evolution of a Virtual Trading Environment and facilitate effective operation.

The project recognises that social relationships and knowledge networks (networks of ideas) established in the virtual marketplace are central to the recognition of opportunities, productive links among existing ideas, and to their exploitation and implementation. The future research will investigate the nature of business networks that can be established in an electronic marketplace, their evolution and possible ways in assisting and guiding their development.

REFERENCES

- Barthelemy F, Dosquet B, Gries S & Magnant X. 2004, Believable Synthetic Characters In A Virtual Emarket. in *proceedings International Conference on Artificial Intelligence And Applications AIA 2004*, February 16-18, 2004, Innsbruck, Austria
- Debenham, JK. 2000, Supporting knowledge-driven processes in a multiagent process management system. In *proceedings Twentieth International Conference on Knowledge Based Systems and Applied Artificial Intelligence, ES'2000: Research and Development in Intelligent Systems XV*, Cambridge UK, December 2000, pp273—286.
- Debenham J K. 2004, Electronic Market Transactions Managed as Knowledge-Driven Processes. In *proceedings The 9th IEEE International Conference on Engineering of Complex Computer Systems*. Florence, Italy, 14-16 April, 2004
- Domingos, P. and Richardson, M. 2001, Mining the Network Value of Customers. *Proc. 7th Int. Conf. on Knowledge Discovery and Data Mining*. San Francisco, CA 57-66
- Dzeroski, S. and Lavrac, N. (eds), 2001, *Relational Data Mining*. Springer-Verlag, Berlin
- Fischer, L. (Ed). 2000, *Workflow Handbook 2001*. Future Strategies
- Hagel III, J. and Armstrong, A.G. 1997, *Net Gain: Expanding Markets Through Virtual Communities*. Harvard Business School Press, Boston, Massachusetts
- Han, J. Lakshmanan, L.V.S. and Ng, R.T. 1999, Constraint-based multidimensional data mining. *IEEE Computer*, 8, 46-50.
- Huhns, M.N. and Singh, M.P. 1998, *Managing heterogeneous transaction workflows with cooperating agents*. In N.R. Jennings and M. Wooldridge, (eds). *Agent Technology: Foundations, Applications and Markets*. Springer-Verlag: Berlin, Germany, pp. 219—239.
- Jennings, N.R., Faratin, P., Norman, T.J., O'Brien, P. and Odgers, B. 2000, Autonomous Agents for Business Process Management. *Int. Journal of Applied Artificial Intelligence* 14 (2) 145-189
- Kersting, K. and De Raedt, L. 2001, “Adaptive Bayesian Logic Programs. In: Rouveirol, C. and Sebag, M. (eds): *Proceedings of the 11th OInternational Conference on Inductive Logic Programming*. Springer-Verlag, Heidelberg (2001), 104-117
- Rao, A.S. and Georgeff, M.P. 1995, “BDI Agents: From Theory to Practice”, in *proceedings First International Conference on Multi-Agent Systems (ICMAS-95)*, San Francisco, USA, pp 312—319.
- Ritter, T., Wilkinson, I.F. and Johnston, W.J. 2002, Measuring network competence: some international evidence. *Journal of Business and Industrial Marketing*. 17 (2/3) (2002) 119-138.
- Van der Aalst, W.M.P & Van Hee, K.M. 2001, *Workflow Management: Models, Methods, and Systems*. MIT Press
- Wasserman, S. and Faust, K. 1994, *Social Network Analysis: Methods and Applications*. Cambridge University Press, Cambridge
- Watkins, M. 2002, *Breakthrough Business Negotiation - A Toolbox for Managers*. Jossey-Bass.

ON THE VERSATILITY OF USE OF BUSINESS GAMES AT THE BEGINNING OF THE NEW CENTURY

Emmanuel Dion
Audencia, Nantes School of Management
8 route de la Joneliere - BP 31222
F-44300 NANTES, FRANCE
E-mail: edion@audencia.com

KEYWORDS

Business, Marketing, Gaming

ASBSTRACT

This short paper describes how business games, originally devised as pedagogical tools, have now gained a new status as public relation instruments, and could in the future develop as a powerful experimental method for the study of managerial behaviour

INTRODUCTION

Two reasons can easily be proposed to justify the growing power of business games:

- First, they develop at the crossroads of simulation and business, which sociologists and essayists have both identified as long-term rising trends for the future of western societies (See for instance Baudrillard (1985) and Lipovetsky (1989))
- Second, because of their digital nature, their capacity develop proportionally to data processing (and hence, on Moore's law), and are therefore not bounded by the cognitive limits of the human mind.

As a consequence of this rising power, business games now tend to be used in situations that do not correspond to their primary use. After their early development stage as a pedagogical tool and their extension as a public relation instrument, new ways of development can be foreseen in the area of management research.

THE FIRST STAGE: BUSINESS GAMES AS A PEDAGOGICAL TOOL

Research seems to teach us that learning activities that reproduce real work situations bring better transfer of learning (Swanson and Holton 1999). Aside from business management, aviation, civil emergency preparedness and medicine all use realistic scenarios to teach or improve complex skills. "When the cost of failure is high and when the performance arena uncertain, simulations are likely to be useful" (Steve Semler,

<http://www.learningsim.com/content/1snews/simdesign.html>,

http://www.learningsim.com/content/sim_enhanced_learning.pdf).

We have known games in general for long as pedagogical tools. As Wolfe explains it, the use of games as a teaching and training technology is hundreds of years old. "The Chinese in about 3,000 B.C. invented the war game Wei-Hai for entertainment purposes as was also the purpose of the Hindu game of Chaturanga. War gaming became serious in 1664 when Weikmann at Ulm developed the King's Game. War Chess by Helwig at the Court of Brunswick in Germany followed in 1780 with the most elaborate "New Kriegspiel" created by George Venturini at Schleswig in 1798".

(<http://www.swcollege.com/management/gbg/history.html>
<http://www.swcollege.com/management/gbg/gbg.html>).

After the first microeconomics simulation attempts by Marie Birshtein in the late-1920s, business gaming really started to develop in the United States in the mid-1950s. Various fields such as computer science, accounting business and operations research were called upon to bring to the emergence of the first business simulations at some of America's most prestigious universities. Wolfe reports the University of Washington as being the first to use a game, called the Top Management Decision Game, in their business policy courses in 1957. "Since that time the adoption of business games has steadily continued. It is now estimated that simulations are used in 97.5% of all AACSB schools in the United States, with the greatest penetration of this teaching method occurring in strategic management and marketing courses" (op. cit.).

In the 1960's mainframe computer systems already allowed the kind of simple calculation which is needed to compute vast arrays of additions, multiplications and ratios. This capacity was enough to edit the accounting and financial statements that constitute the necessary basis of all business games.

In the 1980's, business games migrated on personal computers, but generally remained in the field of management education. Throughout the world, teachers independently discovered their pedagogical power and simplicity of development, and a good number of the early fans started to work on their own software, therefore generating an atomistic and opaque market.

Since the late 1990's, following the main trend, business games have been brought online. This new window open on a wide public in its turn brought some changes. As it is the case in many sectors, one has been able to notice the similarity of many productions, and a lot of "copy/paste" has led to the convergence of some data, including some of the usual description of business games principles and features. However, most of the mathematical models used - the "blackboxes"- remain hidden.

THE NEW FASHION: BUSINESS GAMES AS A PUBLIC RELATION INSTRUMENT

In December 2003, a search on altavista on the phrases "business game" and "business games" brought respectively 14,440 and 19,150 hits.

This is more than 4% of the 150,920 and 634,785 hits observed on the very general terms "business school" and "business schools". From those figures, it seems fair to conclude that at the beginning of the new century, business games are on fashion. And a further look at the main websites on the topic shows that beside their original role as a pedagogical tool, business games now seem to gain a new function in our economy: that of an efficient public relation instrument between large companies and a target of young students who follow a business curriculum.

Following the track of Euromanager (35000 declared participants in 2002) organized since 1979 with a number of international partners such as Price Waterhouse Coopers or Carrefour, we have among others recently noticed on the European market the apparition of:

- E-strat proposed by L'oréal: 30000 participants registered in 2003
- The Global Marketing Game proposed by Spanish business school ESIC under the sponsorship of Efmd: 380 participating teams in 2003.
- Trust proposed by Danone: aiming at 150 participants in 2004

Those business simulation based competitions can be divided into two categories:

- Category A : Pure public relation competitions. Under the control of a unique sponsor, those competitions are characterized by free application for participants, and prizes linked to the sponsor's products or careers opportunities (examples: E-strat, Trust).
- Category B : Profit oriented competitions : Organized by an independent company trying to get multiple sponsorship, those competition are characterized by paying application and valuable prizes (examples: Euromanager, Global Marketing Game).

At the end of year 2003, there already seems to be a small decline in applications for those competitions, showing that the market still is in a process of a rapid change for the years to come. A possible outcome could lead to the

emergence of a small number of widely recognized leaders that would help structure the field.

FUTURE PROSPECTS : BUSINESS GAMES AS AN EXPERIMENTAL DEVICE FOR THE STUDY OF MANAGERIAL BEHAVIOUR

There still exists another possible use of business simulations that in our view has not be fully exploited yet: the experimental method.

Most of the current research in management relies on a fairly standard methodology that aims at comparing various models -always more sophisticated- to empirical data collected in real companies or markets. The experimental design is scarcely used, and whenever it is the case, experiments more often deal with fragmental aspects of consumer behaviour than with cognitive aspects of business decision making.

In our opinion, the width of spread of business games in management schools and the fact that they always imply, almost by definition, the setting of parameters by their author and the control of a number of variables by participants makes them an ideal basis for stimulus-response type experimental testing.

Here is an example of a method that could be used to assess the mimetic aspects of business decision making (We should distinguish here the sociological notion of mimetism popularised by Girard and the new approach proposed by researchers in the new field of memetics (Dawkins, Lynch, Blackmore among others):

Let us first imagine a standard business game for which the model computes -more or less smartly- a number of causal relations between the players' decisions (D) and the outcome of their choices in terms of results, performance or profit (R), which they try to maximize. Let players have an asynchronous access to others decisions and results.

Now what is foreseeable is that many players will tend to mimic the decisions of those who are successful with their results rather than try to understand by themselves the causal links between D and R. This kind of behaviour is frequently observed in actual runs of business simulations, and is usually part of the game. Unfortunately, in the absence of any control device, an outside observer will never be able to attribute clearly this convergence to mimetism or to an authentic improvement in the players understanding of the model.

However, it is experimentally possible to solve the problem. Let us divide D into two subgroups of decisions: DR and DF, DR representing decisions really taken into account by the model and DF being fake variables ignored by the program. If players are not aware of the sub-

grouping, all convergences observed about DF shall be attributed to mimetism.

Such an experimental design will be put in practice in the coming months and results will be presented in a future research paper.

REFERENCES

- Baudrillard, J. 1985. *Simulacres et Simulation*. Galilée, Paris.
- Cohen, K.J., & E. Rhenman 1964. "The role of management games in education and research." *Management Science*, No.7 (July), 131-166.
- Faria, A.J. 1998. "Business simulation games: Current usage levels-An update." *Simulation & Gaming*, No. 29 (3), 295-308.
- Keys, J.B. and J. Wolfe 1990. "The role of management games and simulations in education and research." *Yearly Review of Management*, No. 16(2), 307-336.
- Lipovetsky, G. 1989. *L'Ere du Vide*. Gallimard, Paris.
- Swanson, R.A. and E.F Holton 1999. *Results. How to assess performance, learning, and perceptions in organizations*. Berett-Koehler Publ. Inc., San Francisco.
- Wolfe, J. 1993. "A history of business games in English-speaking and post-Socialist countries: The origination and diffusion of a management education and development technology." *Simulation & Gaming*, 24 (4), 446-463.

AUTHOR BIOGRAPHY

After graduating in business management at HEC Paris in 1986, and a short professional experience in the mail order business, Emmanuel Dion has become a Marketing Professor at Audencia Nantes School of Management in 1989. He has obtained his doctorate in Marketing in 1995. His research interests have led him to develop the Global Simulation Lab, a research and development Center that has produced various business games extensively used in Audencia, especially in MSc and MBA programs.

SYSTEMIC DECISION SUPPORT IN A COMPLEX BUSINESS ENVIRONMENT

Steen Leleur
Decision Modelling Research Group
Centre for Traffic and Transport Research (CTT)
Build. 115, Technical University of Denmark
DK - 2800 Lyngby
Denmark
E-mail: sl@ctt.dtu.dk

KEYWORDS

Decision support systems, general systems theory, policy-making, forecasting, risk analysis

ABSTRACT

This paper presents a so-called systemic decision support system (SDSS) to support decision-making in a complex business environment. The approach can be applied by both private and public enterprises and organisations in their long-term strategic considerations. The approach consists of a framework that is adapted to the particular task by combining appropriate methodologies of different type and nature. Both “soft” and “hard” operations research (OR) methods are applied, for example, Critical Systems Heuristics (CSH) and multi-criteria analysis (MCA) making use of Monte Carlo simulation. As argued in the paper, rationality and optimisation as relating to traditional OR methods become difficult to vindicate and validate in complex decision environments, for which reason it is proposed instead by SDSS to assist the decision-makers by providing decision support based on the contrasting of different insights as successive steps in a wider exploration and learning cycle. The systemic approach is illustrated by a practical example concerning the Øresund Fixed Link. Although this concerns one particular problem type: complex public decision-making, the treatment of the example has made it evident that systemic decision support is well suited for complex strategic decision-making in both public and private settings. Finally some conclusions and a perspective are given.

1. INTRODUCTION

A characteristic of our emerging hypercomplex society is a still increasing degree of uncertainty and complexity in the various societal sectors (Leleur 2004). This makes strategic managerial decision-making in both private and public enterprises and organisations a risky undertaking raising thereby a need to develop types of decision support suited

for complex business environments. On this basis this paper proposes a so-called systemic decision support system (SDSS) approach that is also illustrated by an application example.

The paper is disposed so that Section 2 gives the main ideas behind the SDSS methodology. In this respect the paradigms of Simplicity and Complexity thinking as set out by the French philosopher of science Edgar Morin are presented and a basic framework is formulated. On this basis the following Section 3 sets out SDSS as a multimethodology approach combining “soft” and “hard” operations research (OR) methods to formulate an exploration and learning cycle that can guide and support analysts and decision-makers in the concrete application case. Next in Section 4 an application example is dealt with. It concerns the complex public decision whether or not to build the Øresund Fixed Link at a cost of EURO 3.2 billion for road and rail traffic between Copenhagen and Malmö. The case has been treated with emphasis on the general applicability of SDSS. Section 5 finally gives some conclusions and a perspective.

2. SIMPLICITY, COMPLEXITY AND SYSTEMIC DECISION-MAKING

In accordance with the English organisations- and complexity researcher Ralph Stacey, change processes in business organisations can be categorised into so-called closed change, contained change and open-ended change (Leleur 2004, pp. 16-17):

- Closed change: The key features of closed change are unambiguous problems, opportunities and issues, clear connections between cause and effect, and the possibility of accurately forecasting the consequences of change. Faced with such change, people tend to behave in easily understandable ways. The decision-maker can make use of rational decision-making techniques and the processes of control are formal, analytical and quantitative. There is a clear purpose with clear preferences, and alternative ways of achieving the purpose are known.

- Contained change: The key features of contained change derive from those change situations where it is possible to make probabilistic forecasts based on actions taken now and their most likely consequences. This is made possible because the consequences appear to some degree as repetitions of what has happened in the past or they relate to large numbers of essentially the same event. As a manager looks into the future, accurately predictable closed change declines in relative importance, while less reliably predictable contained change increases in relative importance.
- Open-ended change: Control in open-ended situations in practice means something completely different from what it means in closed and contained situations. In such situations, the future consequences are unknown and forecasting is totally impossible due to an ambiguous purpose and equivocal preferences of the actors involved. The whole situation being confronted is ill structured and accompanied by inadequate information more or less subjective and conditioned by personal ambitions, beliefs and values. There are problems with interpreting data and applying statistical techniques in uniquely uncertain conditions, for which reason forecasting and simulation become problematic.

Decision support systems (DSS) can in many cases be reasonably well specified and developed so they can facilitate managerial decision-making relating to closed and contained change, whereas open-ended change remains a challenge for several reasons. One basic consideration in this respect is that the uncertainties involved in complex decision-making are principally of a generic type that cannot be satisfactorily dealt with by detailing and refining the DSS methods that work well in situations with closed and contained change.

The application of systems science for improving our problem-solving capabilities holds two promises (Ibid., p. 22):

- By seeing our problem or study object as a system we may make use of the systems concepts to make a better representation of it and here capture (and model) various *interrelations among elements*, etc. in a more qualified way.
- By seeing our problem as a system we may be able to focus less on step-by-step approaches and capture more *holistic impressions* which can qualify our study.

The first statement concerns what is sometimes referred to as systems analysis. In an almost generic process we commence by defining our problem and determining the objectives. After this we turn to envisage or model the consequences of various alternatives being relevant. Then we appraise the alternatives to make it possible to select the best one. The final step concerns the implementation of this

alternative and possibly it may be decided to continue the process by monitoring it (Leleur 2000, p. 18). Our ideal in this undertaking is to be rational in our decision-making so the analytical processing of complete information will lead to an optimal result, be it a decision, design, plan, etc. We will see this as a *systematic approach*.

The second statement above relating to the use of systems science, expresses that wholeness matters and it can therefore be seen as a corrective to the first one. With the systematic approach we proceed in our problem-solving by using a step-by-step approach; with the *systemic approach* we are concerned with holistic views.

No doubt the systematic approach is tied to rational-analytical thinking well-known to people educated, for example, as engineers and economists whereas systemic as notion is more difficult to understand and come to grips with. In this situation the so-called paradigms about Simplicity and Complexity are highly relevant, where a paradigm denotes a specific research pattern. On this basis Morin sees classic scientific explanation as based on a Simplicity paradigm. Although he recognises the strength of the Simplicity paradigm in many respects, he also sees certain limitations in its explanatory models. As physics and cosmological thinking have always been major suppliers of ideas to other branches of science, it is interesting that subnuclear physics is the major example that Morin uses to explain the insufficiency of the Simplicity paradigm as it cannot satisfactorily explain new so-called exotic particles for example. Against this background Morin proposes a so-called Complexity paradigm to widen our research pattern by adopting principles that are *complementary* to those contained in the Simplicity paradigm. In Table 1 Morin's two paradigms are indicated by paired keywords (Leleur 2004, p. 24).

Table 1: The Two Paradigms about Simplicity and Complexity

Simplicity paradigm	Complexity paradigm
Universality	Multiplicity
Determinism	Organisation
Dependence	Autonomy
Necessity	Possibility
Lawfulness	Self-organisation
Prediction	Surprise
Separation	Wholeness
Identity	Individuality
The general	The particular
Objects	Subject
Elements	Interactions
Matter	Life
Quantity	Quality
Linear causality	Multi-causality
The automaton	Time
Objectivity	Culture

As concerns the development of systemic decision support we will see systemic thinking as rooted in the Complexity paradigm and systematic thinking as rooted in the Simplicity thinking. The idea is not to replace systematic thinking with systemic thinking but to make wider analysis possible by applying both. As the conventional approach is tied to systematic thinking we adopt the term systemic for such wider analysis which includes both systematic and systemic findings and not just the latter. Another basic complementary relationship behind SDSS concerns

scanning vs. assessment. Dealing with problem-solving, exploration and learning will depend on a kind of alternating between these two modes, i.e. they cannot be problematised at the same time but they will reciprocally influence each other. Cross-referencing these two pairs of complementary relationships we obtain the basic structure behind SDSS, see Figure 1 (Ibid., p. 127). In the figure different methods are indicated to illustrate some possible method choice, see Section 3.

SDSS structure	Systemic	Systematic
Scanning	Example: Critical systems heuristics	Example: Scenario analysis
Assessment	Example: Futures workshop	Example: Multi-criteria analysis & simulation

Figure 1: The SDSS Structure as Four Interrelated Modes of Exploration and Learning

3. THE SDSS APPROACH TO SYSTEMIC DECISION SUPPORT

The SDSS approach is developed by making use of the generic structure shown in Figure 1. Generally this is done by applying appropriate OR methods, see Table 2, in a self-

organising process that embeds conventional optimisation in a wider process of exploration and learning (Ibid., p. 35). The on-going search-learn-debate process moves on by contrasting and interpreting the different findings and insights. The process aims at converging into a satisfactory end result for the decision-makers.

Table 2: The OR Tool Box to be applied within the SDSS

Some current OR methods available for SDSS: Bold types indicate methods applied in the example	
Analytical hierarchy process (AHP) Computer-aided design (CAD) Conflict analysis Cost-benefit analysis (CBA) and cost-effectiveness analysis (CEA) Critical Systems Heuristics (CSH) Critical path method (CPM) Cross-impact analysis Decision analysis (DA) applying SMART and SMARTER Delphi conferencing techniques Environmental impact assessment (EIA) Expert systems Forecasting Futures workshop (FW) Fuzzy set theory Game theory Graph theory Input-output analysis	Interactive planning (IP) Intuitive exploration/brainstorming/metaphor and analogy building Linear programming techniques Multi-criteria analysis (MCA) Multiple perspectives (MP) Network theory Optimisation theory and heuristics Program evaluation and review techniques (PERT) Scenario building Sensitivity analysis Simulation Soft systems methodology (SSM) Statistics, probability and queueing theory Strengths, weaknesses, opportunities and threads analysis (SWOT) Systems dynamics Total systems intervention (TSI)

Generally “hard” OR methods can be seen to provide so-called first-order findings based on calculative rationality whereas second-order findings (or even higher) are associated with “soft” OR methods – based on communicative rationality - that relate to the so-called

subworld created around a complex problem by the various stakeholders and participants in the process (Ibid., pp. 72-73, p. 107). The wider process adapted with SDSS makes it possible to deal with the complex problem in a much more

explicit way. The example below describes some findings relating to the Øresund Fixed Link (Ibid., pp. 132-136).

4. OVERVIEW OF APPLICATION EXAMPLE

The Øresund Fixed Link – open since July 2000 – can be regarded as one of the most complex transport investment decisions made in Scandinavia (Rønnest et al. 1997; Leleur 2004). The case work demonstrates how the huge amount of information produced in studies, etc. over the years could have entered an ex-ante examination applying the SDSS approach. In this respect the case has functioned as a kind of evaluation research methodology laboratory (Decision Modelling Research Group 2004).

The OR methods made use of are, see Table 2: critical systems heuristics (CSH), scenario analysis (SA), futures workshop (FW), multi-criteria analysis (MCA) and simulation (SI). In brief the CSH mapped decision coalitions (“players”) and their motives and different responses at certain stages whereas SA and FW provided a set of interrelating framework- and trend scenarios. These were followed by MCA and SI.

An intermediate result in the exploration and learning cycle was for a central scenario (Scenario 5: Medium integration in the Øresund region & EU-wide regulative economy and sustainability) a total rate of return from the investment based on discounted assessment of all benefits and costs equal to 0.88 (95%-percentile), 1.01 (50%-percentile) and 1.15 (5%-percentile).

This led to various reconsiderations of decision criteria and uncertainties. The overall analysis shows, among other things, that feasibility including strategic concerns about mobility, trade, employment, etc. depends on follow-up issues by the different stakeholders that can ensure integration in the Øresund region in various ways: trade, employment, firms adopting better logistics, etc. (Ibid.). Figure 2 indicates issues that were made explicit by applying SDSS in the complex decision environment characterising the decision-making relating to the Øresund Fixed Link, while Figure 3 as an example presents the Monte Carlo simulation results as concerns the overall feasibility of the investment (narrow and wider issues) for the above mentioned scenario 5 (Ibid.).

Issues made explicit in the decision-making	
Systemic-scanning: Issues of identification and demarcation. General concerns: - Øresund region one of several spheres - The meaning of national barriers - Drivers: market, clusters, culture, etc. - Infrastructure and development Specific concerns: - Limitations of cause-effect model - Interpreting expressed expectations	Systematic-scanning: Issues relating to scenarios. Regional scenarios: - Economy, regulation, transport, etc. - Local integration vs. non-integration - Baltic Sea development: trade, etc. - Competitive transport development EU-wide scenarios: - Economy, regulation, transport - Trends, initiatives, technology
Systemic-assessment: Issues relating to stakeholder preferences. Ex-ante: - Local pro-coalition - Local environmental anti-coalition - National interest - International pro-coalition Ex-post: - National interest - Øresund region citizens - Øresund companies - International interest	Systematic-assessment: Issues relating to multi-attribute decision making. Narrow feasibility (cost-benefit analysis): - Investment - Time savings - Cost savings - Local environment & accidents Wider feasibility (multi-criteria analysis): - Network and mobility - Global emissions (CO₂) - Employment - Logistics and goods effects

Figure 2: Application of SDSS on the Øresund Fixed Link with Indications of Issues made Explicit

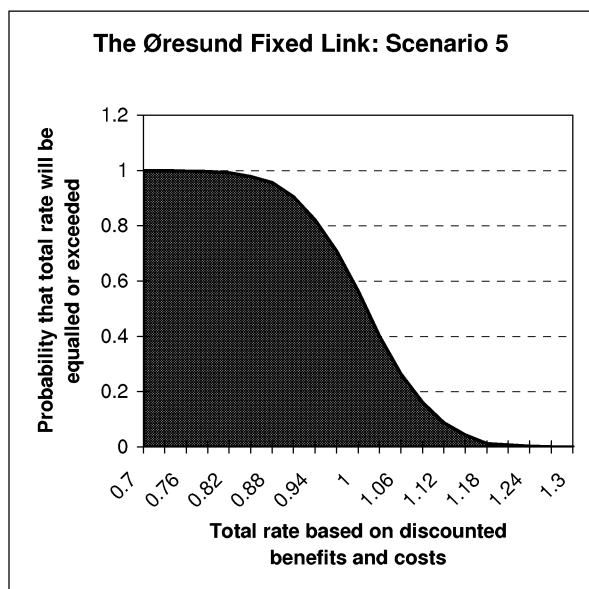


Figure 3: The Results of the Monte Carlo Simulation (SI) carried out as Part of the Systematic-assessment based on Multi-criteria Analysis (MCA) comprising also a Cost-benefit analysis (CBA) Component as One of Its Criteria

The total rate of return (in scenario 5 equal to 1.01) is made up by a cost-benefit analysis (CBA) fraction equal to 0.55 (examining the narrow feasibility) and a multi-criteria analysis (MCA) fraction (examining the wider feasibility) equal to 0.46. Thus going into depth with the MCA & SI results shows that societal feasibility (i.e. total rate > 1) depends on the benefits stemming from wider, strategic issues included into the analysis. Thus applying a traditional cost-benefit analysis concerned only with what is here termed the narrow issues, see Figure 2, would have failed to document the possible feasibility indicated here. Therefore the cost-benefit analysis would not have been appropriate as the sole decision guidance. Furthermore, the linkage of the wider concerns with interpretation of scenarios and stakeholder preferences on the background of reasonable study demarcation underlines the importance of organising complex decision-making by use of interrelated modes of exploration and learning. This is the basic idea of the SDSS approach drawing on both systematic and systemic insights.

5. CONCLUSIONS AND PERSPECTIVE

The limited space only allows an outlining of the major characteristics of SDSS together with an overview of some of the results from an application case. Even on this basis it can, however, be concluded that SDSS holds a potential as guidance for providing decision support of particular relevance for open-ended, complex business or societal

problems. The demand of such methodology with a still increasing complexity in society in general and in the professional spheres of administration and business more specifically, combined with the flexibility and adaptability of SDSS, makes it relevant to pursue a further development of the current SDSS principles, methodology and IT & software. Upcoming events and activities in this respect involving the Decision Modelling Research Group (DMRG) at the Technical University of Denmark is a book to be published about systemic planning and complexity (Leleur 2004) and a download section at the DMRG website www.ctt.dtu.dk/group/bmg with various relevant information. Furthermore, SDSS will be a topic area at an Evaluation Methodology conference arranged by DMRG in November 2004 in Copenhagen as part of the group's research activities in the Danish Centre for Logistics and Goods Transport (CLG).

REFERENCES

- Decision Modelling Research Group. 2004. "Main & Attachment Report, CLG Task 9", Danish Centre for Logistics and Goods Transport (CLG) at the Centre for Traffic and Transport (CTT), Technical University of Denmark. (Feb).
- Leleur, S. 2000. *Road Infrastructure Planning – A Decision-Oriented Approach*, Second Edition, Polyteknisk Forlag, Lyngby, Denmark.
- Leleur, S. (2004). *Systemic Planning*, book manuscript to be published, Centre for Traffic and Transport (CTT), Technical University of Denmark.
- Rønnest, A., Ohm, A. and Leleur, S. (1997). "The Øresund Fixed Link – The Conflicts and the Players", WP7 case report, TEN-ASSESS, The Interdisciplinary Centre for Comparative Research in the Social Sciences (ICCR), Vienna, EU 4th Framework Programme.

AUTHOR BIOGRAPHY

STEEN LELEUR is Professor of Decision Support Systems and Planning at the Centre for Traffic and Transport at the Technical University of Denmark. After having worked for the Danish Road Directorate as a civil servant he joined the Technical University of Denmark in 1986 to lecture and do research on transport related issues with a special concern of project appraisal and evaluation methodology. In addition to these topics he has also lectured at Copenhagen Business School in systems science and planning methodology. Currently he is, among other things, involved in research on topics about systems analysis and evaluation methodology in the Danish Centre for Logistics and Goods Transport. He has participated in a number of EU research projects, has published two textbooks about transport planning and is at present finishing a monograph about systems science and complexity as relating to new ways of planning and providing decision support when confronted with complex societal problems.

ACCESS ANALYSIS OF MONTHLY-CHARGED MOBILE CONTENT PROVISION

Hiroaki Oiso
Codetoys K.K.
Dojima Bldg. 5F, 2-6-8 Nishitenma, Kita-ku,
Osaka, 530-0047, Japan
E-mail: oiso@codetoys.com

Norihisa Komoda
Graduate School of Information Science and
Technology, Osaka University
2-1 Yamada-oka, Suita, 565-0871, Japan
E-mail: komoda@ist.osaka-u.ac.jp

KEYWORDS

Mobile content, Access analysis, Customer retention.

ABSTRACT

Customer retention is a critical challenge for mobile content providers. This paper presents an access analysis aiming at prolonging the subscription period of the users. First, the incentive system adopted by a quiz content site is evaluated based on the trend of users' subscription and gaming activity, through which the target segment for customer retention is identified. Secondly, some of the findings from the user survey were materialized through a change of the rules in the site, which decreased the number of unsubscribers and resulted in a significant increase of subscribers.

INTRODUCTION

The mobile content market, which has been expanding rapidly since around 1998, is said to have become a 300 billion-yen market. NTT DoCoMo's annual turnover from the content fee has exceeded 120 billion yen. There is no other country or region in the world other than Japan where the mobile content market has grown so eminently. We can think of several reasons but the price difference in particular is significant. In other words, (1) the combination of a fixed monthly content fee and a variable packet charge is reasonable for the users and (2) the commission charged by the carriers for collecting the content fee is very low at 10%, which prompts the participation of content providers. However, the users actually have to pay a large sum of packet charge and the share of monthly content fee of the total payment is extremely low. Also, the carriers only collect the monthly content fee as an allocation for the content provider. Content providers must therefore carry out an effective marketing to attract more users and prolong the subscription period per user. In this paper, we will make an access analysis for retaining

the users effectively from the viewpoint of the content provider.

The authors of this paper have been providing quiz game contents (hereinafter referred to as the content) on the official menus of three domestic carriers since 2002. An incentive has been set as part of the rules involving the game, for prolonging the subscription period in the monthly charged system. This paper (1) evaluates the effect of this incentive system through an access analysis and (2) conducts an analysis and survey of the prime reasons for unsubscribing, takes an action based on the findings and evaluates the effect.

OUTLINE OF THE GAME

Game Provision

The game is provided through three mobile carriers. It is provided in the Web system although there are minor differences according to each carrier.

Game Rules

The game is provided in the form of quiz, and one game consists of 15 questions. If the 1st question is answered correctly, the player will receive 10,000 points, and the points double every time the player gets a correct answer. It is possible to acquire 10,000,000 points if the player answers all 15 questions correctly. On the assumption that a number of players will be able to score full points, as it is easy to cheat through a mobile phone, we count not only the points but also the time required to answer the question into the final score.

Game Fee

The content fee is JPY 180 per month. The packet charge is about JPY 3.5 per 1 quiz = 1 page.

The maximum number of games that can be played per month is set at 150. For instance, if the average number of quiz achieved in one game is 8, a total of 10 pages,

including the page before and after the game, will be forwarded per game and therefore the packet charge will amount to about JPY 35. If the player plays the game 150 times a month, it will cost JPY 5,250 per month (JPY 5,430 including the content fee).

INCENTIVE FOR RETAINING LONG-TERM SUBSCRIBERS AND ITS EFFECT

Concept of Incentive Provision

In order to prompt more users to subscribe, we provided a prize to the leading scorers. The cost of the prize is limited by the Law for Preventing Unjustifiable Extra or Unexpected Benefit and Misleading Representation, to under 20 times the cost of the product. If we apply this law to the monthly fee of JPY 180, the limit cost of the prize will be JPY 3,600, which can hardly be considered as an attractive prize. Therefore, we increased the cost of the prize by adding up several months of fees and targeting only the continuous subscribers. Next, we decided how to select the winners – by competition or through random selections such as lottery. Considering the nature of the contents, it was decided that the competition system would suit better and that the incentive system should follow this principle. It was also decided that the ranking should be defined according to the score and the time required to answer the question correctly. Furthermore, four levels of ranking—1st stage (lowest stage) through 4th stage (highest stage)—were set up and the prizes were given to the leading scorers in the highest stage only. The ranking is announced at the end of each month and only the top 20% players of each stage can move up to the next stage. Therefore, it takes at least 3 months to move up to the 4th stage. According to a rough calculation, less than 1% of the users are able to move to the 4th stage in the 4th month.

Analysis of Users and Aspect of Each Stage

In figure 1, the number of users categorized by age and stage is distinguished by according to the length of subscription period (less or more than 4 months). The largest number of users belongs to the 1st stage, and the number of users declines toward the higher stage. If we look at the subscription period, users with the subscription period of 5 months or more are gathered in the 2nd stage and above. The share of subscribers in their 4th month account for less than 3% of the subscribers in the 4th stage. The largest age group is the users in their 20s, followed by teenagers, then those in their 30s. This is especially preeminent in the first stage. However, the age span rises

toward the higher stage, with teenage subscribers disappearing and age groups of 30s and 40s gaining ground. This is a distinctive feature that can be seen in quiz games where knowledge is essential for playing the game.

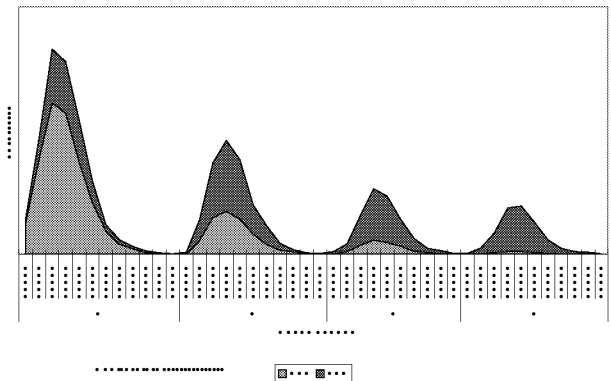


Figure 1: Stage and age distribution of subscriber

Game Count Pattern for Each Stage

Figure 2 indicates the number of games played per month by users in each stage. The game count for 1st and 2nd-stage users in their 1st and 2nd month of subscription is extremely high. Meanwhile, the number of users narrows down as the stage becomes higher and as the period of subscription prolongs. There are, however, heavy users in the higher stages. Therefore, for instance, although the number of users in stage 1 is 10 times that of stage 4, the game count in stage 1 is only 5 times higher than the count in stage 4.

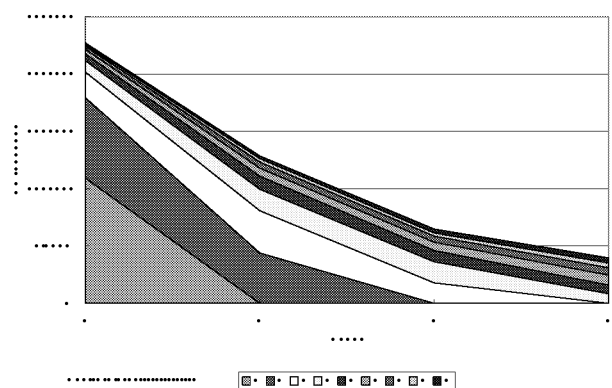


Figure 2: Total game count per stage

Let us look at the game count for each stage. Figure 3 shows the subscription period and the monthly game count of each user with the horizontal axis while distinguishing the game count of users in different stages by colors. As the subscription period becomes longer, the absolute number of users declines. Meanwhile, the total game count increases toward the right side within each subscription period. Therefore, the mountainous curve slanted toward

the right indicates that there are many heavy users that play until the limit in the corresponding subscription period. There are especially many heavy users among the subscribers in their 2nd to 4th month of subscription. Meanwhile, the mountainous curve of the first subscription month is slanted toward the left, meaning that a large number of light users are raising the game count. Also, as the subscription period lengthens, the structural ratios of the upper stages become high. This means that there are dropouts and unsubscribers along the way. From these facts, the following four types of user segment can be acknowledged.

- 1) Quiz mania (competitive users aiming for prizes)
- 2) Enthusiastic/unsubscribers (play many games but give up after 3-4 months)
- 3) Unconstraint users playing to kill time (play less game but continue for long period of time)
- 4) Short-term members (quit without playing much)

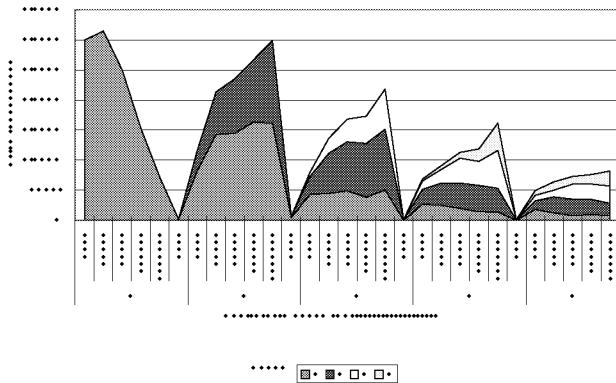


Figure 3: Total game count per subscription period and monthly game count of each user

Players that become one of the leading players in the 4th stage are limited to game manias (refer to 3.4). Players that discover the presence of these manias are the ones that unsubscribe in the 3rd and 4th stage. On the other hand, there are players that remain in the 1st-2nd stages and continue to subscribe for a long period of time—those who we refer to as unconstraint players. Lastly, there are the short-term members who play only a little and unsubscribe within one month. This type of players occupies the most share, and if we can prompt them to become enthusiastic players, the effect would be extremely high.

Evaluation of Incentive System

Figure 4 shows the ratio of (1) average subscription months, (2) game count, (3) age of users, (4) monthly unsubscription rate and (5) number of users in each stage. The values for the 1st stage are set to 100. The curves represent (1) through (5) from top going down. In the 4th

stage, the subscription period (1) increases by 3 times, the game count (2) rises by two times and the unsubscription rate (4) declines by half, indicating the soundness of the heavy users. The age of the players (3) is also about 20% higher in stage 4.

We found out that the trend of (1), (2) and (3) in the higher stages prompts competitive users as we had initially aimed, while contributing to retaining long-term subscriptions. This incentive system can therefore be basically considered as effective.

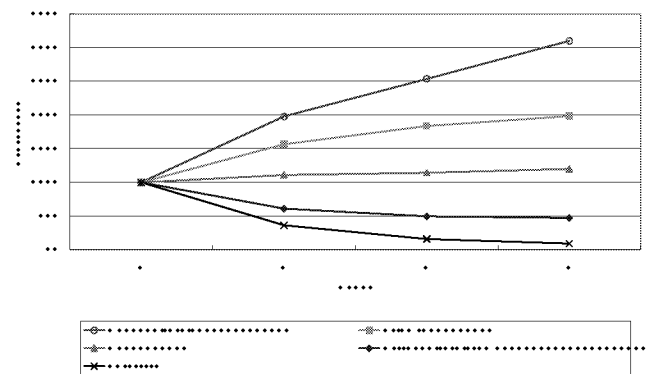


Figure 4: Comparison of stages

If we can increase the ratio of users in the higher stages in the future, we can expect a favorable influence in the overall subscription period. In other words, it will be extremely effective if we can develop a system that can maintain these favorable trends while gradually easing the conditions for players to move up to the next stage in order to increase the ratio of users in the upper stages. To achieve this, we must reduce the number of unsubscription during the 1st-3rd month of subscription in the 1st and 2nd stages that include the largest number of subscribers (to be explained later). We must therefore ease the conditions limiting the players to move upward or add a random ingredient to the incentive. For instance, enabling 1st and 2nd-stage players to move up by good effort or luck even if they do not have the actual ability to move up, or providing a system where some kind of prizes can be obtained in the lower stages as well. However, this will have an adverse effect on the users, although minor in number, that have already moved up to the higher stages. We must therefore make changes in gradual phases.

Adverse Effect of Competition

We have so far explored the quantitative issues but there is another issue that must be considered—the prizes are repeatedly won by certain quiz manias. These quiz manias that move up the stages in a short period of time by answering all the questions correctly in a surprising fast speed, occupy the top rankings, and obtain the prizes most

of the times. As we have chosen to give out the prizes through a competitive system, this is somewhat unavoidable. But it is our constant source of worries as there are many cases where general users unsubscribe when they feel that ‘they cannot compete with the manias (to be explained later).

FACTOR ANALYSIS AND COUNTERMEASURES OF UNSUBSCRIPTION

In order to clarify the cause and reduce the number of unsubscription, we examined the possibility of foreseeing future unsubscription by analyzing the quantitative information such as the shift in game counts. Next, we conducted a questionnaire survey to the unsubscribers and carried out the most effective countermeasures that can be conceived from the result and evaluated its effect.

Quantitative Evaluation of Unsubscription Reasons

By referring to the Churn Analysis Model [Berry] of mobile carriers, we checked if it is possible to predict future unsubscription by analyzing information such as the decline of game count and subscription period.

We tried data mining methods such as the decision tree, but could not obtain a sufficient result and therefore an exploratory data analysis was utilized.

The game count for n-2 month and n-1 month of members at a certain point of time (referred to as n month) were extracted per user, and the difference in the game count of players that unsubscribe in n month or n+ 1 month from that of other players were compared. The analysis carried out in a certain month in 2002 is shown in figure 5.

During the first 3 months of subscription (horizontal axis), the game count increases from the previous month, possibly because none of the players have reached the 4th stage. However, when they reach the 4th month, the game count shifts to a downward trend as the players begin to realize the difficulty of moving up to the 4th stage. The game count continues to decline gradually although the decline rate narrows down after the 4th month. In the months after the 5th month, the decline rate of the game count between the unsubscribers and the others begin to take a different track. Although there is an error of about 10 games due to the fluctuation that occur when changing the n month, it can be considered as one of the indexes that indicate unsubscription in the near future.

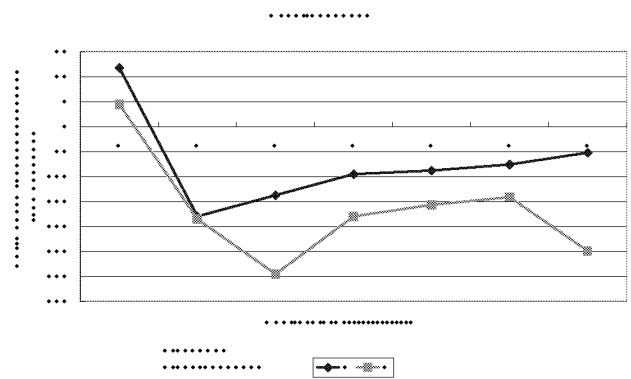


Figure 5: Bigger drop of monthly game count implying higher rate of unsubscription.

Questionnaire Survey of Unsubscribers

It is possible to insert a questionnaire in a series of Web screens that the users go through during the unsubscription procedures. The reason for unsubscription was asked with following multiple choices and a field for entering the text was provided in the questionnaire:

- A) Do not use it anymore
- B) Too few prizes
- C) Unable to win
- D) Discontent with the quiz
- E) Bored with the quiz
- F) Game count limit (5 per day) is too few
- G) Too few Fastest Finger quiz
- H) Have a hard time with the connection
- I) Content fee (JPY 180) is too expensive
- J) Packet charge is too expensive
- K) Poor support service

As a result, four reasons—A, C, F and J—were chosen the most. As mentioned in 3.3, we can expect to achieve the biggest effect by reducing the number of 1st and 2nd-stage unsubscribers in their 1st to 3rd month of subscription. Stratification of the result by focusing on this point is shown in figure 6. The prime reason for users unsubscribing within their 1st month of subscription was “F) Game count limit is too few”, but unsubscribers in their 2nd and 3rd month of subscription indicated “J) Packet charge is too expensive” and “A) Do not use it anymore” as their prime reasons. First-month users tend to have the desire to play the game repeatedly because their scores are still low and complain about the game count limit, but it seems that the complaint soon shifts to the expensive packet charge.

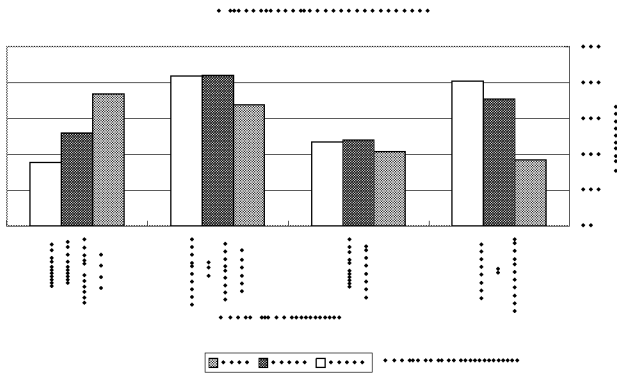


Figure 6: Survey result

As for the limit indicated in the “F) Game count limit is too few”, the limit was set to prompt long-term subscription and so that the players do not get bored with the questions.

Countermeasures and Its Result

The game count limit was changed from 5 plays per day to 150 plays per month. Also, the percentage of players that can move up to the next stage was eased from 20% to 25%. Figure 7 shows the number of subscribers and unsubscribers during 2 months before and 1 month after the change. The number of unsubscribers was reduced to half (over the preceding month) immediately after the change, which possibly had the influence on the increase of subscribers. We believe that the news spread by word of mouth. The number of game count also surged, increasing by as much as 3.5 times (over the previous day) on the day after the change was made. The surge hit the peak in the middle of the month but the cause is still unclear.

When comparing the averages of the two months before the change and the month following the change, the number of subscribers increased by 15% while the number of unsubscribers declined by 10%, increasing the net number of members by 2.83 times. If we look at the result according to each stage, 70% of unsubscription decrease was seen in the 1st stage and 20% in the 2nd stage, which was close to our initial target.

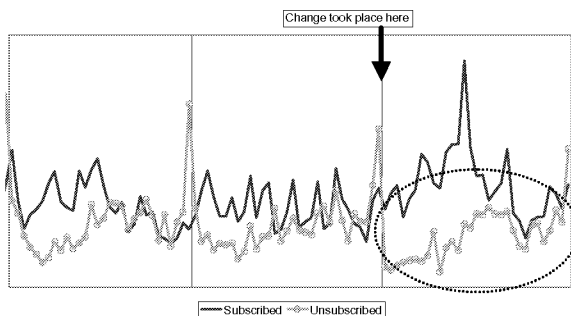


Figure 7: Daily subscribed / unsubscribed

CONCLUSION

As a result of analyzing the incentive system—which we had set up to increase the subscription period—from the attribution of users and from the access trend, we were able to find out that it had an effect of securing a stable heavy-user group and also contributing to prolonging the subscription period. We also discovered that we must target the “short-term unsubscribers”. When we analyzed the main factors for unsubscription, it became clear that a decline in the game count after 5 months of subscription indicates the possibility of future unsubscription. Lastly, as the result of taking actions based on the prime unsubscription factors revealed from conducting the unsubscription questionnaire, we were able to achieve a dramatic effect of 10% decline in the number of unsubscribers over the precedent month.

REFERENCES

M.J.A. Berry and G. Linoff: *Mastering Data Mining*, John Wiley, New York (2000)

SIMULATION MODELS

PERFORMANCE EVALUATION MODEL FOR E-PROCUREMENT SYSTEMS

G.R.Gangadharan
IMIT Class of 2004
Scuola Superiore Sant'Anna
Viale Rinaldo Piaggio, 34 - 56025
Pontedera (PI), Italy
gangadharan_gr@yahoo.com

Rahul Monhot
PGDIE Class of 2004
National Inst. of Indl. Engg. (NITIE)
Vihar Lake, Mumbai
India – 400 087
rahulmohnot@yahoo.com

KEYWORDS

E-Procurement, Performance analysis

ABSTRACT

The basis for the failure of e-procurement implementations lies in the negligence of enterprises in observing the basic principles of measuring business and implementation success. Measuring the success of the project demonstrates the benefits achieved by the business. This paper elaborates the need for measurement and principles of measurement in an e-procurement implementation. Further, this paper describes a case study on performance evaluation for e-procurement systems.

WHY MEASURE?

The advent of e-business and the consequent overhaul of internal procurement and supply chain mechanisms will require enterprises to re-evaluate their measurement criteria. The principal driver for enterprises to adopt e-procurement solutions is the provision of advanced reporting facilities with the aggregation of data in a central location. But the failures of e-procurement implementations arise due to the failure of the enterprises in observing the basic principles of measuring, business and implementation success. In this new and more transparent economy, an enterprise can only exploit comparative information if it has precise and timely information on its own performance.

Measurement is the process (Fenton and Pfleeger 1997) by which numbers or symbols are assigned to attributes of entities in the real world in such a way as to describe them according to clearly defined values and is defines as a mapping from the empirical world to the formal, relational world. Measuring performance is an important role in management. But selection of consistent, acceptable, and meaningful measurements is difficult. Measurement should be viewed as an infrastructure technology (Jose and Joan 2001), which is necessary to achieve systematic improvement.

WHAT SHOULD BE MEASURED?

“E-Procurement projects encompass process change, system change and cultural change, i.e. a very wide and demanding brief, especially for the Project Manager whose

responsibility it is to deliver on each of these three fronts.”
Richard Hounsfield, Biomni

E-Procurement is a web-based procurement process from collection of information through interaction between supplier and buyer to integration of the supply chains of supplier and buyer, and to participation in new business models and markets. Implementation of e-procurement automates the internal and external processes associated with the procurement processes. The success or failure of an e-procurement implementation depends on the understanding and the definition of the goal of the project and the strategies of measurement. The successful e-procurement implementation links directly to the business processes and needs (Stuart and John 2001).

The e-procurement implementation is considered as a success if the project is delivered on time, within budget and the intended business benefits have been achieved. The guidelines for e-procurement implementation are listed below:

- Targeting the areas, where the highest proportion of external expenditure occurs and where, rapid savings can be made.
- Base potential savings at the lower end of the estimates quoted above.
- Considering adoption of new processes rather than simply e-enabling existing ones.
- Allowing sufficient time and effort for re-definition of procurement business processes and rules
- Understanding of emerging standards in content and transaction management.
- The formation of product and services supplier catalogue content may involve bringing in new preferred suppliers accessible via the system. Allow time and effort for negotiations and bringing them on-line.

Measuring the success of the project demonstrates the benefits achieved by the business. The monitoring plan must make use of all possible sources of information to demonstrate achievement of benefits. Measurement criteria might include (BuyIT 2002):

- a target number of users actively using the system;
- a target number of orders processed through the system;
- a target percentage of spend processed through the system;
- the number of suppliers participating;
- a specified reduction in the number of order returns;

- a specified level of data capture for management reporting;
- overall user satisfaction.

The measurement criteria must be both qualitative (user perception questionnaires, supplier perception questionnaires, analysis of help desk calls), and quantitative (reports showing number of orders, value, suppliers, error rate, speed of payment). It must be recorded within a given period of time.

The steps for tracking and monitoring benefits are as follows:

- Clear knowledge of what to measure and the reasons for measuring.
- Mapping expected benefits with Key Performance Indicators (KPIs).
- Keeping KPIs simple, easy-to-measure and correctly defined.
- Considering 'soft' factors.
- Agreeing on the measurement process and measurement baseline.
- Getting sponsors buy-in and sign-off.
- Visibly measure and visibly report.
- Appointing a third party to track and monitor the progress.

MEASURING IMPLEMENTATION PERFORMANCE

The performance of e-procurement implementation is measured by e-procurement scorecard and by Benchmarking. Fields on the scorecard contain E-Procurement transactions to date, Savings recorded to date, Costs to date, Suppliers adopted to date, Number of end users, Improvement in payment terms with key suppliers, Improvement in accounting processes over time and e-tool usage (for e-auctions and e-tenders). The Scorecard should clearly specify what we are measuring and its purpose. It should link KPIs to the benefits expected from e-procurement.

Benchmarking gives a comparison of realised benefits versus targeted benefits. This also decides return on investment (ROI) from business case for the given period and determines progress towards target. Also, Benchmarking analyses integration level of e-procurement with enterprise resources planning (ERP) and other systems.

CASE STUDY ON PERFORMANCE EVALUATION

In line with its strategy of aligning capabilities to meet emerging trends, ABC company recently initiated a mega-transformation process internally to ensure that it emerges as a knowledge-based multinational. The Engineering/Construction (E/C) division of ABC constitutes the largest business segment of the company providing premier services in the areas of process technology, basic and detailed engineering, heavy engineering and skid mounted modular fabrication, procurement, logistics, construction, erection, commissioning and project management.

To thrive in the e-economy, E/C is capitalizing on the strategic advantages that come from speed and information

and are taking steps to address all issues related to e-business. E/C division took a significant step towards its transformation into a clicks and mortar organisation with the launch of its portal, incorporating the first phase of e-procurement initiative. This e-procurement feature enables to carry out procurement transactions with its business partners (suppliers) online using SAP R/3 (E/C Division's ERP system) as the backend. It enables authorized suppliers to view the RFQs along with attachments, submit their offers along with attachments, access purchase order and give order acceptance remotely using this website. All the transactions being ID based, can be viewed and accessed only by the users authorised for the same. This site incorporates robust security features, which ensures confidentiality of data. This launch into the e-procurement era brings in direct benefits to the enterprise and its suppliers who no longer have to handle voluminous documents, make multiple photocopies, experience postal delays and at the same time reduce their overheads. All they need to do is Sign up for e-procurement and start working.

In order to have clear visibility of the aims, objectives, scope and boundaries, Goals/Questions/Metrics (GQM) method (Solingen and Berghout 1999) is adopted for implementing e-procurement systems.

The objectives of measuring e-procurement system are listed as follows:

- Measuring recurring e-procurement activities
- Identifying KPIs
- Helping in the identification of the best practices
- Analyzing performance versus expectations

E-Procurement brings specific risks arising from:

- Dynamic changes in e-procurement technology.
- Consolidation of system providers.
- Dependence of supplier participation.
- IT security implications of linking an external system to the existing systems.
- Dispersed user base across diverse business areas and/or geographical locations.
- Infrequent users.

The following activities are measured:

- Deployment period
- Available user accounts
- Number of users who submitted a purchase transaction during report period
- Number of available suppliers
- Number of suppliers who received a purchase transaction
- Purchase transactions value
- eCatalog spend (percent of total spend)
- Number of purchase transactions released to suppliers
- Average time to approve purchase transactions
- Average number of approvals

A 'Time and Motion' study is conducted within the business to determine existing process activity times to determine transactional benefits monthly.

By estimating levels of maverick spend prior to implementing the e-procurement system, compliance benefit is measured. Once the e-procurement process has been implemented, these measures should be reported monthly. An improvement in these metrics should be seen over time as the e-procurement system becomes fully deployed and accepted throughout the organization.

With the baseline of management information, customer satisfaction and vendor satisfaction levels prior to implementation of e-procurement, management information system benefits are captured periodically occurring as a result of having the procurement scorecard, spend analysis and other e-Intelligence tools available.

To capture any price reductions (price benefits) given by suppliers in contract negotiations as a direct result of the supplier and the business benefiting from e-procurement, pricing offered prior to implementation of e-procurement is used as baseline and measured whenever guaranteed volumes or prompt payment can be used as a leverage in supplier negotiations.

Upon implementation of the e-procurement system, payment benefits are measured on a monthly basis to capture any price reductions given by suppliers in contract negotiations as a direct result of the supplier benefiting from e-procurement and receiving prompt payment of their accounts.

The significant progress achieved by e-procurement implementation and performance evaluated to date are as follows:

- Establishment of preferred suppliers across the enterprise
- Unified platform and process for workflow and approvals
- Increased compliance rate with preferred suppliers
- Improved visibility into expenditure details
- Enhanced focus for purchasing professionals on sourcing specialized items
- Capturing additional spend for future sourcing efforts

The progress is reflected in some common implementation metrics:

- Number of invoices reduced by 1.5%
- Spend capture increased 150% each of past three quarters
- Suppliers available: up 50x since rollout
- Catalogs available: up 3x since rollout

CONCLUSIONS

An enterprise's e-procurement strategy should be conceived as an integrated part of the overall procurement strategy to ensure the business objectives are met. E-Procurement projects require a balanced mix of skills in technical,

business and project management. The enterprise should start out on e-procurement by examining its existing procurement activities in terms of process and spend profile. Knowing in detail how it operates, how it could or should operate and the rate at which it can adapt is absolutely key to the selection and effective implementation of the appropriate solution(s). Achieving the maximum benefit to an enterprise comes from putting the right tools in the right hands to use in the right set of circumstances. To assess the required measurements accurately, information systems are needed to link the various supply chain components that enable the user effectively to collate and analyze information from their e-procurement and other systems.

REFERENCES

- BuyIT e-Procurement Guideline 2002. "Measuring the Benefits: What to measure and how to measure it : Executive Overview", The BuyIT Best Practice Network, 2002.
- Fenton N., Pfleeger S. 1997. "Software Metrics - A rigorous and Practical Approach", PWS Publishing company, 1997.
- Jose M. Esteves, Joan A. Pastor 2001. "Goals/Questions/Metrics Method and SAP Implementation Projects", Technical Research Report, Universitat Politecnica De Catalunya, 2001.
- Solingen R., Berghout E. 1999. "The Goal/Question/Metric Method: a Practical Guide for Quality Improvement of Software Development", McGraw-Hill, 1999.
- Stuart Dodds and John Balchin 2001. "E-Procurement Measurement: It's Not Broken – but It Needs to Be Fixed", Ascet Volume 3, 2001.

BIOGRAPHY

G. R. GANGADHARAN holds M.Sc. degree in Computer Science from Manonmaniam Sundaranar University, Thirunelveli, India. Having 4 years of experience in software industry, his areas of interest include Supply Chain Management, Information Systems and Internet Technologies. He is currently pursuing International Master in Information Technology in Scuola Superiore Sant'Anna, Pisa, Italy.

RAHUL MONHOT is currently a student of National Institute of Industrial Engineering (NITIE), Mumbai, India, pursuing Postgraduate Diploma in Industrial Engineering.

AN INTEGRATED TOOL TO GENERATE AND SELECT CONFIGURATIONS OF SIMULATION MODELS

R. Van Loock, H. Pastijn, F. Van Utterbeeck

Royal Military Academy
Renaissance Avenue, 30
B-1000 Brussels
Belgium
Email: Raf.Van.Loock@rma.ac.be

KEYWORDS

Discrete event simulation, system configuration comparison, outranking methods, multicriteria decision making, performance measures.

ABSTRACT

We created a set of tools helping to assist users with little or no simulation experience using discrete event simulation. These tools can also be useful instruments for training or for educational purposes.

We demonstrate the tools by using them on a simplified model of a call centre. We create different configurations of the model and try to improve them. We look at results obtained from the simulation runs and compare them. In order to select the best model among this finite set of alternative configurations, we use a variant of the Promethee method based on interval evaluations.

INTRODUCTION

We begin with a general description of the simulation model used throughout the paper. It is a model of a call centre which we have implemented in ARENA[®]. Then we describe some performance measures which are also implemented in the model that can be used in the selection of different configurations.

Then we give an overview of the tools written to simplify a simulation study in creating different sets of scenarios. The different available components are: the description of the model itself, which is the input, the simulation run with animation added to it to obtain more insight into the functional components of the model, and the output generated from the simulation runs.

Then we give a more detailed description of a particular stochastic simulation model with values for all the parameters used. We create a set a scenarios where we try to improve the system. This is done by comparing the results of the previous scenarios and removing the bottlenecks. Comparison and improvement are based on four performance measures mentioned in the next section.

The ranking of different configurations to select the best configuration based on a finite set of criteria is a typical multicriteria decision making problem. We use the

Promethee-i method with some of the described performance measures taken as the criteria.

INCIDENT MANAGEMENT MODEL FOR A CALL CENTRE

The simulation model we consider is the incident management process of a call centre. We will describe the elements of the model (the incidents, the resources and the skills matrix) and the process flow (Figure 1). (Van Loock et al., 2003)

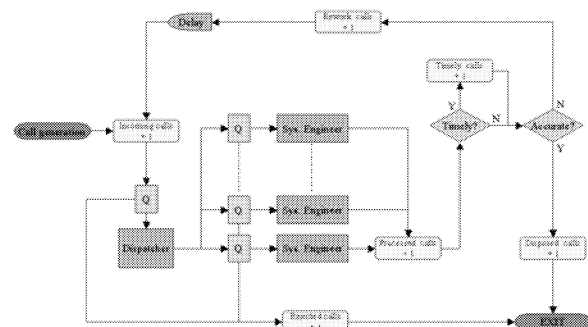


Figure 1: Call Centre Process Flow

The *incidents* are initiated by the customers of the call centre. These incidents are represented by the calls that arrive at the centre. These *incoming calls* follow a stochastic arrival pattern. The calls are subdivided into categories and subcategories, depending on the area of expertise required by the customer. Each category has a specified probability of occurrence, while the subcategories within a certain category are assumed to be equiprobable.

The *resources* in our model are the dispatcher(s) and the system engineers. Each resource has its own weekly working schedule, an hourly cost (based on the number of skills known) and a FIFO queue associated with it. Incoming calls will wait in the FIFO queue if the resource is busy. A call will be *rejected* (and leaves the system immediately) if the time spent waiting in a FIFO queue exceeds a certain fixed threshold.

Every system engineer has his own areas of expertise, which are specified in the *skills matrix*. Every line in the

matrix represents a subcategory, while every column represents a system engineer.

The *process flow* used in our model can be summarized as follows. Every incoming call must pass through a dispatcher. The dispatcher will rout the call to a system engineer whose area of expertise covers the category and subcategory of the call. If multiple system engineers are eligible, the dispatcher will rout the call to the resource with the shortest queue. Ties are broken in favour of the resource located the most to the left in the skills matrix. The *dispatching time* (the time needed by the dispatcher to decide on the routing of the call) follows a stochastic distribution.

The *processing time* (the time the system engineer needs to handle a call) follows a stochastic distribution, regardless of the subcategory. For every *processed call*, there is a fixed probability that the customer is not completely satisfied with the assistance provided. These customers will call back after a stochastic delay. These subsequent *rework calls* will result in a decrease in the performance of the call centre. If the customer is satisfied with the assistance provided, the call is *disposed* and leaves the system.

PERFORMANCE MEASURES

We use four performance measures: productivity, service level, waiting times in queues, and system cost. The productivity is the average of the resource utilisation, which is the busy time divided by the available time, of all the resources. Service level is expressed as the percentage of arriving calls which are finally disposed after a successful handling by the available resources (and as a consequence were not ejected from the system). The waiting time in queues is the average time that a call is waiting in the dispatcher and system engineers queues. The cost of a system engineer depends on his degree of polyvalence (number of skills). The overall system cost is a stochastic entity due to the fact that the resources continue to work at the end of their daily schedule until all calls waiting in their queue at the end of the working day have been processed.

TOOLS

Output analysis and generation of alternatives of simulation models requires knowledge of simulation in general and of the software used to run the model (which is ARENA® in our case). For users without experience in simulation studies we created an integrated tool to assist them in the input and output analysis. ARENA® supports VBA code. We built a tool via VBA.

Input

Comparing different configuration requires input data that describes the configurations. The input can be fed into the model by input forms created via VBA as shown in figure 2. These forms can be completely customized in VBA so

they give access to all relevant parameters and to the number of workstations in the model.

Figure 2: Example of Model Input Forms

Another input option is via Excel. All the input parameters for one scenario are contained in one excel sheet. Several sheets for the simulation of more scenarios are possible. The VBA code can read this and runs the simulation for all the scenarios. The format used for the excel sheet for our model is given in figure 3. The arrows indicate where columns or rows can be added. This format depends of course on the choices made by the developer of the code.

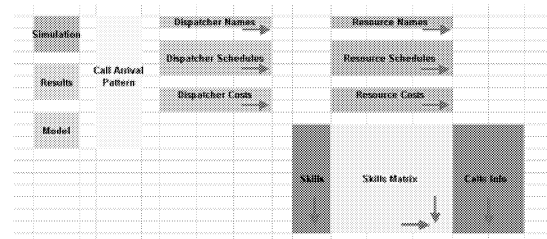


Figure 3: Structure of Excel Sheet of Input of Model

Simulation Animation

There are two options for the simulation run: with or without animation. The runs with the animation can give more insight into the model under consideration. Animation is possible with ARENA® and with Excel which we will show now.

Figure 4 gives an example of a scenario with 1 dispatcher and 5 system engineers. The queues are visible with some calls waiting in it. Counters are added to show the number of different types of calls.

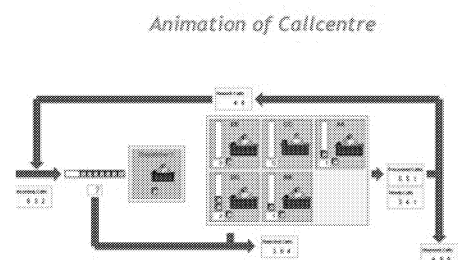


Figure 4: Animation of the Call Centre During the Run

Figure 5 shows an overall description of the run within Arena. We can see the current simulation time, the total replication length, the current replication, the total number of replications and the number of the calls of a given type already handled.

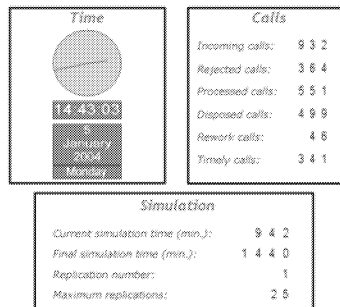


Figure 5: Simulation Time and Call Types Made

Figure 6 plots the evolution of a few measures over time. The top graph shows the total number of calls in the queues and the bottom one the total number of resources busy. The horizontal axis shows the time in minutes. On the right are some numbers that summarize the graphs on the left. From this figure it is clear that there are big fluctuations in these measures.

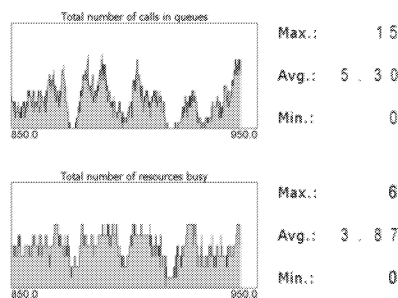


Figure 6: Evolution of Some Measures Over Time

Animation in Excel can happen in parallel with the animation in ARENA®. Figure 7 gives an example of this. Data is immediately exported to Excel and animated graphs can be shown. This can be very helpful because one is not constrained to the animation capabilities available in the ARENA® software. In the figure the scheduled utilization (busy time divided by the available time) is shown for all the resources in the model given. The dark colour of the diagram shows the actual number. The light colour represents the idle time for the resource.

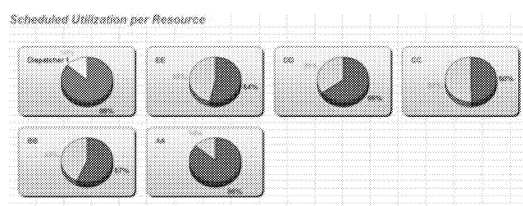


Figure 7: Simulation Animation in Excel

Output

After the simulation run output is generated in Excel. This is one of the advantages of the open architecture of Arena. The reports can be completely customized as a function of the particular model and only data one wants to know are shown. This gives the user a clear overview of the simulation results. Figure 8 shows an example of simulation results written in an Excel file. Results obtained from every replication are also possible and graphs that visualise the data will be shown in a next section.

Simulation Results				
Performance measures				
	Average	Half with	Minimum	Maximum
Timeliness	59.33 %	0.75 %	54.77 %	62.66 %
Productivity	65.83 %	0.58 %	63.41 %	69.26 %
Efficiency	97.09 %	0.18 %	96.12 %	97.88 %
Service Level	61.25 %	8.74 %	67.29 %	64.66 %
Avg. Queue-Waiting Time	2' 08"	0' 02"	1' 52"	2' 17"
Customer Queue	0' 57"	0' 02"	0' 48"	1' 06"
Resource Queue	1' 08"	0' 01"	1' 04"	1' 12"
Financial				
	Average	Half with	Minimum	Maximum
Sales	2984.12	27.60	2888.00	3112.00
Costs	1908.12	6.81	1907.13	1912.33
Profit	1682.00	27.72	1880.38	1802.04
Calls				
	Average	Half with	Minimum	Maximum
Incoming Calls	1267.00	15.00	1211	1393
Rejected Calls	499.16	14.18	443	568
Processed Calls	767.84	7.84	702	828
Rework Calls	80.12	0.68	63	102
Disposed Calls	707.72	6.76	682	738
Timely Calls	487.62	8.10	425	566

Figure 8: Simulation Results in Excel

DESCRIPTION OF CALL CENTRE MODEL

The call centre is open for 8 hours per day with a 1 hour pause during lunch break. Incoming calls arrive at the centre with an exponential distribution with a mean of 150 calls per hour.

System engineers and dispatchers work only during the opening hours of the centre. The cost associated with every resource depends on his expertise. It was chosen to be 25, 38, 50 per hour for a system engineer with 1, 2 or 3 skill(s) respectively. Dispatchers get 25 per hour and are thus equal to a system engineer with 1 skill.

The calls are divided into 2 categories with 3 skills in each category and follow an equiprobable distribution. The duration of a call follows an exponential distribution with a mean of 20 seconds for the dispatching and 2 minutes for the processing calls. Calls that are longer than 3 minutes in a queue are rejected from the system.

The probability that the customers are satisfied with the response of a system engineer (i.e. the accuracy) is 90%. Customers pay a different amount for satisfied calls (4) and unsatisfied calls (2). These unsatisfied calls result in rework calls with a delay which follows an exponential distribution with a mean of 30 minutes.

The replication length is 1 day and we took 25 replications for every configuration.

CONFIGURATIONS

We keep the number of changes from one scenario to the other to a minimum. This is done to clearly see improvements in the system without dramatic changes to it. We only change the number of dispatchers and system engineers and the skills matrix. The rest remains constant.

For the first scenario we have 1 dispatcher and 5 system engineers with a skills matrix shown in table 1. The system engineers are given in the top row with names AA to EE. The 6 possible skills are in the left column. We can see that resource AA expertise consists of skill 1 and 6.

Table 1: Skills Matrix of Scenario 1

	EE	DD	CC	BB	AA
Skill 1					X
Skill 2	X				
Skill 3		X			
Skill 4			X		
Skill 5				X	
Skill 6					X

This case is certainly not optimal because there is one resource with 2 skills while the others have only 1. Some results are given next. The queue waiting time of resource AA (1.6 minutes) is longer compared to the other resources (1 minute on average). This difference results in a higher number of rejected calls which can be seen in figure 9.

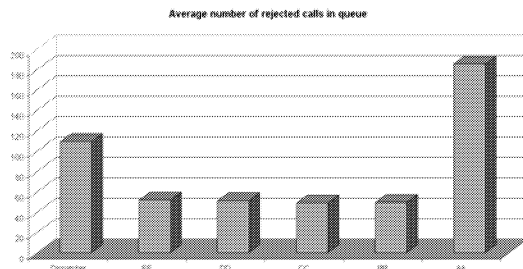


Figure 9: Average Number of Rejected Calls in Queue

In the second scenario we are going to add a system engineer with the same expertise as AA. This will reduce the work on AA and thus improve the performance. Now we have a bottleneck at the dispatcher. The queue there is very long compared to the other resources as can be seen in figure 10. This again results in a high number of rejected calls which we must reduce to reach a higher service level.

The solution to this problem is evident: add a second dispatcher. This is done in the next scenario.

In the third scenario, the average queue waiting time for the dispatchers is reduced enormously by adding a dispatcher. It was around 1 minute for the two first scenarios while it is

now only 6 seconds. This also reduces the total queue waiting time, which is almost halved.

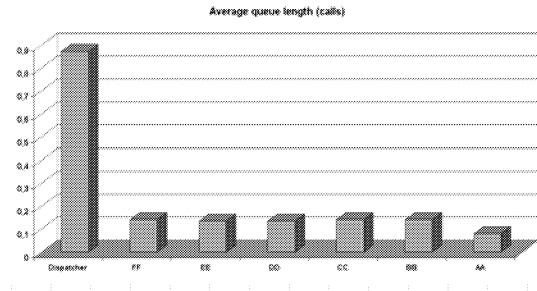


Figure 10: Average Queue Length

In the fourth scenario we add four skills in the skills matrix for the resources with only one skill. We do this by “cross training” of skills between BB and CC and between DD and EE. The skills matrix is shown in table 2. A graph obtained from the results is shown in figure 11. Here one can see that the average number of rejected calls for some resources is higher than the others. This comes from the fact that the dispatcher chooses the left system engineer in the skills matrix when there is more than one available. Our simplistic model does not handle the case of sending the call to another system engineer that becomes available when the call is in a queue of a busy system engineer.

Table 2: Skills Matrix of Scenario 4

	FF	EE	DD	CC	BB	AA
Skill 1	X					X
Skill 2		X	X			
Skill 3		X	X			
Skill 4				X	X	
Skill 5				X	X	
Skill 6	X					X

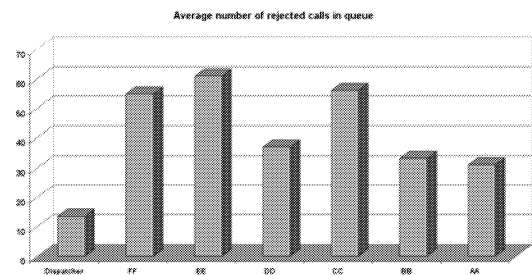


Figure 11: Average Number of Rejected Calls

In scenario 5 we add two system engineers with each 2 skills to further reduce the queues which results in less rejected calls and thus a higher service level. The first system engineer has skills 3 and 4, while the second has skills 2 and 5. Figure 12 shows the average queue waiting time of this scenario. One can see that there is bottleneck at resources AA and FF. Both have the same skills, so we will improve the system by training other resources the same skills.

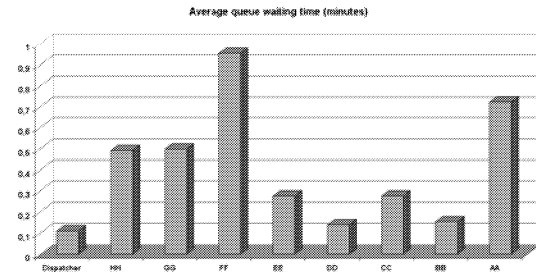


Figure 12: Average Queue Waiting Time of Scenario 5

In the last scenario we consider, we add two skills to reduce the queues of AA and FF. We add the skills to system engineers BB and DD because those have the lowest queue waiting time as can be seen in figure 12. The total skills matrix becomes:

Table 3: Skills Matrix of Scenario 6

	HH	GG	FF	EE	DD	CC	BB	AA
Skill 1			X		X			X
Skill 2		X		X	X			
Skill 3	X			X	X			
Skill 4	X					X	X	
Skill 5		X				X	X	
Skill 6			X				X	X

The model now has 2 dispatchers and 8 system engineers. The queue waiting times for this case are given in figure 13. Compared to figure 12 it is clear that the long queue waiting times are gone. This happened by only adding two skills.

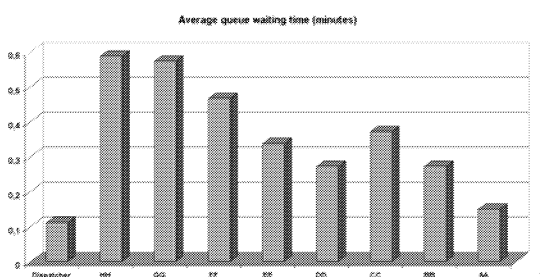


Figure 13: Average Queue Waiting Time of Scenario 6

In order to create new alternative configurations we are looking primarily at the queues and rejected calls. We try to reduce them by sharing skills over several system engineers. We also look at the work performed by every system engineer. A system engineer which is more busy than the others can create a bottleneck. Therefore it is necessary to distribute the work more equally among all the system engineers.

SIMULATION RESULTS

Figure 14 shows the performance measures productivity and service level for the 6 scenarios. The service level is increasing in each scenario and reached 90% in the last. The productivity is more difficult to explain but in general it reduces when we are adding resources because the same amount of work is divided among more people. Handling calls that were otherwise rejected without additional skills results in more work done thus a higher productivity can be seen in that case.

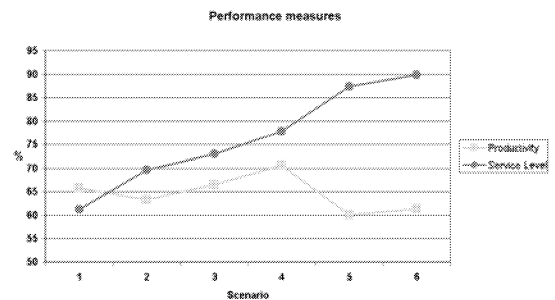


Figure 14: Performance Measures

Figure 15 shows the average queue waiting times. We consider the waiting times of the queues of the dispatcher and the system engineers (resources). The average queue waiting time (the top most curve) is the sum of the other two. It can be clearly seen that adding the dispatcher has a huge impact on the waiting time of the dispatcher queue which happened in scenario 3.

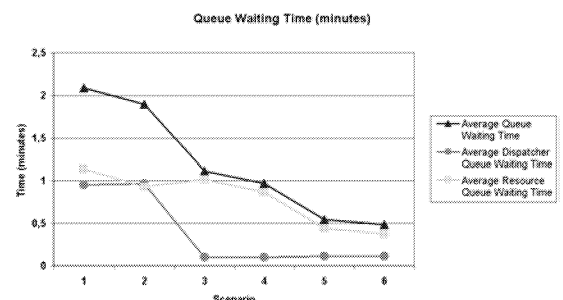


Figure 15: Queue Waiting Time (minutes)

The financial information of the model is given by the sales, costs and profit. This is shown in figure 16. The profit is the difference of the sales and the costs. Both sales and costs seem to increase over the scenarios. The addition of resources increases the cost and more processed calls increases the sales. For the scenarios considered, the maximum profit is reached in the second scenario.

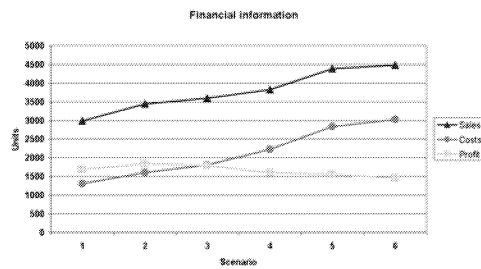


Figure 16: Financial Information

PROMETHEE METHOD

PROMETHEE is a multicriteria decision making method to rank different configurations on certain criteria. This lets us select the best system from a finite set of configurations. Due to the stochastic nature of the data we take an extension to the original PROMETHEE method, called the PROMETHEE-i method as described in (Pastijn et al., 2003). This extension allows the use of interval data instead of crisp data. Several possibilities exist here but we take the interquartile intervals of the data.

The parameters used for this example are shown in table 4. As criteria we take the productivity, the service level, the average queue waiting time and the cost. The type parameter describes the type of preference function used defined by the parameters q and p . The best candidate corresponds to the largest ψ . Figure 17 shows the evolution of the ψ operator. After 25 replications the ranking is scenario 4, 3, 6, 1, 5 and 2. This means that scenario 4 is the best in this case.

Table 4: PROMETHEE-i parameters

Criteria	weight	min/max	type	q	p
Productivity	7	max.	4	2	10
Service Level	5	max.	3	0	25
Queue W. Time	4	min.	6	0.5	0
Cost	6	min.	5	100	1500

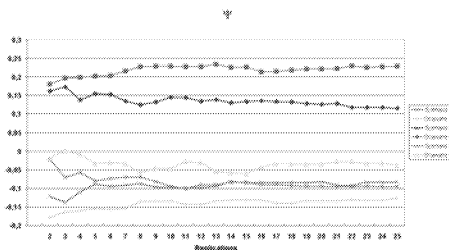


Figure 17: ψ Operator of the PROMETHEE-i Method

CONCLUSION

We began by a description of a call centre which was used as an example throughout the paper.

We described a set of tools used to generate, run and compare different configurations of the same model. These

tools are certainly very helpful for people with no or only little experience in simulation modelling. It can also be very useful in training or education situations for system managers.

We used the simulation model of the call centre to create a set of different scenarios. We tried to improve this model by removing the existing bottlenecks. This was done by only changing the number of dispatchers or system engineers and changes made to the skills matrix. We described the results obtained from the simulation runs and saw that in some cases serious improvements were possible with small changes to the model. The best configuration among the whole set of different configuration was selected using the PROMETHEE -i selection method.

Current research focuses on the implementation of a neighbourhood search procedure that tries to find the local optimum in the vicinity of the selected candidate.

REFERENCES

- Pastijn, H. Van Utterbeeck, F. and Van Loock, R. 2003. *PROMETHEE-i selecting the best simulation model configuration based on multiple performance measures*, ESS2003, Delft, The Netherlands, Oct 26-29 2003
- Van Loock, R. and Pastijn, H. 2003. *Performance measures in simulating incident management processes of a call centre*. Orbel 17, Brussels, Belgium, Jan 23-24 2003
- Van Loock, R. Van Utterbeeck, F. and Pastijn, H. 2003. *TAPE Performance measures implemented in an incident management model*, ISC2003, Valencia, Spain

SIMULATION MODEL OF DECENTRALIZED COORDINATION STRATEGIES IN A LINEAR SUPPLY CHAIN

Francesco Costantino

Giulio Di Gravio

Massimo Tronci

Department of Mechanical and Aeronautical Engineering

University of Rome "La Sapienza"

Via Eudossiana 18, 00184 - Rome, Italy

E-mail: giulio.digravio@uniroma1.it

KEYWORDS

Supply chain, Beer Game, bullwhip effect, kanban, responsibility tokens, simulation.

ABSTRACT

To increase competitiveness in the actual market, enterprises look for integration with customers and suppliers, developing more and more organized framework of physical and informative flows to improve general performances at all levels. One of the most famous tools that underline what aspects are mainly to care about is MIT Beer Game. This business game recognized in bullwhip effect a characteristic and critical problem in supply chain management.

This paper presents a simulation research that approached a linear productive chain integration system with kanban and responsibility tokens, in comparison with Beer Game standard results, to suggest a way to reduce the orders variability and amplification through agents, together with a correspondent reduction of stocks. Four strategies were studied, with two different level of integration and two different targets, either focused on total management cost or on stabilization period of the supply chain. Important results were brought out opening the way to future analysis.

1. INTRODUCTION

In the traditional conception of supply chain, every activity by the different agents is independently managed without considering inter-companies connections and pursuing exclusively the business targets of a single entity. The management is so focused on a single knot of the logistic chain trying to increase its efficiency.

Nevertheless, from the second half of the '70s we assisted to a deep change in the approach: the production system, organized in hierarchical and rigid schemes, showed lack of balances progressively more pronounced, caused by the economic environment and the competition terms.

If in the '60s and in the first half of the '70s the market was in strong expansion and the dominant strategic factor was the production capacity, nowadays the most important drivers became quality of the product and service offered and time to market (from design to delivery): with the enormous increase of the

competition and the globalization of the market, spatial and temporal boundaries, that in the past limited and regulate the competition among the economic agents, fell down.

The enterprises need to re-engineer their internal processes and external relationships, in a way to replace the isolate vision of the firm with an "extended" one, where the informative flows that relate the functions (production, purchases, marketing and sales) are shared with all the actors of the supply chain to connect them and synchronize the operations.

Only with collaboration and coordination of the knots an enterprise can reach high level of efficiency, effectiveness, flexibility and reactivity, all necessary elements to guarantee:

- increase of competitiveness;
- improvement of innovative capacity;
- improvement in following demand changes;
- increase of profitability;
- increase of customer satisfaction.

2. BULLWHIP EFFECT

The bullwhip effect can be defined with its three main characteristic effects:

- amplification of the orders variability by the different agents of the supply chain and consequent amplification of demand;
- tendency of these oscillations to increase as going back through the chain, from the customer to the producer;
- presence of a certain delay between the peaks of the demand and the order levels in the different agents.

The main reason of these fluctuation is a lack of transparency in the stages of the supply chain that creates an erroneous perception of the information progressively arriving at the agents.

The first analysis of the phenomena was arranged by Jay. W. Forrester in 1958 in a paper on the Harvard Business Review, even if a later study of experts in logistics from Procter & Gamble gave visibility to the problem and, for the first time, used the name of bullwhip effect. The effects of this propagation of oscillation in the order quantities were completely classified by many studies and can be identified in a general decrease of the business levels and in an

increase of the management costs. In particular, it's to face:

- higher level of stock to limit the unexpected variation in the demand and avoid out of stocks, with an increase in storage costs and capital expenditure;
- lower service level to the customer due to out of stocks that can cause damages from loss of incomes to image on market;
- unsatisfying quality of the product due to a necessary increase in production rhythm to cover the demand peaks, with relative costs for frequent programming and scheduling or replacement of scraps.

To quantify the effects of these consequences, basing on a study of Kurt Solomon Associates, we can underline that the average inefficiency costs derived from an incorrect supply chain management are between 9% and 20% of the total product value on a six month monitor for a general enterprise. About stocks, it is calculated that the bullwhip effect causes an average of 100 days of surplus storage for the entire chain. These data are sufficient to justify the great interest on the study and comprehension of the phenomena, without forgetting how the reduction of the lifecycle time and the new lean production model, gave more tools and focused the attention in fitting the market requirements. Many systematic studies (Jay. W. Forrester, J. Sterman. H.L. Lee, V. Padmanabhan and S. Whang) demonstrated how a production and distribution system, due to its policies, organizations and lead times, it's naturally oscillatory and any minimum disturbance in inputs can show this nature for misperceptions of feedback by the agents in correctly interpreting the laws of the causal structure and the connections among any action of the agents and the environment.

Synthesizing the results, it was evident how this is a natural consequence of the rational behaviour of the different level agents and, to control the bullwhip effect, it's necessary to intervene principally on the infrastructure of the supply chain and on the behavioural model of the single actors to eliminate the four main causes (demand forecast updating, order batching, price fluctuation and rationing and shortage gaming).

3. THE TOOL FOR ANALYSIS

On these theoretic basis, few years after Forrester's publications, at Boston's MIT Beer Game was born. The games simulates the management of a four level, totally independent, linear supply chain – fig. 1 – to produce and distribute beer (retailer, wholesaler, distributor, factory) where each member of the chain knows exclusively the orders from the upper level; only the retailer has the real demand information that discovers as the game proceed.

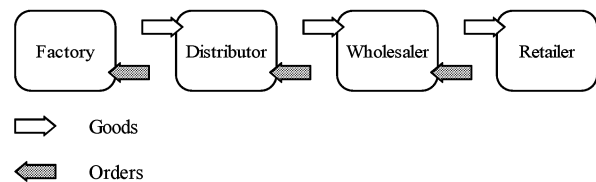


Fig. 1 – Beer Game supply chain

The characteristics of the model can be summarized in the following points:

- the flow of the order between two agents has a lead time of one period, while the material lead time is of two period;
- every agent has to satisfy the whole order of its customer. If this is not possible, the remaining quantities are put in queue (backlog orders) and filled as soon as possible;
- the management cost of each agent are calculated at the end of each period: 0,5\$ for each unit in stock or travelling and 1\$ for each unit in backlog;
- production capacity and stock capacity of each agent are unlimited.

The target of the agent is to minimize the total management cost, respecting the rule of not communicating among them excepting on leaving and arriving orders.

Notwithstanding the extreme simplicity of the described supply chain, compared to the real productive situations, the bullwhip effect is well highlighted by the variability of the stock and the orders at the end of each game.

This because any agent doesn't consider the consequences that his proper actions have on the other players. In particular, the hardest difficulties are in a correct evaluation of the lead times effect and, in general, in the non linearity of the system.

In 1989 Sterman explained in details the causes of the bullwhip effect in relation to the players behaviour. He verified that a simple unexpected increase in the demand generates a decrease of the retailer stock level that continues until the quantity of supplies is the one requested by the market. Having to face an increasing number of backlog, the orders become bigger than the demand level, trying to solve the problem as soon as possible. But, due to lead times, the increase of the order by the retailer causes an out of stock at wholesaler, not satisfying retailer's total demand. The normal reaction of the retailer is a further increase in the orders quantity even if the beer cases that are moving along the chain are more than sufficient. To this facts, the amplification along the supply chain has started and an explosion of storages or many out of stock are expected.

After a long campaign of experiments in the following years, with the scope of finding an average behaviour of the players and all their possible strategies – fig. 2 and 3 – MIT defined some characteristic value for the

parameters of the described supply chain (average costs and period of stabilization – the period when the order curve shape fits back the market demand) for a deterministic step 4 to 8 demand – tab. 1.

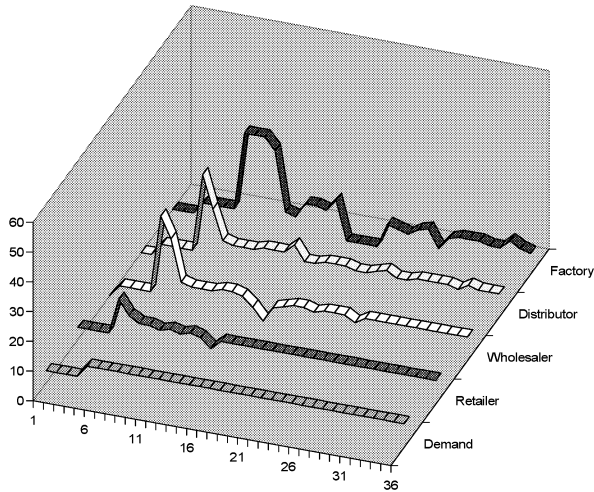


Fig. 2 – Average order state for step demand

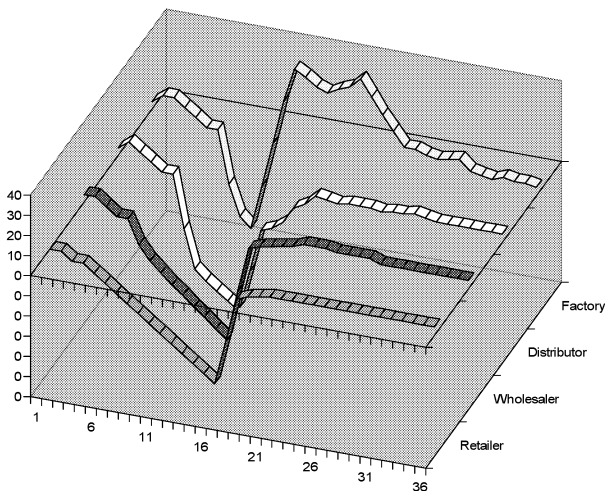


Fig. 3 – Average stock state for step demand

Tab. 1–Deterministic demand parameters

	Average costs	Average period of stabilization
<i>Retailer</i>	350	21
<i>Wholesaler</i>	609	18
<i>Distributor</i>	703	17
<i>Factory</i>	1028	16
<i>Supply chain</i>	2690	18

4. THE MODEL

The model takes as starting point the various studies in logistic management based on kanban JIT tool to optimize the functioning and interactions among different production cells and on “Responsibility Tokens” model by Evan L. Porteus that works at supply chain level.

The use of tokens has been known for long time in production programming field and the idea of synchronization and control of the informative and materials flows has been developed in different ways, according to the specific characteristics of the plants.

The innovation carried out by Porteus is: if a supplier doesn't satisfy the entire order of his customer, he sends responsibility tokens instead of product units not delivered. The customer deals with them as if they were a material unit of product and the financial consequences are for real, they are not material units are on suppliers' shoulders who is so pushed to reach the best possible service level.

The principle of our model, starting from both experiences, proposes to synchronize the behaviour of the companies in the supply chain dividing the information flows of the order placed in two distinguished parts, to reduce the bullwhip effect and its consequences. A generic order, in this way, travels on parallel channels for the following components:

- the first part (X) represents the need of the agents based on market demand. More precisely, when the final customer demand changes, X follows the variation. This part hasn't to be modified by any agent during his course in supply chain as it represents the real market requirement;
- the second part (Y_i – where i refers to the agent) is constituted by the tokens that represent the necessary quantities to manage the market demand changes and oscillations.

The role of X is to inform on the real tendency of the market demand, while the one of tokens is to manage the effects of the demand change that, due to lead times materials and information, causes backorder or oscillations of storages, with a consequent high variability of the order quantities. Obviously Y_i depends on the entity of the demand variation of each agent who, furthermore, has to add to the stock the amount of tokens coming from his customer that wasn't able to satisfy.

The use of this management tool permits to cover the variation of the demand with a single quantity of product that, as the tokens arrive at the factory, will leave from the producer to the retailer to refill the stock of all the agents of the chain, avoiding the unjustified fluctuation of different orders when the market has again become stable.

The fundamental element of the model is the determination of the correct number of tokens to send, related to the supply chain structure. First, it's necessary to evaluate the different lead times that characterize the information and materials flows,

choose the entity of safety stocks and define the service level or customer satisfaction degree to obtain. On these parameters it's possible to find the optimal quantity of tokens to control a generic ΔX , maintaining a sufficient low level of stock while contemporarily avoiding backorders. At this point, we can calculate the number of kanban to sum to the market demand and define the final order to place:

$$Kanban_i = LeftKanban_{i-1} + \alpha \Delta X$$

$$Order = Market\ demand + Kanban$$

where:

α = parameter dependent on the strategy adopted;
 ΔX = increase in market demand.

The described model – fig. 4 – can be defined as a decentralized coordination of supply chain: each productive entity keeps its complete decisional and operative autonomy. In regards to centralized coordination models (complete demand and stock information sharing), that generally reduce freedom and privacy of the entities, kanban model is less “invasive” however needing a certain degree of integration.

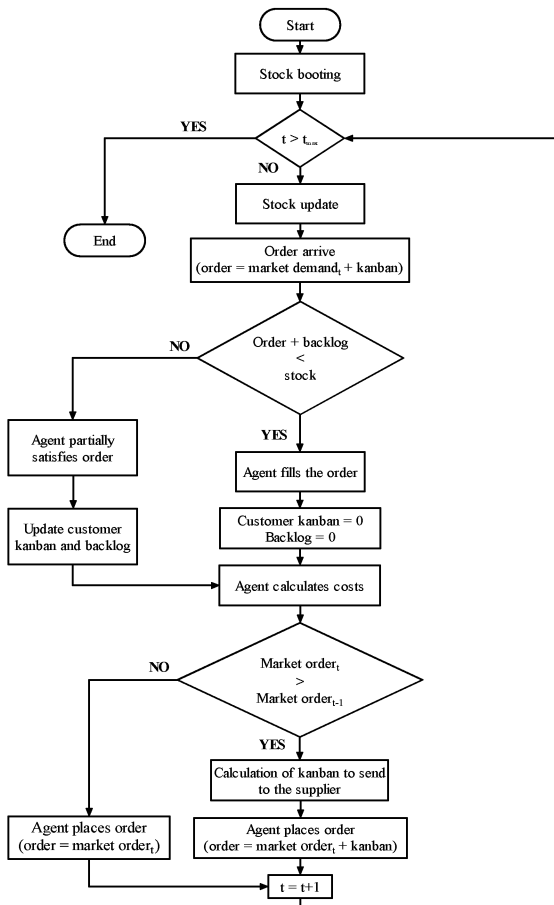


Fig. 4 – Kanban model

The technique has been validated finding an optimal solution for a deterministic step function demand that doubles (4 to 8) its value when the perturbation occurs. In this case, the supply chain is, during the first periods of simulation, in perfect balance and the step puts immediately the agents in backlog.

In this case, the model gives an optimal solution with kanban: every agent places an order equal to market demand added to his customer kanban that couldn't satisfy with his stock and three times the difference between the actual demand and the previous period demand ($\alpha = 3$). The choice of the constant value to multiply the demand increase it's established on the three periods of lead times that pass from the order placement to the delivery of the materials.

This strategy allows to clear any agent backorder in an average of about three or four periods before than normal Beer Game plays (fig. 6) and, furthermore, the chain comes back in balance: all the companies, even not having stock, have never to bear backorders (fig. 5). The general supply chain costs register a saving of about 20% in comparison to the average costs of MIT experimental campaign. It's interesting to notice that, using kanban model, a unique perturbation in demand curve reflects in a single perturbation in each agent order curve, avoiding unwanted oscillation after the stabilization of the market, characteristic of a non integrated supply chain (see fig. 2)

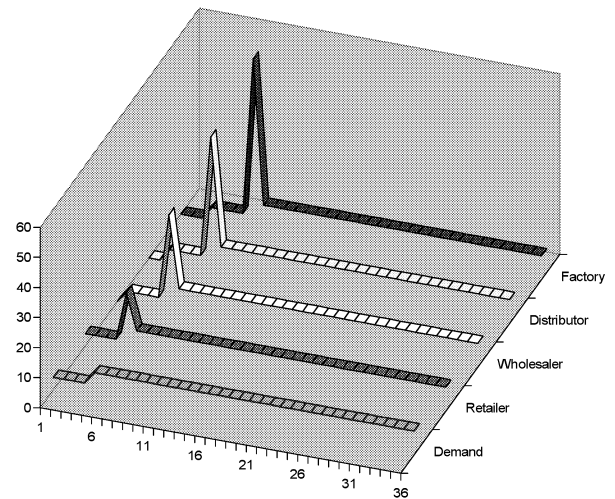


Fig. 5 – Optimal order state for step demand (kanban solution)

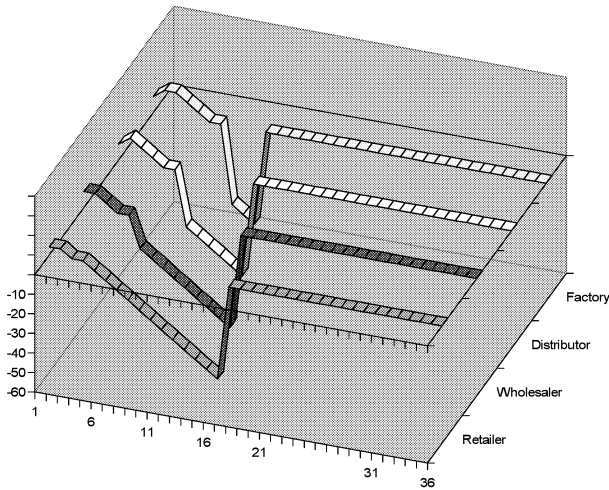


Fig. 6 – Optimal stock state for step demand (kanban solution)

5. THE EXPERIMENTATION

The model was then tested to show its advantages on an experimentation with stochastic demand (average 6, variance 2), considering two levels of integration:

- agents informed only on their own stock;
- agents informed on arriving quantities in the next period.

For both integration models two different strategies were developed, focusing either on average total management cost or average period of stabilization of the supply chain. Defining as follows the characteristic parameters of the different strategies:

k = number of kanban;
 Q_{mag} = quantity in stock;
 Q_{arr} = quantity arriving next period;
 ΔX = increase of market demand in two subsequent periods;

the possible choices are so framed:

Basic level of integration (I1)

kanban sending condition:

- $Q_{mag} \leq 2$ (retailer)
- $Q_{mag} \leq 1$ (other agents)

Advanced level of integration (I2):

kanban sending condition:

- $Q_{mag} + Q_{arr} \leq 6$ (any agent)
- $backlog \geq 12$ (any agent)

Period of stabilization strategy (S1):

$k_{ret} = 0,5 \Delta X$
 $k_{who} = k_{ret} + \Delta X$
 $k_{dis} = k_{who} + \Delta X$
 $k_{fac} = k_{dis} + \Delta X$

al management costs strategy(S2):

$= 0,5 \Delta X$
 $= k_{ret} + 0,5 \Delta X$
 $= k_{who} + \Delta X$
 $= k_{dis} + \Delta X$

ailer choice to send a kanban quantity equal to half increase of market demand ($\alpha = 0,5$) reflects the necessity not to propagate wide oscillations and to limit whip effect. Nevertheless, the quantity of tokens 't be too little because, in this case, there could be a ng risk of delay for stabilization period with sequent high management costs.

lso wholesaler (S2) has the same constant value for ket demand, there is a further reduction in order llation but, above all, a general lower level of stocks after the stabilization period.

Distributor and factory can not use the same value because the supply chain becomes too much instable and slow in absorbing market fluctuations

The inferior stock value to send kanban it's to allow a sufficient and fast recovering of backorder, while limiting the average one. The higher value for retailer (I1), the first ring of the chain, gives more stability to all the other agents: even in this case it's a compromise with the average stock level.

In the advanced level of integration strategy (I2), it was reasonable to impose a threshold value equal to demand average, adding a control on backlog units to avoid the rare case of stabilization with still backorder on.

In the following tab are summarized, for Beer Game and the four kanban strategies, the average total costs and period of stabilization t_s for each agent and the whole supply chain, together with an evaluation of bullwhip effect. This is calculated by an index that relates the variance of market demand to the variance of the agents orders:

$$B.E. = \sum_{j=1}^n \frac{Var(q_j)}{Var(D)}$$

where:

$Var(q_j)$ = agent j orders variance;
 $Var(D)$ = market demand variance;
 n = number of agents.

Comparing the results obtained with the four different strategies to the Beer Game campaign – tab. 2 – we can show a general improvement in average costs, stabilization period and bullwhip effect caused by the different level of integration that characterized our model.

Among kanban strategies, it's to notice that a limitation in tokens sent by wholesaler has always an improvement in total costs but a worsening, even if not stressed, in stabilization period: a reduction of average stock lowers costs but increases risk of backorder. So, the solution in case of highest level of integration is

pseudo-optimal, leaving to the company policy the decision of which target to reach.

About stabilization period it's also to explain that some strategies and some agents in particular strategies have high value of the parameter even with integration. That's because not always a quick arrangement means in any case lower costs, especially when a rapid increase in demand and related backorder are covered by peak of orders that creates huge stock levels, in particular in the agents far from the market.

Tab. 2 – Results of the experimentation

	B.G.	I1S1	I1S2	I2S1	I2S2
t_s Retailer	24	25	26	23	23
t_s Wholesaler	23	22	26	21	25
t_s Distributor	27	24	24	23	22
t_s Factory	24	26	25	26	25
t_s Supply Chain	25	24	25	23	24
Average Costs	3678	3392	3174	3226	3026
Average B.E.	19,75	10,57	9,64	9,66	9,04

Even the misalignment of the stabilization period (that can also increase for particular agents) has to be seen in an integrated sight. If the companies continue to concentrate on their own individual objectives, the kanban strategies are difficult and seem not worthy to apply at all levels of the supply chain, while, if the common scope of the involved enterprises becomes the final customer satisfaction, an harmonized distribution of dues and benefits allows a greater increase in cost savings and market position for all the agents – tab. 3.

Tab. 3 – Percentage savings

	Beer Game	I1S1	I1S2	I2S1
I1S1	7,77%	—	—	—
I1S2	13,70%	6,42%	—	—
I2S1	12,29%	4,89%	—	—
I2S2	17,73%	10,79%	4,66%	6,20%

For further explanation, in the following diagrams – fig. 7, 8 and 9 – are represented the optimal state of orders and stock with the higher level of integration (I2S2) and a comparison between the different strategies and the Beer Game stock level.

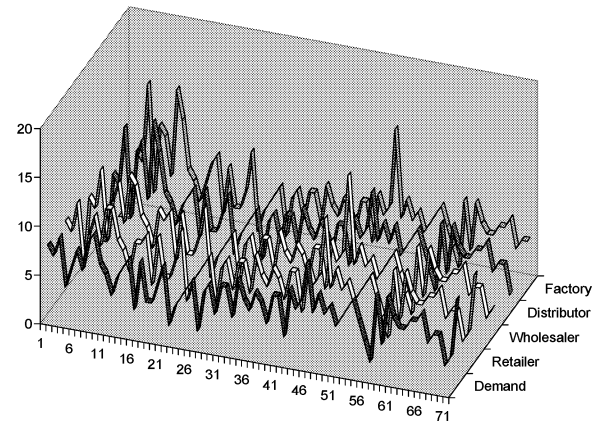


Fig. 7 – Optimal order state for stochastic demand (kanban solution, I2S2 strategy)

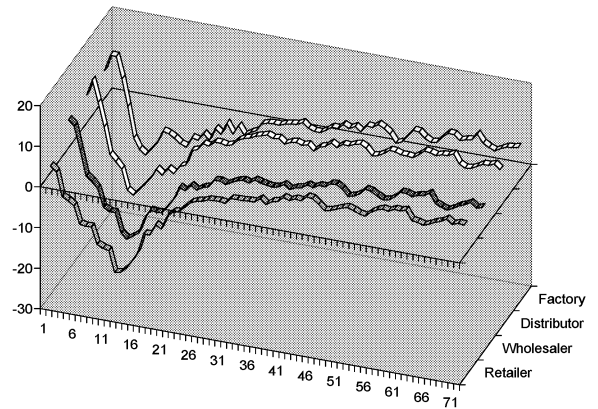


Fig. 8 – Optimal stock state for stochastic demand (kanban solution, I2S2 strategy)

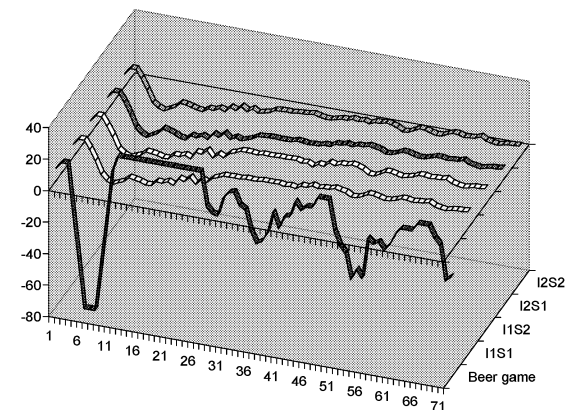


Fig. 9 – Stock level comparison

6. CONCLUSION

In this paper, the application of kanban and responsibility tokens in a linear supply chain model has

been investigated and discussed. The MIT Beer Game is used as benchmark to test the effectiveness of this management tool. A simulation study was realized to achieve case analysis, the results obtained gave the research a pseudo-optimal solution to reduce and contain general management costs and bullwhip effect. Using responsibility tokens involves the supply chain in an integrated working, with different strategy linked to the demand growth value. Comparing to average Beer Game plays, the complete kanban model strategy with the higher level of integration saves more than 20% in costs and cuts average period of stabilization eliminating backorder problems. Nevertheless, the benefits reached with a lower level of integration face up with the stabilization time of the single agents, so a supply chain less overloaded by high level of storage pays out with more time to fit the demand profile. Future work will aim at different improvement of the study: the actual model is linear as in Beer Game framework, but it would be interesting to test the kanban model in a non linear supply chain. Next step in the research will deal with sensitiveness study of different strategies approaches, to gain more information about the supply chain response. For example, the natural further investigation could test the effect of negative demand growth on kanban level to reduce the average stocks or study dynamic responses to the fluctuation introducing function like $\alpha = f(\Delta X)$ and implementing forecasting demand tools based on the same principle, to avoid the natural state of backorder for the first increase of demand.

Bibliografia

- [1] Chen F., Drezner Z., Ryan J., Simchi-Levi D., (2000), *Quantifying the Bullwhip Effect in a Simple Supply Chain: The Impact of Forecasting, Lead Times, and Information*, Management Science, v. 26, n. 3, p. 436-443.
- [2] Chen F., (1999), *Decentralized supply chains subject to information delays*, Management Science, v. 45, n. 8, p. 1076– 1090.
- [3] Chen F., Drezner Z., Ryan J. K., Simchi-Levi D., (1998), *The Bullwhip Effect: Managerial Insights on the Impact of Forecasting and Information on the Variability in a Supply Chain*, Chapter 14 in Quantitative Models for Supply Chain Management, International Series in Operations Research & Management Science, v. 17, Kluwer Academic Pub., Boston, Massachusetts.
- [4] Chen F., (1997), *Decentralized supply chains: subject to information delays* Management Science, v. 45, n.8, p. 1076-1090.
- [5] Durfee E. H., (2001), *Scaling up agent coordination strategies*, Computer, v.34, n.7.
- [6] Forrester J. W., (1960), *Industrial Dynamics*, MIT Press Cambridge, MA.
- [7] Hariharan R., Zipkin P., (1995), *Customer-order Information, Leadtimes, And Inventories*, Management Science, v. 41, n. 10, p. 1599-1608.
- [8] Kimbrough S.O., Wu D.J., Zhong F., (2000), *Computers Play the Beer Game: Can Artificial Agents Manage Supply Chain?*, Proceedings of the 34th Hawaii International Conference on System Science.
- [9] Lee H. L., So K. C., Tang C. S., (2000), *The Value of Information Sharing in a Two-Level Supply Chain*, Management Science, v. 46, n. 5, p. 626-643.
- [10] Lee H.L. and Whang S. (1999), *Decentralized multi-echelon supply chains: Incentives and information*, Management Science, 45(5):633-640.
- [11] Lee H. L., Padmanabhan V., Whang S., (1997), *Information Distortion in a Supply Chain: The Bullwhip Effect*, Management Science, v. 43, n. 4, p.546-558,
- [12] Lee H., Padmanabhan P., Whang S., (1997), *The Paralyzing Curse of the Bullwhip Effect in a Supply Chain*, Sloan Management Review, v. 38, p.93–102.
- [13] Moyaux T., Chaib-draa B., D'Amours S., (2003), *Multi-Agent Coordination Based on Tokens: Reduction of the Bullwhip Effect in a Forest Supply Chain*, Proceedings of the Second International Joint Conference on Autonomous Agents and Multiagent Systems.
- [14] Porteus E. L., (2000), *Responsibility tokens in supply chain management*, Manufacturing & Service Operations, v. 2, n. 2, p. 203-279,.
- [15] Raghunathan S., (2001), *Information Sharing in a Supply Chain: A Note on its Value when Demand is Non-stationary*, Management Science v. 47, n. 4, p. 605-610.
- [16] Sikora R., Shaw M., (1998), *A multi-agent framework for the coordination and integration of information systems*, Management Science, v. 40, n. 11, p. 65–78.
- [17] Serman J.D., (1989), *Modeling managerial behavior: Misperceptions of feedback in a dynamic decision making experiment*, Management Science, 35(3):321-339.
- [18] Taylor D., (1999), *Measurement and analysis of demand amplification across the supply chain*, The International Journal of Logistics Management, 10(2):55-70.

SYSTEM DYNAMICS MODELLING OF THE BUSINESS PRODUCTION OF THE SHIPBUILDING PROCESS

Ante Munitic
Frane Mitrovic
Josko Dvornik
University of Split
Maritime Faculty of Split
Zrinsko-Frankopanska 38, 21000 Split, Croatia
e-mail: amunitic@pfst.hr, josko@pfst.hr

KEYWORDS

System dynamics modelling, continuous computer simulation, business-production shipbuilding process

ABSTRACT

Simulation Modelling, together with System Dynamics and intensive usage of modern digital computers, which means application on massive scale, nowadays very cheap, and in the same time very powerful personal computer (PC-a), is one of the most suitable and effective scientific way for investigation the dynamics of non-linear and very complex: natural, technical and organization systems. The methodology of System Dynamics (prof. dr. J. Forrester-MIT), i.e. relatively new scientific discipline, showed its efficiency in practice as very suitable means for scientific investigation of the problems of management, of behaviour, of flexibility and sensibility of behaviour dynamics of large numbers of systems and processes. The aim of this paper is to present efficiency of application of System Dynamics simulation modelling in investigation of complex business production shipbuilding process – building the ship.

1. CHARACTERISTIC OF BUSINESS – PRODUCTION SHIPBUILDING PROCESS

System Dynamic Computer Simulation Model of The Shipbuilding Process is a continuous model, which comprise: 1. qualitative (mental, verbal and structural model), and 2. quantitative (mathematical and computer model, behaviour dynamics model) model presented using DYNAMO SYMOBOLICS (diagram of material and informational flows) and as well in DYNAMO higher programming language, which is very compatible for continuous use of business-production managing structure because of the:

1. Tracking and controlling of behaviour dynamics of PPBP (Abbreviation for term: Business-Production Shipbuilding Process (in Croatian: Poslovno proizvodnog Brodograđevnog Proces-a-PPBP)),
2. Forecasting the future behaviour of PPBP and

3. Optimization of parameters of PPBP.

The authors are using three main principles, which are: approximation and aggregation; System thinking philosophy, in order to present the Organization Business-Production Shipbuilding Process (PPBP) in following dynamic phases (flow of SP process) and in discrete control events (discrete event DD):

1. SP: SUPPLY OF SHIPBUILDING CAPACITY – SSC (Abbreviation for Term Supply Of Shipbuilding Capacity (in Croatian Ponuda brodograđevnog kapaciteta –PBK)) is now a part of global world marine market which requires: fresh market information on supply/demand fluctuation, obligate supply documentation, as well as information on market competition and information on client's financial standing (debts or credits).
2. D.D.: EFFICIENT COMPETITION TERMINATION, i.e. signing the agreement with consignee – SA (Abbreviation. Signing the Agreement), what initiate the sub process of PREPARING PROJECT DOCUMENTATION AND PREPARATION FOR OTHER PHASES. Also, material specification and production materials are needed, concluding contract with clients and cooperates, finishing deadlines and concluding deliveries, payments for debts and credits, and determination dynamics of sub-contractor employment.
3. Concluding the preparation of project documentation, which includes finishing the complete technological project documentation, begins the process of preparing sections and part of equipment out of slipway which requires the adequate documentation for reception warehouse that's completing is the basis for DD: setting (PUTTING DOWN THE KEEL -PK.)
4. Putting down the keel – PK is the basis for beginning of the shipbuilding on the slipway - PGBNN, beginning of the construction of the hull of the ship and continuously beginning of the other phases of ship constructor. In this phase it is very important to ensure the adequate documentation for worker's warehouses.
5. Finishing the shipbuilding on the slipway begins the process of DD: launching of the ship - PB, and then begins the process of SP: FINAL EQUIPPING OF THE SHIP-

ZOB which requires the adequate equipment, and preparation for the controlling and documentation for transfer.

6. After phase of the equipping of the ship - ZOB follow DD: TRANSFER OF THE SHIP – PPB, which required collaudation documentation in case of possible finishing due to complaint of ship owner.

7. After transfer of the ship – PPB start SP: COMERCIAL USE OF THE SHIP IN THE GUARANTEE PERIOD – KKBGR, which end in DD: when guarantee period run out – IGR.

8. At the end of guarantee period – IGR start SP: COMERCIAL USE OF THE SHIP – KKB. This phase comprise agreed time period (from 8 to 10 years), in which ship owner is obliged to fulfil all financial obligation (sum due) i.e. full cost price of ship toward shipyard. «Merchandise credit» shipyard => ship owner determine agreed time period.

9. DD: Speed (rate) in which shipyard collection debts - BNPB influence on decrease of SP: state of unpaid debts of shipyards – SNPOTB, whit final sum i.e. state equal zero.

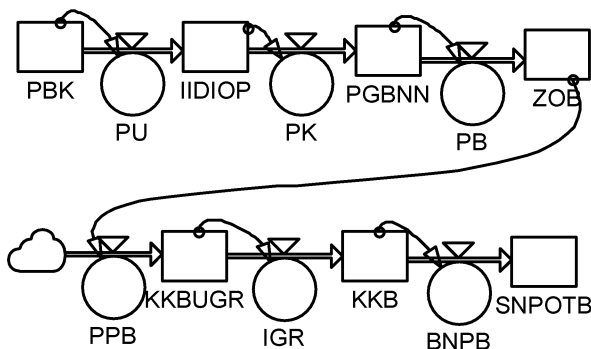


Fig. 1. The Rudimentarily diagram of shipbuilding process

2. BUSINESS – PRODUCTION PROCESS (PPBP) WITH SYSTEM DYNAMICS APPROACH

System Dynamics methodology of system modelling is very suitable for computer simulation of behaviour dynamics of the most complex organizational system. Business-production process (PPBP), without doubt, belongs to that group of system. Furthermore, we will observed PPBP as a whole, in accordance with System dynamics methodology, i.e. system is consist of nine relevant sub systems:

1. Planned process of shipbuilding as a whole (CPPIB)
2. Co-operation, i.e. external flexible labour capacities (KEKRRK)
3. Internal performer of labour duty, i.e. labour unit (IIRZ)
4. Supplying of materials, production materials, machinery and equipment which will be build in to the ship (NM)
5. State of outstanding debts and debt (SDP)
6. State of transfer account, i.e. incoming and outgoing money (SNZR)

7. Total income, income, profit, expenses, penalty, stimulations (UPDT)

8. Investment in basic and permanent working capital (IOTOS)

9. Short terms and long terms loans, i.e. sub system of financing shipbuilding

Structural informative intersectional model of construction of Skiff No.356 (Amorella) shown on Figure 2.

3. CONCLUSION

System dynamic simulating modelling is one of the most appropriate and successful scientific dynamics modelling methods of the complex, non-linear, natural, technical and organizational systems.

Implementation of System dynamic continuous computer simulation sub model of the business-production shipbuilding process – PPBP, which is a part of discrete digital computer simulated models of continuous nonlinear realities, in BrodoSplit, has allowed it's managing structure the continuous application of «active simulation process» in archiving following effects:

1. Quantitative grading of historical equals of PPBP,
2. Narrowing the future uncertainty,
3. Completing the list of the costs of shipbuilding process of the actual process,
4. Heuristic optimization of PPBP and enlarging the financial stability,
5. Improvement of organization of the process and complete automation of PSBP.

REFERENCES

- Forrester, Jay W. 1973/1971. *Principles of Systems*, MIT Press, Cambridge Massachusetts, USA, 1973.
- Richardson, George P. and Pugh III Aleksander L. 1981, *Introduction to System Dynamics Modelling with Dynamo* MIT Press, Cambridge, Massachusetts, USA,
- Munitic, A. 1989. *Computer Simulation with Help of System Dynamics*, BIS, Croatia, 1989.
- Ante Munitic, Slavko Simundic, System dynamics continuous cpmputer simulation model of ship construction documentation elaboration subprocess, *IMAM'95*, 23-27 April, 1995., Croatia,
- Slavko Simundic, Ante Munitic, The knowledge basis model in the expert system structure for the tehnological preparation phase in shipbuilding, *IMAM'95*, 23-27 April, 1995., Croatia
- Ante Munitic, Slavko Simundic, Josko Dvornik, System Dynamics Continuous Computer Simulation Model of Shipbuilding process, *ISC 2003*, 9-12 June, 2003, Spain, 336-339.
- Ante Munitic, Slavko Simundic, Josko Dvornik, Computing Simulation Model of Shipbuilding Organization Process, *SCSC 2003*, 20-24 July, 2003, Canada, 191-194.
- Ante Munitic, Slavko Simundic, Josko Dvornik, Shipbuilding Organization Simulation Modelling, *ISDC 2003*, 20-24 July, 2003, New York, USA,

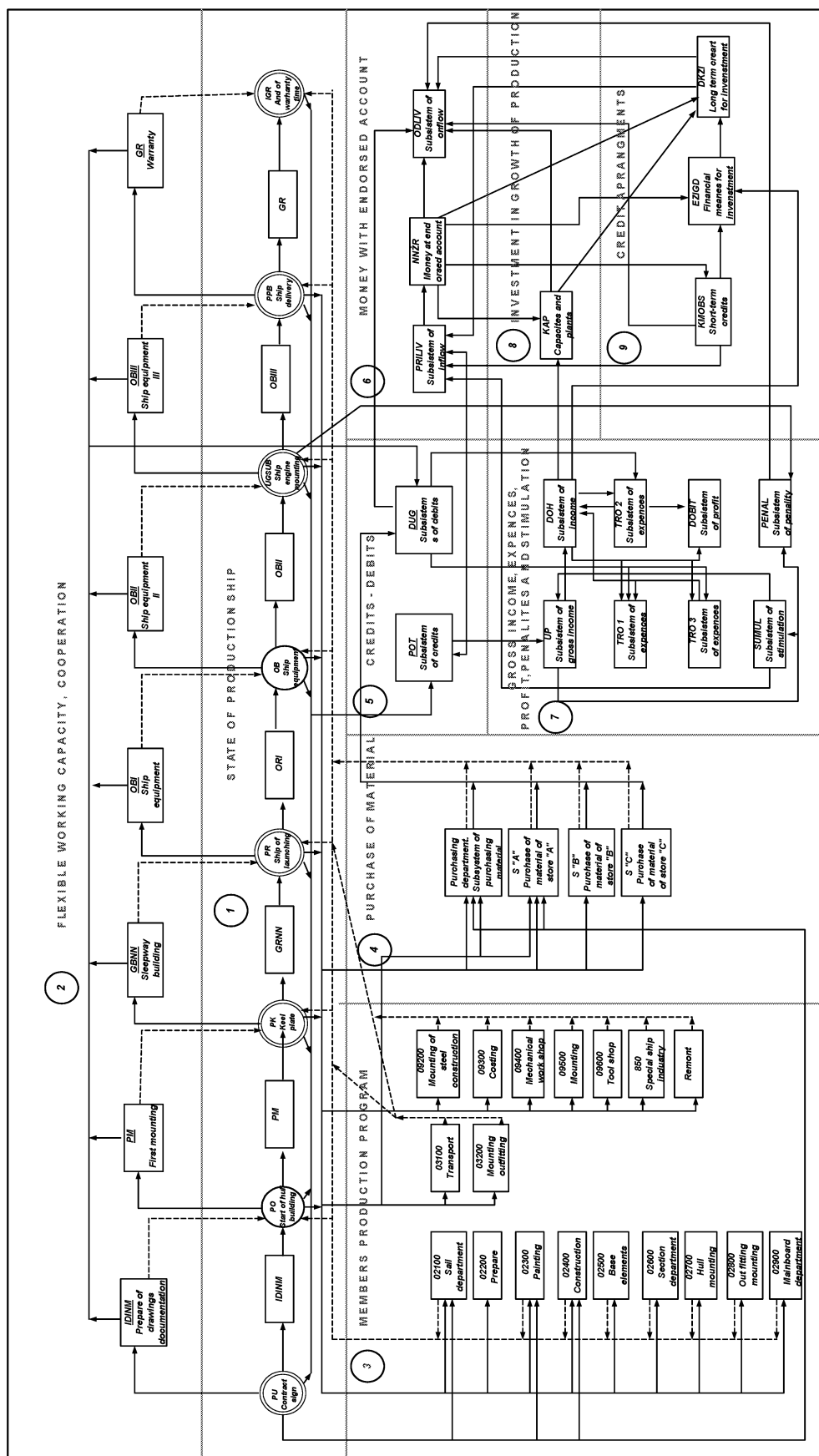


Figure 2. Structural informative intersectional model of construction of Skiff No.356 (Amorella)

AUTHOR LISTING

AUTHOR LISTING

Baron C.	28	Mitrovic F.	99
Boukachour H.	35	Monhot R.	83
Cakici K.	5	Munitic A.	99
Costantino F.	92	Oiso H.	76
De Moor B.	15	Passerini K.	5
Di Gravio G.	92	Pastijn H.	86
Diegel O.	20	Reschke C.-H.	9
Dion E.	68	Simon G.	35
Duflou J.	15	Thawonmas R.	47
Dvornik J.	99	Tronci M.	92
Esteve D.	28	Van Dromme D.	15
Gangadharan G.R.	83	Van Loock R.	86
Green L.	63	Van Utterbeeck F.	86
Ho J.-H.	47	Van Welden D.	55
Komoda N.	76	Vandermeulen B.	15
Leleur S.	71	Vertommen J.	15
		Yoshida Y.	40